

The Economics of Industrial Cyberespionage



Oleh Stupak

Kellogg College

University of Oxford

Word Count: 73,846

Thesis submitted in partial fulfilment
of the requirements for the degree of

DPhil in Cyber Security
in the Oxford Internet Institute at the University of Oxford

Trinity Term 2022

Acknowledgements

Completing this DPhil thesis was an exciting, unpredictable, and life-changing adventure which would not be possible without the help and support of numerous people.

First, I am deeply indebted to my two supervisors, Dr Greg Taylor and Dr Alexei Parakhonyak. I have benefited greatly from their wealth of knowledge, guidance, and continuous support. Greg's patience synergised with Alexei's strictness was crucial in proceeding through the program and completing this thesis.

Besides my supervisors, my sincere gratitude goes to my milestone assessors Dr Joss Wright, Dr Péter Eső, Dr Daniel Arce, and Dr Andrew Simpson, for their insightful comments and questions, which significantly strengthened my work.

I am fortunate to have simultaneously been a part of two departments, the CDT in Cyber Security and Oxford Internet Institute, which provided a healthy working environment and supported my research in every conceivable way. Its members will forever remain dear to me.

I am also very grateful to Dr David Steinsaltz, Dr Grant Blank, and Dr Stephanie Solywoda for helping me to take the first steps in teaching at a graduate level. Teaching was a significant part of my DPhil journey. Thank you for making it enjoyable.

My special thanks go to Janet Sadler and David Hobbs, who managed to lift most of the administrative load from our shoulders and let us focus on the research.

My friends Andrew, Oleksii, Marcel, Vlad, Lamuel, Vitaly, Nikita, and Roman helped me to maintain a somewhat healthy work-life balance. You all are truly wonderful.

I shared the final stages of my DPhil with my partner Chenzi, whose care and support helped me to maintain a clear mind and healthy diet.

Finally, I am genuinely grateful to my parents, Alina and Oleksandr Stupak, who made my journey possible and supported me at every step I took.

Abstract

Industrial cyberespionage is a widespread phenomenon that costs the global economy up to \$6 trillion annually. It influences how we invest in innovation, store data, and create laws, and ultimately the welfare of society. Yet, it has received limited attention from the economics community. This thesis contributes to the emerging interdisciplinary field of economics of information security, bringing an economic perspective to the matter and considering industrial cyberespionage from three different angles.

The first chapter examines a dynamic R&D race in which competitors can conduct cyberespionage against each other. We develop a framework that analyses the influence of cyberespionage on innovative incentives, companies' payoffs and the quality of the end product. We demonstrate that industrial espionage has an ambiguous influence on the overall investments exerted in the race and companies' expected payoffs and might even be beneficial for the quality of innovative end-products under certain circumstances.

The second chapter provides new empirical evidence that research-intensive industries are particularly susceptible to information leakage attacks. Based on Eurostat aggregated data on enterprises' innovative and digital activity, we construct a tailored data set that allows us to achieve robust statistical inference and study the relationship between information leakage attack rate, research-intensity, and companies' data reliance. The study uses multivariate fractional regression analysis to distinguish two industry-specific associations: high-tech manufacturing industries are prone to experience targeted attacks, while knowledge-intensive service companies are more likely to fall victim to opportunistic attacks.

The third chapter aims to understand efficient network formation and optimal defensive resource distribution in the presence of an intelligent attacker. We present a two-player dynamic framework in which the Defender and the Attacker compete in a network formation and defence game with heterogeneous vertices' values. Such a model allows for studying the trade-off between network efficiency and security. Contrary to the literature, we find that a centrally protected star network does not yield the maximum payoff for the defending side in most circumstances, even being the most secure network formation. Additionally, it reveals a new type of network that often arises in an equilibrium of the games with limited defensive resources—a maxi-core network.

Contents

List of Figures	xi
1 General Introduction	1
1.1 Motivation	1
1.2 Thesis outline and contributions	6
1.3 Background and literature	9
1.3.1 Industrial espionage in competitive settings	10
1.3.2 Data breach risks	14
1.3.3 Organisational cybersecurity	18
1.4 Future research	25
1.4.1 Chapter 2: patent design, multiple innovations, and endogenous attack function	25
1.4.2 Chapter 3: causality and ecological fallacy	26
1.4.3 Chapter 4: multiple attack resources, probabilistic defence function, and reinforcement learning	27
1.5 Dissertation structure	28
2 Industrial Cyberespionage in R&D Races	29
2.1 Introduction	29
2.1.1 Related literature	32
2.2 The model of economic cyberespionage	37
2.2.1 Development stage	41
2.2.2 Research stage	47
2.3 Social welfare	55
2.3.1 Espionage and welfare	56
2.4 Conclusion and discussion	57
3 No Research, No Risk?	61
3.1 Introduction	61
3.1.1 Related literature	65
3.2 Hypotheses	74
3.3 Empirical strategy	75

3.3.1	Data	75
3.3.2	Methodology	83
3.3.3	Methodology discussion	90
3.4	Results	94
3.4.1	Fractional regressions results	94
3.4.2	Marginal effects analysis	96
3.4.3	Targeted and opportunistic attacks	97
3.4.4	Alternative explanations and robustness checks	103
3.5	Conclusion	106
4	Secure and Efficient Networks	109
4.1	Introduction	109
4.1.1	Related literature	113
4.2	The model of network design and defence	116
4.2.1	Model setup	116
4.2.2	Analysis: encounter stage	118
4.3	Equilibrium	125
4.3.1	Limiting properties of equilibrium	130
4.3.2	Relaxed problem solution, $\delta = 1$	134
4.3.3	Feasibility of solutions	140
4.3.4	Numerical verification	145
4.4	Weighted version of the model	148
4.5	Limitations and future research	150
4.5.1	Limitations	151
4.5.2	Future research	152
4.6	Conclusion	154
5	General Conclusion	155
Appendices		
A	Appendices for Chapter 2	161
A.1	Mathematical Appendix	162
A.1.1	Equilibrium derivation and analysis	162
A.1.2	Social welfare function derivation and analysis	170

B	Appendices for Chapter 3	175
B.1	List of industries' aggregates	177
B.2	Independent variables histograms	178
B.3	Multiple Imputation conditions	179
B.4	Assumptions for GLM class models	181
B.5	Eurostat indicators on high-tech industry and knowledge-intensive services	181
B.6	Collinearity in multivariate specifications	182
B.7	RESET	184
B.8	Alternative specifications test	185
B.9	Kernel density plots of residuals of imputation models	186
B.10	Complete case regression analysis results for r^{per}	187
B.11	Multiple imputation regression analysis results for r^{per}	190
B.12	Complete case regression analysis results for r^{res}	193
B.13	Multiple imputation regression analysis results results for r^{res}	195
B.14	Complete case regression analysis results for r^{sup}	198
B.15	Multiple imputation regression analysis results results for r^{sup}	201
B.16	Complete case regression analysis results for models with research- intensity proxies and CRM variable	204
B.17	Multiple imputation regression analysis results results for models with research-intensity proxies and CRM variable	206
B.18	Multiple imputation regression analysis results results for models with research-intensity proxies and CRM variable for manufacturing industries	208
B.19	Multiple imputation regression analysis results results for models with research-intensity proxies and CRM variable for service industries	210
B.20	Robustness checks	212
C	Appendices for Chapter 4	215
C.1	Mathematical appendix	215
C.2	Supplementary material	234
C.3	Extreme points characteristics and the convex hull	235
C.3.1	Full support	235
C.3.2	Convex hull construction for the full support case	245
C.3.3	Partial support	249
C.3.4	Convex hull construction for the partial support case	261
C.3.5	Extreme points of convex regions	263
C.4	Feasibility of intermediate family sequences	268
C.5	Supplementary material for Subsection 4.3.4	273
	References	275

List of Figures

2.1	Timings of the R&D race. Circles represent the timings for firms' investment decisions	38
2.2	The uniqueness of Nash equilibrium depending on the slopes of PDF and cost functions.	43
2.3	Influence of attack power on the overall investment.	51
2.4	Influence of technological edge on firm's payoff.	53
3.1	Dependent variable histogram.	77
3.2	Geographical distribution of the countries included in the study. The countries are colour-coded according to the average share of enterprises that experienced information leakage attacks.	77
3.3	Industrial distribution of the research-intensity proxy variables. . .	80
3.4	Industrial distribution of the CRM usage rate.	82
3.5	Point estimates of the marginal effects of the research proxy variables. .	97
3.6	Marginal effects of the share of R&D support personnel and CRM usage rate point estimates.	100
4.1	Examples of secure and efficient networks on nine nodes. Each node is valued according to the number of connections it has.	110
4.2	The Defender's expected value from a game on five vertices ($n = 5$) with two defensive resources ($\delta = 2$) as a function of a single link cost, c , built for every feasible connected network structure. Highlighted lines are spawned by the networks, which appear as equilibrium solution for some given c . Boundaries on the figure are: $c_l = 1.7$ is the lower cost boundary, $c_u = 1.78$ is the upper boundary, and $c_i = 1.71$ is an intermediate boundary.	131
4.3	All optimal networks for a game with five vertices and two defensive resources. The colors of networks correspond to the colors of lines demonstrated in Figure 4.2.	132
4.4	The set of feasible values, Υ^4 , for nodes 1, 2, and 3 along the edge v_4 , (a) and (b); and the convex hull of set of feasible values, $\text{Conv}(\Upsilon^4)$, for nodes 1, 2, and 3 along the edge v_4 , (c) and (d). More details about this example can be found in Appendix C.3.2.	136

4.5	M-maxi-core network on nine vertices with two vertices in the core.	143
4.6	Results of numerical optimisation. See Appendix C.5 for precise values of cost boundaries and more details on optimal maxi-core networks.	147
B.1	Share of research personnel histogram	178
B.2	Share of researchers histogram	178
B.3	Share of R&D support personnel histogram	179
B.4	CRM usage rate histogram	179
B.5	Residuals distributions for univariate imputation models	186
B.6	Residuals distributions for multivariate imputation models	186

1

General Introduction

Cybercrime constitutes the greatest
transfer of wealth in history.

Gen. Keith Alexander

Contents

1.1	Motivation	1
1.2	Thesis outline and contributions	6
1.3	Background and literature	9
1.3.1	Industrial espionage in competitive settings	10
1.3.2	Data breach risks	14
1.3.3	Organisational cybersecurity	18
1.4	Future research	25
1.4.1	Chapter 2: patent design, multiple innovations, and endogenous attack function	25
1.4.2	Chapter 3: causality and ecological fallacy	26
1.4.3	Chapter 4: multiple attack resources, probabilistic de- fence function, and reinforcement learning	27
1.5	Dissertation structure	28

1.1 Motivation

Humanity has witnessed a giant technological leap since the last century's second half. The shift from mechanical and analogue systems to digital computing marked

the beginning of the digital industrial revolution (Paul Markillie, 2012).

The revolution was fuelled by the development of the Internet—another powerful technology which has reshaped established conventions, led to scientific advances, and accelerated progress. Humanity entered the information era with an entirely new set of rules and became heavily dependent on interconnected digital systems.

However, the Internet has not been designed with security in mind (Sweeney, 2015). Unlawful practices have quickly adapted to the information environment alongside the activities of law-abiding citizens. Crime has gone cyber.

Cybercrime issues are addressed by a newly emerged cybersecurity field, which absorbs knowledge from a wide range of natural and social disciplines: mathematics and statistics, computer science and psychology, sociology and economics. Still, most cybersecurity research is conducted on a technical side, with a disproportionally small share of studies written from a social science perspective. Only in the mid-2000s was it recognised that cybersecurity failures are caused at least as often by social factors as bad technological design (Anderson & Moore, 2006). Ultimately, cyberspace is just a mediator for social interactions. The authors argued that the “tools of game theory and microeconomic theory are becoming just as important to the security engineer as the mathematics of cryptography”. Their seminal paper paved the way for an important branch of cybersecurity research—the economics of cybersecurity.

There is currently no consensus on how to define the economics of cybersecurity. Kianpour et al. (2021) offered a comprehensive definition based on 628 scholarly articles related to the subfield, which we adopt for the thesis.

Definition 1.1 (Economics of Cybersecurity). *Economics of cybersecurity is a research field that offers a socio-technical perspective on the economic aspects of cybersecurity, including cybersecurity investment and budgeting, information asymmetry, cybersecurity governance and motivation, tools, and consequences of cybercrime. Its main goal is to provide sustainable cybersecurity policy recommendations, regulatory options, and practical solutions that improve the cybersecurity posture of the interacting agents in socio-technical systems.*

The research presented here is devoted to issues of economics of cybersecurity. It examines unfair practices in competition in the information environment and, specifically, industrial cyberespionage and information leakage attacks¹. With the estimated global costs of cybercrime to the global economy reaching up to \$6 trillion (a significant share of which is allegedly associated with data breaches), phenomena have become increasingly influential (Morgan, 2020). The average data breach cost to a single company has reached a historical height of \$4.24 million (IBM Security, 2021), making information leakage attacks one of the primary concerns of CEOs and CISOs (Dhillon et al., 2021; PwC, 2020). The overall annual cost of cybercrime to the UK is estimated to be £27 billion, of which £9.2 billion is affiliated with the theft of intellectual property (HOSAC, 2018). It was also reported by Oxford Economics (2014) that at least every fifth company in the UK suffered from a data breach in 2013. Additionally, 61% and 59% of breached companies claimed the loss of competitive advantage due to the loss of intellectual property (IP) or sensitive information, respectively. However, any estimate in this domain should be taken with scepticism. Due to the secrecy of the phenomenon, there is no reliable evidence to claim any number with certainty (Anderson et al., 2013).

Industrial cyberespionage is a multidimensional problem which embeds both technical and social aspects. The technical side of the phenomenon is extensively addressed by computer science. Computer science literature about industrial cyberespionage generally focuses on preventive technological measures—cybersecurity hardware and software, their interaction and management. It studies operational security, including user authentication techniques, identity management, antivirus software, firewalls, and encryption, as well as technological measures enabling organisations to detect intrusions, monitor traffic, discover vulnerabilities, and manage patching². This strand of literature is generally focused on the applications of computer science and information systems security knowledge to tackle specific

¹An information leakage attack (or data breach attack) is a cyberattack performed by an unauthorised third party in which confidential, sensitive data (e.g. trade secrets or clients' databases) are stolen, copied, viewed, or transmitted.

²See Eckhart et al. (2019), Ahmad et al. (2019), Khraisat et al. (2019), and Yuan and Wu (2021) for comprehensive surveys about technical approaches to organisations' cybersecurity.

cyberattacks and system failures that might lead to sensitive data disclosure. It answers the question “how”, often adopting an industry-specific lens.

Moreover, cyber-enabled industrial espionage is widely acknowledged as a socio-technical problem (Khan et al., 2021). Sadok et al. (2020) argued that technical disciplines tend to disregard a human as the cornerstone of any interconnected system. This gap is addressed by various social disciplines, including management and marketing, criminology, psychology, law, finance, and economics. For instance, management science is concerned with identifying antecedents of data breaches, developing managerial procedures and routines to control organisational cyber risks, and operational crisis management strategies that minimise the damage in the event of a cyberincident³. Crisis management is often supported by marketing science, which is devoted to reputational damage control and customer retention⁴. Psychology and criminology address the phenomenon on the individual level, assessing the role of human error and motivation behind illegal behaviour⁵. Finance is predominantly focused on the consequences of data breaches, studying their influence on the market and stock price⁶. Scholarly articles on law discuss the legal approaches to mitigate cyber risks for organisations and individuals (e.g. compulsory breach disclosure, cybersecurity audits, and enforcement)⁷.

Both technical and social research strands discussed above are abstracted from another influential factor without which the picture would be incomplete—a strategic

³For instance, see T. Wang and Hsu (2010), who developed guidelines for making board structure and composition cyber-resilient, or Say and Vasudeva (2020) who analysed organisational learning from digital failures.

⁴One could cite Choi et al. (2016), who analysed an effective recovery of customer relationships after privacy breaches, and Syed (2019) who investigated reputation threats on social media and their mitigation.

⁵For instance, Kim and Kwon (2019), Liginlal et al. (2009), and McLeod and Dolezel (2018) analysed human errors and their prevention in the context of the “Swiss cheese” model and advocated a multilayer cybersecurity approach; or Sen and Borle (2015) who attempted to explain the motivation behind information leakage attacks with exaggerated emphasis on economic success in western society.

⁶For instance, see Hovav et al. (2017) and Malliouris and Simpson (2020), who studied the influence of data breaches on the market’s reaction and systematic cyber risk in the form of cost of equity.

⁷See Png et al. (2008), who studied governmental enforcement and its deterrence effects on information security attacks, or Kolevski et al. (2021) who studies pros and cons of the obligatory breach disclosure.

element inherent in every conflict. Industrial cyberespionage involves at least two competitive sides: the offensive side, which aims to compromise the Information and Communications Technology (ICT) systems of the opponent and retrieve valuable information, and the defensive side, which tries to prevent this from happening or adjusting its strategy to deal with the circumstances. Competitive situations in which outcomes of two or more sides with conflicting interests depend not only on their own actions but also the actions of the opponent are generally addressed by economics and, particularly, microeconomics and game theory (Osborne, 2003). Game theory research revolves around the concept of (Nash) equilibrium—a situation in which two or more competitive sides (called “players”) choose actions to achieve the best possible outcome for themselves, accounting for the actions of others. The approach allows us to predict the optimal choices of multiple decision-makers, which would be impossible if we consider a single player in isolation. Thus, the economics perspective applied to cybersecurity problems can support better decision-making in the management and governance or even be used to inform a better security technology design. However, the game-theoretical perspective is rarely adopted in cybersecurity studies on industrial cyberespionage and information leakage attacks, with only a few notable exceptions (e.g. Higgs et al., 2016; Laszka et al., 2014).

This dissertation addresses the lack of decision support for organisations and policymakers facing industrial cyberespionage. Our overall research question is *what are the implications of industrial cyberespionage for organisational competitive and security incentives?* To tackle the problem, we approach it from three different perspectives. In Chapter 2, we develop a game-theoretical framework to analyse competitive incentives of R&D companies in the presence of industrial cyberespionage. In Chapter 3, we empirically investigate the drivers behind information leakage attacks at the industry level and provide supporting evidence for the results of our game-theoretical analysis about industrial cyberespionage incentives. Chapter 4 introduces a new theoretical framework which considers trade-offs between security and efficiency in organisational network optimisation. The specific research questions in each chapter can be presented as follows:

1. How does industrial cyberespionage influence innovative incentives and social welfare in research and development races?
2. Are research-intensive industries particularly susceptible to industrial cyberespionage?
3. How do we design and protect a network in the presence of an intelligent external attacker?

While these three questions all address the problem from different perspectives, they share a focus on organisation-level risks associated with information security breaches and strategies in response to them.

1.2 Thesis outline and contributions

This dissertation endeavours to extend—within reasonable limits—our understanding of industrial cyberespionage. By considering the economics of industrial cyberespionage, this work contributes to the emerging interdisciplinary field of *economics of information security*. We adopt mixed methods. Chapter 2 situates the research in the context of R&D race by developing a game-theoretical model in which innovative companies compete for better products and can perform industrial cyberespionage against their rivals. The model studies the influence of espionage on competitive incentives of companies, the quality of end products, as well as the motivation behind information leakage attacks. It demonstrates that the difference in relative positions of the opponents might be the main driver behind cyberespionage, and that the phenomenon might influence innovative incentives of companies and the quality of end products either negatively or positively—depending on cyberespionage intensity. These results are then used to inform the empirical research presented in Chapter 3, which studies drivers of information leakage attacks and provides supporting evidence for the theoretical results about cyberespionage incentives. The empirical study demonstrates that research-intensive industries are susceptible to targeted information leakage attacks as, in those industries, companies engage

in R&D races more. Given that in the vast majority of circumstances, industrial espionage is harmful to an innovative company, in Chapter 4, we offer a new theoretical framework, which allows us to optimise the network structure and defensive resource allocation in the presence of an intelligent attacker with respect to both efficiency and security. More details on the contributions and methods of each chapter can be found below.

Chapter 2 is based on Stupak (2022a). We use a modified version of Lazear and Rosen’s (1981) rank-order tournament model to investigate investment incentives in R&D races when firms can conduct industrial espionage against their rivals. The study demonstrates that industrial espionage has an ambiguous influence on the overall investments exerted in a race and companies’ expected payoffs. This is because espionage discourages early investment, whose fruits might be stolen, but also weakens the discouragement effect that arises if one firm gets too far ahead. Moreover, industrial cyberespionage might benefit the quality of innovation brought to the market by the R&D race winner. We demonstrate that aggressive enforcement of cybersecurity laws can have knock-on effects on competition and innovation that may be negative rather than positive. These concerns may be included in the cost-benefit analysis of public cybersecurity, patenting policies, and their interaction.

Chapter 3 is based on Stupak (2022b). Here we empirically test the plausibility of some of the results obtained in Chapter 2. Specifically, we test whether cyberespionage incentives arise from the difference in relative positions of innovation race participants. If it is the case, research-intensive industries must be particularly susceptible to industrial cyberespionage and information leakage attacks. This statement is the central hypothesis of the empirical research presented in the chapter. Based on Eurostat aggregated data on enterprises’ innovative and digital activity, we construct a tailored data set that allows us to achieve robust statistical inference, and study the relationship between the information leakage attack rate, research-intensity, and companies’ data reliance. The study uses multivariate fractional regression analysis to distinguish two industry-specific associations: high-tech manufacturing industries are prone to experiencing sophisticated targeted

attacks, while knowledge-intensive service companies are more likely to fall victim to opportunistic financially-driven attacks. These findings can be utilised to develop better cybersecurity policies or insurance pricing models.

Chapter 4 is based on Stupak (2022c). The last question we raise in the thesis is how an enterprise can protect itself against information leakage attacks and industrial cyberespionage. The scholarship of economics of networks generally focuses on a fundamental aspect: either network efficiency (e.g. Jackson, 2010) or network security (e.g. Dziubiński & Goyal, 2013; Goyal & Vigier, 2010). We contribute to the literature by considering both of these aspects of network formation simultaneously. Particularly, this chapter aims to understand efficient network formation and optimal defensive resource distribution in the presence of an intelligent attacker. We present a framework in which two asymmetric players, the Defender and the Attacker, compete over the two stages in a network formation and defence game with heterogeneous vertices' values. In a baseline model, nodes are valued according to their degree centrality, allowing the Defender to manipulate the network value and expected losses from an attack by creating costly links. Additionally, players are not perfectly informed about the allocation of defensive and offensive resources. Such a model allows for studying the trade-off between network efficiency and security. The chapter produces three main results. First, contrary to the literature, it demonstrates that a centrally protected star network does not yield the maximum payoff for the defending side in most circumstances—even though it is the most secure network formation, it is not efficient, as any two peripheral nodes need to overcome at least two links to communicate. Secondly, it reveals a new type of network that often arises in an equilibrium of the game with limited defensive resources—a maxi-core network with a completely connected core and sparse periphery. Thirdly, the model demonstrates that the density of optimal network formations decreases in the cost of connection. The theoretical framework can be used to study the optimal organisation of employees and ICT devices in an enterprise or aid the development of optimal cybersecurity awareness training strategies.

1.3 Background and literature

Industrial espionage is a complex term covering a range of illegal activities performed to attain competitive advantage or economic gain, resulting in considerable financial losses. It influences research efforts and, as a result, the quality of innovations and social welfare (Hou & Wang, 2020). The phenomenon of industrial espionage can be traced back to the mid 18th century. On the brink of the industrial revolution, the French Bureau of Commerce was systematically hiring British and French spies to retrieve trade secrets, know-how and military technology. Already back then, it was recognised that espionage could be an efficient and relatively cheap method of gaining competitive advantage (Harris, 1985, 2017).

With the development of information communication technologies, the exponential growth of processing powers available to the industry, and intensified competitiveness in the newly emerged information environments, intelligence gathering has become an integral part of modern competition (Crane, 2005). However, espionage has undergone a substantial change in nature. According to the classification proposed by Anderson et al. (2013), industrial espionage can be attributed to the “transitional cybercrime” category as the modus operandi of espionage has substantially changed with the introduction of the Internet. Modern espionage is predominantly conducted by means of the Internet and cyberattacks. Thus, the thesis focuses on the evolutionary successor of industrial espionage—industrial cyberespionage. At this very moment, there is no unified definition of industrial cyberespionage (the same is true about the definition of conventional espionage). Therefore, we offer the following definition based on J. A. Lewis and Baker (2013) and K. Baker (2022).

Definition 1.2 (Industrial Cyberespionage). *Industrial cyberespionage is the illegal practice of obtaining competitors’ trade secrets or other intellectual property stored in digital formats without permission to gain a competitive edge or economic advantage. It can be performed using various hacking techniques, malicious software or social*

engineering methods. Cyberespionage may be wholly conducted online or involve infiltration by computer-trained spies and moles.

The issue of industrial cyberespionage has been extensively addressed by technical and social disciplines. In reviewing this literature here we focus on strands that either provide us with tools and methodologies to analyse the problem or inform our theoretical and empirical frameworks. Additionally, we include a concise survey of economics literature, which is not generally considered to be part of the cybersecurity field but could be helpful in developing economic thinking about industrial cyberespionage. More specifically focused literature reviews can be found at the beginning of Chapters 2–4.

Thus, we divide the literature about industrial cyberespionage in three strands: (i) a first strand, which relies on microeconomic theory, and focuses on the influence of espionage on innovative incentives and social welfare, (ii) a second strand, driven by actuarial, statistical, and qualitative research, studies determinants of cyber risk and ways to quantify it, and (iii) a third strand, stemming from computer science and more recent economics of networks, examines the ways to protect organisational structures and ICT networks from cyberattacks.

1.3.1 Industrial espionage in competitive settings

Given the strategic nature of industrial espionage, it is natural to study it in a competitive setting (e.g. R&D races in which firms can target competitors' trade secrets and IP). The literature about espionage incentives in competitive situations can be divided into two sub-strands which differ by way of depicting an espionage mechanism: (i) a first sub-strand, which draws ideas from information economics, assumes that espionage is a tool capable of acquiring information about the opponent's actions, and (ii) a second sub-strand, which relies on classic models of competition and contest theory, depicts espionage as a mechanism of stealing or duplicating competitors' technologies or efforts. Most of the papers in this literature were written about conventional espionage. However, given the level of

abstraction of economic models, the insights from these papers are still relevant to the issues of modern cyberespionage.

Information Economics

The first sub-strand of literature is related to the economics of information, which stems from a seminal paper, “Economics of Information” by Stigler (1961). Economics of information generally focuses on how information affects economic decisions. For instance, it studies situations in which decision-makers are not perfectly informed about the utility or actions of their opponents.

Several papers modelled espionage as Bayesian games in which firms are not perfectly informed about their opponent’s actions and can choose to spy on him to retrieve this information. The early theoretical work by Matsui (1989) studied a two-person repeated game in which there is an exogenously defined chance of espionage, i.e. one or both firms will be informed about the other’s firm strategy and will be provided with an opportunity to revise its own strategy based on this information before the start of the next period. The author concluded that if the chance of espionage is sufficiently small, any pair of Subgame Perfect Nash Equilibrium payoffs is Pareto efficient⁸—contrary to the baseline case without espionage. Solan and Yariv (2004) studied a normal 2×2 game in which players choose their strategies before the start of the play, and one of the players can perform costly espionage and retrieve a noisy signal about his opponent’s actions. The paper recognised that a player chooses to spy only if his rival has a first-mover advantage.

Barrachina et al. (2014) studied the effects of industrial espionage on entry deterrence. The paper considered a sequential game in which the incumbent monopoly decides whether to invest in new capacity or keep the current production level, while the potential entrant decides whether or not to enter the market. The entrant is not perfectly informed about the decisions of the monopoly and receives a positive payoff only if the monopoly maintains its current capacity level. The entrant has an espionage device that gives him a noisy signal about the monopoly’s

⁸An outcome of the game is said to be Pareto efficient if there does not exist any other outcome of the game in which some player is better off and no other individuals are worse off.

action. The authors concluded that depending on the accuracy of the espionage device, it could benefit either the incumbent monopoly (if it is not sufficiently precise) or the potential entrant (if the device is sufficiently accurate).

Competition and Contests

An alternative approach to industrial espionage modelling was proposed by Whitney and Gaisford (1996, 1999), who developed a Cournot game in which companies could perform industrial espionage against their opponents in order to gain access to a competitor's technology and decrease their marginal costs. The papers provided evidence that industrial espionage inevitably leads to increased competitive incentives, higher payoffs and social welfare.

Billand et al. (2012) and Chen (2016) took a similar route, employing a modified Cournot setting to study competition with espionage. Billand et al. (2012) studied multi-market oligopolistic competition, in which companies can spy on their opponents. In Billand et al.'s paper, espionage was modelled as an exogenously defined boost awarded to the company if it decides to spy on its competitor. Echoing the results of Whitney and Gaisford (1996, 1999), they found that espionage can be beneficial for social welfare under certain circumstances.

Chen (2016) studied a game in which one of the companies possessed a superior technology, which allowed it to produce products with lower marginal costs. The author assumed that espionage incentives arise from relative differences in opponents' technological levels. Contrary to the predecessors, Chen (2016) found that espionage deteriorates competitive incentives.

More recently, Send (2022) adopted a Tullock contest framework to study the one-shot R&D race in which companies can steal each other's innovations by copying the research effort of their opponents and adding them to their own. The paper demonstrated that industrial espionage strictly decreases innovative incentives and the leading company's expected payoffs.

The study presented in Chapter 2 is closely related to the second sub-strand of the literature but absorbs some ideas and findings from the first one. Similarly to Solan and Yariv (2004) and Chen (2016), we assume that espionage incentives arise from the difference in their relative positions. However, contrary to the main body of literature, we study cyberespionage in a dynamic setting. Unlike Whitney and Gaisford (1999), Chen (2016) and Send (2022), we employ an endogenous technology award mechanism, meaning that companies must first compete for better technology before implementing it into a final product and competing for market profits. The dynamic setting allows us to distinguish two opposite effects of industrial cyberespionage on competitive incentives: the negative effect on research incentives and the positive effect on product development incentives. Thus, our framework captures most of the known results from the literature and explains the ambiguity of industrial cyberespionage effect on R&D race incentives, payoffs, and social welfare.

Additionally, we draw ideas from tournament theory initially introduced by Lazear and Rosen (1981) and Nalebuff and Stiglitz (1983). The original papers can be affiliated with personnel economics and presented a rank-order tournament framework which allowed the study of situations in which employees' wage differences arise from relative differences between individuals rather than their marginal productivity. It was assumed that an employer could not perfectly observe employees' effort levels and could only observe a noisy signal based on their real effort levels. The signal was modelled as a sum of the employee's effort level and a random component. Given that research always involves an element of luck (random component) and patents and advanced innovative products are often achieved based on a relative basis, this framework is well suited for studying R&D races.

The rank-order tournament framework was also utilised to study cheating in tournaments (e.g. Gilpatric, 2011; Kräkel, 2007; Stowe & Gilpatric, 2010) and head starts (e.g. Denter & Sisak, 2015, 2016; Siegel, 2014)—notions that we believe are fundamental to studying the influence of espionage in R&D race setting. Refer to the corresponding chapter for a more detailed literature review.

1.3.2 Data breach risks

This strand of literature is devoted to studies of cyber risks⁹. Due to the exponential growth of cybersecurity events that we observe in the digital era, there are numerous established strands of literature studying cyber risks¹⁰: from international relations studies focusing on the influence of large-scale cyberincidents on global welfare (e.g. RMS, 2016; WEF, 2015) to theoretical studies searching for methods to model cyberinsurance markets (e.g. Biener et al., 2015). Thus, we narrow our scope to studies that focus on the quantification of risks behind data breaches—cyberincidents that often happen as a result of industrial espionage or information leakage attack. We additionally divide this literature into two sub-strands: (i) a first sub-strand relies on various identification and analysis techniques to recognise and evaluate data breach risks and aid the development of better managerial practices, and (ii) a second sub-strand, which draws from actuarial science and econometrics in an attempt to quantify the probability and potential losses from various cyberincidents and data breaches. The sub-strands are highly intertwined, meaning that results from both are generally inter-applicable (e.g. the cyber risk factors initially identified in empirical risk management studies might be used to build better insurance pricing models).

Risk Management

A considerable amount of research was aimed at identifying organisational attributes connected with cyber risks to support better decision-making within an organisation and aid the development of particular management solutions to tackle cybersecurity problems. Most of the papers in this group were empirical and considered one or several factors that were believed to be associated with data breach risks.

Ransbotham (2010) and Moore and Anderson (2011) analysed the association between *the characteristics of software utilised by the enterprise* and data breaches.

⁹Cyber risk is any potential financial loss, disruption or damage to a company or its reputation due to a failure of its ICT systems (IRM, 2014).

¹⁰See Eling and Schnell (2016) for a comprehensive survey.

Ransbotham (2010) analysed the durability of open- and closed-source software against cyberattacks. The paper offered empirical evidence that open-source software is generally more vulnerable. The phenomenon was explained by the fact that public access to source code aids hackers in recognising security holes and exploiting them. Moore and Anderson (2011) hypothesised that hackers often employ relatively trivial techniques—for instance, they use publicly available search engines to identify websites with common security holes. The authors concluded that the presence of a commonly known (and unpatched) vulnerability in an enterprise’s software must be strongly associated with data breach risks. Thus, the type of software used by the organisation and its resilience were identified as critical attributes connected with organisational cyber risks.

T. Wang and Hsu (2010), Kwon et al. (2013), Hsu and Wang (2021), and Haislip et al. (2021) investigated the association between *a company’s board characteristics* and the likelihood of data breaches. The papers discovered that factors such as board diversity, the level of computer literacy of its members, and the presence of a Chief Information Officer (CIO) are all negatively correlated with the chance of being breached. Thus, the works suggest a strategic approach to board composition.

The link between *organisational structural complexity* and cyber risks was analysed by Tanriverdi et al. (2019) and Say and Vasudeva (2020). The authors found that structurally complex organisations are more likely to experience data breaches by analysing cyberincident occurrence after recent mergers, acquisitions, and subsidiary divestitures.

Additionally, it was demonstrated by C.-W. Liu et al. (2020) that *the degree of centralisation of IT systems* is another critical factor behind data breach risks. The author presented empirical evidence that organisations with a higher degree of centralisation of IT systems are less likely to experience security breaches.

A considerable body of research focused on analysing cyber risks in the healthcare industry. Angst et al. (2017) demonstrated that a hospital’s *size* is a statistically significant prediction in the empirical model of cyber risks. McLeod and Dolezel

(2018) and Kim and Kwon (2019) demonstrated that medical institutions' *reliance on electronic patient management systems* is positively associated with the data breach rate.

There is also a prolonged debate around the relationship between cybersecurity investments and data breaches. For instance, W. Liu et al. (2007) utilised regression analysis to verify the association between computer virus incidents and cybersecurity measures. The analysis identified that continuous investment in cybersecurity measures is associated with a decrease in data breach risks. The opposite results were demonstrated by Sen and Borle (2015), who demonstrated that cybersecurity investments are associated with an increase in the rate of data breaches. However, cybersecurity investment was not the main factor studied in Sen and Borle (2015). The paper has also demonstrated that *industry and country of a company's residence* play an essential role in modelling data breach rate.

The contradicting results around cybersecurity investments are most likely caused by the low quality and heterogeneity of data, which did not allow for causal inference (a common problem for empirical cybersecurity research). Only recently, those issues were addressed by Gandal et al. (2022), who managed to obtain higher quality data on Israeli companies and achieve a causal inference demonstrating that cybersecurity investment decreases the probability of being breached.

Actuarial science

Cyberinsurance is not per se studied in this thesis. However, in a race for better cyberinsurance models, a significant body of actuarial research is devoted to identifying insurance risk and rating factors¹¹ and exposure measures¹²—variables that are used for insurance premium pricing. Here we briefly review the main variables identified in actuarial science without going into details about insurance modelling.

¹¹Insurance risk factors (determinants of the level of cyber risks) and rating factors (risk factors or their proxies that are observable, measurable, and acceptable for marker use) used in cyberinsurance to reduce heterogeneity and allow for differentiating insurance rates according to the level of risk the insured company faces (Bardopoulos, 2020).

¹²Exposure measure represents a quantity that can serve as the basic unit of risk used for insurance premium estimation (Bardopoulos, 2020).

Early work by Hoo (2000) employed utility theory to analyse companies' decision trees based on computer security surveys. The author proposed to use *companies' wealth indices* (e.g. revenue) as rating factors and utilise *the customer database size* as a primary measure of exposure.

Böhme and Kataria (2006), H. S. Herath and Herath (2011), and Barracchini and Addessi (2014) built analytical models for cyberinsurance premium pricing. While the approaches varied among all papers (from regression analysis to mixed binomial models and Weibull distribution copulas), all presented frameworks agreed on a choice of the main exposure measure—*number of computers within the organisation* and the probability of their infection (often accelerated by correlated network effects).

Interesting qualitative work in this field was presented by Innerhofer-Oberperfler and Breu (2010). The authors interviewed 36 cybersecurity field experts to identify the main indicators for cyberinsurance models. The study concluded that “*the level of dependency on IT systems*”, “*the existence of know-how, patents, and other valuable information*”, and “*possessing sensitive, confidential data*” are the highest-ranked factors mentioned by most interviewed experts.

A large body of cyberinsurance research stems from actual insurance companies. Romanosky et al. (2017) and Bardopoulos (2020) performed a systematic qualitative analysis of cyberinsurance policies employed by large-scale insurers. Romanosky et al. (2017) identified four main categories of insurance risks and rating factors, including *organisational factors*, *technical factors*, *legal and compliance factors*, and *level of cybersecurity policies and procedures*. Bardopoulos (2020) further revealed that companies often utilise *internet revenue*, *the size of the workforce*, and *the number of confidential customer records* for cyberinsurance modelling.

The research presented in Chapter 3 is related to both sub-strands described above. It contributes to the literature by identifying another factor that might be used in cyber risk management or cyberinsurance modelling—research-intensity. We demonstrate that research-intensive industries are more likely to fall victim to information leakage attacks. Moreover, we provide new empirical evidence that this

association is industry-specific. High-tech manufacturing industries are more likely to experience targeted information leakage attacks, while knowledge-intensive service industries are more likely to fall victim to opportunistic financially-driven attacks.

Additionally, we draw on advanced fractional outcome modelling and partial data analysis techniques from the applied statistics literature to achieve robust statistical inference on low-quality cybersecurity data (Rubin, 1987; Wooldridge, 2010). A more detailed literature survey can be found in Chapter 3.

1.3.3 Organisational cybersecurity

While the previous literature mainly focused on identifying and quantifying cyber risks, the literature about organisational cybersecurity attempted to develop practical methods and strategies for defending from cyber threats, information leakage attacks, and industrial espionage. We distinguish two strands of literature relevant to our work: (i) a first, more technical strand draws ideas from computer science and machine learning to develop practical tools, mechanisms, and strategies to protect the organisation from cyber threats, and (ii) a second strand, which emerged at the junction of game theory and graph theory, is about optimal network design and protection—a recurring topic in cybersecurity.

Cybersecurity and Protection

Similarly to cyberinsurance, technical aspects of cybersecurity research are out of this work’s scope. However, some technical papers have informed the assumptions employed in the network design and defence framework presented in Chapter 4. Here we present a concise survey of papers relevant to our study.

A large body of computer science research is devoted to studying technical approaches and techniques to protecting sensitive data and intellectual property. Papers on the topic are focused on either the offensive side or the defensive side of the problem.

A deeper understanding of offensive techniques and methods is necessary to inform better solutions in defensive software and hardware design. A comprehensive

review of offensive means used for modern industrial espionage can be found in Androulidakis and Kioupakis (2016). The authors focused mainly on equipment and techniques which enable insider threats and eavesdropping (e.g. GSM interception systems, wall and parabolic microphones, and electronic interception devices) and malware that would allow a spy to access the necessary information from a distance (e.g. malware¹³ and man-in-the-middle attack¹⁴). The study pointed out the importance of the initial choice of the target regardless of a preferred method of espionage. The authors also discussed potential countermeasures to prevalent information leakage attacks and emphasised the significance of encryption and IT network monitoring and threat detection systems. Nevertheless, the study recognised that efficient protection from industrial espionage must not rely only on technical measures, suggesting a complex solution involving (1) legal aspects (e.g. patents), (2) both IT and physical security, (3) governance (e.g. enforcement mechanisms and audits), (4) finance (e.g. cybersecurity budgeting), and (5) management (e.g. employee cybersecurity awareness training).

Wangen (2015) and Ahmad et al. (2019) studied advanced persistent threats (APTs)¹⁵. Wangen (2015) offered a taxonomy of significant cyber-enabled espionage and information leakage incidents (e.g. GhostNet¹⁶, Dragonfly/Energetic Bear¹⁷, Careto¹⁸), describing their mechanisms and consequences. Ahmad et al. (2019)

¹³A malware attack is any cyberattack supported by malicious software that is designed to harm or damage an ICT system of an unaware end-user (Nahari, 2018).

¹⁴A man-in-the-middle (MITM) attack is a type of cyberattack in which the hacker secretly intercepts and relays the communications between two parties who continue to believe that they are in direct communication (Callegati et al., 2009).

¹⁵The advanced persistent threat is a method of cyberattack that utilises sophisticated, continuous and stealthy hacking techniques to gain prolonged access to the ICT systems of a target to retrieve valuable information or disrupt its operation.

¹⁶GhostNet was a large-scale targeted cyberespionage operation that relied on spear-phishing (crafted contextual emails) aimed at governmental, business, and academic institutions. The operation was conducted in 2009 and is believed to have been orchestrated by Chinese government-backed hacking groups.

¹⁷Dragonfly/Energetic Bear was a large-scale cyberespionage attack aimed at strategically important western organisations related to the aviation, energy, and defence sectors. The cyberattack happened in 2014 and was affiliated with Eastern-European hacking groups.

¹⁸Careto is espionage malware, which was used in targeted cyberespionage attacks (enabled by spear-phishing) against a relatively low number of industrial and governmental victims. The malware was operational for seven years in Spanish-speaking countries. However, neither the perpetrator nor the motivation behind the cyberattacks is known.

attempted to study the phenomenon from the perspective of military science and the science of organised crime to develop a theoretical model of the stages of advanced APTs. Both studies discovered that the large-scale cyberespionage and information leakage attacks rely on social engineering and spear-phishing. Thus, any targeted attack must include the “reconnaissance phase” during which the foe collects all the necessary information about the particular victim in the organisation (e.g. their role in the organisation, names of colleagues and managers) and information about the organisation’s technical systems vulnerabilities. The studies recognised that cyber-enabled spies often employ a strategic approach to espionage attacks and exploit socio-technical characteristics of an organisation’s networked structure.

Ahmed et al. (2020) analysed the impact of socio-technical factors in cloud data breaches. The study demonstrated that breaches happen due to aligning both technical and human factors. Additionally, it was discovered that often attacks are enabled due to human errors made by the professionals who manage the operation and security of the IT network. For instance, some data breaches (e.g. the Equifax data breach¹⁹) happened because of outdated server software, which remained unpatched for over five years.

More recently, Siddiqi et al. (2022) conducted a study of cyberattacks involving social engineering techniques, which exploit the human factors embedded in any organisation’s architecture. The study recognised that attacks of this type are extremely difficult to counter as they do not follow a specific pattern and are often crafted following a specific target’s characteristics. The authors claimed that due to an increasing number of cyberattacks, the technical countermeasures are also improving, making purely technical attacks difficult to perform. Thus, more and more attacks rely on social engineering as the most efficient technique to infiltrate an organisation’s networks and bypass technical measures. Echoing the results of Wangen (2015), the authors found that a successful social engineering attack requires an in-depth knowledge of the targeted organisation’s social and ICT structures. The authors also briefly analysed technical countermeasures which can

¹⁹The data breach occurred in 2017 and resulted in the disclosure of over 170 million private records of the clients of the American credit bureau Equifax.

mitigate the risks from social engineering attacks and suggested using comprehensive solutions reliant on firewalls, encryption, secure data management systems, and machine learning-driven spam filters.

The importance of machine learning in protection against information leakage attacks was also recognised by Parrend et al. (2018) and Henningsen et al. (2018). In their works, the authors argued that modern cybersecurity attacks are specifically tailored to the very network of a targeted organisation. Thus, static defensive measures are inefficient against novel threats, particularly zero-day²⁰ and advanced multi-step attacks. Therefore, Parrend et al. (2018) identified that statistics and machine learning could be used for real-time network data analysis—outliers and behavioural anomaly detection. Henningsen et al. (2018) suggested using similar techniques for insider attack detection in industrial ICT networks. The study suggested that the monitoring of networked data (e.g. nodes' communication patterns) could significantly improve the security of ICT systems of the organisation.

A machine learning approach to cybersecurity was also explored by Darvari et al. (2021), who utilised a reinforcement learning approach for a network design and defence problem. The researchers demonstrated that a reinforcement learning algorithm can be utilised to optimise real-world network structures (e.g. power grids) to be more resilient to both random and targeted (but not strategic) attacks.

While the above-mentioned research focused on either the attack vectors and techniques and recommended potential defensive techniques, a large body of research focused on implementing those techniques in particular software or hardware solutions.

Network monitoring and intrusion detection are reoccurring topics in computer science literature about data breaches. Mukkamala et al. (2005), Cheng et al. (2017), Khraisat et al. (2019), and Yuan and Wu (2021) offered comprehensive studies on the development, implementation, and limitations of Signature-based Intrusion Detection Systems (SIDS), Anomaly-based Intrusion Detection Systems (AIDS), and commercial antivirus software. Such systems utilise statistical and machine

²⁰Zero-day attack is a type of cyberattack that exploits recently discovered (and often not publicly disclosed) vulnerabilities.

learning algorithms to detect irregular agents' behaviour or patterns associated with known cyberattacks and insider threats. It was also recognised that apart from monitoring network traffic and traces, one can monitor human communication to filter out phishing and malicious mail. A review of natural language processing methods, techniques, and their implementation as well as the analysis of their effectiveness against social engineering attacks can be found in Stone (2007), Sawa et al. (2016), and Kidmose et al. (2018).

Apart from developing and implementing specialised cybersecurity tools, scholars have researched ways to improve the security of an organisation's automation systems. Eckhart et al. (2019) offered a review of good practices for software security testing and system integration and presented a framework for the semi-automated security analysis of cyber-physical systems. An interesting approach to secure software development was presented by Heitzenrater (2016). The author offered an economically-inspired decision-making framework for software security investment based on estimating returns on (security) investment. Thus, it was demonstrated that economic considerations could be one of the main drivers of secure software engineering processes.

Economics of Networks

Any organisational structure, an enterprise's hierarchy, or an interconnected system of ICT devices, can be depicted as a network. This fact was recently explored by a new strand of economics literature—the economics of networks, concerned with the understanding of economic phenomena by using concepts of game theory and graph theory.

A large body of literature considered an efficient network formation. For instance, early studies by Jackson and Wolinsky (1996) and Jackson and Rogers (2005) considered games in which each player can form costly links or sever them to extract the maximum value, which is a function of the number of direct and indirect connections that the player has. The authors analysed the stability²¹ and efficiency²²

²¹A network is stable when neither of the players wants to form or sever any links.

²²An efficient network is a network with the largest possible total value.

of such games. They demonstrated that stable networks, which are also efficient, do not always exist. In subsequent work, Jackson and Rogers (2005) expanded the base framework by allowing a heterogeneous link cost function (e.g. it is cheaper to connect to players that are geographically closer to you). They discovered that stable networks that emerge in this setting possess “small-world” features—that is, short average path length between two nodes and high clustering. These studies paved the way for a plethora of economics studies about efficient network formation. Refer to Jackson (2010) and Goyal (2007) for comprehensive surveys.

The economics of networks is also recognised to be helpful in security studies. Goyal and Vigier (2010, 2014) in their seminal papers analysed optimal secure network formation in the presence of an intelligent adversary. In their framework, two players, a designer and an adversary, compete over two stages in the network formation and defence game. First, the defender connects nodes in some network and distributes the defensive resources among them. Then, the adversary observes the network and the distribution of defensive resources, allocates contagious offensive resources among unprotected nodes, and chooses how those resources navigate the network. The papers focused on the trade-off between connectivity and the increased vulnerability it implies. It was discovered that in most circumstances, a star network with a protected central node emerges in the equilibrium. The authors offered to apply this model to guide secure computer network design or for the defence of Peer-to-Peer (P2P) networks.

Cerdeiro et al. (2015), Goyal et al. (2016), and Dziubiński and Goyal (2017) developed extensions to the previous framework in which the design and defence decisions were decentralised, thus, in the fashion of Jackson and Wolinsky (1996) and Jackson and Rogers (2005) allowing the nodes to choose their connections and defence efforts (“immunisation” plan). Cerdeiro et al. (2015) compared a decentralised approach to cybersecurity against the centralised approach demonstrated in Goyal and Vigier (2010). The authors concluded that in a decentralised framework, two main problems arise depending on the cost of defensive effort: over-protection—for the intermediate cost of defence, or under-protection if the cost of defence

is sufficiently high. Goyal et al. (2016) studied a similar decentralised setting and its applications to internet security. They discovered that any equilibrium network that arises from this setting will have at most $2n - 4$ edges for any $n \geq 4$. Dziubiński and Goyal (2017) studied the conflict intensity (defined as the minimum costs spent by the attacker and the defender) in a decentralised network defence setting. The paper discovered a new star-like formation that minimises conflict intensity—a windmill graph.

Hoyer and Jaegher (2016) studied the optimal network design against two types of external attacks: link-deletion and node-deletion attacks. Contrary to the general convention, the authors discovered that a star formation is the worst possible choice against a vertex deletion attack, as an adversary can deal enormous damage to the network by simply removing the central node. However, the designer did not have any defensive options.

Additionally, the economic problem of network design and defence was approached with reinforcement learning techniques. Y. Wang et al. (2021) used a game-theoretic specification of reward function to analyse optimal attack strategies. The paper discovered that the attacker never chooses a pure strategy in the equilibrium and always attacks some number of top valued nodes with positive probability (thus, choosing a mixed strategy).

Chapter 4 studies the design and defence of networks utilising the framework, which is closely related to the second strand of the literature. Instead of focusing on either network efficiency (e.g. Jackson & Rogers, 2005; Jackson & Wolinsky, 1996) or network security (e.g. Goyal & Vigier, 2010, 2014) we present a new framework which considers both of those aspects simultaneously. Our framework utilises a more sophisticated network value function, which accounts for graphs' connectivity and allows us to capture the tension between connectivity and efficiency. Moreover, we model the actual conflict on a network as a simultaneous move game in which neither the defender nor the attacker is aware of the actions of the other—which is more in line with real-world cybersecurity scenarios. This setting allowed us to

discover a new family of networks that arise when defensive resources are scarce—maxi-core networks with completely connected core and sparse periphery. The theoretical framework can be utilised to guide the secure design of an organisation’s ICT network or the development of algorithms behind intrusions detection systems.

Additionally, we draw ideas from contest theory. For instance, we utilise insights from Roberson (2006) and Kovenock and Roberson (2008, 2012, 2021), who studied optimal resource allocation over multiple battlefields (Colonel Blotto and General Lotto games). We also rely on computer science literature to inform the assumption utilised in our framework. For instance, we assume that if the attacker targets an unprotected node, he compromises it with probability one. The assumption is in line with real-world cyberincidents described in Wangen (2015), Ahmed et al. (2020), and Siddiqi et al. (2022). A more focused literature survey can be found at the beginning of Chapter 4.

1.4 Future research

The thesis offers an economic perspective on the issues of industrial cyberespionage and information leakage attacks. Each of the main chapters is a self-contained research project—however, the thesis just scratches the surface of a significant and sophisticated problem. In this section, we discuss the vectors of future research and opportunities which were discovered but have not been explicitly addressed in this work.

1.4.1 Chapter 2: patent design, multiple innovations, and endogenous attack function

Optimal patent design The chapter focuses on the influence of industrial cyberespionage on competitive incentives and social welfare. While the immediate attention of the study was on the influence of espionage power on competitive incentives, payoffs, and social welfare, we also identified that the framework might be linked to innovation economics and optimal patent design studies. Two exogenous variables utilised in the model, technological step and technological edge, can

be alternatively interpreted as patentability threshold and patent length, which correspond to patent design dimensions described by Scotchmer (2004). This interpretation and corresponding adjustments can lead to a framework capable of analysing optimal patent designs in an aggressive information environment (where competitors violate intellectual property laws for their good). This modification would also allow for the development of indirect policy measures that might foster higher R&D investments, payoffs, and better quality products through patent law regulations.

Multiple innovations Currently, the model assumes that each company invests in a single innovation. However, in reality, the product rarely relies on a single innovation. Allowing the players in the model to invest in multiple innovations might yield insights into optimal research budget allocation or strategic data organisation for increased security.

Attack mechanism The attack mechanism that we utilise in the model is relatively simple. Instead of using an exogenously defined industrial espionage power, it might be possible to employ a more sophisticated endogenously defined function of espionage power. For instance, we can employ an espionage function defined by monitoring and enforcement efforts of the social planner in the fashion of Stowe and Gilpatric (2010). In turn, this would allow us to study optimal enforcement policy design.

1.4.2 Chapter 3: causality and ecological fallacy

Causality In Chapter 3 we study the link between information leakage attack rate and research-intensity and demonstrate that research-intensive industries are particularly susceptible to information leakage attacks. We hypothesise that this association is caused by increased industrial cyberespionage activity in those industries and offer supporting evidence to support this causal mechanism. However, similar to most empirical cybersecurity studies, our work did not establish the actual

causality. The causal inference would require a set of instrumental variables or a much better data set with a time dimension for discontinuity study.

Ecological fallacy The conclusions drawn from our empirical study applied to the industry level. However, those conclusions do not necessarily apply to an individual company as industrial indices are aggregated at a high level, leaving considerable within-sector heterogeneity. The study would require more detailed microdata to overcome this issue and achieve conclusions applicable on an individual firm level.

1.4.3 Chapter 4: multiple attack resources, probabilistic defence function, and reinforcement learning

Multiple attack resources Currently, the Attacker in the model is assumed to have only a single offensive resource. While being adequate for cyberattack modelling, the assumption might not capture other real-world situations (e.g. military supply-chain defence). Thus, allowing the Attacker to have multiple attack resources will make the model much more general. However, this modification would require finding a method of determining the Attacker's support (the set of nodes the adversary attacks with positive probability) as the support is no longer monotonic when the assumption is relaxed.

Probabilistic defence The study employs the deterministic defence function. If the Defender protects some node and the Attacker chooses to attack it, the adversary always loses, and the node remains intact. While this assumption corresponds to multiple cybersecurity scenarios²³, it might not correspond well to other real-life situations in which the success of both parties depends on their efforts (e.g. air traffic security against terrorist attacks). Thus, the model can be generalised further with a probabilistic success function in a rent-seeking contest (Gradstein & Konrad, 1999).

²³For instance, the hacker attempting to exploit some common vulnerability in an organisation's ICT system, which was already patched, is unlikely to succeed (Ahmed et al., 2020; Moore & Anderson, 2011).

Reinforcement learning Network defence and design problems are also studied in reinforcement learning literature. Given the successful implementation of game-theory inspired reward function in Y. Wang et al. (2021), it might be feasible to adopt our framework for a similar purpose. This would allow for studying the problem numerically and achieving solutions for specific modifications which are beyond analytical approach capabilities.

1.5 Dissertation structure

The rest of the dissertation is built as follows. Chapter 2 considers issues of industrial cyberespionage in research and development races. The empirical model that studies the link between research-intensity and information leakage attacks is presented in Chapter 3. Chapter 4 offers a theoretical framework of network design and defence. Chapter 5 concludes.

2

Industrial Cyberespionage in R&D Races

Contents

2.1	Introduction	29
2.1.1	Related literature	32
2.2	The model of economic cyberespionage	37
2.2.1	Development stage	41
2.2.2	Research stage	47
2.3	Social welfare	55
2.3.1	Espionage and welfare	56
2.4	Conclusion and discussion	57

2.1 Introduction

This chapter is devoted to the issues of modern research and development (R&D) races. Specifically, we focus on the incentives behind industrial cyberespionage and how it influences R&D investments, research companies' profits, and the end-product quality. We define industrial cyberespionage as the practice of obtaining competitors' secrets and information stored in digital formats without permission in order to gain an economic advantage. It can be performed using various hacking techniques, malicious software or social engineering methods. Cyberespionage may be wholly conducted online or involve infiltration by computer-trained spies and

moles. Seemingly any modern espionage is cyberespionage since it is improbable that it would not involve at least some information communication technology (ICT).

It is straightforward to find examples of industrial cyberespionage in the news feed. For example, in October 2018, the supply chain of technological giant Super Micro was infiltrated by Chinese adversaries. They created a backdoor in some of Super Micro’s server equipment that would allow them to steal corporate secrets in the future (Lynch, 2018). In November 2018, the US justice department filed criminal charges against Fujian Jinhua, a Chinese state-owned company, for the alleged theft of trade secrets worth up to \$8.75 billion from Micron, an American semiconductor manufacturer (Shubber, 2018). Industrial espionage has caused an enormous loss for the global economy. For instance, The Commission on the Theft of American Intellectual Property (2017) estimated that espionage via hacking costs between \$180 billion and \$540 billion annually to the US economy.

The situation aggravates every year with an increasing number of organisations and individuals being affected. The current COVID-19 crisis makes this even more apparent. During the crisis, national security agencies reported a significant increase in cyberespionage activity over the globe. People stay at home during lockdowns, stressed, and often violate good cybersecurity practices. It creates an excellent attack surface for a series of social-engineering attacks performed by state-backed hacking groups (Warrell & Manson, 2020).

Cyberespionage has become an inalienable part of contemporary competition. Nonetheless, the phenomenon has not been rigorously studied under the competitive framework. We fill this gap by developing a multi-stage R&D race model with endogenous espionage decisions. The model is built on a rank-order tournament framework by Lazear and Rosen (1981), which is consistent with innovative activity as it embeds a degree of uncertainty in companies’ decisions and has a relativistic award mechanism that resembles the patenting process.

We extend Lazear and Rosen’s framework to consider two ex-ante identical companies that compete in an R&D race over two stages for a better innovation to acquire the market. During the initial *research* stage, rivals invest in research to

obtain superior technology—a technological edge. In order to obtain a technological edge, a company must outperform its opponent by a substantial margin, which we call a technological step. During the following *development* stage, companies invest again to perfect their technologies and release them to the market.

Espionage incentives may arise from the difference in relative positions of competitors who race to innovate. If the gap between rivals is sufficiently large, it will create a strong incentive for the underdog company to steal the leader’s intellectual property. Thus, cyberespionage might increase an underdog’s chances of success and its expected profits. For example, an attacker can obtain classified information to manufacture copy-cat products and gain market share.

If one of the companies gains a technological edge in the research stage, it utilises this advantage in the development stage to achieve better innovation and increase its chances of winning the market. A company that has not gained a technological edge (an underdog) can choose to perform a cyberattack at a fixed cost to steal a fraction of the competitor’s technological edge in an attempt to equalise the chances in the development stage. If neither company gains a technological edge in the research stage, the race continues with the symmetric development stage, where both companies have the same available technologies, and nobody performs industrial espionage. By the end of the game, the company with a better innovative product wins the race and obtains all the market to itself. The framework is stated formally and in greater detail at the beginning of Section 2.2.

The study makes two main contributions. First, we demonstrate that industrial espionage has an ambiguous influence on the overall investments exerted in the race. Espionage negatively influences the investments in the research stage and positively influences the investments in the development stage. Intuitively, none of the companies would be willing to invest in the technological edge during the initial *research* stage, knowing that it is most likely to be stolen. However, if one firm gets far ahead in the research stage, its rival might “give up” in the development stage. On the other hand, if it catches up with the opponent (through espionage or otherwise), it will want to compete because it has a realistic chance of winning.

Therefore, by regulating the power of industrial espionage through legal means, a policymaker can control competitive incentives in the R&D race and maximise investments. Cyberespionage can maximise firms' payoffs from participating in the R&D race and make the race more attractive for them.

Second, we investigate the welfare implications of cyberespionage. Social welfare is measured by the quality of the innovative product delivered to the market by the winner of the race. We demonstrate that, under specific conditions, dosed espionage might be beneficial for society. It generally happens when the technological edge awarded at the initial stage is too large, which decreases competitive incentives at the second development stage. In turn, diminished efforts at the development stage decrease the overall quality of an end-user product. Thus, industrial cyberespionage might counter the inefficient (from a market perspective) technological edge, motivate competition, and foster higher quality products. It also means that by choosing an optimal value of technological edge (e.g. by setting an optimal duration of patent protection), policymakers can make industrial cyberespionage obsolete while encouraging competition.

The chapter is built as follows. This section provides an introduction and a review of related literature. Section 2.2 formalises the model and provides the results. The social welfare analysis is presented in 2.3. Section 2.4 concludes. For those interested in computing equilibria and related proofs, see Appendix A.1.1. The derivation of the Social Welfare Function and related proofs can be found in Appendix A.1.2.

2.1.1 Related literature

This chapter contributes to a rich strand of literature about industrial espionage in R&D races. The early works by Matsui (1989) and Solan and Yariv (2004) could be attributed to information economics as the espionage mechanism was modelled as a tool which ex-ante uninformed players could use to obtain information about the actions of their opponents. Matsui (1989) studied a two-person repeated game where espionage happened with some exogenously set probability and provided players with an opportunity to learn about their opponents' actions. It was demonstrated

that dosed espionage could lead to a Pareto-efficient subgame perfect equilibrium in which both players increase their payoffs. Solan and Yariv (2004) modelled strategic espionage in a sequential non-zero-sum game and demonstrated that a firm is incentivised to conduct espionage only if its rival has a first-mover advantage. Even though we use a very different espionage mechanism, the findings of those papers are echoed in the presented framework. Similarly to Matsui (1989), we find that espionage can work as a payoff-increasing collusion device. The assumptions behind the espionage mechanism in our framework absorbed the findings of Solan and Yariv (2004)—in our model, espionage incentives arise when one of the companies is sufficiently behind in the innovative process.

More recent work presented by Barrachina et al. (2014) studied the effects of espionage on market entry deterrence. They studied an asymmetric two-player game in which a monopoly incumbent may choose to extend its capacity in order to deter the entrance of a potential entrant. The entrant is unaware of the incumbent's choice. However, he can choose to use "The Intelligence System"—a device that provides a noisy signal about the incumbent's actions. Similarly to Matsui (1989), they found that espionage makes the market more competitive and leads to higher payoffs, which could also be observed in our framework under certain circumstances.

Another popular way of modelling espionage was initially presented by Whitney and Gaisford (1996, 1999). The authors utilised a modified Cournot game in which espionage could decrease the marginal cost of a firm by providing access to a better production technology retrieved from the competitor. They primarily focused on the influence of espionage on social welfare. It was argued that espionage could be considered as a technology transfer, which inevitably leads to increased competitive incentives, higher payoffs, and higher social welfare. Our framework similarly depicts espionage, allowing companies to steal technologies from each other.

Unlike the papers mentioned above, we study espionage in a dynamic setting. In our model, the espionage decision is made in the middle of the race and only if one of the ex-ante identical companies has achieved superior technology. This allows for a more flexible model that distinguishes the effects of espionage on innovative

incentives in a race for a technological edge and competitive incentives in a race for a market share. Additionally, our approach reveals that, depending on its power, espionage can influence social welfare both positively and negatively.

Billand et al. (2012) and Chen (2016) have also studied espionage in modified Cournot settings. Billand et al. (2012) studied an oligopolistic competition with espionage in a multi-market setting. Espionage was modelled as an exogenous boost that companies attain if they choose to spy on their competitors. The magnitude of a boost did not depend on the effort levels of the opponents and simply increased in the number of compromised rivals. This framework did not allow for the investigation of the relationship between espionage and competitive incentives but, similarly to Whitney and Gaisford (1996, 1999), demonstrated that spying is beneficial for social welfare. Similarly to Solan and Yariv (2004) and to our framework, Chen (2016) assumed that espionage incentives arise from relative differences in opponents' technological levels. In Chen's framework, one of the companies was assumed to possess a more sophisticated technology and be capable of producing products at a lower marginal cost. They demonstrated that espionage reduces rivals' competitive incentives, which contradicts the findings of our study. This difference arises from a static setting of the model. By allowing for the endogenous technological edge award mechanism, we demonstrate that the influence of espionage is ambiguous: it decreases innovative incentives in the initial race for better technology and increases the incentives in the following stage where the technologies are implemented in a final product.

Send (2022) also considered a one-shot R&D race with espionage in which one of the companies is assumed to be more technologically advanced. The author, however, drew ideas from a contest theory and employed a modified version of a Tullock contest (Tullock, 2001). Similarly to Chen (2016), a more advanced company is assumed to have a smaller marginal cost of innovative efforts. Espionage was modelled as a device which copies the opponent's research effort and adds it to the company's own. The authors found that the disadvantaged player is more likely to win the race if espionage is possible. It was also demonstrated that espionage strictly

decreases overall R&D investment and the leading company's expected payoff. On the contrary, our model demonstrates that espionage's influence on overall efforts and payoffs is ambiguous. The difference can be once more attributed to the static nature of Send's model. This model can be considered a variation of the second stage of the R&D race presented in the chapter and does not capture our main results.

As mentioned before, this study builds on the ideas of tournament theory originally introduced in seminal articles by Lazear and Rosen (1981) and Nalebuff and Stiglitz (1983). The original papers studied wage differentiation based on relative performance differences between employees with a framework widely recognised as a rank-order tournament. They analysed optimal prize (wage) structure in a static setting to find methods to increase employees' competitive incentives and productivity. Since then, their model has found numerous applications in various fields unrelated to personnel economics. The random element of the original rank-order tournament makes it a good fit to model an innovation process. Additionally, we modify an originally static setting into a dynamic one. The dynamic structure of the framework is essential for the investigation as incentives behind cyberespionage attacks arise when competitors find themselves in uneven positions in the middle of the R&D race.

Several other papers studied dynamics in tournaments. Denter and Sisak (2015, 2016) and Siegel (2014) have analysed head starts in multi-stage contests under symmetric and asymmetric assumptions. In their papers, the authors discovered that awarding a small head start to one of the players whenever they interact over multiple stages is optimal, even in the case of symmetric players. This research contributes to the literature by introducing an endogenous award mechanism, which allows players to obtain a head start by performing a cyberespionage attack.

There were a small number of papers in tournament theory which focused on cheating in contests—a phenomenon which resembles cyberespionage. Berentsen and Lengwiler (2005) in their article on fraudulent accounting and doping in sports utilised a modified rank-order tournament model in which they implemented a cheating mechanism—the doping game. They investigated the evolutionary

heterogeneous players’ framework where contestants can improve their performance by using performance-enhancing drugs. They described cases where high-ability contestants are more likely to cheat than low-ability ones. The players could boost their performance by adding an exogenously defined variable to their effort levels. However, the model assumes ex-ante heterogeneous players and does not account for uncertainty in a race: the “doped” player wins over any clean player with certainty. Moreover, the cheating mechanism in the framework does not depend on the opponent’s performance—any player can boost her effort levels regardless of how successful the opponent is. This makes the model unsuitable for studying cyberespionage in R&D races.

The doping game was also studied by Kräkel (2007). The author analysed the dynamics of two asymmetric players in the tournament with cheating. The paper discovered three determinants of a player’s cheating decision: the likelihood effect (cheating increases the winning probability), the cost effect (cheating influences the effort costs) and the windfall-profit or base wages effect (the decrease in base salary due to the possibility of getting caught). However, the cheating mechanism broadly remained similar to the predecessors.

Rather than study cheating, Curry and Mongrain (2000) focused on ways to prevent it. They demonstrated that prize structure plays a crucial role in determining incentives for cheating. It is possible to decrease cheating solely by setting up a correct prize structure instead of investing in sophisticated monitoring systems. Unlike the model presented in the chapter, they assumed that both effort and cheating decisions are simultaneous, which might not be the case in many real-life situations (e.g. R&D).

The series of papers on cheating, enforcement and audit in contests by Stowe and Gilpatric (2010) and Gilpatric (2011) is closely related to our research as well. In the 2010 paper, the authors analysed the incentives behind prohibited behaviour in rank-order tournaments while the base effort was considered as given. Their framework should be considered the final stage of the tournament, where effort levels are already known, and the players only make a dichotomous choice on whether to cheat or not.

The latter paper described how the prize structure, uncertainty, monitoring efficiency, number of contestants, and associated penalty impact cheating. They found that greater enforcement might harm players' productivity, which partially echoes our results. However, due to the limitations of a static framework, the authors did not capture the cases in which cheating positively influences contestants' effort levels.

Therefore, this study contributes to the literature by studying the dynamic interaction between contestants, their relative effort levels and payoffs, and how they are affected by the presence of espionage (or cheating). It investigates the influence of espionage on social welfare by converting the relative outcomes of rank-order tournament into absolute ones, which, to our best knowledge, has not been done before.

2.2 The model of economic cyberespionage

Two players, firm 1 and firm 2, engage in an R&D race over two stages: the initial **research stage** and the following **development stage**. Participating in the research stage and the development stage, the companies invest x_i and $\hat{x}_i$ (where $i = 1, 2$), respectively, which increases their chances of gaining a superior innovation by the end of the race. If one of the companies pulls ahead after the research stage and achieves an advantage (a technological edge), its opponent might choose to carry out a cyberespionage attack to steal a fraction of the advantage and increase the chances of winning the race. During the race, the competitors experience random shocks, which account for the uncertain nature of R&D (N. R. Baker et al., 1986). In line with the classical framework of price competition by Bertrand (1989), the company which delivers superior technology by the end of the development stage remains the only one active on the market and collects all of the profits to itself. Society benefits from the quality of the new technology delivered to the market by the winner of the race. The social welfare function is formally defined in Section 2.3. Figure 2.1 outlines the timings of the game.

Similar to Lazear and Rosen (1981), shocks (ϵ_i and $\hat{\epsilon}_i$) are drawn from a known distribution with zero mean, infinite support, and variance $\frac{1}{2}\sigma^2$. It is assumed

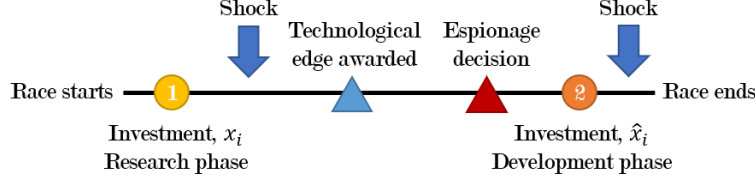


Figure 2.1: Timings of the R&D race. Circles represent the timings for firms' investment decisions

that shocks are independent and identically distributed (i.i.d.) across firms and stages. They represent the luck of competitors and can influence the research either positively or negatively. The shocks for the research and development stages are denoted ϵ_i and $\hat{\epsilon}_i$, respectively. Throughout the chapter, we use the notations $f(\cdot)$ and $F(\cdot)$ to denote the probability density function (PDF) and the cumulative distribution function (CDF) of the difference between two shocks.

A firm's output of the **research stage** is expressed as the sum of investment and the random shock:

$$q_i = x_i + \epsilon_i. \quad (2.1)$$

The structure of the research stage closely resembles the original model of Lazear and Rosen (1981). However, it uses an alternative award mechanism to account for two possibilities in an innovation race: concurrent failure and simultaneous invention, both of which can be treated as a similar event from a relative standpoint (Lemley, 2011). A company is said to be the winner of the research stage if it outperforms its opponents by a substantial margin k , e.g. firm i wins in the research stage if $q_i - q_{-i} > k$. We call variable k a technological step. Firm i is then awarded a technological edge h , which adds on to the output of the development stage and provides an advantage over its opponent. Conversely, if the output difference is not significant, $|q_i - q_{-i}| < k$, participants achieve parity, and nobody gains a technological edge h . Note that technological step k and technological edge h are assumed to be discrete. While one could argue that innovation is a continuous phenomenon, innovation often happens in discrete steps. For instance, microprocessor manufacturers compete to achieve a better fabrication process,

which enables smaller silicon transistors on a printed circuit board (PCB). The advancement of the process is measured in nanometers and usually happens in discrete steps (e.g. 14 nm, 10 nm, 7 nm). In turn, the advanced fabrication process allows for microprocessors with more cores and threads—which are also discrete characteristics (Shrout, 2018).

The race then proceeds to the final stage of the R&D process—the **development stage**. During this stage, the companies implement their innovations into end-user commercial technologies. A company with a better end-user product wins the market, the value of which is normalised to 1, $w = 1$.

The development stage depends on the outcome of the research phase. If neither of the companies has gained a technological edge, the game follows a standard rank-order tournament pattern. Each firm then produces a product of quality:

$$\hat{q}_i = \hat{x}_i + \hat{\epsilon}_i. \quad (2.2)$$

If one of the firms obtains a technological edge h , the underdog firm might choose to carry out a cyberespionage attack to retrieve a fraction $\alpha \in [0, 1]$ of the opponent's head start at a fixed cost A and boost her chances of winning the race. This relatively simple cyberattack mechanism captures the fact that the cyberattacks are mainly performed by hired third-party hacking groups as companies rarely participate in the offensive operations themselves.

Without loss of generality, assume that firm 1 is the leader while firm 2 is the underdog. Then their output levels are:

$$\hat{q}_1^a = \hat{x}_1 + \hat{\epsilon}_1 + h, \quad (2.3)$$

$$\hat{q}_2^a = \hat{x}_2 + \hat{\epsilon}_2 + (\alpha h), \quad (2.4)$$

where $\alpha h = 0$ if the underdog decides not to attack. The company that achieves a higher output in the development stage wins the main prize.

The R&D effort x (either x_i or $\hat{x}_i$), is costly and is produced at cost $c(x)$. The cost function is assumed to be convex in x : $c'(x) > 0$ and $c''(x) > 0$ for all $x > 0$. As the study focuses on the process dynamics, sunk costs are omitted from the

research: $c(0) = 0$. The cost function's convexity indicates that the marginal cost of the R&D effort increases, and every additional investment unit becomes less efficient. For instance, the effectiveness of the investment in developing a new matrix for a computer screen tends to zero as soon as the pixels become indistinguishable to a human eye. Despite the convexity of the cost function, competitors' payoff functions do not need to be quasiconcave in investments. We also impose the following assumption:

Assumption 2.1. *The second derivative of the cost function is greater than the marginal benefits for any $x \in \mathbb{R}$:*

$$\max f'(x) < \min c''(x).$$

The assumption guarantees that $c'(\cdot)$ crosses $f(\cdot)$ only once from below (see Figure 2.2), which ensures the existence of interior pure strategy equilibrium in every stage.

The overall payoffs of the game are equal to:

$$E[\Pi_i] = P^S \hat{P}_i^s + P_1^A \hat{P}_1^a + P_2^A (\hat{P}_2^a - (A)) - c(\hat{x}_i) - c(x_i), \quad (2.5)$$

where P^S is the probability that neither of the companies has gained a technological edge in the research stage and $\hat{P}_i^s$ is the probability of achieving a superior innovation in the symmetric development stage from an even position; P_1^A is the probability of gaining a technological edge in the research stage and $\hat{P}_1^a$ is the probability of achieving a superior innovation in the development stage from a leading position; P_2^A is the probability of losing a technological edge to the opponent in the research stage and $\hat{P}_2^a$ is the probability of achieving a superior innovation in the development stage from an underdog position; A is a cost of cheating, which equals to 0 if the underdog decides not to engage in cyberespionage. Note that cheating might only occur when the development stage of the game is asymmetric and, therefore, incorporated in terms $\hat{P}_1^a$ and $\hat{P}_2^a$.

The study uses the Subgame Perfect Nash Equilibrium (SPNE) as a solution concept and assumes that each company rationalises against its opponent's optimal

investment level, representing a firm's inability to influence the market. The game is solved by backward induction beginning at the development stage.

2.2.1 Development stage

Depending on the outcome of the research stage, there are three possible cases: (i) symmetric case if nobody has a head start (technological edge) in the development stage, $HS = 0$, (ii) asymmetric case with cyberespionage and the leader's head start, $HS = h(1 - \alpha)$, and (iii) asymmetric case without cyberespionage and the leader's head start, $HS = h$.

Symmetric case, HS=0

In the symmetric case both companies have the same marginal incentives. Thus, their behaviour should also be identical if Assumption 2.1 is satisfied. The firms' expected payoff functions are then:

$$\hat{\Pi}_i^s = \hat{P}_i^s - c(\hat{x}_i), \quad (2.6)$$

where $\hat{P}_i^s$ is the probability of winning the game. It can be expressed as:

$$\begin{aligned} \hat{P}_i^s &= P(\hat{q}_i > \hat{q}_{-i}) \\ &= P(\hat{x}_i - \hat{x}_{-i} > \hat{\epsilon}_{-i} - \hat{\epsilon}_i) \\ &= P(\hat{x}_i - \hat{x}_{-i} > \hat{\zeta}) \\ &= F(\hat{x}_i - \hat{x}_{-i}), \end{aligned} \quad (2.7)$$

where $\hat{\zeta} = \hat{\epsilon}_{-i} - \hat{\epsilon}_i$ is the convolution of two random variables. As shocks are i.i.d., $E(\hat{\zeta}) = 0$ and $E(\hat{\zeta}^2) = \sigma^2$. Substituting (2.7) into (2.6) yields:

$$\hat{\Pi}_i^s = F(\hat{x}_i - \hat{x}_{-i}) - c(\hat{x}_i). \quad (2.8)$$

Each player then chooses $\hat{x}_i$ to maximise (2.8), taking its opponent's strategy as given. Given that Assumption 2.1 holds and that the objective function is twice differentiable, maximisation requires:

$$\frac{\partial \hat{\Pi}_i^s}{\partial \hat{x}_i} = f(\hat{x}_i - \hat{x}_{-i}) - c'(\hat{x}_i) = 0, \quad (2.9)$$

$$\frac{\partial^2 \hat{\Pi}_i^s}{\partial \hat{x}_i^2} = f'(\hat{x}_i - \hat{x}_{-i}) - c''(\hat{x}_i) < 0, \quad (2.10)$$

where (2.9) is the set of first-order conditions (FOCs) and (2.10) is the set of second-order conditions (SOCs).

We find that if the solution exists, both firms choose identical investment strategies. The result should be intuitive, considering the symmetry of marginal incentives of firms. Note that the marginal increase in the probability of winning the market by a firm is precisely the density of the uncertainty shock evaluated at the relative investment level. Since participants also have the same well-behaved cost functions, their marginal costs are also identical. Consequently, given any prior distribution $F(\cdot)$, firms have identical marginal incentives and choose the same strategies. Thus, the equilibrium must be symmetric. This finding is summarised in the proposition below.

Lemma 2.1. *The equilibrium of the symmetric case of the development stage exists and is unique if Assumption 2.1 is satisfied. In the equilibrium of the subgame, both companies choose identical investments:*

$$\hat{x}_i^s = c'^{-1} [f(0)] . \quad (2.11)$$

Proof of Lemma 2.1. See Appendix A.1.1. ■

The intersection of marginal benefits $f(\cdot)$ and marginal costs $c'(\cdot)$ might not be unique. In Figure 2.2, graph (a) demonstrates a case of the unique intersection, while graph (b) shows a case where the intersection is not unique.

Assumption 2.1 ensures that marginal costs intersect with marginal benefits only once “from below”, as illustrated in Figure 2.2 (a). It rules out the possibility of multiple intersections, as shown in Figure 2.2 (b), and guarantees that the solution for the system of the FOCs is indeed an equilibrium.

Substituting optimal investment levels (2.11) into (2.8) returns the equilibrium payoffs for the symmetric case of the stage:

$$\hat{\Pi}_i^S = \frac{1}{2} - c \left[c'^{-1} [f(0)] \right] . \quad (2.12)$$

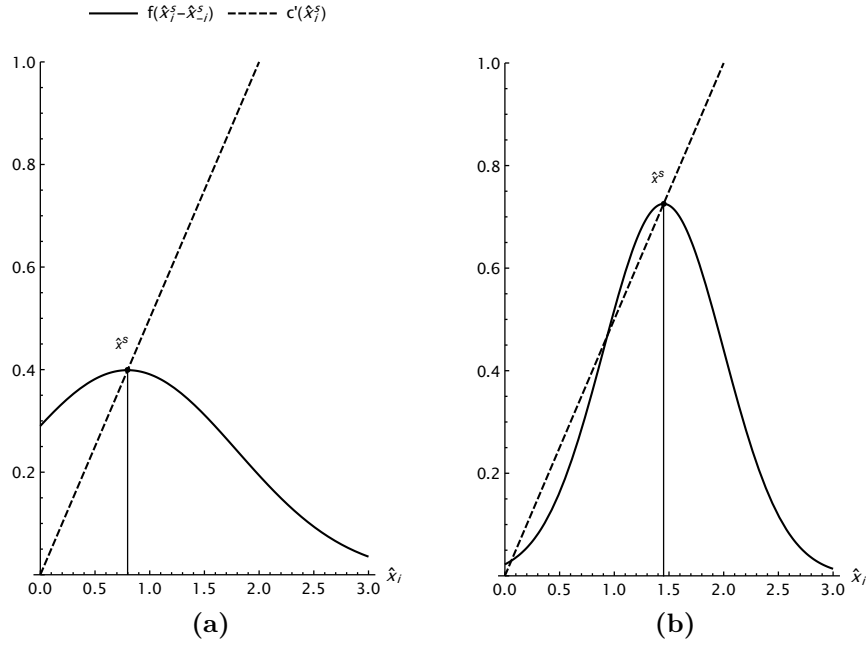


Figure 2.2: The uniqueness of Nash equilibrium depending on the slopes of PDF and cost functions.

It is clear that if neither of the companies gains a technological advantage in the research stage, the winner is decided by a coin toss. Thus, in equilibrium, the development stage has zero influence on the expected winner of the race: both choose the same strategies, have the same marginal incentives, and zero net effect.

Asymmetric case involving cyberespionage, $HS=h(1-\alpha)$

If one of the competitors gains a technological edge in the research stage, the contenders continue the race from uneven positions. Without loss of generality, we assume that firm 1 is the leader who has gained the technological edge h in the previous stage. Then, firm 2 is the underdog and may choose to carry out a cyberespionage attack and retrieve a portion α of the competitor's head start if the attack costs A are sufficiently small. Assuming sufficiently small attack costs, the firms produce innovative products of qualities $\hat{q}_1^a$ and $\hat{q}_2^a$ respectively:

$$\hat{q}_1^a = \hat{x}_1 + \hat{\epsilon}_1 + h,$$

$$\hat{q}_2^a = \hat{x}_2 + \hat{\epsilon}_2 + \alpha h.$$

The firms' payoffs for the asymmetric case are then:

$$\begin{aligned}\hat{\Pi}_1^a &= \hat{P}_1^a - c(\hat{x}_1), \\ \hat{\Pi}_2^a &= \hat{P}_2^a - c(\hat{x}_2) - A,\end{aligned}\tag{2.13}$$

where A is the cyberattack costs and $\hat{P}_i^a$ are the probabilities of winning the stage. The leader then wins with the probability:

$$\begin{aligned}\hat{P}_1^a &= P(\hat{q}_1 > \hat{q}_2) \\ &= P(\hat{x}_1 - \hat{x}_2 + h - \alpha h > \hat{e}_1 - \hat{e}_2) \\ &= P(\Delta\hat{x} + h(1 - \alpha) > \hat{\zeta}) \\ &= F(\Delta\hat{x} + h(1 - \alpha)),\end{aligned}\tag{2.14}$$

while the underdog wins with the probability:

$$\begin{aligned}\hat{P}_2^a &= 1 - \hat{P}_1^a \\ &= 1 - F(\Delta\hat{x} + h(1 - \alpha)).\end{aligned}\tag{2.15}$$

Substituting equations (2.14) and (2.15) into payoffs (2.13) yields:

$$\begin{aligned}\hat{\Pi}_1^a &= F(\Delta\hat{x} + h(1 - \alpha)) - c(\hat{x}_1), \\ \hat{\Pi}_2^a &= 1 - F(\Delta\hat{x} + h(1 - \alpha)) - c(\hat{x}_2) - A.\end{aligned}\tag{2.16}$$

The firms then maximise their expected payoffs by choosing optimal investment levels $\hat{x}_i^a$, which must satisfy the following FOCs:

$$\begin{aligned}\frac{\partial \hat{\Pi}_1^a}{\partial \hat{x}_1} &= f(\Delta\hat{x} + h(1 - \alpha)) - c'(\hat{x}_1) = 0, \\ \frac{\partial \hat{\Pi}_2^a}{\partial \hat{x}_2} &= f(\Delta\hat{x} + h(1 - \alpha)) - c'(\hat{x}_2) = 0,\end{aligned}\tag{2.17}$$

and SOC:

$$\begin{aligned}\frac{\partial^2 \hat{\Pi}_1^a}{\partial \hat{x}_1^2} &= f'(\Delta\hat{x} + h(1 - \alpha)) - c''(\hat{x}_1) < 0, \\ \frac{\partial^2 \hat{\Pi}_2^a}{\partial \hat{x}_2^2} &= -f'(\Delta\hat{x} + h(1 - \alpha)) - c''(\hat{x}_2) < 0.\end{aligned}\tag{2.18}$$

Even though the firms' profit functions are now asymmetric, marginal incentives are still identical. The marginal increase in the probability of winning the game is

the density function evaluated at the relative investment level, which is influenced by technological edge h and espionage attack power α . Given that the shock is drawn out of a distribution symmetric around zero, marginal benefits are precisely the same. The cost functions are symmetric by assumption. Therefore, it must be the case that both firms again choose identical investments. The finding is summarised below.

Lemma 2.2. *The equilibrium of the asymmetric case of the development stage involving cyberespionage exists and is unique if Assumption 2.1 is satisfied. In the equilibrium of the subgame, both companies choose identical investments:*

$$\hat{x}_i^a = c'^{-1} \left[f(h(1 - \alpha)) \right]. \quad (2.19)$$

Proof of Lemma 2.2. See Appendix A.1.1. ■

Substituting optimal investment levels (2.19) into (2.16) yields equilibrium payoff levels:

$$\begin{aligned} \hat{\Pi}_1^A &= F(h(1 - \alpha)) - c \left[c'^{-1} \left[f(h(1 - \alpha)) \right] \right], \\ \hat{\Pi}_2^A &= 1 - F(h(1 - \alpha)) - c \left[c'^{-1} \left[f(h(1 - \alpha)) \right] \right] - A. \end{aligned} \quad (2.20)$$

In the asymmetric case of the race, the leader has more chances to win the market due to the technological edge, while the underdog might balance the game by cyberattack means. The underdog's binary decision about espionage is investigated in the subsection below.

It also follows from (2.11) and (2.19) that companies are exerting less effort in the asymmetric case than in the symmetric one. Thus, uncertainty encourages R&D companies to invest more.

Asymmetric case not involving cyberespionage, HS=h

It is straightforward to obtain the equilibrium payoffs for the fair instance of the asymmetric case (without cyberespionage). Substituting $\alpha = 0$ and $A = 0$ to (2.17) yields a solution to the game with no espionage:

$$\hat{x}_i^f = c'^{-1} [f(h)]. \quad (2.21)$$

Proposition 2.1. *The equilibrium investments in the asymmetric case with cyberespionage are higher than those in the asymmetric case without cyberespionage.*

Proposition 2.1 immediately follows from (2.19), (2.21) and the fact that $\hat{\zeta}$ has a single-picked distribution with zero mean and infinite support.

The intuition behind the proposition should be straightforward. When the gap between the leader and the underdog is too large after the research stage, it decreases investment incentives for both of them: the leader does not need to invest a lot as she is already far ahead, and the underdog is not motivated to invest in a surely lost game. By diminishing the gap, cyberespionage encourages players to invest more and stimulates competition.

Optimal payoffs for the asymmetric case without cyberespionage are:

$$\begin{aligned}\hat{\Pi}_1^F &= F(h) - c \left[c'^{-1} [f(h)] \right], \\ \hat{\Pi}_2^F &= 1 - F(h) - c \left[c'^{-1} [f(h)] \right].\end{aligned}\tag{2.22}$$

The existence and uniqueness of the equilibrium are guaranteed if:

$$f'(\hat{x}_i - \hat{x}_i^f) < c''(\hat{x}_i^f),$$

which is always true by Assumption 2.1.

Cyberespionage decision

It is evident that the underdog will choose to perform a cyberespionage attack only if:

$$\hat{\Pi}_2^A > \hat{\Pi}_2^F,$$

or, alternatively, if:

$$A < F(h) - F(h(1 - \alpha)) + c \left[c'^{-1} [f(h)] \right] - c \left[c'^{-1} [f(h(1 - \alpha))] \right] = A^T. \tag{2.23}$$

The sign of the right-hand side (RHS) of the inequality above is always positive if Assumption 2.1 is satisfied. It implies that the underdog always chooses to perform an attack when it is sufficiently cheap. This finding is summarised in the proposition below.

Proposition 2.2. *Cyberespionage decision threshold, A^T , (2.23) is always positive if Assumption 2.1 is satisfied.*

Proof of Proposition 2.2. See Appendix A.1.1. ■

Note that Proposition 2.2 may not hold if Assumption 2.1 is relaxed. This case is worth investigating in future research as it might reveal scenarios in which the underdog chooses not to perform an attack even if it is free. Following the result of Proposition 2.1, cyberespionage may substantially increase the investments of the leader, which may decrease the expected profits of the underdog even further. In this case, cyberespionage may not be performed even if the underdog is paid for it.

From now on, we assume zero cyberattack costs. The assumption allows us to focus on the particular case of the R&D race with industrial espionage and neglects the less interesting case where cyberespionage never occurs in equilibrium (which is also formally equivalent to the standard rank-order tournament). The assumption is summarised below and imposed until the end of the chapter.

Assumption 2.2. *The cheating is costless, $A = 0$.*

2.2.2 Research stage

The study now has all the necessary ingredients to solve the research stage of the game. Since both companies are assumed to start their research simultaneously without prior knowledge, the stage is symmetric. It is said that the firm gains a technological edge h if its output is at least k higher than the opponent's:

$$q_i - q_{-i} > k,$$

or:

$$\Delta x + \zeta > k,$$

where $\Delta x = x_i - x_{-i}$ and $\zeta = \epsilon_i - \epsilon_{-i}$. There are three possible outcomes of the stage for each player:

(1) neither firm gains a technological edge and the game continues symmetrically with probability:

$$\begin{aligned} P^S &= P(-k < \Delta x + \zeta < k) \\ &= P(-\Delta x - k < \zeta < -\Delta x + k) \\ &= F(-\Delta x + k) - F(-\Delta x - k); \end{aligned}$$

(2) firm i gains a technological edge and becomes the leader of the race, while firm $-i$ becomes the underdog with probability:

$$\begin{aligned} P_1^A &= P(k < \Delta x + \zeta) \\ &= P(-\zeta < \Delta x - k) \\ &= F(\Delta x - k); \end{aligned}$$

(3) firm i loses the technological edge to firm $-i$ and continues the race as the underdog with probability:

$$\begin{aligned} P_2^A &= P(\Delta x + \zeta < -k) \\ &= P(\zeta < -\Delta x - k) \\ &= F(-\Delta x - k). \end{aligned}$$

The distribution of the new random variable ζ is symmetric around zero $E(\zeta) = 0$ and has a variance $E(\zeta^2) = \sigma^2$ as shocks are assumed to be i.i.d.

Combining outcome possibilities with equilibrium payoffs (2.12), (2.20) yields the following expected payoffs at the beginning of the initial stage:

$$\begin{aligned} \Pi_i &= P^S \Pi^S + P_1^A \Pi_1^A + P_2^A \Pi_2^A - c(x_i) \\ &= (F(-\Delta x + k) - F(-\Delta x - k)) \Pi^S + \\ &\quad + F(\Delta x - k) \Pi_1^A + F(-\Delta x - k) \Pi_2^A - c(x_i). \end{aligned} \tag{2.24}$$

The firms then choose the optimal investment levels x_i^* , which, assuming differentiability of the objective function, must satisfy two sets of conditions:

(1) FOCs:

$$\frac{\partial \Pi_i}{\partial x_i} = f(\Delta x + k) (\Pi^S - \Pi_2^A) + f(\Delta x - k) (\Pi_1^A - \Pi^S) - c'(x_i) = 0 \tag{2.25}$$

(2) and SOC:

$$\begin{aligned} \frac{\partial^2 \Pi_i}{\partial x_i^2} = & f'(\Delta x + k) (\Pi^S - \Pi_2^A) + \\ & + f'(\Delta x - k) (\Pi_1^A - \Pi^S) - c''(x_i) < 0. \end{aligned} \quad (2.26)$$

Thus, we demonstrate that the research stage of the game also has a unique symmetric equilibrium. The intuition remains the same: marginal incentives are equal for both firms. Proposition 2.3 summarises this claim.

Proposition 2.3. *The equilibrium of the research stage exists and is unique if Assumption 2.1 is satisfied. In the equilibrium the firms choose identical investments:*

$$x_i^* = c'^{-1} \left[f(k) (2F(h(1 - \alpha)) - 1) \right]. \quad (2.27)$$

Moreover, investments decrease in cyberespionage power $\frac{\partial x_i^*}{\partial \alpha} < 0$.

Proof of Proposition 2.3. See Appendix A.1.1. ■

The intuition behind the fact that the research stage investments decrease in cyberespionage power is that the firms are unwilling to invest in the research, a large share of which will be stolen. Therefore, cyberespionage diminishes competitive incentives in the initial stage. The combination of the results of Propositions 2.1 and 2.3 implies that cyberespionage has opposite effects on the investments in the research and development stages.

It is now possible to write down the optimal expected payoff at the beginning of the race. Substituting optimal effort levels (2.27) and payoffs of the latter stage into (2.26) and simplifying yields:

$$\begin{aligned} \Pi_i = & \frac{1}{2} + 2F(-k) \left(c \left[c'^{-1} [f(0)] \right] - c \left[c'^{-1} [f(h(1 - \alpha))] \right] \right) - \\ & - c \left[c'^{-1} [f(0)] \right] - c \left[c'^{-1} \left[f(k) (2F(h(1 - \alpha)) - 1) \right] \right]. \end{aligned} \quad (2.28)$$

Thus, both players initially have the same chance to win the race. However, it is unclear how cyberespionage attacks influence the firms' expected payoffs and investment levels. The following subsection investigates this matter.

Espionage externalities

The results suggest that cyberespionage power has opposite effects on the investment levels in the research and the development stages. A policymaker can optimise the overall R&D race investments by influencing cyberespionage power α . It can be regulated by legal instruments that impose obligations to businesses to implement cybersecurity measures to protect their sensitive information. For instance, the EU has two primary regulative documents that cover cybersecurity: General Data Protection Regulation (GDPR) and Network and Information Security Regulations (DCMS, 2018; Van Alsenoy, 2019). Both documents provide guidelines that aid companies in coping with common cyberattacks and ensuring the security of private and sensitive information. It is relatively hard to quantify the effects of the new cybersecurity law. Still, there is evidence that those laws are beneficial. Cisco (2019) reports that GDPR compliant companies were 15% less likely to have experienced a breach in the last 12 months, had 3 hours shorter average service downtimes, and lost significantly less money to breaches. Moreover, \$63 million in fines were issued in the first year after the release of GDPR, which was insignificant compared to overall losses due to cyberattacks. This supports Assumption 2.2 about negligible cyberespionage costs (Sobers, 2018).

Therefore, by changing the severity of the cybersecurity landscape, policymakers can regulate espionage power α and influence the competitive incentives in the R&D market. Well-designed cybersecurity law may then be used as a balancing device that erodes asymmetries and intensifies competition.

The research demonstrates that there exists a non-zero cyberespionage power $\alpha^I \in (0, 1)$ that maximises the overall expected investments if the technological edge h is too large. The result implies that “dosed” industrial espionage can promote competition by balancing out the effects of espionage on investments in the research and development stages. Proposition 2.4 summarises the results and Figure 2.3 illustrates them.

To allow for in-depth inference and simplify the notation, we assume the normality of shocks and impose a quadratic cost function.

Assumption 2.3. *Shocks are Normally Distributed: $\epsilon_i, \hat{\epsilon}_i \sim N(0, \frac{1}{2}\sigma^2)$ and cost function is quadratic $c(x) = \frac{x^2}{2}$.*

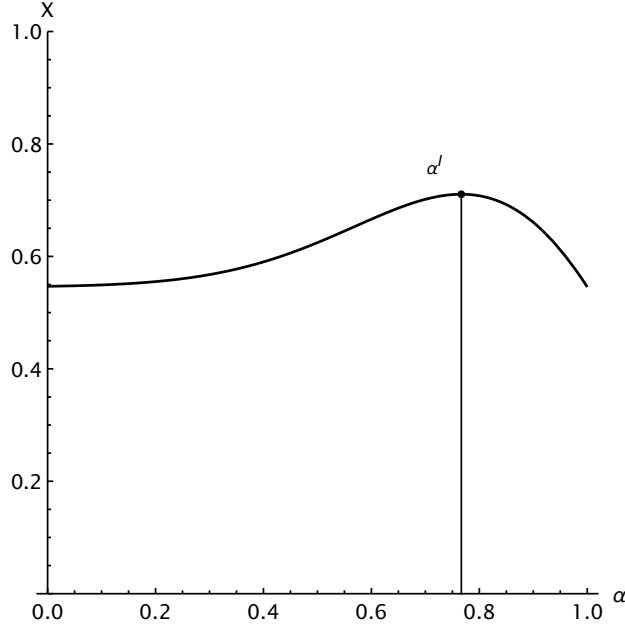


Figure 2.3: Influence of attack power on the overall investment.

Proposition 2.4. *There exists a technological edge $\bar{h}$ such that overall effort is maximised by*

$$\alpha^I \equiv 1 - \frac{1}{h} \sqrt{\frac{1}{2\pi}} \frac{\sigma}{\Phi(-k)} \exp\left(-\frac{k^2}{2\sigma^2}\right),$$

if $h > \bar{h}$ and $\alpha^I = 0$ otherwise. $\Phi(\cdot)$ is a CDF of normal distribution with zero mean and standard deviation σ .

Proof of Proposition 2.4. See Appendix A.1.1. ■

We now demonstrate that effort-maximising espionage power, α^I , (i) weakly increases in technological edge h and (ii) decreases in k . The first relationship should be straightforward as the effort-maximising espionage power α^I should increase in the distance between the actual h and $\bar{h}$. It happens because if the awarded technological edge is too large, it severely diminishes investment incentives in the development stage, as the underdog has little chance of winning. In this case, cyberespionage allows the underdog to level with the leader and continue

the race. It balances out the depressing effect of the large technological edge in the development stage. The second relationship is a consequence of the first. A larger technological step k means a higher “price” that companies need to pay to be awarded a technological edge h . It makes the companies more reluctant to invest in the research stage and reduces the weight of the technological edge, driving α^I down. Those results are summarised below in Proposition 2.5.

Proposition 2.5. *The effort-maximising espionage power α^I increases in technological edge, h , and decreases in technological step, k .*

Proof of Proposition 2.5. See Appendix A.1.1. ■

The influence of cyberespionage power on a company’s expected payoff in the race is not straightforward either. The research investigates the issue under three different espionage regimes: (i) ultimate espionage power $\alpha^* = 1$, (ii) interim power with $\alpha^* \in (0, 1)$, and (iii) absolute fairness with $\alpha = 0$.

The ultimate espionage regime completely diminishes the importance of the initial research stage. No firm would invest in the research stage, knowing that every single bit of the developed innovation will be stolen. The whole competition is then decided in the development stage of the game. We now demonstrate that such a high α can never maximise a firm’s payoffs. The result is formalised below.

Lemma 2.3. *Choosing the ultimate espionage power, $\alpha^* = 1$, is never optimal for payoff maximisation.*

Proof of Lemma 2.3. See Appendix A.1.1. ■

In order to simplify further derivations, let $H = (1 - \alpha)h$. Proposition 2.6 demonstrates that there always exists a unique optimal value of head start, H^* , that maximises the firms’ expected payoffs. If the technological edge is too small, $h \leq H^*$, the payoffs are maximised under the regime of absolute fairness as there is no point in diminishing it even further from its optimal value by choosing a positive

cyberespionage power. This finding echoes the known result from the competition theory that it is always optimal to reward one of the contestants with a small head start to intensify competition (Denter & Sisak, 2016). On the contrary, if the technological edge is sufficiently large, $h > H^*$, it negatively influences the firms' payoffs by diminishing the competition in the development stage. Then choosing an interim espionage power regime (with $\alpha^* \in (0, 1)$) can reduce the technological edge to the payoff-maximising value H^* . The results are illustrated by Figure 2.4.

It follows that policymakers can influence firms' decisions to enter the R&D race by imposing stimulating cybersecurity laws. This would be especially interesting to consider in an N-players setting, which will be considered in future research.

Observe also that cyberespionage might serve as a collusion device. It diminishes the investment levels in the initial research stage, but does not influence the prize the winner of the race receives. Therefore, under certain conditions, espionage can reduce the competitive incentives of firms while increasing their expected payoffs at the beginning of the game—imposing a self-policing cartel.

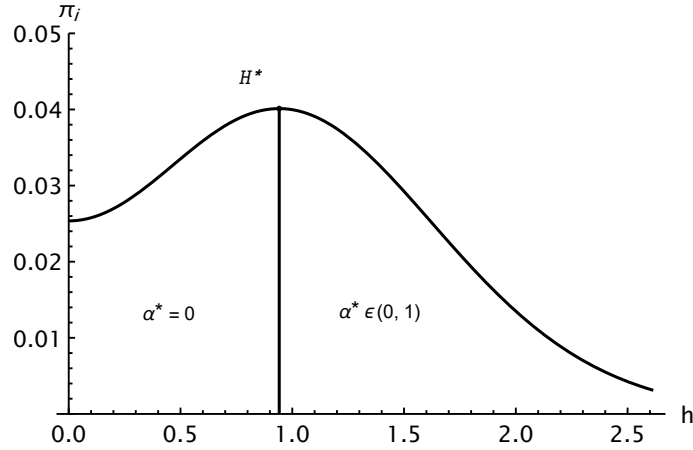


Figure 2.4: Influence of technological edge on firm's payoff.

Proposition 2.6. *Joint profit maximising espionage power is equal to 0, $\alpha^* = 0$, if technological edge is less or equal than optimal head start value, $h \in (0, H^*]$, and obtains its intermediate value $\alpha^* \in (0, 1)$ if $h \in (H^*, \infty)$. Moreover, α^* always exists regardless of h and k and is a global maximum.*

Proof of Proposition 2.6. See Appendix A.1.1. ■

Depending on the size of payoff maximising H^* , k can influence α^* either positively or negatively. For instance, if H^* is sufficiently small, then optimal espionage power increases in the technological step, $\frac{\partial \alpha^*}{\partial k} > 0$. The phenomenon occurs because the extra investments needed to gain a technological edge surpass the added benefits, making the symmetric case more appealing to the competitors.

If the optimal technological edge is already sufficiently large, the relationship becomes complicated. However, it still follows a logical pattern: the α^* increases in distance between h and its optimal value H^* and decreases otherwise. Proposition 2.7 formalises the findings and provides more details.

Proposition 2.7. *The technological step k has an ambiguous influence on payoff maximising cyberespionage power α^* . It does not influence the optimal attack level if $h < H^*$. If $h > H^*$ then the sign depends on the values of k and H^* :*

1. *if H^* is sufficiently small and $q(H^*) \in \left(0, \frac{\sqrt{2}}{\sqrt{e\pi}\sigma^2}\right)$ then $\frac{\partial \alpha^*}{\partial k} > 0$;*
2. *if H^* is sufficiently large and $q(H^*) \in \left(0, \frac{\sqrt{2}}{\sqrt{e\pi}\sigma^2}\right)$ then (1) $\frac{\partial \alpha^*}{\partial k} < 0$ if $k \in (k_1^*, k_1^*)$, (2) $\frac{\partial \alpha^*}{\partial k} > 0$ for any $k \in (0, k_1^*) \cup (k_2^*, \infty)$,*

where $q(H) = \frac{f'(H)}{1-2F(H)}$.

Proof of Proposition 2.7. See Appendix A.1.1. ■

There may also exist an optimal value of the technological step k^* , which maximises the firm's expected payoff. The presented framework can be adopted to study cyberespionage in relation to the patent regimes, with h being interpreted as the patent length and k as the patent breadth (or patentability threshold)—the dimensions of a patent identified by Scotchmer (2004). We will pursue this line of logic in future research.

2.3 Social welfare

The research adopts the expected quality (output) function to measure the social benefits of the race. We assume that society benefits from the quality of innovation delivered to the market by the winner of the R&D race and loses the cumulative costs that both companies incurred in the R&D race (which could be spent on something else otherwise).

While the model measures the performance of the companies in relative terms, absolute investments and quality of the end-product may be crucial from a social perspective. For instance, it is arguably of minor importance for a car manufacturer to reduce the gasoline its vehicles consume, as far as the car of its competitor uses more. However, any additional economised litre of gasoline matters for the environment and influences social welfare.

The social welfare function (SWF) is then expressed as expected output of the winner of the race accounting for all exogenous shocks that might happen during the R&D race subtracting expected overall costs of both competitors:

$$W = X_i + E^S + E^L + E^U - C, \quad (2.29)$$

where X_i is the overall effort exerted by the winner of the race, E^S is a contribution of symmetric case shocks to the output of the game, E^L is a contribution of asymmetric case shocks when a leader wins the race, E^U is a contribution of asymmetric case shocks when an underdog wins the race, and C is the overall expected costs of the R&D race. See Appendix A.1.2 for a complete SWF derivation.

Note that technological edge h influences the SWF (2.29) only indirectly through the overall expected effort exerted by the winner of the race. It is not, however, present in the function as a separate term. The technological edge is a measure of the relative advantage of one firm over another. Given that society is interested in the absolute quality of innovation, the inclusion of the technological edge into the function is unnecessary.

2.3.1 Espionage and welfare

Now we focus on the relationship between the end-product quality and cyberespionage power. The function (2.29) then can be simplified by excluding terms that are not functions of α .

Let,

$$\begin{aligned} g(k) &= \int_{-\infty}^{\infty} \int_{\epsilon_{-i}+k}^{\infty} \epsilon_i f(\epsilon_i) f(\epsilon_{-i}) d\epsilon_i d\epsilon_{-i} \\ &= - \int_{-\infty}^{\infty} \int_{-\infty}^{\epsilon_i-k} \epsilon_{-i} f(\epsilon_{-i}) f(\epsilon_i) d\epsilon_{-i} d\epsilon_i > 0, \end{aligned}$$

and

$$\begin{aligned} i(H) &= \int_{-\infty}^{\infty} \int_{\hat{\epsilon}_2-H}^{\infty} \hat{\epsilon}_1 f(\hat{\epsilon}_1) f(\hat{\epsilon}_2) d\hat{\epsilon}_1 d\hat{\epsilon}_2 \\ &= \int_{-\infty}^{\infty} \int_{\hat{\epsilon}_1+H}^{\infty} \hat{\epsilon}_2 f(\hat{\epsilon}_2) f(\hat{\epsilon}_1) d\hat{\epsilon}_2 d\hat{\epsilon}_1 > 0. \end{aligned}$$

It is then sufficient to study the properties of the following function:

$$\begin{aligned} \psi &= F(H) \left(2g(k) + f(k) + F(-H) f(k)^2 \right) + \\ &\quad + F(-k) \left(2i(H) + f(H) - f(H)^2 \right). \end{aligned} \tag{2.30}$$

The study now demonstrates that the welfare-maximising value of cyberespionage power, α^W , depends on the size of the technological edge: $\alpha^W \in (0, 1)$ if the technological edge is sufficiently large and $\alpha^W = 0$ if it is sufficiently small. The finding is summarised in a subsequent proposition.

Proposition 2.8. *Social welfare maximising espionage power is equal to $\alpha^W = 0$, if technological edge is less or equal to the socially optimal head start value, $h \in (0, H^W]$, and attains its intermediate value $\alpha^W \in (0, 1)$ if $h \in (H^W, \infty)$.*

Proof of Proposition 2.8. See Appendix A.1.2. ■

The result of Proposition 2.8 echoes the results about effort and payoff maximisation derived in the previous section. The regime of absolute fairness ($\alpha = 0$) can benefit social welfare only if the technological edge is sufficiently low $h \leq H^W$. However, it achieves the opposite result if the technological edge is sufficiently large. In this case, choosing an interminable espionage power regime ($\alpha^W \in (0, 1)$)

leads to higher social welfare. As with payoff maximisation, choosing the regime of absolute espionage power ($\alpha = 1$) is never optimal as it completely deteriorates the investments in the initial research stage leading to low-quality end-user products. The findings align with the empirical results described by Busso and Galiani (2014).

Of course, increasing the power of cyberespionage in R&D competition if the technological edge is too big might not be an option for policymakers. Still, if we interpret a technological edge (h) as a patent length—the dimension that policymakers can control—the model yields some insights into patent law regulations. By choosing an optimal patent regime and IP laws (e.g. trade secret laws) for a specific R&D industry, policymakers might maximise investment levels, firms' payoffs, and, as a result, the quality of the end-products while making cyberespionage worthless. Those ideas will be investigated in more detail in future research.

2.4 Conclusion and discussion

In this chapter, we explored the influence of industrial cyberespionage on the R&D incentives, the R&D race outcome, and social welfare. The study finds that cyberespionage power has an ambiguous effect on the R&D investments in the dynamic race. Espionage diminishes the competitive incentives in the initial research stage while increasing them in the development stage. Companies will be reluctant to invest in the research if their technology will most likely be stolen. On the contrary, espionage can serve as a balancing device in the development stage: by decreasing the gap between the leader and the underdog, providing the latter with a fighting chance. Therefore, by regulating the cyberespionage power via legal means, policymakers can maximise the overall R&D investments in the race.

Additionally, cyberespionage can positively influence competitors' payoffs. By choosing an appropriate espionage regime, policymakers can make the R&D race more attractive to the participants and motivate competition. Moreover, cyberespionage can lead to higher quality innovation and produce social benefits. The social welfare function is defined as the absolute measure of the quality of the innovation delivered to the market by a winner of the race. Thus, imposing an

appropriate espionage regime can sufficiently increase overall competitive incentives, positively influencing the expected quality of the innovation.

It is rather challenging to retrieve empirical evidence favouring the proposed theory due to the secrecy of the industrial espionage phenomenon. It would require a detailed data set that includes variables on individual companies' R&D efforts (e.g. R&D expenses, research human capital, base stock of knowledge), R&D profit breakdowns, and the information on industrial espionage decisions, which is probably the most difficult to obtain.

However, it might be possible to study the dynamics of social welfare function expressed in economic growth terms (e.g. TFP or GDP in per capita terms) and research it against indices of industrial espionage. This would also require a data set containing information on the flows of the stolen sensitive data, which is not currently available.

The study also identified three possible vectors for future research. Firstly, it is worth exploring the link between the R&D process, innovation economics, and optimal patent design. The technological step k might be alternatively interpreted as a patentability threshold, and the technological edge h can be regarded as a function of patent length. The framework then can be easily modified to study the optimal patent design in the aggressive digital environments (Scotchmer, 2004).

Secondly, the framework can be extended by allowing firms to diversify investment into various technologies during the research stage. The product design is rarely built on a single technology and often incorporates multiple innovations that synergise together in a unique design. Thus, the expansion can lead to insights about optimal research budget allocation and strategic data organisation for enhanced cybersecurity.

Thirdly, the model can be extended with a more sophisticated cyberattack mechanism. Instead of using exogenously defined attack power, one could use a more sophisticated endogenously defined function that might capture the offensive and defensive efforts of participants or affiliation mechanisms, and social planners'

investments into industrial espionage enforcement, in the fashion of Stowe and Gilpatric (2010).

3

No Research, No Risk?

Contents

3.1	Introduction	61
3.1.1	Related literature	65
3.2	Hypotheses	74
3.3	Empirical strategy	75
3.3.1	Data	75
3.3.2	Methodology	83
3.3.3	Methodology discussion	90
3.4	Results	94
3.4.1	Fractional regressions results	94
3.4.2	Marginal effects analysis	96
3.4.3	Targeted and opportunistic attacks	97
3.4.4	Alternative explanations and robustness checks	103
3.5	Conclusion	106

3.1 Introduction

This chapter investigates the association between the research and development (R&D) intensity in an industry and the rate of information leakage cyberattacks. The relationship is studied under several statistical settings. The chapter's primary goal is to answer a simple question: are research-intensive industries more likely to become victims of information leakage attacks?

Cybersecurity has become one of the major global threats. Thousands of businesses and millions of people are becoming victims of cybercrime every year. The cost estimate of cybercrime to the global economy surpassed \$1 trillion in 2018 and almost doubled to \$2 trillion in 2020, making it roughly 2% of global GDP. With a predicted annual growth rate of 15%, this number is expected to reach \$4–5 trillion by 2025 (Juniper Research, 2019). Recent reports suggest that more than two-thirds of global cybercrime costs were associated with data breaches: intellectual property theft, cyberespionage and financially motivated crimes (McAfee, 2020).

Around 32% of companies experienced data breaches in the United Kingdom in 2019 with total economic losses estimated around £27 billion, £16.7 billion of which was associated with industrial espionage and intellectual property theft (IBM Security, 2020). The average cost of a data breach incident in the UK reached £3.5 million (DCMS, 2022).

Data breaches are generally divided into two big groups: targeted and opportunistic (or untargeted) (National Cyber Security Centre, 2015). The former is closely associated with intellectual property theft and industrial espionage, while the latter is often associated with more financially-driven attacks that take advantage of unidentified victims' vulnerabilities to obtain easily monetised data. However, it should not be treated as a rule: some targeted attacks could be aimed at consumer records, while some opportunistic attacks might extract trade secrets from the victims.

This study focuses on targeted attacks and theorises that they might be adopted as means of unfair competition. An underdog in an R&D race can choose to obtain competitors' know-how or corporate secrets utilising a cyberattack. Similarly, client databases can be targeted by cyberattacks in sales driven competition. According to Oxford Economics (2014), in 2013 every fifth company in the UK experienced a data breach, with 61% and 59% of companies reporting the loss of competitive advantage due to the loss of intellectual property (IP) and commercially sensitive data, respectively. Stealing know-how or client records could be a cheap, efficient

and relatively safe method of gaining a competitive edge in an R&D race (Crete-Nishihata et al., 2018).

Our study employs a tailored data set based on aggregated survey data from Eurostat to investigate this claim. The data contain cyberattack rates and plausible proxies for research-intensity and customer data use. -

The research estimates fixed effects fractional outcome regression models with the country and industry controls. This estimation technique is superior to alternative methods (e.g. ordinary least squares (OLS) or Tobit models), as the chosen dependent variable is the information leakage attack rate in a given country in a given industry $\in [0, 1]$. However, the coefficients estimated by the fractional outcome model are not directly interpretable. The study utilises marginal analysis to estimate average marginal effects, which can be interpreted similarly to classic OLS coefficients.

The dependent variable is estimated separately against three research-intensity proxy variables that capture distinct groups of R&D personnel: (1) a share of all R&D personnel in overall employment in a given country in a given industry, (2) a share of researchers in overall employment in a given country in a given industry, and (3) a share of research support personnel in a given country in a given industry. The choice of the research-intensity proxies follows conventional economic growth theory, where research personnel is often employed as an essential part of the ideas output function (e.g. Porter & Stern, 2000).

Due to the secrecy of information leakage attacks, the available data possess two major challenges. Firstly, the data set has a considerable share of partial observations, leading to biased estimates if not treated carefully. Therefore, fractional outcome regression analysis is backed up by the multiple imputation (MI) technique, a hybrid Bayesian-frequentists approach that allows for unbiased estimates in the presence of partial observations. Secondly, the dependent variable aggregates both targeted and untargeted attacks.

Verizon's Data Breach Investigation Report (2011) suggests that targeted attacks are generally aimed at intellectual property and trade secrets, which are the results of research activity (Verizon Business, 2011). On the contrary,

untargeted (opportunistic) attacks are responsible for 90% of compromised records, customer financial data, and credentials. Thus, the relationship between the rate of information leakage attacks and research-intensity cannot be studied without controlling for opportunistic attacks. To tackle the issue, we use an additional Client Resource Management (CRM) proxy variable representing the end-user data reliance in a given country in a given industry to absorb the variance associated with opportunistic attacks. The additional variable provides an opportunity to consider another (secondary) research question: are data-reliant industries susceptible to information leakage attacks?

The theory is tested with univariate and multivariate fractional outcome models estimated on Eurostat data. Univariate models demonstrate a strong association between information leakage attack rate and research-intensity regardless of a chosen proxy. The share of research support personnel demonstrates a more than fourfold larger marginal effect than the other two proxies. Moreover, it is the only proxy that maintains a statistically significant coefficient after controlling for opportunistic attacks. Therefore, the share of R&D support personnel might be a better proxy to capture the value of innovative ideas. A discussion on this matter can be found in Subsection 3.3.1.

The multivariate estimation reveals that the CRM proxy variable has a statistically significant coefficient. By employing the multivariate model specified with both research-intensity and end-user data reliance proxies, the study distinguished two positive associations potentially linked to targeted and opportunistic attacks.

To back up the findings, the study additionally performs multivariate estimations for six subsamples of different technological environments: all manufacturing industries, high-tech manufacturing industries, low-tech manufacturing industries, all services, knowledge-intensive services, and less knowledge-intensive services. Consistent with the proposed causal mechanism, information leakage attacks in high-tech manufacturing industries are strongly associated with the research-intensity proxy while demonstrating no association with the data reliance proxy. On the contrary, the information leakage attacks rate in knowledge-intensive service

industries does not demonstrate any association with the research-intensity proxy while showing a strong association with the data reliance proxy.

Therefore, the results suggest that research-oriented manufacturing industries are prone to fall victim to targeted information leakage attacks (potentially industrial espionage). Meanwhile, knowledge-intensive service industries are more likely to experience opportunistic financially-driven attacks. However, this can be considered solely as supporting evidence favouring the proposed causal mechanism. Establishing an actual causality will require a better and more complete data set, which is currently unavailable.

The rest of the chapter is built as follows. The next section briefly discusses related literature. Section 3.2 summarises main the hypotheses. Detailed information on the chosen methodology and data can be found in Section 3.3. Section 3.4 reports the results and Section 3.5 offers conclusions and discussion.

3.1.1 Related literature

This chapter contributes to a rich strand of the economics of cybersecurity literature focused on data breach risk quantification. A majority of work is conducted empirically and stems from managerial and actuarial sciences (see the comprehensive survey in Schlackl et al., 2022). The research in this area focuses on discovering organisational attributes and characteristics that could serve as predictors for cyber risk models. Those models can be used for insurance premium pricing, supporting better decision-making within the organisation, or guiding the development of cybersecurity policies. Here we offer a review of both actuarial and managerial literature, grouping the papers by the primary data breach factors discussed.

Cyber risk management

Cyber risk management studies generally attempt to identify antecedents of data breaches, which the company can control. Identifying those factors aids the company in the development of better managerial practices and mitigating data breach risks.

Software characteristics Ransbotham (2010) and Moore and Clayton (2011) analysed how the characteristics of software employed by an enterprise are linked with data breaches.

Ransbotham (2010) compared the durability of open- and closed-source software. Open-source code is developed and tested by many programmers and software testers, which, in theory, should lead to the discovery of more bugs, glitches, and vulnerabilities compared to closed-source software. However, by analysing two years of data on intrusion detection, Ransbotham found that open-source software was targeted more frequently, and exploits were being developed faster.

Moore and Clayton (2011) researched how cybercriminals choose their targets. The authors recognised that foes often use relatively trivial techniques to exploit their target. For instance, they utilise publicly available search engines to identify possible victims—websites with common (and usually obvious) security holes. The authors claimed that the simplicity of compromising the website is one of the main drivers of victim selection. We believe that those findings can be used to explain the nature of opportunistic financially-driven attacks.

Therefore, several studies identified a link between data breach risk and software employed by the company. We contribute to this line of research by demonstrating that opportunistic attacks described by Ransbotham (2010) and Moore and Clayton (2011) are industry-specific and generally happen more often in service industries.

Board composition and characteristics A significant body of literature represented by T. Wang and Hsu (2010), Kwon et al. (2013), Hsu and Wang (2021), and Haislip et al. (2021) analysed the link between a company's board composition and characteristics with exposure to data breach risks.

T. Wang and Hsu (2010) investigated the association between a company's board structure and a data breach occurrence. They demonstrated that board size and the number of independent directors in a company have a strong positive correlation with the likelihood of a company being breached. Additionally, it was demonstrated that a board's diversity (e.g. age-wise) has a strong negative

correlation with the likelihood of information security breaches. Kwon et al. (2013) continued this line of work and examined the association between the presence of an IT executive on a board and the likelihood of information security breaches. The authors demonstrated that companies with executive IT officers on the board are less prone to experience data breaches. The compensation of such a director has also demonstrated a strong negative correlation with cyber risk exposure.

In a more recent study, the influence of a company's board characteristics on the likelihood of a data breach was investigated by Hsu and Wang (2021). The authors empirically tested the hypothesis that companies in which board members are "too busy" (i.e. holding multiple positions within the board) are more likely to fall victim to data breaches. The phenomenon is explained by the inability of a "busy" board member to commit to IT security duties, which results in a company's increased vulnerability.

Haislip et al. (2021) analysed the influence of a board's level of IT expertise on data breaches. They have demonstrated that companies which have a CEO with IT expertise are less likely to experience a data breach. Similarly to Kwon et al. (2013), the paper has also demonstrated that the presence of a Chief Information Officer on the board is associated with a decrease in the likelihood of a data breach.

Thus, several studies have demonstrated a strong link between exposure to cyber risks and the structure of a board. We believe this is only the part of the picture, and the composition and characteristics of the body of regular employees can also be an influential factor in modelling cyber risks. We contribute to the literature by testing the hypothesis that the share of research personnel in a given industry is positively correlated with the rate of information leakage attacks in the industry.

Attack surface factors Angst et al. (2017), McLeod and Dolezel (2018), and Kim and Kwon (2019) analysed the link between attack surface (e.g. size of the workforce and the presence of advanced electronic managerial systems) and exposure to data breach risks, utilising the data on information leakage attacks in the US healthcare industry.

Angst et al. (2017) have demonstrated that the likelihood of data breach has a strong positive association with a hospital's workforce size. Moreover, it was demonstrated that the deeper integration of security into IT-related routines and processes decreases the risks of being hacked. Similar results were obtained by McLeod and Dolezel (2018). The authors attempted to model the occurrence of consumer data breaches against two sets of explanatory variables: the size of medical institution proxies and proxies for computerisation of the medical personnel (e.g. the share of physicians that utilise Electronic Medical Records (EMR) software in a given hospital). The study concluded that both the size of a medical institution and the share of personnel with EMR access have a strong positive association with data breach occurrence.

Kim and Kwon (2019) studied the effects of meaningful use initiative implementation on data breaches. The program sets governmental standards for EMR system implementation in US hospitals. The study found that hospitals which joined the initiative and implemented EMR software are three times more likely to experience a data breach than hospitals which did not switch to digital patient management systems. It was also demonstrated that larger hospitals are more likely to be breached.

We contribute to this line of research by investigating the link between electronic client resource management (CRM) system usage rate and information leakage attack rate in 19 European manufacturing and service industries. Moreover, we demonstrate that CRM usage rate has a strong positive association with information leakage attack rate only in service industries, while not showing any statistically significant correlation in manufacturing industries.

Institutional factors A significant body of research has been devoted to investigating the link between the characteristics of the environment in which an organisation operates (e.g. industry and country) and data breaches. Sen and Borle (2015) empirically tested the opportunistic theory of crime and the institutional

anomie theory¹ by estimating the risk of data breach incidents within the industry and state. It was demonstrated that industry plays an essential role in determining the risk of a data breach. Similarly to our findings, they found that industries which are more reliant on customer (or patient) data are more prone to data breaches. Supporting the opportunistic theory of crime, they have also found that the data breach rate is state-specific. The states with stricter data breach disclosure rules demonstrated higher resilience to data breaches.

Another empirical paper was presented by Q. H. Wang and Kim (2009). The authors used time-series regression analysis to study cyberattacks on a country level. For the most part, the authors studied the spatial autocorrelation of cyberattacks and the effects of joining IT conventions. They found evidence of autocorrelated cyberattacks and concluded that physical boundaries must be considered an essential factor in designing cybersecurity regulations.

Similar conclusions could be found in Romanosky (2016), who analysed 12,000 cyberincidents (including data breaches, phishing, and privacy violations) in 20 different US industries. They echoed the result that the cyberincident rate is industry-specific. Moreover, they confirmed that more data-reliant industries are particularly susceptible to cyberincidents. However, the author claimed that the average damage from a cyberincident is minor, constituting on average only 0.4% of annual company revenues—roughly the same as an average share of a typical firm’s annual IT security budget. Thus, they concluded that public concerns about cyberincidents are excessive.

We use these findings to inform our choice of control variables (e.g. we utilise country and industry control variables). Additionally, we contribute to this literature by distinguishing between two associations potentially linked to different types of information leakage attacks. Thus, we confirm that more data-reliant service industries are more likely to experience data breaches. However, those data breaches are more likely to be caused by opportunistic financially-driven attacks. On the

¹Institutional anomie theory hypothesises that the crime is motivated by an exaggerated emphasis on economic success within society.

contrary, we demonstrate that research-intensive manufacturing industries are more likely to experience sophisticated targeted attacks.

Structural factors Tanriverdi et al. (2019) and Say and Vasudeva (2020) studied the association between data breach risks and a company’s structural complexity. Both studies found that recent subsidiary divestiture has a strong negative association with the likelihood of a data breach. Additionally, Say and Vasudeva (2020) demonstrated that mergers and acquisitions are, on the contrary, positively correlated with the risk of a data breach. Both papers hypothesise that data breach risk is a function of a company’s complexity—by selling its subsidiary, the company’s complexity decreases, decreasing the risk, while mergers and acquisitions have the opposite effect.

C.-W. Liu et al. (2020) analysed the link between the structural characteristics of organisations’ IT systems and cyber risks. Utilising data on security breaches in higher education institutions in the US, they demonstrated that educational organisations with a higher degree of centralisation are less likely to experience a data breach. Moreover, the authors recognised (with the help of correlation tables) that R&D activities carried out by a faculty are a particularly attractive target for information leakage attacks. It follows that public US universities with a larger amount of funding are more likely to experience information leakage attacks. However, this finding was utilised solely to construct an additional control variable.

Instead of focusing on structural characteristics of organisations as main predictors of data breach risks, we primarily focus on research-intensity. We contribute to this literature by extending the results of C.-W. Liu et al. (2020) to the private sector and, particularly, high-tech manufacturing industries.

Cybersecurity investments An ongoing debate in empirical cybersecurity literature revolves around the influence of cybersecurity investments on cyber risks. The results found in the literature are generally data specific and often contradict each other. The issue is caused by the prevalence of low-quality data in

the field, which, in turn, does not allow scholars to determine causal relationships between variables.

For instance, W. Liu et al. (2007) used a regression analysis to investigate the relation between computer virus incidents and cybersecurity measures in Japanese companies (e.g. information security policy, human cultivation). They found evidence suggesting that continuous investment in cybersecurity can significantly mitigate cyber risks related to data breaches. Similar results were also obtained by McLeod and Dolezel (2018), who demonstrated that cybersecurity investments in medical institutions reduce the exposure to information leakage attacks. The opposite results were found by Sen and Borle (2015) and Angst et al. (2017) who demonstrated that higher cybersecurity investment in an industry is associated with a higher rate of data breaches.

The first causal inference empirical model was achieved only recently by Gandal et al. (2022). The authors have investigated cyberincidents in Israeli companies and demonstrated a strong negative influence of investment in four basic security precautions (strong password policy, strict patching policy, backup systems, and monitoring systems) on the likelihood of experiencing a data breach. To establish causality, the authors utilised instrumental variables estimation. The study instrumented cybersecurity precaution proxies with an indicator of whether the company is alert about cybersecurity risks (having experienced them in past periods) and intentionally vigilant towards cyber risk exposure.

Cybersecurity investments are out of the scope of our study. Nevertheless, the evidence described above could inform organisational risk management procedures, mitigating data breach risk in research-intensive and data-reliant companies.

Actuarial studies

A significant body of research is devoted to quantifying cyber risks in a race for better cybersecurity insurance models. Such quantification is a fundamental actuarial challenge that has been addressed by scholars, insurers, and businesses (see the comprehensive survey in Bardopoulos, 2020). Even though we do not

consider cyberinsurance in the thesis, multiple actuarial science papers are devoted to identifying insurance risks, rating, and exposure measures—variables that are used to model cyber risks of the insured. Here we briefly review actuarial literature, grouping articles by the main variable of interest studied in the paper.

Customer base An early attempt to quantify cyber risk was made by Hoo (2000). The author employed utility theory to analyse companies’ decision trees based on computer security surveys. The study identified that the main factors behind cyber risks are the company’s wealth indices and probabilities of the security events associated with decision trees. It proposed to use the customer base as a main measure of exposure. However, the main factors were defined following the anecdotal evidence from a small sample survey analysis.

Quantity of ICT devices, platforms, and correlated risks Barracchini and Addressi (2014), Böhme (2005), Böhme and Kataria (2006), and H. Herath and Herath (2007) presented analytical side cyberinsurance models that relied on the number of infected computers and their interconnectedness as the main insurance exposure measures. Böhme (2005) argued that the interconnected nature of industrial IT infrastructure might intervene in developing a proper cyberinsurance market. The author recognised that cyber risks of organisations are correlated due to the dominance of a few IT platforms, which, in turn, lead to correlated losses. Böhme and Kataria (2006) expanded on those observations by investigating the correlation of cyber risks within and across firms based on “honeypot” data (Leita et al., 2008). Thus, Böhme and Kataria (2006) echoed the results of Ransbotham (2010) and Moore and Clayton (2011) that the type of software used by the company might be strongly correlated with cyber risk exposure. Barracchini and Addressi (2014) in their work took a similar route and offered an alternative epidemiological model of cyberinsurance premium calculation, which took the number of terminals (computers) of organisations and its network structure as main inputs.

H. Herath and Herath (2007) presented a premium pricing framework for assessing the potential financial losses based on the observed distribution of the

number of compromised computers. The empirical work was based on survey data which was analysed employing a copula. Data defects were identified as one of the main limitations of the approach (the data had only 15 observations).

Other factors Innerhofer-Oberperfler and Breu (2010) conducted a qualitative study to identify the main risk indicators for cyberinsurance models. They interviewed 36 field experts and reduced their statements to the set of 94 indicators ranked according to their relative importance (subjectively chosen by experts during the post-interview questionnaire). In their study, “processing sensitive data with high confidentiality requirements” and “existence of worth-protecting know-how, patents and otherwise valuable information” are ranked third and fourth, respectively—right after the indices of dependency on IT systems. The empirical model presented in this chapter attempts to quantify those qualitative findings.

However, a significant share of cyberinsurance research happens behind closed doors of insurance companies rather than in academia. Publicly available pricing policies of large-scale cyberinsurance companies were analysed by Romanosky et al. (2017) and Bardopoulos (2020).

Romanosky et al. (2017) conducted a systematic qualitative review of policies of large-scale insurance companies. The authors identified four main categories of risk and rating factors: (1) organisational factors, which include customer data collection handling, IT security budget, and previous incident history; (2) technical factors, including access controls and overall IT security level indices; (3) cybersecurity policies and procedures, including IT security and privacy policy maturity and cybersecurity awareness of personnel; (4) legal and compliance factors, which include the level of organisational compliance to governmental cybersecurity standards and laws. In a similar study, Bardopoulos (2020) analysed the principal exposure measures affiliated with data breach risks. The author further revealed that companies often utilise internet revenue, the size of the workforce, and the number of confidential customer records for cyberinsurance modelling.

Thus, while insurance companies identify the importance of organisations' reliance on customer data as a significant factor for cyberinsurance modelling, they are generally agnostic of the company's research-intensity indicators. Our work provides evidence that this factor might also be influential in modelling cyber risks.

To conclude, risk management and actuarial literature about cyber risk quantification revealed a number of factors that can be used for data breach risk modelling. It identified that an organisation's attack surface, reliance on customer data, and institutional environment factors play an essential role in those models. While it was recognised by Innerhofer-Oberperfler and Breu (2010) and C.-W. Liu et al. (2020) that research-intensity might also play a significant role in the mixture, the literature provides little empirical evidence around it. Following our observations from the analytical R&D race model with espionage presented in Chapter 2, we investigate the link between the research-intensity of a given industry and the information leakage attack rate in this industry. We present new empirical evidence that this association is industry-specific. Our study demonstrates that high-tech manufacturing industries are particularly susceptible to targeted information leakage attacks, while knowledge-intensive service industries are more likely to fall victim to opportunistic financially-driven attacks. Those findings can be used to develop better risk management procedures and routines or inform more sophisticated cyberinsurance premium pricing models.

3.2 Hypotheses

We construct two hypotheses to test for the possible associations of the rate of information leakage cyberattacks² with (1) research-intensity and (2) industries' reliance on the end-user data. The first hypothesis concerns the association of research-intensity with the rate of information leakage attacks. The association is believed to be linked to targeted attacks that are predominantly aimed at IP and corporate secrets.

²Defined as the rate of cyberattacks that result in the theft of information.

Hypothesis 1. *Higher research-intensity in the industry (independent variable) is associated with an increase in the information leakage cyberattack rate in the industry (dependent variable).*

Simultaneously the study examines the association between information leakage attacks and end-user data reliance.

Hypothesis 2. *Higher reliance on end-user data in the industry (independent variable) is associated with an increase in the information leakage cyberattack rate in the industry (dependent variable).*

3.3 Empirical strategy

3.3.1 Data

The research uses a tailored data set that is based on three publicly available sets from Eurostat:

1. Structural business statistics (SBS) data set that covers the structure, conduct and performance of the industry, construction, services and trade (Eurostat, 2010c);
2. Science, technology and innovation data set which covers statistics in the fields of science, technology and innovation (Eurostat, 2010b);
3. Digital economy and society data set which covers all aspects of the usage of Information and Communication Technologies (ICT) by individual households and businesses (Eurostat, 2010a).

All selected data sets are survey-based, and aggregate responses from 148,000 enterprises categorised by year, country, and economic activity according to the second revision of NACE³ (European Commission, 2008). However, data sets have a different granularity of aggregation. For instance, the SBS data set has detailed information on every economic activity present in the NACE classification, while

³NACE stands for “Nomenclature Statistique des Activités économiques dans la Communauté Européenne”. It is a standard European system for classifying firms by the industrial sector.

Table 3.1: List of countries included in the study

1 Belgium	9 Spain	17 Hungary	25 Slovakia
2 Bulgaria	10 France	18 Malta	26 Finland
3 Czechia	11 Croatia	19 Netherlands	27 Sweden
4 Denmark	12 Italy	20 Turkey	28 UK
5 Germany	13 Cyprus	21 Poland	29 Iceland
6 Estonia	14 Latvia	22 Portugal	30 Norway
7 Ireland	15 Lithuania	23 Romania	31 Macedonia
8 Greece	16 Luxembourg	24 Slovenia	32 Serbia

the two others utilise industry aggregates that group several similar economic activities. Therefore, to make merging possible, data sets were aggregated to the highest common level, that of the “Digital economy and society” data set. Countries that are included in the research are summarised in Table 3.1. Figure 3.2 illustrates the geographical distribution of the countries included in the study and the corresponding average shares of companies that experienced cyberincidents in each country. A list of industry aggregates can be found in Appendix B.1. Country and industry dummies constitute two sets of control variables.

Note that even the finalised data set includes only variables for 2010 due to the fact that “Security Incidents and Consequences” surveys included in “Digital economy and society data set” was conducted only for 2010 and 2019; thus, 2010 was the only year where all three data sets overlap. The partial observations constitute 36% of data, which we address by utilising a multiple imputation approach described in Subsection 3.3.2.

Dependent variable

The dependent variable—the *rate of information leakage attacks* (B_{ij})—represents the share of the enterprises which experienced cybersecurity incidents that resulted in the disclosure of confidential data due to intrusion, pharming or phishing attacks in a given country (i) in a given industry (j).

The variable has a right-skewed distribution, which suggests the usage of the logarithm to normalise it (see Figure 3.1). However, zero observations constitute

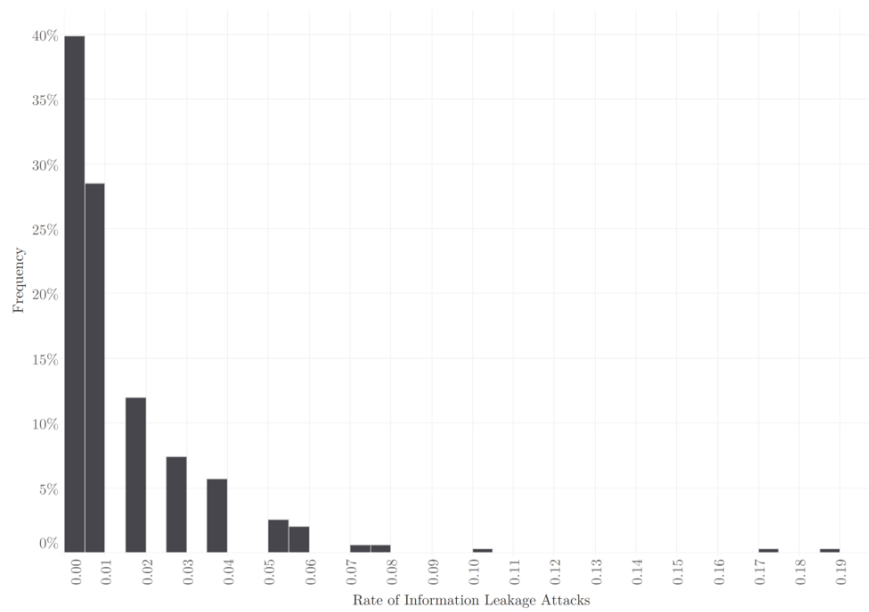


Figure 3.1: Dependent variable histogram.

39.9% of observations, which means that log-transformation might significantly affect the results.

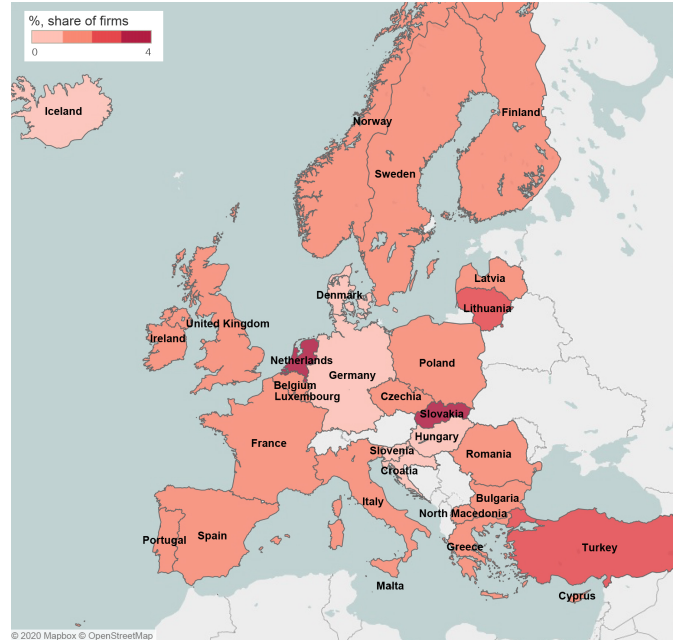


Figure 3.2: Geographical distribution of the countries included in the study. The countries are colour-coded according to the average share of enterprises that experienced information leakage attacks.

There is no reasonable way of determining the stolen type of data (e.g. personal records, financial information, blueprints). The variable aggregates all the possible

cases regardless of the original motive of the breach. To shed some light on the issue, the research refers to 2011 Data Breach Investigation Report (based on the Veris Community Database (VCDB)), which analyses the majority of known data compromise incidents in 2010 (Verizon Business, 2011).

The report divided cyberincidents into two big groups: (1) targeted, and (2) opportunistic attacks. Around 20% of all discovered breaches can be affiliated with targeted attacks, which, however, correspond to less than 10% of all compromised records (e.g. credentials, card details, personal data) in 2010. Those attacks are disproportionally aimed at more specific types of data that are not generally stolen in bulk: IP and corporate secrets. As later reports suggest, the primary motivation behind targeted attacks is industrial espionage—that is, the stolen data are used to achieve a competitive edge on the market (Verizon Business, 2020).

Opportunistic breaches constitute around 80% of cases. Victims of these attacks are chosen due to the exploitable vulnerabilities they exhibit (e.g. SQL Injection-vulnerable websites) rather than their business activity. Opportunistic attacks are predominantly financially motivated and are responsible for 90% of compromised records, including customers' financial data, personal information, and credentials that might be quickly sold on the black market.

It should be recognised that the actual rate of intellectual property and corporate secret theft is likely larger than reported. Financial fraud detection remains the most effective method of discovering a breach. However, IP and corporate secrets are not generally involved in financial identity theft. It is then reasonable to assume that a substantial share of targeted attacks has never been discovered.

It is worth mentioning that the research focuses on the cyberattacks performed by external actors. While there is no doubt that an insider threat is a common problem in modern competition, external actors are the predominant source of information leakage attacks. External actors were responsible for 92% of all reported data breaches in 2010 (Verizon Business, 2011). We leave the discussion of insider threats for another occasion.

Independent variables

As neither the exact form of the targeting mechanism nor the actual form of the ideas output function is known, the research tests the central hypothesis using three proxy variables. Proxy variables represent shares of interlayers of organisations' R&D personnel and resemble the variables used in macroeconomic and economic growth literature to model ideas output (Coe & Helpman, 1995; Jones, 2003; Porter & Stern, 2000). The variables are derived according to the Frascati Manual⁴ and are defined as follows (OECD, 2015):

1. *R&D personnel* (r_{ij}^{per}) is the share of all persons in overall employment in a given country in a given industry who engaged in R&D as well as those who provide direct support for the R&D activities (e.g. R&D managers, technicians);
2. *Researchers* (r_{ij}^{res}) is the share of professionals and researchers in overall employment in a given country in a given industry who are engaged in the conception or creation of new knowledge, conduct research and improve or develop theories, models, techniques instrumentation, software or operational methods;
3. *Technicians and other R&D supporting personnel* (r_{ij}^{sup}) represents the share of persons in overall employment in a given country in a given industry who provide direct services for the R&D activities (e.g. technicians, R&D managers, and administrators).

Note that the R&D personnel variable can be expressed as a sum of the latter two variables, $r_{ij}^{per} = r_{ij}^{res} + r_{ij}^{sup}$. The industrial distribution of all three variables is demonstrated in Figure 3.3. The figure demonstrates that manufacturing industries (codes start with C) possess a much larger share of research personnel than service industries. As shown in the figure, the manufacture of computer, electronic, and

⁴The Frascati Manual is a document setting forth the methodology for the collection of analyses of statistics about research and development.

optical products (C26) has the largest share of research personnel among all the industries. Services are generally less research-oriented, with only two groups demonstrating a relatively large share of R&D personnel: publishing of books, software, and media (J55–J63) and professional, scientific, and technical activities (M69–M74). See Appendix B.1 for a complete list of industry aggregates.

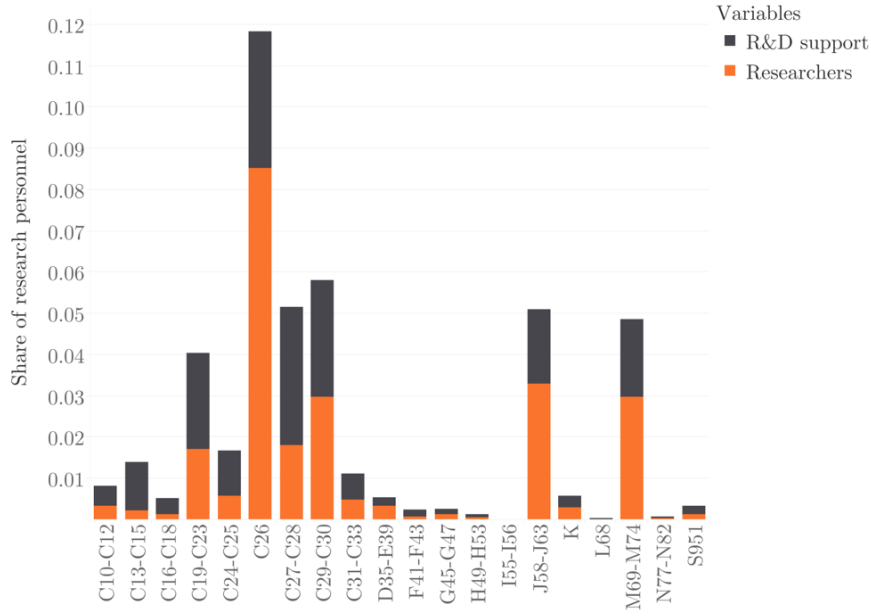


Figure 3.3: Industrial distribution of the research-intensity proxy variables.

Due to collinearity the independent variables cannot be included in the specification simultaneously. Therefore, we use three separate sets of specifications to analyse the association between each independent variable with information leakage attack rate.

The specifications with the share of research personnel as a primary independent variable are built following the assumption that researchers and R&D support personnel cooperatively contribute to innovative ideas. It assumes that industries with the largest share of research personnel are the most research-intensive and produce the most valuable results. However, it also implies that research-intensity is proxied by the sum of the researchers and research support personnel, which might not be the case (e.g. the cooperation might be synergistic, or researchers might contribute significantly more than the research support personnel).

The second set of specifications assumes that research-intensity can be attributed mainly to the input of researchers and research professionals. While the assumption is plausible as the research support personnel might not contribute to the research directly, it also neglects the latter's potential effects on the output of the ideas.

Finally, the last set of tests is performed with the share of technicians, R&D managers and other support personnel as the primary independent variable. A company with highly productive researchers might hire a larger R&D support body. If researchers and research support personnel are complements in the production function of the ideas, then they would allow the researchers to focus on their work and increase their productivity.

P. A. Lewis (2017) suggested that a greater proportion of technicians and R&D managers indicates mature high-tech organisations specialising in process development. Those organisations have a more elaborate division of labour which allows them to efficiently commercialise innovation and carry out abundant manufacturing in the pilot plants (Steger & James, 2011). Similarly, a significant fraction of R&D managers, administrators, and equivalent personnel in the service industries might indicate the high commercial potential of research output. In other words, industries with a large share of R&D support personnel are typically those with the most immediately marketable ideas and thus are often attractive targets for industrial espionage. Moreover, the share of research support personnel has the most significant positive correlation with the R&D expenses and overall turnover among all three proxies in our data.

On the other hand, a larger body of R&D support personnel might disproportionately increase the company's attack surface. It increases the number of people with access to sensitive data without a substantial contribution to idea generation. Hence, the variable might capture both research-intensity and extended attack surface. However, industry control variables can rule out the attack surface effect. Additionally, we test the plausibility of the attack surface hypothesis using the extended specification (see Subsection 3.4.4 for details), which includes the auxiliary attack surface variable defined as the average amount of R&D support personnel in

an enterprise in a given industry in a given country. Being an absolute measure, it allows for comparing the average attack surfaces among industries while neglecting the research-intensity itself⁵ (Comanor, 1965; Gruber et al., 1967; Terleckyj, 1960).

All research-intensity proxy variables have a right-skewed distribution (see Appendix B.2). Still, as with the dependent variable, there are no solid theoretical grounds to use any transformation to normalise them.

We use an additional proxy variable to control the variance associated with opportunistic attacks and test Hypothesis 2. Note that opportunistic attacks are responsible for most compromised consumer data and credentials—the information commonly stored in Customer Relationship Management (CRM) systems. Therefore, the last proxy variable is defined as *CRM usage rate* (d_{ij}) in a given country in a given industry. We believe that the rate of CRM usage will be higher in industries with more valuable customer data. The industrial distribution of CRM usage rate is demonstrated in Figure 3.4. The histogram of this variable can be found in Appendix B.2.

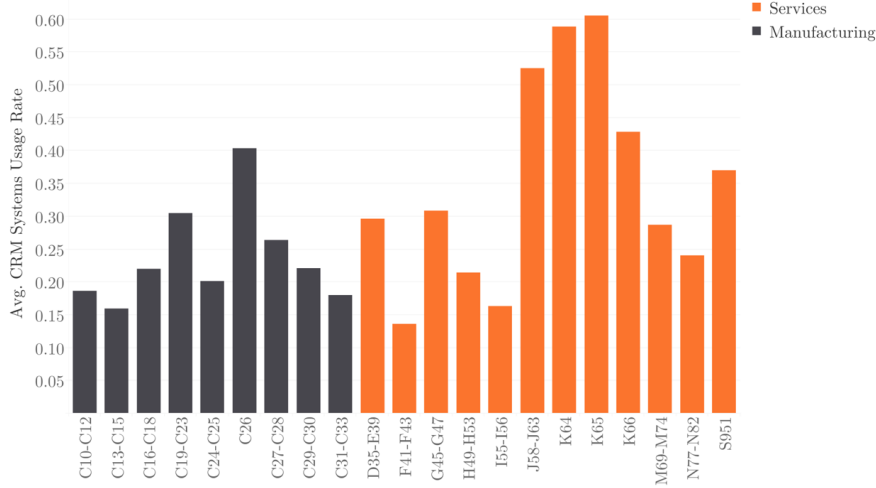


Figure 3.4: Industrial distribution of the CRM usage rate.

The basic descriptive statistics for the dependent and independent variables can be found in Table 3.2.

⁵The attack surface variable does not capture research-intensity as it does not allow for comparability of R&D support bodies among industries. As the total average employment is not used in the robustness setups, it only captures the attack surface’s variance.

Table 3.2: Descriptive statistics for all main variables. Indices i and j denote a country and an industry respectively

Variable	Label	Obs.	Avg.	Std. Dev.	Min	Max
Information breaches rate	B_{ij}	351	.014	.021	0	.19
Share of R&D personnel	r_{ij}^{per}	224	.027	.059	0	.649
Share of Researchers	r_{ij}^{res}	224	.015	.042	0	.506
Share R&D support personnel	r_{ij}^{sup}	222	.012	.020	0	.143
CRM usage rate	d_{ij}	315	.285	.171	0	1

Categorical variables

The study utilises two sets of categorical variables to control for heterogeneity among countries and industries. They account for unobserved differences across countries and industries by allowing individual intersects for each category.

Alongside with base sets of industry and country control variables, we use two additional sets of Eurostat indicators on high-tech industry and knowledge-intensive services (Eurostat, 2016). The manufacturing-specific variable is based on the categories of manufacturing industries, including high-technology, medium-high-technology, medium-low-technology, and low-technology. They were divided according to technological intensity. Table B.2 of Appendix B.5 presents aggregated categories. Following a similar approach, the study utilises service-specific variables based on Eurostat's definitions of knowledge-intensive services (KIS) and less knowledge-intensive services (LKIS). The aggregation is demonstrated in Table B.3 of Appendix B.5.

3.3.2 Methodology

To investigate the relationship between information leakage attack rate and the research-intensity based on the available data set, this study adopts the Bayesian-frequentists hybrid method: multiple imputation to handle partial observations followed by fractional outcome regression analysis. The procedure is performed in three steps:

1. Multiple imputation, where M completed data sets (imputations) are simulated under the chosen imputation specification;
2. Complete-data analysis, where fractional regression analysis is performed on each imputation $m = 1, \dots, M$;
3. Pooling, where the results from M completed data sets are combined.

This subsection offers detailed information on the underlying statistical methods used in each step.

Multiple imputation

The major challenge imposed by the data is partially missing observations: around 36% of observations are missing for research-intensity variables and 10% for the *CRM* variable. The issue is mainly caused by the non-response in Eurostat's surveys. If not handled carefully, it could reduce statistical power, lead to biased estimates, and larger standard errors due to the smaller sample size.

There are two generally accepted approaches to handling missing data: (1) omission, by which partially incomplete observations are discarded from the analysis, and (2) imputation, by which missing observations are filled in.

The size of the data set, the share of partially missing observations, and the lack of robustness for standard statistical inference techniques render the omission option rather inappropriate. On the other hand, imputation techniques are designed to extract all available information from partially observed data and minimise bias.

The research adopts the multiple imputation (MI) technique introduced by Rubin (1987). The multiple imputation approach is a Bayesian simulation-based statistical technique designed to handle partially observed data. The method was initially used to tackle non-response in medical surveys and handle data sets with up to 90% of observations missing (Madley-Dowd et al., 2019).

The method is superior to the single imputation counterpart (e.g. last observation carried forward, the sample mean imputation) as it accounts for the uncertainty of missing values. The single imputation approach treats imputed values as genuinely

observed, which causes standard errors to be too small and causes a biased statistical inference. In reality, it is not possible to know the exact values of the missing data, and the MI technique accounts for all inherent uncertainty by injecting appropriate variability into the multiple imputed values (Sterne et al., 2009).

The research utilises two MI procedures designed for univariate and multivariate models.

The *univariate imputation* of research-intensity independent variables utilises the Gaussian normal regression model. The procedure is described below.

Consider a 351×1 vector $R^L = (r_{1j}^L, \dots, r_{ij}^L; r_{i1}^L, \dots, r_{ij}^L)'$, which records values for the research-intensity variables: shares of research personnel ($L = per$), researchers specifically ($L = res$), or R&D support personnel ($L = sup$) in an industry i and country j . The variable then follows a Gaussian normal regression model:

$$r_{ij}^L | z_{ij} \sim N(z_{ij}', \sigma^2), \quad (3.1)$$

$z_{ij} = (z_{ij1}, z_{ij2}, \dots, z_{ijq})$ is a vector of values of predictors of r_{ij}^L for observation ij , β is the $q \times 1$ vector of unknown regression coefficients, and σ^2 is the unknown scalar variance.

Consider then the division of $R^L = (R_o^L, R_m^L)$ into $n_0 \times 1$ and $n_1 \times 1$ vectors containing the complete and incomplete observations respectively. Similarly $Z = (Z_o, Z_m)$ is divided into $n_0 \times q$ and $n_1 \times 1$ submatrices.

The imputation then proceeds as follows to fill in R_m^L :

1. Specification (3.1) is fitted to the observed data (R_o^L, Z_o) to obtain estimates $\hat{\beta}$ and $\hat{\sigma}^2$ of the model parameters.
2. New parameters $\tilde{\beta}$ and $\tilde{\sigma}^2$ are simulated from their joint posterior distribution under the noninformative improper prior for a scaling parameter $\Pr(\beta, \sigma^2) \propto \frac{1}{\sigma^2}$. This is done in two steps:

$$\begin{aligned} \tilde{\sigma}^2 &\sim \hat{\sigma}^2 (n_0 - q) / \chi_{n_0 - q}^2, \\ \tilde{\beta} | \tilde{\sigma}^2 &\sim N\left(\hat{\beta}, \tilde{\sigma}^2 (Z_o' Z_o)^{-1}\right); \end{aligned}$$

3. The new set of imputed values is obtained by simulating missing values of R_{m1}^L from $N(Z_m\tilde{\beta}, \tilde{\sigma}^2 I_{n_1 \times n_1})$;
4. Steps 2 and 3 are repeated to obtain M sets of imputed values:

$$R_{m2}^L, R_{m3}^L, \dots, R_{mM}^L.$$

Steps 2 and 3 follow the procedure of simulating the missing data from the posterior predictive distribution, described by Gelman et al. (2013).

The univariate imputation models utilise the rate of information leakage attacks in industry i , country j as the explanatory variable and both sets of categorical variables. Specification 3.1 can then be represented in the linear form:

$$r_{ij}^L = \beta_1 + \beta_2 B_{ij} + a_i + c_j + \epsilon_{ij}.$$

where B_{ij} is the share of enterprises that experienced information leakage attacks in an industry i and country j , and ϵ is the error term.

The multivariate models utilise the modified version of the MI technique to allow the *imputation of multiple variables* with missing observations simultaneously: multiple imputation by chained equations (MICE) developed by van Buuren et al. (1999). MICE is more suitable than other multivariate imputation methods as it can handle arbitrary missing-data patterns (contrary to a standard multivariate imputation, which requires missing data patterns to be monotone).

The MICE procedure imputes multiple variables iteratively following a sequence of univariate multiple imputation models (one model for one partially observed variable). The procedure can be described as follows. For imputation variables $X_1, X_2, \dots, X_q$ and complete predictors, Z , imputed values are drawn from:

$$\begin{aligned} X_1^{(t+1)} &\sim g_1 \left(X_1 | X_2^{(t)}, \dots, X_p^{(t)}, Z, \phi_1 \right), \\ X_2^{(t+1)} &\sim g_2 \left(X_1 | X_1^{(t+1)}, X_3^{(t)}, \dots, X_p^{(t)}, Z, \phi_2 \right), \\ &\vdots \\ X_q^{(t+1)} &\sim g_p \left(X_1 | X_1^{(t+1)}, \dots, X_{p-1}^{(t+1)}, Z, \phi_q \right). \end{aligned}$$

for iterations $t = 0, 1, \dots, T$ until the imputed values converge at doomsday $t = T$. Here, ϕ_q is the vector of corresponding model parameters with a uniform prior and, g_k is the corresponding univariate imputation model appropriate for imputing corresponding independent variables, X_k , where $k = (1, \dots, q)$. As for specification models, the research utilises multivariate normal Gaussian regressions' specifications:

$$\begin{aligned} X_1 &= \beta_0 + \beta_2 X_2 + \dots + \beta_q X_q + \beta_{q+1} B_{ij} + c_j + a_i + \epsilon_{ij} \\ X_2 &= \beta_0 + \beta_1 X_1 + \beta_3 X_3 + \dots + \beta_q X_q + \beta_{q+1} B_{ij} + c_j + a_i + \epsilon_{ij} \\ &\vdots \\ X_q &= \beta_0 + \beta_1 X_1 + \dots + \beta_{q-1} X_{q-1} + \beta_{q+1} B_{ij} + c_j + a_i + \epsilon_{ij}, \end{aligned}$$

For the described procedure to lead to valid statistical inference, four more conditions must be satisfied: (1) the chosen imputation models should be *appropriate* for the further statistical analysis; (2) MI procedure must be *proper*; (3) MI procedure requires the missing data mechanism to be ignorable; and (4) a sufficient number of imputations, M , must be chosen (Rubin, 1987). These conditions are discussed in more detail in Appendix B.3.

Thus, the research adopts the multiple imputation procedure to handle the relatively big share of partially observed data. The procedure is designed to account for the uncertainty embedded in the analysis of partially observed data instead of simply generating a data set with no missing data. Given that the above-mentioned conditions are satisfied, MI procedure should lead to the unbiased statistical analysis described in the following subsection.

Complete-data analysis

We use a fractional outcome regression developed by Wedderburn (1974). The method belongs to the generalised linear models (GLM) class and is designed to fit outcome variables $\in [0, 1]$.

Fractional outcome regression is specified as GLM of the binomial family with robust standard errors and a link function $G(\cdot)$ which satisfies $G(z) \in [0, 1]$

for all $z \in \mathbb{R}$:

$$E(Y_i|X_i) = G(X_i\beta). \quad (3.2)$$

Cumulative distribution functions (CDF) are often utilised as the link function for fractional regression, with the two most common options being the logistic function $G(z) = \Lambda(z) = \frac{\exp(z)}{1+\exp(z)}$, and the probit function $G(z) = \Phi(z)$, where $\Phi(z)$ is the CDF of a standard normal distribution. However, it can be essentially any function bounded between 0 and 1, not necessarily a CDF.

Additionally, GLM class models require fewer assumptions to be satisfied to achieve unbiased results (in comparison with the OLS regression)—see Appendix B.4 for the discussion. Given that data cases are independently distributed, it is merely enough to specify the conditional mean of the variable of interest correctly to achieve an unbiased statistical inference.

The research utilises two linear specifications that will be estimated on the imputed data: S^U (univariate specification) and S^M (multivariate specification).

Univariate specification:

$$E(B_{ij}|r_{ij}^L, c_j, a_i) = G(\beta_0 + \beta_1 r_{ij}^L + c_j + a_i) = G(S^U), \quad (3.3)$$

which is used to assess individual effects of the research-intensity variables on the rate of information leakage attacks.

Multivariate specification:

$$E(B_{ij}|r_{ij}^L, d_{ij}, c_j, a_j) = G(\beta_0 + \beta_1 r_{ij}^L + \beta_2 d_{ij} + \dots + c_j) = G(S^M), \quad (3.4)$$

which is used to assess the effects of the research-intensity variables on the rate of information leakage attacks in the presence of an additional CRM usage rate variable (in the base multivariate specification). The variable is added to the model to control for the variance induced by opportunistic attacks in which the leakage of personal records of clients is highly likely to predominate. Note that the multivariate specification lacks the industry controls. They were removed from the

specification due to collinearity: the share of enterprises that use CRM systems is industry-dependent. Appendix B.6 provides the variance inflation factor (VIF) and estimated within-industry variance of the CRM variable.

Following good practice procedures proposed by Villadsen and Wulff (2021) it is now necessary to (1) verify the conditional mean specifications and (2) choose an appropriate link function.

The conditional mean specification can be tested using the Regression Equation Specification Error Test (RESET) developed by Ramsey (1969). It test whether non-linear combinations of the fitted values (e.g. $\hat{B}_{ij}^P$, where P is the maximum power of the non-linear combination) are useful to explain the outcome variable. If non-linear combinations have any influence on the outcome variable, then the tests suggest that the model is misspecified and it may be better approximated via some non-linear function.

Tests results for both univariate and multivariate specifications under different link functions are provided in Appendix B.7. All four tests fail to reject the null hypothesis that the fractional outcome model is specified correctly.

We follow the canonical approach and discriminate between *logit* and *probit* link functions. To do so, the research utilises the test of econometric specification in the presence of alternative specification (Davidson & MacKinnon, 1981). Test results for univariate and multivariate specifications are provided in Appendix B.8. Both performed tests fail to reject the null hypothesis implying the usage of *probit* link function. However, both link functions lead to extremely close estimates, and the choice of link function would not influence the results.

The binomial GLM with probit link function then fits coefficients via maximising the following quasi-likelihood functions for every imputed data set:

$$\log L = \sum_{i=1}^{N^{in}} \sum_{j=1}^{N^c} \left(B_{ij} \log\{\Phi(S)\} + (1 - B_{ij}) \log\{1 - \Phi(S)\} \right), \quad (3.5)$$

where $N^{in/c}$ is the number of industries and countries represented in the research, S is the linear specification with $S = S^U$ for univariate specification and $S = S^M$ for multivariate specification, and Φ is a CDF of standard normal distribution. The

maximisation is performed numerically following the Newton-Raphson iterative method (McMullen, 1987).

The statistical significance of a single coefficient is determined based on the standard score (z) statistic. However, the coefficients of fractional regression outcome with probit link function have limited interpretability. They can be interpreted as the difference in the standard scores associated with a unit change in the regressor. Therefore, raw models' results can only reveal statistically significant effects and signs.

To achieve interpretable and comparable results, we estimate the marginal effects of each variable. For univariate-specification models, the marginal effect estimation demonstrates how the outcome (the conditional mean of the rate of information leakage attacks) variable changes when a specific explanatory variable changes, $\frac{\partial Y}{\partial X}$. For models with a probit link function, the marginal effect is estimated as follows:

$$\frac{\partial Y}{\partial X_i} = \beta_i \phi(S), \quad (3.6)$$

where $\phi(\cdot)$ is a standard normal probability density function.

The procedures described above are performed for each imputed data set, and then the results are combined in the following pooling step.

Pooling step

Given that multiple imputations are proper, unbiased estimates can be achieved by a simple averaging of the estimates derived by fitting fractional outcome regression models on each of the 40 imputed data sets separately. Pooled estimates are then interpreted similarly as the coefficients of a classic data set model. All results reported in the next section are based on the pooled coefficients.

3.3.3 Methodology discussion

This subsection provides a discussion on the chosen methodology.

Controlling for unobserved effects

The study utilises the conventional fixed effects method to account for unobserved heterogeneity of industries and countries. Alternatively, it is possible to model heterogeneity using the mixed-effects method (also called multi-level or hierarchical models). This family of models is especially attractive for research designs in which individuals are organised hierarchically at more than one level (e.g. at the country and industry level in this study).

Unfortunately, the data adopted for this study do not allow for hierarchical mixed-effects models. This happens due to the absence of replicated measurements at the industry-country level (i, j) . The hierarchical model is not identified for singleton pattern data because the random intercepts estimated for the lower level (which is the “industry” level in the case of this research) are confounded with the overall error terms (Stata Corporation, 2013).

Handling missing data

The unbiasedness of the results depends on the validity of the imputation model. The imputation method might not lead to an unbiased result if not handled correctly. Several pitfalls might lead to a biased statistical inference.

Misspecification of regressors mainly exists due to the exclusion of the dependent variable from a linear specification of the imputation equation. An outcome variable might possess additional information about the missing values of the explanatory variables. Failure to include the outcome variable into the imputation model weakens the association between the dependent variable and imputed regressors (von Hippel, 2013). The problem was avoided in the proposed MI procedure, and every imputation model includes the dependent variable.

Imputation of non-normally distributed variables is another concern that is discussed in the literature. The normality assumption of Gaussian regression, in general, applies to the error terms. The assumption is satisfied, as demonstrated in Appendix B.9, so it is not an issue in this particular research.

More debatable is whether the assumed imputation model is appropriate to handle the *imputation of bounded data* (all imputed variables in this study are $\in [0, 1]$). The main concern with the chosen specification is that the predicted values may fall outside the range of the imputed variable. It is essential to remember that the MI procedure is not utilised to recreate missing observations. Instead, it achieves an unbiased statistical inference by accounting for uncertainty embedded in the missing observations. Imputed values are, therefore, not required to be bounded.

Moreover, the study by Rodwell et al. (2014), which investigated the efficiency of imputation methods of bounded variables based on 1000 incomplete simulated data sets, suggested that the Gaussian MI procedure is superior to any conventional alternatives. Post-imputation rounding, truncated normal regression, zero-skewness log-transformation, and predictive mean matching are highly sensitive to the skewness of the imputed variables and prone to reducing the variance of imputed values.

To conclude, the potential bias may arise from misspecification of MI procedure or imputing data with a non-random missing pattern. Following Rubin's recommendations, the proposed MI procedure is carefully crafted to achieve proper multiple imputation. The Gaussian specification is chosen based on a comprehensive empirical study of alternative imputation methods of bounded variables. The inference based on complete cases is reported in Appendix B.

Analytical model

Four main analytical models are used in statistic/econometrics literature to model fractions (from the most popular to the least popular): (1) OLS regression, (2) OLS regression with transformed dependent variable, (3) Tobit models, and (4) fractional regression (Villadsen & Wulff, 2021). Additionally, the excess share of zero observations might suggest using zero-inflated models. However, all the alternative methods have theoretical and practical drawbacks discussed below.

The main problem with *OLS regressions* is that it ignores the boundaries of the outcome variable. Modelling an outcome variable $\in [0, 1]$ with Gaussian normal regression will most certainly require a violation of the linearity assumption. Another

problem is the error term's heteroskedasticity, which requires special treatment (e.g. estimation of robust standard errors). Furthermore, OLS is highly sensitive to skewed variables (which is not the case with fractional outcome regressions with non-linear link functions). However, this model can still be used to investigate the signs and sizes of the relationships, though the procedure must be performed carefully.

The most popular *transformation of the fractional dependent variable* is a *log-transformation*, which might normalise the data and solve the problem with outcome boundaries. The main problem with this method is that it does not perform well with highly skewed dependent variables with a substantial share of zero observations. It requires some ad-hoc solution to make the log-transformation of zero numbers possible (e.g. addition of some small negligible number), which may lead to a biased inference. Moreover, the model requires substantial efforts to verify the underlying assumption of the OLS regression.

The Tobit model was initially developed to model censored outcomes (e.g. outcomes whose values are only partially known) (McDonald & Moffitt, 1980). Fractional outcomes are not censored in any way but instead defined only on the interval of $\in [0, 1]$. Censored outcomes and corner solution responses (like fractional outcomes) are often confused. The application of the Tobit model to the fractional outcomes implies that there exist some observations of the dependent variable outside the unit range, which is not the case (Wooldridge, 2010). This leads to difficulties in the interpretation of the Tobit model coefficients. The marginal effect is only reflected on the coefficients of the latent (unobserved) variable, which is non-existent. Moreover, the model assumes homoskedasticity and normality of error terms—conditions likely to be violated when dealing with fractional outcomes.

Zero-inflated models might be another reasonable choice given the share of zero observations. However, those models are generally harder to set due to the necessity of setting the predictor for zero observations appearance, which, conditional on the nature of the data, might not be feasible. Moreover, those models are generally used to tackle the problem of overdispersion, which is not the case for a given data (see Table 3.2: the zero-inflated Poisson model would be appropriate if mean of the

dependent variable is much smaller than the variance). Additionally, they cannot be applied to fractional variables directly and would require an additional estimation step (e.g. the two-part composite model in W. Liu & Xin, 2014), which would further complicate the interpretation and tracking of assumptions (Williams, 2019).

Fractional outcome regression is most suitable for our data. It performs extraordinarily well in the presence of heteroskedasticity, provides valid statistical inference for the bounded and skewed outcomes, and produces interpretable and adequate estimates. Moreover, it is agnostic about the moments of the outcome and requires only the correct specification of the conditional mean.

3.4 Results

This section presents the results of the study. The study adopts the univariate and multivariate specifications described in the previous section. Both sets of specifications are tested with three proposed research-intensity proxy variables.

3.4.1 Fractional regressions results

Table 3.3 presents results of MI univariate fractional outcome regressions with all three research-intensity proxy variables. Column (I) of the table reports base regression results with no control variables, columns (II) and (III) report regression results with the country and industry controls included respectively, and column (IV) reports the results with both control variables sets. Complete case regression analysis tables can be found in Appendices B.10, B.12, and B.14.

The upper section of Table 3.3 contains results of univariate regression models with the share of research personnel as the main explanatory variable. Research personnel's coefficients are positive in all four specifications and statistically significant at 5% level in three out of four specifications (with the specification with the country-only controls as an exception).

The middle section of Table 3.3 reports the results with an alternative measure of research-intensity: the share of researchers. As previously, models report positive coefficients next to the primary explanatory variable. However, only coefficients for

the specifications with industry-only controls, or both controls are statistically significant at 5% level.

Finally, the bottom section reports the results of models with the share of research support personnel as the main explanatory variables. The coefficients are positive and statistically significant in all four specifications at 1% level.

It is evident that regardless of the included research-intensity variable, models specified with both country and industry controls display positive and statistically significant coefficients.

Raw fractional probit models' coefficients, however, are not further interpretable. Results only indicated the positive sign of the correlation between the alternative research-intensity proxy variables and the information leakage attacks rate. The interpretation of results is included in the following analysis of coefficients' marginal effects.

Table 3.3: Regression results: fractional probit regressions with share of research personnel as an independent variable

	I Base model	II Country controls	III Industry controls	IV All controls
Research personnel MI models				
r^{per}	1.377** (2.23)	1.011* (1.95)	1.820** (2.44)	1.285** (2.43)
Constant	-2.233*** (-67.86)	-2.100*** (-14.56)	-2.252*** (-24.76)	-2.110*** (-13.69)
Researchers MI models				
r^{res}	1.366* (1.68)	1.052 (1.53)	1.747** (1.97)	1.315** (2.07)
Constant	-2.212*** (-70.84)	-2.086*** (-14.27)	-2.239*** (-24.96)	-2.101*** (-13.38)
R&D support personnel MI models				
r^{sup}	6.225*** (3.81)	3.838*** (2.80)	9.106*** (4.38)	5.521*** (3.67)
Constant	-2.273*** (-68.83)	-2.110*** (-15.78)	-2.279*** (-23.31)	-2.118*** (-14.35)
Observations	351	351	351	351

z statistics in parentheses; Significance levels: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$

C10–C12 and BE are chosen as base levels; See Appendices B.11, B.13, B.15 for complete tables

Table 3.4: Average marginal effects of research-intensity proxy variables in models with all controls

	$\partial B / \partial r^{per}$	$\partial B / \partial r^{res}$	$\partial B / \partial r^{sup}$
Marginal effect	0.0432**	0.0442 **	0.185***
	(2.38)	(2.04)	(3.55)

z statistics in parentheses; Significance levels: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$

3.4.2 Marginal effects analysis

Table 3.4 presents average marginal effects estimates (as shown by equation (3.6)) for all three alternative research-intensity proxy variables estimated for models with all controls included. As expected, all presented coefficients are positive and statistically significant at 5% level (or even at 1% level for the share of research support personnel). It is now possible to interpret the coefficients in the same fashion as coefficients of a classic OLS regression.

In the model with the share of research personnel as the main independent variable, a 10 percentage points increase in the share of research personnel is associated with a 0.432 percentage points increase in the information leakage attack rate. This marginal effect is relatively similar to that demonstrated by the model with the share of researchers as the primary independent variable: a 10 percentage points increase in the share of researchers is associated with a 0.442 percentage points increase in the information leakage attack rate.

The model with the share of research support personnel delivers much more prominent results. A 10 percentage points increase in the share of research support personnel is associated with a 1.85 percentage points increase in the information leakage attacks rate.

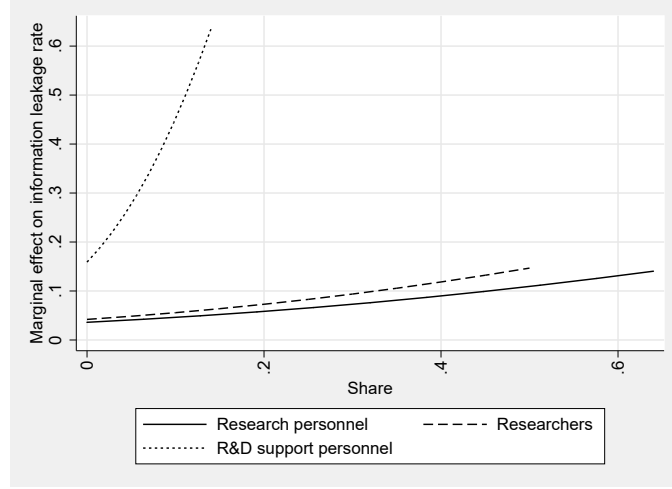


Figure 3.5: Point estimates of the marginal effects of the research proxy variables.

The substantial difference in the magnitude is also evident from the shape of the partial derivative function. As illustrated in Figure 3.5, all three marginal effects follow quadratic growth as the share of the corresponding research body increases. In particular, the marginal effect of the share of research support personnel has a much steeper shape than the marginal effects of the other two proxy variables. In its peak ($r^{sup} = 14\%$), a 1 percentage point increase in the share of research support personnel is associated with a 0.634 percentage points increase in the information leakage attack rate. Meanwhile, the marginal effects of the shares of research personnel and researchers have a similar, yet more gradual incline ($\partial B / \partial r^{per} = 0.14$ at $\max r^{per} = 64\%$ and $\partial B / \partial r^{res} = 0.15$ at $\max r^{res} = 50\%$).

3.4.3 Targeted and opportunistic attacks

As stated above, most data breaches are affiliated with opportunistic, financially motivated breaches responsible for the majority of compromised customers' data.

It is possible that the chosen dependent variable accounts mostly for opportunistic attacks, and the discovered relationships are spurious and not related to the targeted information leakage attacks. To control for this possibility, the study estimates multivariate specifications, which include a proxy variable that represents a measure of how much end-user data are available to be stolen in an industry—the CRM usage rate. Results for the extended multivariate specification for all three research

Table 3.5: Regression results: fractional probit regressions with research proxies, CRM variable and country controls

	I R&D personnel model	II Researchers model	III R&D support model
MI models			
r^{per}	0.562 (1.22)		
r^{res}		0.410 (0.53)	
r^{sup}			3.041** (2.35)
d	0.715*** (3.63)	0.762*** (3.78)	0.641*** (3.48)
Constant	-2.409*** (-14.62)	-2.416*** (-14.64)	-2.395*** (-15.37)
Observations	351	351	351

z statistics in parentheses; Significance levels: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$

C10–C12 is chosen as a base level; See Appendices B.16 and B.17 for complete tables

proxy variables (as demonstrated by equation (3.4)) are presented in Table 3.5. The table for complete case analysis can be found in Appendix B.16.

One should approach the interpretation of the coefficients estimated for the CRM variable with caution. As industry control variables were excluded from the following sets of models due to collinearity, the CRM variable might absorb the variance generated by industrial heterogeneity unrelated to the CRM usage rate.

In all three presented models, the coefficients for the CRM usage rate are positive and statistically significant at 1% level. All three models report positive signs of coefficients for the research-intensity variables. Nonetheless, only the share of research support personnel has a statistically significant coefficient at 5% level.

As suggested earlier, the variable might indicate mature high-tech manufacturing or service industries with a substantial body of commercialisable research. We now consider the model presented in column (III) of Table 3.5 in further detail. The average marginal effects of the share of R&D support personnel and the CRM usage rate on the information leakage attacks rate are presented in Table 3.6.

Table 3.6: Average marginal effects in extended models

	$\partial B/\partial r^{sup}$	$\partial B/\partial d$
Marginal effect	0.103** (2.28)	0.0217*** (3.34)
<i>z</i> statistics in parentheses; Significance levels: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$		

The average marginal effect of the share of R&D support personnel in the MI models is positive and statistically significant at 5% level: a 10 percentage points increase in the share of R&D support personnel is associated with a 1.03 percentage points increase in the information leakage attack rate. Likewise, the raw coefficient, the average marginal effect of the CRM usage rate on the information leakage attack rate, is positive and significant at 1% level. A 10 percentage points increase in the CRM usage rate is associated with a 0.22 percentage points increase in the information leakage attack rate.

The average marginal effect of the share of R&D support personnel is larger than that of the CRM system usage rate. Still, direct comparison of the effects is case-specific due to the industrial heterogeneity of variables' means. Point estimates of the marginal effects of both variables are demonstrated in Figure 3.6. As with the univariate models, the measured marginal effects of the share of R&D support personnel and the CRM usage rate exhibit quadratic growth. The marginal effect function for the research-intensity proxy is much steeper than the marginal effect function of CRM usage rate. It means that more research-intensive industries experience a larger effect of an additional percentage point of R&D support personnel.

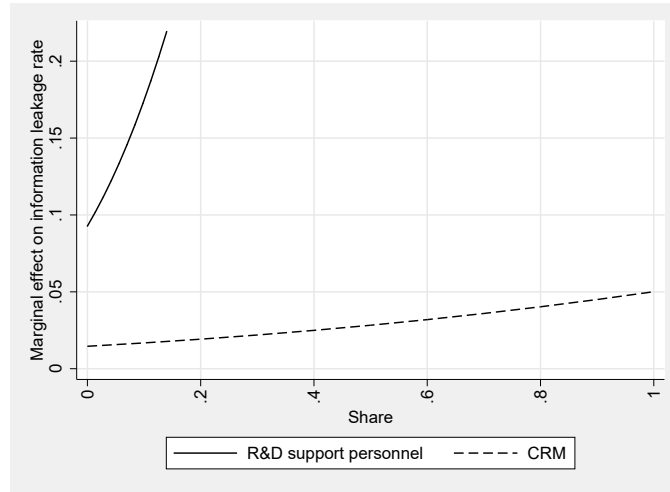


Figure 3.6: Marginal effects of the share of R&D support personnel and CRM usage rate point estimates.

Technological environments

This subsection investigates the relationships between the primary dependent variable, the research-intensity proxy, and the CRM usage rate in diverse technological environments. The multivariate specification is estimated separately for manufacturing, services, and corresponding technological and knowledge-intensity levels. Examining coefficients' magnitudes in different technological environments might provide supportive evidence for the underlying causal mechanism.

The manufacturing and service industries are divided according to the Eurostat indicators on high-tech industry and knowledge-intensive services provided in Appendix B.5. Note that to maintain a healthy amount of observations, the manufacturing industries are grouped into two groups instead of four: the high-tech group absorbs high-technology and medium-high-technology groups. In contrast, the low-tech group includes medium-low-technology and low-technology industries.

If the causal mechanism is plausible, two associations identified in Subsection 3.4.3 should exhibit different patterns depending on the technological environment. It is expected that the information leakage attacks in manufacturing industries are closely associated with research-intensity, and the association would be stronger in high-tech manufacturing industries. On the other hand, relatively less R&D intensive and more end-user data-reliant service industries would experience information

leakage attacks more closely related to CRM usage. Capturing this pattern might suggest that manufacturing industries are prone to experience targeted attacks, while services suffer from opportunistic, financially-driven attacks.

The specifications are estimated in technological clusters. This should reduce the potential heterogeneity captured by the CRM usage rate variable in the full-sample estimation and reveal less biased coefficients. Fractional regression results for manufacturing industries are presented in Table 3.7.

Table 3.7: Regression results: fractional probit regressions on imputed data with research proxy and CRM variable for manufacturing industries

	I All manufacturing	II Low-tech manufacturing	III High-tech manufacturing
r^{sup}	4.488*** (2.77)	3.200 (0.66)	7.206*** (3.13)
d	-0.776 (-1.41)	-0.651 (-0.87)	-1.337 (-1.29)
Constant	-1.549*** (-5.43)	-1.705*** (-4.98)	-1.223** (-2.24)
Observations	152	81	71

z statistics in parentheses; Significance levels: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$

C10–C12 is chosen as a base level; See Appendix B.18 for the full table

Coefficients of the CRM system usage are not significant in all demonstrated specifications. It suggests that opportunistic attacks do not contribute to the information leakage attack rate variance in manufacturing industries. Coefficients for the share of R&D support personnel are positive and significant at 1% for high-tech and all manufacturing industries on average. Neither the CRM system usage nor the research-intensity proxy demonstrate statistically significant coefficients for low-tech industries. Results suggest that information leakage attacks in the manufacturing industries are more likely to be targeted, as they are strongly associated with the R&D intensity proxy while demonstrating no association with the data reliance proxy.

Table 3.8 presents results of the average marginal effects of the share of R&D support personnel. As expected, the marginal effect is not significant in low-tech

industries. The average marginal effect of the share of R&D support personnel in the high-tech industries is substantially larger than the effect in all manufacturing industries. A 10 percentage points increase in the share of research support personnel is associated with a 1.26 percentage points increase in the information leakage attack rate in all manufacturing industries on average. In high-tech industries, a 10 percentage points increase in the share of research support personnel is associated with a 1.92 percentage points increase in the information leakage attack rate. The results align with the previous finding that the marginal effect function of the share of R&D support personnel is convex and steep.

Table 3.8: Average marginal effect in manufacturing models

	I All manufacturing	II Low-tech manufacturing	III High-tech manufacturing
$\partial B / \partial r^{sup}$	0.126*** (2.70)	0.091 (0.66)	0.192*** (3.38)

z statistics in parentheses; Significance levels: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$

Results of fractional outcome regressions for the subsamples of service industries are presented in Table 3.9. The models demonstrate the opposite outcomes to those seen in the manufacturing environment. The coefficients for the share of R&D support personnel are insignificant in all demonstrated specifications. In contrast, the CRM usage rate coefficients are significant at the 1% level for knowledge-intensive and all knowledge-based services. Table 3.10 offers average marginal effects estimates for the CRM variable: a 10 percentage points increase in the CRM usage rate is associated with a 0.274 percentage points increase in the information leakage attacks rate in all knowledge-based service industries. The magnitude of the average marginal effect is more pronounced in the knowledge-intensive service industries, with a 10 percentage points increase in the CRM usage rate associated with a 0.420 percentage points increase in the rate of information leakage attacks. Those results suggest that, contrary to manufacturing industries, information leakage attacks in the service industries are more likely to be opportunistic, as the dependent variable demonstrates a strong association with the data reliance proxy and no association

with the research-intensity proxy. Moreover, knowledge-intensive service industries (which are generally more data-reliant) tend to be affected by the phenomenon more than their less knowledge-intensive counterparts.

Table 3.9: Regression results: fractional probit regressions on imputed data with research proxy and CRM variable for service industries

	I All knowledge- based services	II Less knowledge- intensive services	III Knowledge- intensive services
r^{sup}	3.698* (1.82)	10.70* (1.89)	1.862 (1.01)
d	0.729*** (3.48)	0.238 (0.43)	0.985*** (4.55)
Constant	-2.573*** (-23.41)	-2.345*** (-8.96)	-2.698*** (-27.33)
Observations	199	105	94

z statistics in parentheses; Significance levels: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$
C10–C12 is chosen as a base level; See Appendix B.19 for the full table

Table 3.10: Average marginal effect in service models

	I All knowledge- based services	II Less knowledge- intensive services	III Knowledge- intensive services
d	0.0274*** (3.36)	0.00776 (0.42)	0.0420*** (4.45)

z statistics in parentheses; Significance levels: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$

Therefore, as suggested in the above analysis, the information leakage attack rate in manufacturing industries is mainly attributed to targeted attacks (potentially cyberespionage), while the attacks in service industries are primarily attributed to opportunistic attacks. However, the tests only provide supporting evidence for causal relationships instead of establishing them.

3.4.4 Alternative explanations and robustness checks

The study focuses on the share of research support personnel as the primary research-intensity proxy. The variable demonstrates statistically significant and positive coefficients in all univariate and multivariate specifications and possesses

the largest average marginal effect among the alternatives. The findings align with the proposed causal mechanism, but they do not rule out other potential causal paths. In search of additional supporting evidence, we perform several robustness tests to verify the behaviour of the primary variable in the presence of auxiliary variables that stand for alternative causal mechanisms.

Table 3.11: Regression results: robustness checks

	I	II	III	IV	V
	A.Surface controls	Turnover controls	E-comm controls	Sysadmin controls	All controls
r^{sup}	3.277*** (2.68)	3.235*** (2.72)	3.783*** (3.25)	2.999** (2.01)	3.981** (2.53)
d	0.585*** (3.02)	0.581*** (2.96)	0.285 (1.34)	0.540*** (2.88)	0.310 (1.45)
A.Surface	-18E-7 (-0.17)				703E-9 (0.07)
Turnover		0.000120 (0.32)			0.0000710 (0.19)
E-comm			0.00768** (2.22)		0.00787** (2.30)
Sysadmin				0.0822 (0.33)	-0.0711 (-0.29)
Constant	-2.370*** (-15.16)	-2.371*** (-15.08)	-2.464*** (-14.51)	-2.380*** (-14.61)	-2.458*** (-14.41)
Observations	351	351	351	351	351

z statistics in parentheses; Significance levels: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$

C10–C12 is chosen as a base level; See Appendix B.20 for the full table

Observe the R&D support personnel variable might capture two effects simultaneously: (1) the values of the research ideas, and (2) the extended attack surface. Column (I) of Table 3.11 shows a specification with an additional attack surface variable that represents an average quantity of the R&D support employees in the enterprise in a given industry in a given country. The additional variable hardly has any effect on the original variables' coefficients, and does not support the claim.

Another alternative scenario is that cyberattacks are generally performed on rich industries, which, in turn, can afford a larger R&D support body. It means that both information leakage attacks and research-intensity might be confounded by the company's average turnover in a given industry in a given country. This

alternative is tested in Column (II) of Table 3.11. The newly included variable does not have any effect on the results.

Column (III) of Table 3.11 demonstrates the results of the specification with the e-commerce variable included, which shows the share of enterprises in a given country in a given industry that have any electronic sales. It tests for the alternative scenario in which the most attacked industries are the ones that participate in e-commerce. Again, the additional variable does not significantly affect the research-intensity variable. However, it renders the CRM variable insignificant. This phenomenon can be easily explained: the CRM and e-commerce variables are highly correlated. The e-commerce variable should not significantly change the results for the research-intensity proxy as its magnitude is mainly attributed to manufacturing industries, which have relatively less engagement in e-commerce.

It might also be the case that research-intensity proxies are highly correlated with the number of system administrators in a given country in a given industry. It is a known fact that system administrators are the most common threat actor in internal industrial espionage cases (Verizon Business, 2020). Assuming that the same might be true for attacks performed by the external actors, it could be the case that the effects of system administrators might confound independent and dependent variables. Column (IV) displays results under the specification, which utilises an additional variable, *sysadmin*, which stands for the share of enterprises with an in-house system administrator in a given country in a given industry. The coefficient next to this variable is insignificant and does not affect the original variables' coefficients.

Column (V) displays results under the specification with all additional auxiliary variables included. Again, the main result of the research-intensity variable demonstrates robustness in a setting with all additional controls and even a slight increase in its magnitude.

Another potential source of endogeneity is research personnel engaged in cybersecurity research. However, cybersecurity researchers constitute only a tiny fraction of the overall employment. For example, only 43,000 employees worked

in the cybersecurity sector in the UK in 2019, which was less than 0.25% of the overall workforce. Moreover, only a fraction of them engaged in research and development (Donaldson et al., 2020).

The additional tests suggest that the main result is robust, even when controlling for alternative scenarios. The tests provide no evidence suggesting alternative causal mechanisms. However, those possibilities cannot be ruled out entirely until the causality between the information leakage attack rate and the research-intensity proxy is established: a challenge to overcome in future research.

3.5 Conclusion

This chapter studies the possible determinants of industrial information leakage attacks. Such attacks are often used as a means of unfair competition and are associated with industrial cyberespionage.

The study theorises that a fair share of targeted information leakage attacks happen as a part of the R&D race and, therefore, research-intensive industries are expected to be particularly susceptible to industrial cyberespionage. However, targeted attacks and industrial espionage constitute only a part of the picture.

The study implements a multi-stage approach in which the leading hypothesis is tested in a series of univariate and multivariate specifications. The multi-stage design allows assessing the performance of three different research proxy variables, both solely and in the presence of auxiliary variables that control for opportunistic attacks and alternative causal mechanisms.

The results of the univariate model of Eurostat data are consistent with the proposed causal mechanism. Regardless of the chosen research-intensity proxy, the models demonstrate positive and statistically significant coefficients after controlling for industry and country unobserved effects. Nevertheless, the variables differ considerably in their magnitudes. The share of research support personnel has a four times larger average marginal effect than the other two research-intensity proxies. Furthermore, this is the only research-intensity proxy that maintains a significant coefficient in the multivariate model.

The multivariate model reveals that user data reliance is another major contributor to the information leakage attack rate variance. The study distinguished two significant associations, potentially connected with targeted and opportunistic attacks respectively.

Estimating the multivariate specification on the subsamples of manufacturing, services, and their respective technological levels reveals that the discovered associations are industry-specific. Manufacturing industries (which are generally more R&D oriented) experience cyberattacks associated with research-intensity. The correlation is more prominent in the subsample of high-tech manufacturing industries, which can be considered additional supporting evidence favouring the proposed theory. The coefficients for user data reliance are insignificant in all models estimated for manufacturing industries. The situation is the opposite in the service industries, which experience information leakage attacks associated solely with end-user data reliance. Similarly, the effect becomes more salient in the subsample of knowledge-based services.

The association between the information leakage attack rate and research-intensity in manufacturing industries suggests that a major share of attacks might be targetted (potentially cyberespionage). Moreover, the attack rate is sensitive to the manufacturing industry's research-intensity and technological level. Those results provide supporting evidence in favour of the proposed causal mechanism. However, the causality is still to be established in future research, and alternative scenarios should not be ruled out completely.

To conclude, our study provides supporting evidence that research-intensive industries are more likely to become victims of information leakage attacks. This phenomenon could be associated with unfair competition and industrial espionage.

4

Secure and Efficient Networks

Contents

4.1	Introduction	109
4.1.1	Related literature	113
4.2	The model of network design and defence	116
4.2.1	Model setup	116
4.2.2	Analysis: encounter stage	118
4.3	Equilibrium	125
4.3.1	Limiting properties of equilibrium	130
4.3.2	Relaxed problem solution, $\delta = 1$	134
4.3.3	Feasibility of solutions	140
4.3.4	Numerical verification	145
4.4	Weighted version of the model	148
4.5	Limitations and future research	150
4.5.1	Limitations	151
4.5.2	Future research	152
4.6	Conclusion	154

4.1 Introduction

The scholarship of *economics of networks* generally focuses on a fundamental aspect: either network efficiency (Jackson, 2010) or network security (Dziubiński & Goyal, 2013; Goyal & Vigier, 2010). In most real-life scenarios, however, the issues of efficiency and security cannot be considered alone or at the expense of the other.

The fastest military supply chain is worthless during a war if it is not adequately protected with anti-aircraft defences, and can be easily disturbed by cheap unmanned combat aerial vehicles (UCAVs). Harsh social distancing measures can mitigate the spread of a virus but also reduce the density of social connections and severely damage business processes. Creating overconnected information communication technology (ICT) networks can significantly increase the data flow but also facilitate the hacker's navigation inside the enterprise's file system.

To illustrate the concepts of efficiency and security in a networked environment, we refer the reader to Figure 4.1. A star network (Figure 4.1a) is generally considered the most secure network: a significant share of its value is concentrated in the central node, making it easy to protect. However, any two peripheral nodes must overcome two links in order to communicate with each other, which makes it inefficient in some scenarios.

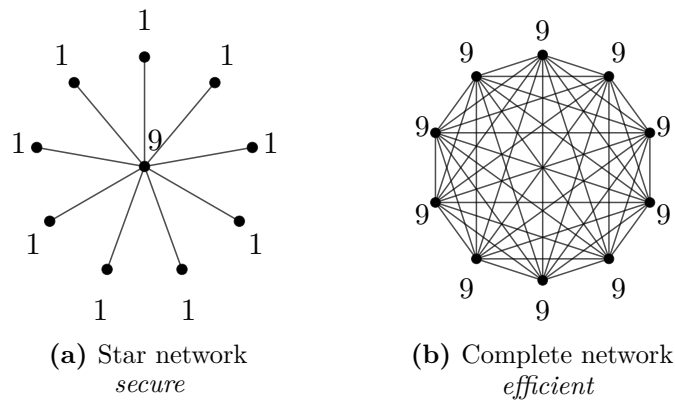


Figure 4.1: Examples of secure and efficient networks on nine nodes. Each node is valued according to the number of connections it has.

On the contrary, the complete network in Figure 4.1b is highly efficient. Each pair of nodes in this network are immediate neighbours, which allows for direct communication. However, complete networks are also the most insecure as it is hard to protect all high-value nodes when the defensive resources are scarce. The trade-off between security and efficiency is the central topic of this research.

More specifically, this study aims to understand how to construct an efficient network in the presence of an intelligent attacker. To tackle this problem, we

develop a game in which two asymmetric players, the Defender and the Attacker, compete in network design and defence over two stages.

In the initial *wiring stage*, the Defender is endowed with some fixed number of nodes and is free to arrange them in any connected network structure by creating costly links. For instance, the Defender can represent an IT department, who chooses an optimal information technology structure and hierarchical protocols, and optimises the data and information exchange efficiency against the expected security incidents rate.

At the beginning of the following *encounter stage*, the Attacker observes the network structure. For example, he can perform port scanning to retrieve the organisation of Internet Protocol (IP) address or use a specialised search engine to discover the underlying server structure of the organisation (e.g. Shodan Search Engine¹). The Attacker and the Defender then engage in a simultaneous move game in which the Attacker attempts to compromise one of the vertices while the Defender tries to protect the network by distributing some scarce defensive resources among the vertices. If the Defender does not protect the vertex that the opponent attacked, the Attacker performs a successful attack and compromises the vertex, receiving its value, while the Defender receives the residual value of the network. On the contrary, the attack fails if the Defender protects a correct node, in which case the Defender keeps the full value of the network only to herself. For instance, if a hacker is trying to compromise the enterprise's network using a particular exploit, the cybersecurity specialist's success depends on whether she/he correctly predicted the attack vector and patched the vulnerable machines. The framework is stated formally and in greater detail at the beginning of Section 4.2.

The game utilises a Subgame Perfect Nash Equilibrium as a solution concept and is solved via the backward induction beginning at the encounter stage. The main technical challenge in the solution occurs in the wiring stage of the game. The stage can be stated as the Defender's maximisation problem in which she chooses a network which yields her the maximum expected payoff accounting for the overall

¹Shodan is a search engine that allows users to search various ICT devices connected to the Internet.

network value and expected loss from an attack in the encounter stage. Formally it can be written as a constrained integer maximisation problem defined over a non-convex set. Problems of this type, in general, do not allow for direct analytical treatment. Thus, to characterise an equilibrium solution, we rely on the convex relaxation technique, which implicates the relaxation of integrality constraints and linearisation of non-convex constraints. This approach yields a simpler maximisation problem defined on a convex hull of the original set while maintaining the original form of the objective function. We then analytically identify the equilibrium solution to the relaxed maximisation problem and verify numerically whether this solution is indeed the equilibrium solution to the original game. The solution to the relaxed maximisation problem and numerical verification can be found in Section 4.3. Convex relaxation procedure is described in great detail in Appendix C.3.

The model produces three main results. First, contrary to the literature, it demonstrates that the centrally protected star does not lead to the highest payoff for the defending side in most circumstances. Instead, it is an extreme formation, which yields the highest possible protection and the lowest possible efficiency among all connected graphs. Second, it reveals the new type of network that is often optimal in games with scarce defensive resources—a *maxi-core* network with a completely connected core and sparse periphery. A *maxi-core* network is an intermediate option that optimally balances between a star network’s security and the efficiency of a complete network. This network allows the Defender to enjoy the efficiency of a completely connected core while extracting the maximum possible value from a periphery without an increase in expected loss from an attack. Third, it shows that the density of the optimal network decreases in the cost of a single link. Additionally, we have developed a weighted model modification, which allows us to account for the possibility that the damage received by the Defender extends beyond the value of the compromised vertex (e.g. reputational losses) or constitute only a share of the value of a compromised vertex (e.g. partial or temporary loss of a node).

The chapter is built as follows. This section provides an introduction and reviews related literature. Section 4.2 formalises the model and offers the analysis of the

final encounter stage of the game. Section 4.3 analyses the overall equilibrium of the game. In Section 4.4 we present the weighted versions of the model. The discussion of limitations and future research is presented in Section 4.5. For those interested in computing the equilibria, related proofs, and supplementary materials, see Appendix C.

4.1.1 Related literature

This study contributes to the literature on network design and defence, which emerged on two pillars of theoretical economics: contest theory and network theory.

Contest theory is concerned with strategic resource allocation in conflict situations; see Konrad (2009) for a comprehensive survey. A considerable share of literature addresses optimal resource allocation over multiple battlefields. Roberson (2006), Hart (2008), Washburn (2013), and Kovenock and Roberson (2008, 2012, 2021) researched the family of Colonel Blotto and General Lotto games, in which players simultaneously distribute limited resources over several battlefields. Each player then obtains the payoff equal to the sum of the valuations of the battlefields she won. Battlefields' valuations in Colonel Blotto games are assumed to be exogenous. On the contrary, our model allows a defender to choose network structure and, consequently, manipulate the values of battlefields (named *vertices* in our setting). An additional strategic element allows us to study the optimal interconnected structures of the battlefields and how the choice of a network structure interacts with contest incentives.

Network theory focuses on strategic network formation, optimal structures, and performance of social and economic networks in diverse circumstances; Jackson (2010), and Goyal (2007) provided perhaps the most comprehensive surveys on the matter. While the network theory takes its origins in the early 17th century, studies on network design and defence in the presence of an intelligent adversary appeared only recently (Naumowicz, 2014).

Arguably, the most famous series of papers on the topic was started by Goyal and Vigier (2010, 2014). In their original framework, a designer and an adversary

compete over several stages in a network formation and defence game. In the first stage, the designer chooses a network formation and distributes defensive resources. The adversary then observes the network and the distribution of defensive resources, allocates contagious attack resources on nodes, and chooses how successful resources should navigate the network. The network in their framework represented the interconnected system of computers, while the attack was modelled after a virus, which navigated through it. Similarly to our research, the authors then studied the trade-off between the connectivity and the increased vulnerability that the connectivity implies. Still, our study has two crucial differences.

Firstly, in the original papers, the authors assumed that all nodes were homogeneous and used a simple cardinality² function to evaluate the network. While useful for studying optimal defensive networks, the approach neglected the influence of network formation on overall graph efficiency, suggesting that a centrally-protected star is an optimal structure in most scenarios. On the contrary, the vertices in our framework are heterogeneous and valued according to their degree centrality—the number of immediate neighbours. The alternative setting confirms that a centrally-protected star is indeed the most secure structure but also demonstrates that it is the least efficient connected network. It can only be optimal for a defending side whenever the cost of a single connection is sufficiently high. Secondly, the original model assumed perfect information, meaning that the attacker knew both the network structure and the defensive resource allocation. On the contrary, we employ a simultaneous game approach in the encounter stage that represents the players' uncertainty about the actions of their corresponding opponents, which is more in line with real-world cybersecurity incidents (Anderson, 2001).

Cerdeiro et al. (2015), Goyal et al. (2016), Dziubiński and Goyal (2017) utilised a very similar setting to the original Goyal and Vigier papers. The authors explored the defence of the edges and decentralised approaches to network protection, employing the same cardinality-based value function as described above. Cerdeiro et al. (2015) and Goyal et al. (2016) assumed every node to be a player, allowing them to

²Cardinality refers to the number of vertices in the network.

unilaterally create connections and choose their own “immunisation” plan. The strategic capabilities of the adversary in the paper by Goyal et al. (2016), however, were limited, enabling him only to choose the type of nodes over which he mixes uniformly at random. Dziubiński and Goyal (2017) studied conflict intensity, denoted as the minimum sum of costs spent by the defender and the attacker. Even though the network was treated as given, the author discovered a new efficient star-like formation—the windmill graph. The decentralised defence is out of the scope of this thesis, but we expand on the ideas of those papers by introducing more sophisticated network value functions and uncertainty about the opponents’ actions.

Gueye et al. (2011) proposed an alternative approach to the problem. In their simultaneous network topology game, a defender chooses a spanning tree of some connected undirected graph while an attacker chooses an edge to attack to disrupt the communication between vertices. The authors assumed that any spanning tree has the same cost (normalised at one) and that the attacker must guess which edges are present in the chosen tree. If the attacker does not guess the correct edge, the game stops, and the network wins. The authors concluded that the attacker mixes uniformly at random over the set of potential critical edges (defined as edges that yield the largest payoffs for the attacker). Similarly to Goyal and Vigier (2014), networks in this setting are differentiated only by defensive capabilities, neglecting the efficiency that the defending side might extract from the structure itself. Moreover, any connected network structure with the same cardinality has the same price for the network designer. Our framework utilises a more sophisticated network value function and assumes that connections are costly. This enables us to study the influence of an edge cost on the optimal choice of the network.

Acemoglu et al. (2016) studied the model of security investments in a network of interconnected agents. The network structure was treated as given, and the main focus was shifted towards the possibility of cascading failures following the exogenous or endogenous attacks, thus not covering our main findings.

The sequential game presented by Hoyer and Jaegher (2016) studied the optimal network structure against two modes of external attack: a link-deletion attack

and a vertex-deletion attack. Contrary to the general literature, they found that star formation is the worst choice against a vertex-deletion attack. Their result is caused by the network designer's inability to protect (or immunise) vertices. The authors, however, discovered that the cost of a link is one of the leading influencers on optimal network formation, which echoes our results.

Studies of optimal network design and defence are not limited to economics literature. For instance, recent papers by Makridis (2021) and Y. Wang et al. (2021) utilised a reinforcement learning approach to solve a very similar design and defence problem. However, due to the limitations of the methodology, the action sets of both the defender and the attacker were limited in both studies. In the former paper, a defender had an initial network and was only allowed to add links, while the attacker could attack only the nodes of the largest value. In the latter paper, the authors used a game-theoretic specification to set up a reward function and recognised the existence of mixed-strategy equilibrium, but did not provide any insights about the optimal network formation.

To conclude, our study contributes to the literature along three dimensions: (1) we introduce a more sophisticated network value function, which recognises the influence of connectivity on overall graph value and allows us to study the tension between security and efficiency; (2) we study the simultaneous version of the game that models the uncertainty of both defending and attacking sides, which is more in line with real-world cybersecurity events; (3) the costly links assumption allows us to study the optimal network formation under different cost regimes.

4.2 The model of network design and defence

4.2.1 Model setup

Two players, the Defender (She) and the Attacker (He), compete in the network defence game over two stages: the initial **wiring stage** and the following **encounter stage**.

In the **wiring stage** the Defender is endowed with a fixed quantity $n > 1$ of vertices (or nodes), which she arranges in some connected graph³ G by creating costly edges (or links) between them. For instance, the vertices can represent the employees of some organisation, while edges are organisational links between them. Similarly, we could think of vertices as computers belonging to the organisational network and edges as literal connections between ICT devices. In our baseline model, we assume that every vertex in the newly arranged graph is valued according to its degree of centrality.

Definition 4.1 (Degree Centrality). *Degree centrality is defined as the number of links incident upon a vertex (i.e. the number of ties that a vertex has).*

The usage of degree centrality represents the idea that more connected vertices are more valuable for organisational integrity and allows us to study the trade-off between the security and efficiency extracted from a network structure. The partially ordered set of vertices' values is represented by vector $\vec{v} = (v_1, \dots, v_n)$. The network's overall value is assumed to be the sum of the value of all vertices minus the cost of all edges. As each additional edge increases the sum of values of all vertices by two, each network must have exactly $\frac{1}{2} \sum_{i=1}^n v_i$ edges. Multiplying the number of edges by the cost of a single link, c , yields the overall cost of the edges.

Assumption 4.1. *The overall value of the networks G is the difference between the sum of values of all vertices and the cost of all edges:*

$$V(\vec{v}) = \left(1 - \frac{c}{2}\right) \sum_{i=1}^n v_i, \quad (4.1)$$

where c is the cost of a single edge.

At the beginning of the **encounter stage**, the Attacker observes the network structure G created by the Defender in the wiring stage. The Attacker can obtain information on the enterprise's organisational structure via social media (e.g. LinkedIn) or perform port scanning to determine the organisation of IP addresses,

³A connected graph is a graph in which there is a path from any vertex to any other vertex.

hosts, ports and their networked structure. The Attacker and the Defender then engage in a simultaneous zero-sum game: the Attacker attempts to compromise one of the vertices of his choosing, while the Defender attempts to prevent that by protecting $\delta \geq 1$ vertices. For example, the Attacker chooses an entry point to exploit in the network, and the Defender chooses an appropriate immunisation plan (e.g. device patching policies). Alternatively, the Attacker might choose an individual employee for a spear-phishing attack, while the Defender decides which employees to prioritise for cybersecurity awareness training.

The attack is successful if the Attacker targets an unprotected vertex, receiving a payoff equal to this vertex's value, while the Defender obtains the overall value of the network minus the value of the compromised node⁴. If the Attacker targets a protected node, the attack fails, and the game finishes. In this case, the Attacker obtains nothing, and the Defender loses nothing and obtains the value of the network in full. The payoff of the Attacker is derived in the following subsection, while the Defender's expected payoff at the beginning of the game is formally stated at the beginning of Section 4.3.

The game utilises Subgame Perfect Nash Equilibrium (SPNE) as a solution concept and assumes that each player rationalises against its opponent's optimal choices. The game is solved via backward induction, beginning at the encounter stage.

4.2.2 Analysis: encounter stage

In the encounter stage **the Defender** already possesses a set of nodes arranged in a connected network G and $\delta \geq 1$ defensive resources, which she uses to protect δ nodes. If some node i is defended, then any attack on this node is unsuccessful. Vertices' values are written as a finite set $\vec{v} = (v_1, \dots, v_n)$ such that $v_1 \geq v_2 \geq \dots \geq v_n$ without loss of generality.

The Attacker insidiously attempts to compromise one of the vertices from the network G . The strategic form of the zero-sum game when $\delta = 1$ is illustrated by Table 4.1.

⁴In Section 4.4 we consider the case where the damage is allowed to be larger or smaller than the value of the compromised node.

Table 4.1: Strategic form game on the network with n vertices and $\delta = 1$, where the number in each cell gives the Attacker's payoff (or the loss of the Defender).

		Defender				
		1	2	3	$\dots$	n
Attacker	1	0	v_1	v_1	$\dots$	v_1
	2	v_2	0	v_2	$\dots$	v_2
	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$
	n	v_n	v_n	v_n	$\dots$	0

Observe from Table 4.1 that the subgame played in the encounter stage does not have an equilibrium with pure strategies. For instance, suppose vertex i is attacked with probability 1. The best response of the Defender is to defend i with probability 1, leaving the Attacker with a payoff of zero. The Attacker then has a profitable deviation to attack a different node. A similar situation occurs when the Defender has multiple defensive resources, $\delta > 1$. In this case, if the Attacker attacks vertex i with probability 1, the Defender's best response is to allocate one of the defensive resources to node i with probability 1. Again, this leaves the Attacker with zero payoff and a profitable deviation to attack a different unprotected node. It follows that the equilibrium strategies of the Defender and the Attacker must be mixed.

Claim 4.1. *The encounter stage subgame does not have an equilibrium with pure strategies.*

Additionally, we assume the scarcity of defensive resources. Observe that if $\delta \geq n$, the Defender's best response is to protect every vertex, even if it is not attacked with a positive probability. In this case, the Attacker is indifferent to any strategy, as he knows that the attack will fail regardless of his actions. From now on, we focus on the more interesting case with $\delta < n$.

Assumption 4.2. *The Defender has $\delta < n$ defensive resources.*

It is now essential to establish the set of vertices attacked with positive probability, referred to later in the chapter as *the Attacker's support*.

Definition 4.2 (Attacker's Support). *The Attacker's support (supp_A) is the set of nodes attacked with positive probability.*

Similarly, we refer to the set of nodes defended with positive probability as the Defender's support, supp_D .

First, observe that in the mixed equilibrium $|\text{supp}_A| = k > \delta$. If it is not the case (and $k \leq \delta$), the Defender can protect all k vertices included in the Attacker's support leaving the latter with zero payoff.

Claim 4.2. *The Attacker always attacks $k > \delta$ vertices with positive probability:*

$$|\text{supp}_A| > \delta.$$

Second, observe that if some vertex is not attacked with positive probability, the Defender never protects it. There is no incentive to protect the vertex, which will certainly not be attacked. Thus, the Attacker's support is always larger or equal to the Defender's.

Claim 4.3. *The Defender protects a vertex with positive probability only if it is attacked with positive probability:*

$$\text{supp}_D \subseteq \text{supp}_A.$$

Now we demonstrate that in any equilibrium, the Attacker mixes over $k > 1$ highest value vertices of the network G . If some vertex with index $m > 1$ in $\vec{v}$ is attacked with positive probability, then every vertex of a higher value is attacked with positive probability as well. The finding is summarised below.

Claim 4.4. *In any equilibrium if node m is attacked with positive probability, so is any node l such that $v_l > v_m$.*

Proof of Claim 4.4. See Appendix C.1 ■

An immediate corollary of Claim 4.4 is that in any organisational structure, the most valuable vertices should always be defended with positive probability in equilibrium. Ensuring the protection of the most valuable employees (e.g. board members, research personnel) and the most connected ICT devices (e.g. servers and routers) can minimise the expected loss from an attack.

Consider now some equilibrium candidate with the top k highest value vertices attacked with positive probability. Suppose that node j is defended with probability q_j and not defended with probability $x_j = 1 - q_j$. Then, the Attacker is indifferent between attacking vertex i and vertex j if:

$$v_i x_i = v_j x_j.$$

Observe that $\sum_{j=1}^k x_j = \sum_{j=1}^k (1 - q_j)$ and $\sum_{j=1}^k q_j = \delta$. It follows that $\sum_{j=1}^k x_j = k - \delta$. Using $x_i = \frac{v_k}{v_i} x_k$ we get:

$$\left(\frac{v_k}{v_1} + \frac{v_k}{v_2} + \cdots + \frac{v_k}{v_k} \right) x_k = k - \delta. \quad (4.2)$$

Rearranging equation (4.2) yields the probability of node k to remain undefended:

$$x_k = \frac{(k - \delta) \prod_{j=1}^k v_j}{v_k \sum_{j=1}^k \prod_{i \neq j}^k v_i}. \quad (4.3)$$

Since x_k is a probability, we must have $x_k \leq 1$, which yields the following necessary condition for k to be in the equilibrium Attacker's support:

$$\frac{(k - \delta) \prod_{j=1}^k v_j}{v_k \sum_{j=1}^k \prod_{i \neq j}^k v_i} \leq 1. \quad (4.4)$$

It is now necessary to check for the monotonicity in k of condition (4.4), such that the condition for k vertices implies the condition for top $k - 1$ vertices. This would guarantee that any pure strategy in the Attacker's support yields him an equal expected payoff and, consequently, that there are no nodes $z < k$ excluded from the Attacker's support and that there is no profitable deviation to smaller support for the Attacker.

Claim 4.5. *The condition (4.4) for the top k vertices implies the condition for the top $k - 1$ vertices.*

Proof of Claim 4.5. See Appendix C.1. ■

Moreover, condition (4.4) is also sufficient for there to exist an equilibrium with k nodes included in the support, as it also guarantees that any pure strategy in

the Attacker's support yields him an equal or higher payoff than the value of any node excluded from the Attacker's support.

The encounter stage of the game follows the “no soft-spots” principle for Colonel Blotto games with multiple targets described by Dresher (1961). The Defender mixes over high-value nodes such that each becomes equally attractive (in expectation) to the Attacker, implying that there are no “soft-spots” among all the protected nodes. Furthermore, the nodes that are never protected with positive probability are of lesser value to the Attacker than the expected gain from an attack on any of the protected nodes. Thus, the Attacker mixes only over nodes protected with positive probability or nodes with a value equal to the expected gain from the attack on any of the protected nodes.

Now we denote the ordered set of values of nodes that are included in the Attacker's support as $\vec{v}_k$. Finally, we can write down the expected loss function for the Defender $D(\vec{v}_k)$, or the expected gain function for the Attacker $A(\vec{v}_k)$ on a given graph G knowing the Attacker mixes over top k highest value nodes:

$$D(\vec{v}_k) = -A(\vec{v}_k) = -\frac{(k - \delta) \prod_{j=1}^k v_j}{\sum_{j=1}^k \prod_{i \neq j} v_i} = -x_k v_k. \quad (4.5)$$

Observe that if condition (4.4) holds strictly for some node k and does not hold for node $k + 1$, the game has a unique equilibrium in which the Attacker and the Defender have identical support $\text{supp}_D = \text{supp}_A$. The equilibrium is unique, as if the Attacker chooses smaller support mixing over $k - z$ nodes, where $z \in \mathbb{Z}$ and $z \in [1, k - \delta)$, then by Lemma 4.2, the Defender's optimal strategy is to protect only $k - z$ nodes with positive probability. In this case, if both players mix over the $k - z$ nodes such that the indifference principle is satisfied for both of them, the Attacker receives a strictly lower expected payoff than from mixing over the top k nodes.

Claim 4.6. *If condition (4.4) holds strictly for some node k , while it does not hold for node $k + 1$, then the encounter stage has a unique equilibrium in which the Attacker and the Defender mix over the top k nodes $|\text{supp}_D| = |\text{supp}_A| = k$ and $\text{supp}_D = \text{supp}_A$.*

Proof of Claim 4.6. See Appendix C.1 ■

However, the equilibrium support might not be unique. If condition (4.4) holds strictly for some node y and with equality for nodes indexed $\nu \in (y, y + z]$, where $z \in \mathbb{Z}$ and $z \in [1, n - y]$, then the Attacker has multiple equilibrium supports, all of which are payoff-equivalent. This echoes a known result about the payoff-equivalence in zero-sum games by von Neumann (1928). We summarise this observation in Claim 4.7.

Claim 4.7. *If condition (4.4) holds strictly for some node y and with equality for nodes indexed $\nu \in (y, y + z]$ then the Attacker has multiple equilibrium supports $|\text{supp}_A| \in [y, y + z]$, all of which are payoff-equivalent.*

Proof of Claim 4.7. See Appendix C.1. ■

Therefore, the size of the Attacker's equilibrium support can be written down as follows:

$$|\text{supp}_A| \in [k, g],$$

where

$$k = \max \left\{ z : \frac{(z - \delta) \prod_{j=1}^z v_j}{v_z \sum_{j=1}^z \prod_{i \neq j} v_i} < 1 \right\},$$

and

$$g = \max \left\{ z : \frac{(z - \delta) \prod_{j=1}^z v_j}{v_z \sum_{j=1}^z \prod_{i \neq j} v_i} \leq 1 \right\}.$$

Now observe that the expected gain from an attack on top k nodes (4.5) is strictly decreasing in δ . It immediately follows that under Assumption 4.2, the Defender utilises all the available defensive resources in any equilibrium.

Claim 4.8. *The Defender utilises all the available defensive resources in any equilibrium.*

Observe that the Defender cannot influence k in the encounter stage. However, as the size of the Attacker's support is a function of vertices' values, it follows that the Defender can manipulate it in the wiring stage when choosing the network formation.

We now show that the Attacker's expected gain in the encounter stage is minimised when the game is played on a star network⁵. At the same time, it attains its maximum when the game is played on a complete network⁶.

Observe first that the Defender's expected payoff is convex over $\vec{v}_k$, or, alternatively, that the Attacker's expected payoff is concave over $\vec{v}_k$, where index k corresponds to the number of nodes included in the Attacker's support.

Lemma 4.1. *For any strategy of the Attacker satisfying Claims 4.5–4.7 and condition (4.4), the expected loss of the Defender/payoff of the Attacker in the encounter stage is convex/concave over $\vec{v}_k = \{(v_1, \dots, v_k) \in \mathbb{R}_+^k \mid 1 \leq v_i \leq n-1\}$.*

Proof of Lemma 4.1. See Appendix C.1. ■

We now establish the results about which networks yield the highest and the lowest expected gain for the Attacker. We demonstrate that the Attacker obtains the highest expected gain if the game is played on a complete network, while he obtains the lowest possible gain if the game is played on a star network.

Proposition 4.1. *The Attacker's expected gain attains its minimum when the game is played on a star network, while it attains its maximum when the game is played on a complete network.*

Proof of Proposition 4.1. See Appendix C.1. ■

The results of Proposition 4.1 should be intuitive. For a fixed number of nodes n , star and path⁷ networks have the smallest possible number of links, limiting the

⁵A star network has a degree distribution in which one node is completely connected, and the rest of the nodes have one connection each, $G_* = (n-1, 1, \dots, 1)$.

⁶A complete network has a degree distribution in which each node is completely connected, $G_C = (n-1, \dots, n-1)$.

⁷The path graph (or network) has two nodes of degree 1, and all other nodes of degree 2.

expected gain of the Attacker. However, unlike in a path network, half of the value of a star is concentrated in a single core node. Thus, the Defender is always better off choosing a star network than a path network since protecting the core node with high probability limits her expected loss while keeping the same overall network value⁸. On the contrary, a complete network has the largest possible quantity of links, meaning that every single node attains the highest possible value. Given that defensive resources are scarce, it provides the Attacker with an opportunity to compromise a completely connected node with high probability, yielding the maximum possible gain. Since the encounter stage is a zero-sum game, it also implies that Defender's expected loss is maximised when the encounter stage is played on a complete network and minimised when it is played on a star network.

These findings echo the results obtained by Goyal and Vigier (2014), who also recognised that a centrally protected star leads to the highest possible payoff for the defending side in a strategic contest versus the intelligent Attacker. The star network is, however, not necessarily an optimal choice for the Defender if we also account for the overall network value—the research proceeds with the analysis of the initial **wiring stage** of the game.

4.3 Equilibrium

In the **wiring stage** the Defender chooses a network formation. She optimises the graph structure by considering the graph's overall value and the equilibrium expected payoff in the encounter stage.

Consider first the expected value of the Defender at the beginning of the game if she chooses some network G .

$$U(\vec{v}) = V(\vec{v}) + D(\vec{v}_k), \quad (4.6)$$

where:

⁸In fact, the expected loss of the Defender in the encounter stage played on a star network never exceeds 1 in the equilibrium.

$V(\vec{v})$ is the overall value of the network, $V(\vec{v}) = \left(1 - \frac{c}{2}\right) \sum_{i=1}^n v_i$ by Assumption 4.1;

$D(\vec{v}_k)$ is the expected loss of the Defender from an attack on k highest value nodes in the encounter stage.

We restrict the cost of an edge to be $c \leq 2$. If $c > 2$, the Defender's payoff is always negative, and she always chooses the network with the least possible number of edges and best defensive capabilities—a star network. The following assumption rules out this possibility.

Assumption 4.3. *The cost of a single link is restricted to $c \leq 2$.*

The Defender's goal in the wiring stage is to maximise the expected value of the game (4.6). In order to do so, she must choose some degree sequence⁹ which maximises her payoff. We call this sequence *a maximising degree sequence*. Note that two networks with identical degree sequences yield the same payoff for the Defender.

Definition 4.3 (Maximising Degree Sequence). *A maximising degree sequence is a degree sequence which maximises the Defender's expected value function.*

It immediately follows from Lemma 4.1 that the Defender's expected utility at the beginning of the game (4.6) is convex over the set of node values' $\vec{v}$ for a given size of the Attacker's support k as the sum of a linear and convex functions.

Claim 4.9. *The Defender's expected value function at the beginning of the game, $U(\vec{v})$, is convex over $\vec{v} = \{(v_1, v_2, \dots, v_n) \in \mathbb{R}_+^n \mid 1 \leq v_i \leq n - 1\}$ for a given size of the Attacker's support k .*

Observe that the objective function of the Defender (4.6) is discontinuous as the expected loss from an attack has a different form depending on how many nodes are included in the Attacker's support. To proceed, we exploit the fact that any network with a given degree distribution can be attributed to one or several

⁹A degree sequence of a graph is some monotonic nonincreasing sequence of positive integer numbers, which represents all of its vertex degrees.

classes with all networks in a given class sharing the same size of the Attacker's support, $|\text{supp}_A|$. Given that the minimum Attacker's support must be $k > \delta$, there exist $n - \delta$ classes of networks, each of which has the same form of the expected loss from an attack. If the encounter stage on some network has multiple equilibrium supports, then this network belongs to several classes with respect to the Attacker's support. However, by Claim 4.7, regardless of which equilibrium support the Attacker chooses (and to which class the corresponding network is attributed), both players always obtain equivalent payoffs.

Thus, we divide the Defender's problem in the wiring stage into two sub-problems:

1. The full support case ($k = n$) in which all nodes pass the Attacker's support condition (4.4), meaning that all nodes will be attacked with positive probability in equilibrium.
2. The partial support case ($k \leq n$) in which k nodes pass the Attacker's support condition (4.4) and the other $n - k$ nodes do not. In this case, top k nodes will be attacked with positive probability in equilibrium, and other $n - k$ nodes will not. The partial support case yields $n - \delta - 1$ similar maximisation problems, which we consider simultaneously.

We can now state the Defender's maximisation problem in a standard form for a given class of networks k . Finding the solutions to each of $n - \delta$ maximisation problems and then finding among them the one that yields the maximum expected payoff for the Defender yields the equilibrium solution for the game.

$$\max_{v_i} \quad \left(1 - \frac{c}{2}\right) \sum_{i=1}^n v_i - \frac{(k - \delta)}{\sum_{i=1}^k \frac{1}{v_i}} \quad (4.7)$$

$$\text{s.t.} \quad v_i \in \mathbb{Z} \quad \forall i \in [1, n], \quad (4.8)$$

$$1 \leq v_i \leq n - 1 \quad \forall i \in [1, n], \quad (4.9)$$

$$\sum_{i=1}^y v_i \leq y(y - 1) + \sum_{i=y+1}^n \min(v_i, y) \quad \forall y \in [1, n], \quad (4.10)$$

$$\frac{k - \delta - 1}{\sum_{j \in K \setminus \{i\}} \frac{1}{v_j}} - v_i \leq 0 \quad \forall i \in [1, k], \quad (4.11)$$

$$- \frac{(k - \delta)}{\sum_{i=1}^k \frac{1}{v_i}} + v_j \leq 0 \quad \forall j \in (k, n], \quad (4.12)$$

where:

K is the set of all vertices in the Attacker's support endogenously defined in the encounter stage;

(4.7) is the objective function of the Defender;

(4.8) is the integer constraint on the values of the nodes;

(4.9) is the set of degree constraints¹⁰;

(4.10) is the set of Erdos-Gallai's sufficient and necessary conditions for the graphicality of a sequence (Tripathi & Vijay, 2003);

(4.11) is the lower boundary for the degrees of vertices inside the Attacker's support derived from condition (4.4);

(4.12) is the upper boundary for the degrees of vertices outside the Attacker's support derived from condition (4.4).

We denote the set of all degree sequences that satisfy constraints (4.8)–(4.12) as Γ .

Erdos-Gallai conditions (4.10) ensure that the degree of some vertex i in degree sequence G does not exceed the maximum possible value given the degrees of all

¹⁰Note that any vertex in a connected network must have at least one edge and at most $n - 1$ edges.

other vertices. If the conditions are satisfied and the sum of all degrees in a sequence is even, the sequence is guaranteed to be graphic. A sequence is said to be graphic if it is possible to construct a graph having the sequence as its degree sequence. For more information about Erdos-Gallai conditions, refer to Appendix C.2.

Constraints (4.11) and (4.12) are imposed by the Attacker's support condition (4.4) for a given class of networks with k top nodes attacked with positive probability. From (4.4) we have that the vertices included in the Attacker's support must be:

$$\frac{(k - \delta)}{\sum_{j=1}^k \frac{1}{v_i}} \leq v_i,$$

or rearranged and simplified:

$$\frac{k - \delta - 1}{\sum_{j \in K \setminus \{i\}} \frac{1}{v_j}} \leq v_i. \quad (4.13)$$

Then the vertices outside the Attacker's support must be:

$$v_j \leq \frac{(k - \delta)}{\sum_{i=1}^k \frac{1}{v_i}}. \quad (4.14)$$

It follows from Lemma 4.1 that conditions (4.12) are convex in standard form and, therefore, must span a convex set. However, by the same lemma, conditions (4.11) are defined by a sum of the concave and linear functions, which might result in non-convexity of Γ even assuming the continuity of nodes' values. Therefore, the maximisation problem (4.7) is an integer maximisation problem defined over a non-convex set. Problems of this kind can be attributed to the family of Mixed-Integer Non-linear Problems (MINLP), which, in general, are known to be non-deterministic polynomial-time (NP) hard problems (Sahinidis, 2019). Therefore, to approach the maximisation problem (4.7), integrability and non-convex constraints must be relaxed.

Roadmap for the analysis We perform the equilibrium analysis in three steps. First, in Subsection 4.3.1, we establish general analytical results about the equilibrium of the game. Second, in Subsection 4.3.2, we analyse the base case of the game with $\delta = 1$ utilising the convex relaxation technique to obtain candidate

solutions for the original maximisation problem (4.7). Third, in Subsection 4.3.3, we check the feasibility of relaxed problem solutions for the original maximisation problem. Finally, in Subsection 4.3.4, we perform a numerical analysis of the Defender's problem to verify that the solutions to the relaxed problem are indeed the solutions to the original problem.

4.3.1 Limiting properties of equilibrium

First, observe that the condition for some network G to be optimal is a function of the cost of a single link c . This is because the Defender's expected utility function at the beginning of the game (4.6) is linear in c . Assuming that we can write down the set of all possible connected networks on n nodes, we then can compare them for a given value c and choose the particular network which yields the maximum utility for the Defender for this c . Performing this exercise for some arbitrary large network is not tractable due to the sheer quantity of possible connected networks. However, we demonstrate that two network structures always occur in equilibrium for extreme values of c regardless of the number of vertices and defensive resources in the Defender's possession: complete and star networks. A complete network emerges as a maximising degree sequence if the cost of a single edge is sufficiently small, $c < c_l$, while a star is an optimal choice if the cost of a single edge is sufficiently high, $c > c_u$. We denote c_u and c_l as higher and lower cost boundaries. This result is summarised in Proposition 4.2.

Definition 4.4. *A cost boundary is a threshold which demarcates two different maximising degree sequences.*

Proposition 4.2. *A complete network is optimal for the Defender if the cost of a single edge is sufficiently small, $c < c_l$. A star network is optimal whenever the cost of an edge is sufficiently high, $c > c_u$.*

Proof of Proposition 4.2. See Appendix C.1. ■

Figure 4.2 illustrates the principles behind Proposition 4.2. We take $n = 5$ and $\delta = 2$ and draw 21 lines that correspond to the Defender's utility functions (w.r.t. c) calculated for all of the possible connected networks. For the relatively small quantity of vertices and defensive resources, it is fairly easy to determine all of the potential maximising degree sequences and then verify which ones can serve as maximisers. Observe that in the illustrated case, there are four maximising degree sequences and three corresponding cost boundaries. The quantity of the intermediate boundaries might be significantly increased with an increase in n and δ , but the lower-cost boundary (c_l) and the upper boundary (c_u) will always exist.

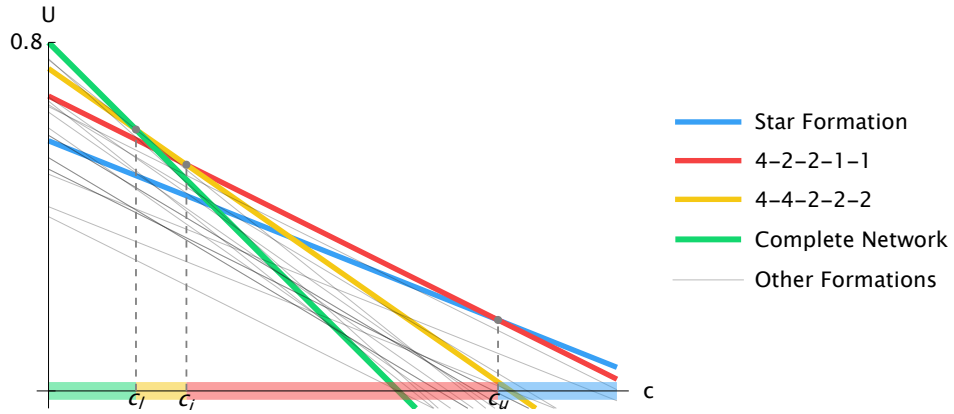


Figure 4.2: The Defender's expected value from a game on five vertices ($n = 5$) with two defensive resources ($\delta = 2$) as a function of a single link cost, c , built for every feasible connected network structure. Highlighted lines are spawned by the networks, which appear as equilibrium solution for some given c . Boundaries on the figure are: $c_l = 1.7$ is the lower cost boundary, $c_u = 1.78$ is the upper boundary, and $c_i = 1.71$ is an intermediate boundary.

In Figure 4.3 we demonstrate the graphs which correspond to maximising degree sequences for $n = 5$ and $\delta = 2$.

Thus, Proposition 4.2 implies that whenever the cost of connection is sufficiently low, the Defender prioritises network connectivity over the potential security concerns—choosing a complete network. This is because low connection costs allow the Defender to extract large benefits from every additional link, which vastly outweigh potential losses from an attack. On the contrary, if the connection cost is

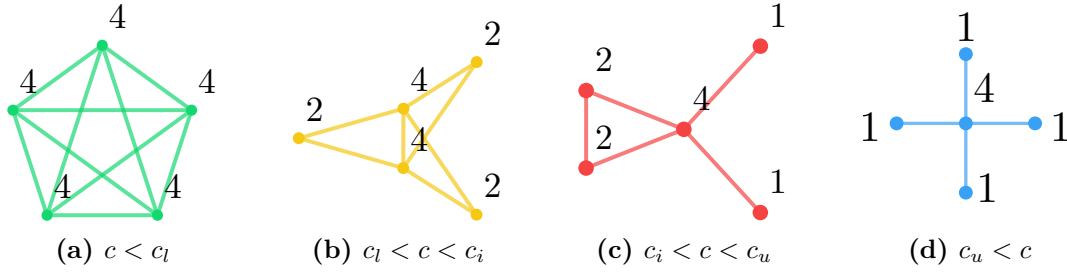


Figure 4.3: All optimal networks for a game with five vertices and two defensive resources. The colors of networks correspond to the colors of lines demonstrated in Figure 4.2.

sufficiently high, the Defender's benefits from additional links are limited and can be easily washed out by an external attack. In this case, the Defender prioritises network security over its connectivity. By choosing a star network and protecting a central node with high probability, the Defender can significantly reduce the potential losses from an attack while still extracting some value from limited connectivity.

Additionally, observe from Figures 4.2 and 4.3 that the higher cost of a single edge is associated with the optimal network structures of a lower density¹¹, while the lower cost of an edge is associated with higher density structures. In Proposition 4.3, we demonstrate that this observation is true in the general case.

Proposition 4.3. *The density of equilibrium degree sequence is a monotonically decreasing function of the cost of a single edge, c .*

Proof of Proposition 4.3. See Appendix C.1. ■

Propositions 4.2 and 4.3 suggest that the Defender evaluates the added benefits of connectivity against an increase in potential damage from an attack that the connectivity brings. Moreover, the density of an optimal network is a decreasing function of the cost of a connection, with the star and complete networks being extremes of this function. Thus, an organisation can optimise its network structure by balancing additional connectivity benefits against security concerns. Those results

¹¹Network density is the portion of the *potential edges* in a graph that are actual edges. *Potential edge* is the edge that might exist between two nodes—regardless of whether it actually exists or not. Any network on n vertices has $\frac{n(n-1)}{2}$ potential edges. A network with n vertices and e edges has a density of $\frac{2e}{n(n-1)}$.

echo the findings of Slikker and van den Nouweland (2000), who have demonstrated that, in the general case, the higher cost of a connection in communication networks is associated with structures of lower density.

We now demonstrate that when the Defender possesses some arbitrary finite number of defensive resources, $\delta \geq 1$ and a large number of vertices $n \rightarrow \infty$, the Attacker can only compromise some infinitesimal share of the network. Therefore, if the Defender possesses a large number of nodes, $n \rightarrow \infty$, network effects outweigh the potential loss from an expected attack, and it is always optimal for her to choose a complete network. We summarise this claim below.

Proposition 4.4. *If the Defender possesses $n \rightarrow \infty$ vertices and some finite number of defensive resources $\delta \geq 1$, the game has a unique equilibrium solution in which the Defender chooses a complete network.*

Proof of Proposition 4.4. See Appendix C.1. ■

Thus, if some organisation has a very large number of nodes (e.g. a large international corporation) and has limited defensive resources, an attack is successful with a probability which tends to 1. However, the benefits extracted from the connectivity of a large network are disproportionally larger than any potential damage the organisation might encounter. In this case, a connected network yields the largest possible payoff for the organisation. However, note that this result might not hold if the upper boundary of a loss from an attack exceeds the value of the largest node. In Section 4.4, we offer a modified version of the model, which can account for this possibility.

To conclude, we obtained three general properties of the equilibrium: (1) if the Defender has some finite number of nodes, a star network is always optimal for a sufficiently high cost of a single edge, while a complete network is always optimal for a sufficiently small cost of a single edge; (2) the density of maximising degree sequence monotonically decreases with respect to a cost of a single edge and

(3) a complete network always yields the maximum payoff for the Defender if $n \rightarrow \infty$.

To analyse the equilibrium further we now consider the base case of the game in which $\delta = 1$.

4.3.2 Relaxed problem solution, $\delta = 1$

Since the original maximisation problem of the Defender in the wiring stage (4.7) is defined over a non-convex set, it cannot be solved utilising standard convex optimisation methods. Thus, to proceed, we utilise a convex relaxation method widely used to find approximate solutions to MINLP problems (Pulleyblank, 1989; Tuy & Van Thuong, 1988).

Convex relaxation is a modelling technique in which some of the constraints of the original problem are relaxed, extending the objective function to the larger convex space. We perform convex relaxation in two steps: (1) we lift the integrality constraints (4.8) and Erdos-Gallai constraints (4.10) which are inherently integer, and (2) we linearise non-convex constraints (4.11).

1. Integer relaxation Lifting integrality constraints (4.8) and Erdos-Gallai graphicality constraints results in the following relaxed maximisation problem:

$$\max_{v_i} \quad \left(1 - \frac{c}{2}\right) \sum_{i=1}^n v_i - \frac{(k - \delta)}{\sum_{i=1}^k \frac{1}{v_i}} \quad (4.15)$$

$$\text{s.t.} \quad 1 \leq v_i \leq n - 1 \quad \forall i \in [1, n], \quad (4.16)$$

$$\frac{k - 2}{\sum_{j \in K \setminus \{i\}} \frac{1}{v_j}} - v_i \leq 0 \quad \forall i \in [1, k], \quad (4.17)$$

$$- \frac{(k - 1)}{\sum_{i=1}^k \frac{1}{v_i}} + v_j \leq 0 \quad \forall j \in (k, n]. \quad (4.18)$$

We denote the set over which the maximisation problem (4.15) is defined as Υ , $\Gamma \subseteq \Upsilon$.

2. Convex hull relaxation Observe that set Υ might still be non-convex due to the presence of non-convex Attacker's support conditions (4.17). In order to “convexify” the problem and preserve maximum information about the original set, we construct a convex hull of the set over which optimisation problem (4.15) is defined. Convex hull of Υ is defined as the smallest convex compact set $\text{Conv}(\Upsilon)$ such that $\Upsilon \subseteq \text{Conv}(\Upsilon)$. We then maximise the objective function (4.15) over the convex hull of set Υ .

Since non-convex regions Υ are defined by smooth concave functions, the convex hull can be constructed by linearising constraints (4.17) (Boyd & Vandenberghe, 2004). Thus, in order to convexify Υ , we perform the following procedure: (1) we find all extreme points of non-convex regions of set Υ , (2) we find corresponding hyperplane equations spanned by the extreme points of non-convex regions of Υ , and (3) we replace non-convex constraints with linear constraints defined by hyperplane equations. The procedure is described in great detail in Appendix C.3. Figure 4.4 illustrates an example where $n = 4$ and the Attacker has full support.

The convexification of the set of feasible values yields an optimisation problem which requires a maximisation of a convex function over a convex compact set. Therefore, convex relaxation allows us to utilise the Bauer maximum principle (Bauer, 1958).

Theorem 4.1 (Bauer maximum principle). *Any function that is convex and continuous, and defined on a compact convex set, attains its maximum at some extreme point of that set.*

Since the objective function of the original maximisation problem is convex by Claim 4.9 and continuous under integer relaxation for a given class of networks, and $\text{Conv}(\Upsilon)$ is convex and compact by definition, it must then be maximised at some extreme point of a convex hull. If any extreme point of $\text{Conv}(\Upsilon)$ is feasible for the original problem (4.7), then this point is a maximising degree sequence for the original problem (Geoffrion, 1971).

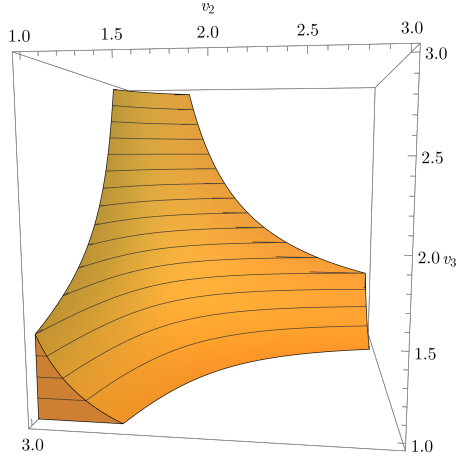
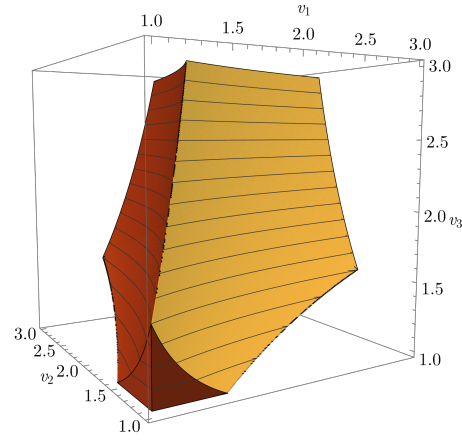
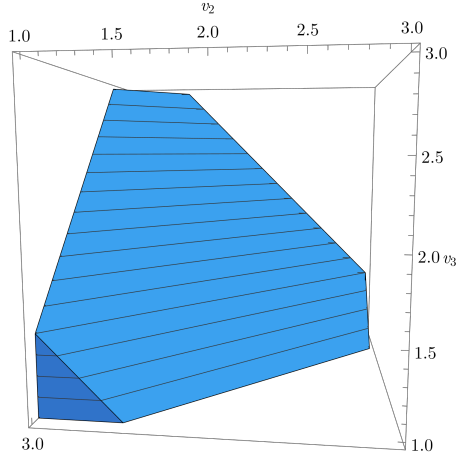
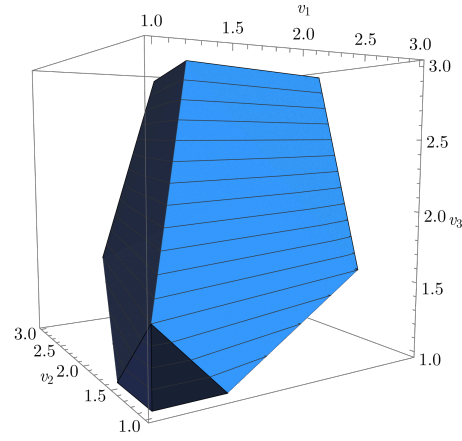
(a) $\Upsilon^4, v_4 = 1$ (b) $\Upsilon^4, v_4 = 3$ (c) $\text{Conv}(\Upsilon^4), v_4 = 1$ (d) $\text{Conv}(\Upsilon^4), v_4 = 3$

Figure 4.4: The set of feasible values, Υ^4 , for nodes 1, 2, and 3 along the edge v_4 , (a) and (b); and the convex hull of set of feasible values, $\text{Conv}(\Upsilon^4)$, for nodes 1, 2, and 3 along the edge v_4 , (c) and (d). More details about this example can be found in Appendix C.3.2.

However, relaxation techniques might lead to sub-optimal maximising degree sequences, as by relaxing integer and non-convex constraints, we might lose some of the solutions that appear on non-convex regions of the original set. Thus, to back up our analysis, we numerically verify that the solution to the relaxed problem is indeed the optimal (equilibrium) solution for networks of up to 20 nodes in Subsection 4.3.4. The discussion on the limitations of relaxation techniques can be found in Subsection 4.5.

Also, observe that the original sets for full and partial support cases have different

sets of constraints: in the partial support case, constraints (4.18) are active, while in the full support case, they are not. This implies that sets of extreme points for partial and full support cases might differ. Thus, we consider relaxed problems for both cases separately. To distinguish the cases, we denote the sets of feasible values for the full and partial support cases as Υ^C and Υ^P , respectively.

Characterising extreme points for the full support case and every partial support case and then finding the extreme points which yield the maximum payoff for the Defender for a given c yields a set of all maximisers for the relaxed maximisation problem.

Full support We state the relaxed maximisation problem for the full support case as:

$$\max_{v_i} \quad \left(1 - \frac{c}{2}\right) \sum_{i=1}^n v_i - \frac{(k - \delta)}{\sum_{i=1}^k \frac{1}{v_i}} \quad (4.19)$$

$$\text{s.t.} \quad v_i \in \text{Conv}(\Upsilon^C). \quad (4.20)$$

The set of constraints which define $\text{Conv}(\Upsilon^C)$ can be found in Appendix C.3.2.

In Appendix C.3.1 we show that the set $\text{Conv}(\Upsilon^P)$ has extreme points that fall into one of four families:

1. Star network family:

$$\left(n - 1, \underbrace{1, \dots, 1}_{\times(n-1)} \right); \quad (4.21)$$

2. Complete network family:

$$\left(\underbrace{n - 1, \dots, n - 1}_{\times n} \right); \quad (4.22)$$

3. Intermediate family:

$$\left(\underbrace{n - 1, \dots, n - 1}_{\times q}, \underbrace{\frac{(n - 1)(q - 1)}{q}, \dots, \frac{(n - 1)(q - 1)}{q}}_{\times(n-q)} \right), \quad (4.23)$$

where $q \in [2, n - 1]$;

4. Lower family:

$$\left(n-1, \frac{n-1}{n-2}, \underbrace{1, \dots, 1}_{\times(n-2)} \right). \quad (4.24)$$

Partial support We first demonstrate that any node excluded from the Attacker's support must attain the maximum possible value that satisfies constraints (4.18). To demonstrate that, we restate the objective function of the original maximisation problem to reflect that, in the partial support case, the network has two types of nodes: (i) nodes that are always attacked with a positive probability denoted v^k and (ii) nodes which are not attacked with positive probability v^s :

$$\max_{v_i^k, v_j^s} U^P = \sum_{i=1}^k v_i^k + \sum_{j=k+1}^n v_j^s - \frac{c}{2} \left(\sum_{i=1}^k v_i^k + \sum_{j=k+1}^n v_j^s \right) - \frac{(k-1)}{\sum_{i=1}^k \frac{1}{v_i^k}}, \quad (4.25)$$

Since in the partial support case, nodes indexed $j > k$ are not included in the Attacker's support, the Defender is always better off choosing the maximum possible value for them. Since the upper boundary for nodes outside the Attacker's support is defined by convex constraint (4.18), it follows that it is always optimal for the Defender to choose the value for nodes indexed $j > k$ such that (4.18) binds. In this case, the Attacker is indifferent between adding nodes of value v^s to his support or not, as by Claim 4.7 in both cases, his expected payoffs are equivalent. Therefore, by choosing the maximum possible value for the nodes outside the Attacker's support, the Defender increases the overall value of the network without increasing her expected loss from an attack.

Lemma 4.2. *Consider the subset of networks with partial Attacker's support, $k < n$. In equilibrium, any node outside the Attacker's support must attain a value equal to the expected gain of the Attacker from an attack on top k nodes.*

Proof of Lemma 4.2. See Appendix C.1. ■

We now state the relaxed maximisation problem for the partial support case as:

$$\max_{v_i} \quad \left(1 - \frac{c}{2}\right) \sum_{i=1}^n v_i - \frac{(k - \delta)}{\sum_{i=1}^k \frac{1}{v_i}} \quad (4.26)$$

$$\text{s.t.} \quad v_i \in \text{Conv}(\Upsilon^P). \quad (4.27)$$

The set of constraints which define $\text{Conv}(\Upsilon^P)$ can be found in Appendix C.3.4.

In Appendix C.3.3 we demonstrate that in addition to extreme point families (4.21)–(4.24), set $\text{Conv}(\Upsilon^C)$ yields three more:

1. k -family:

$$\left(\underbrace{\frac{k}{k-1}, \dots, \frac{k}{k-1}}_{\times k}, \underbrace{1, \dots, 1}_{\times (n-k)} \right); \quad (4.28)$$

where $k \in [2, n-1]$.

2. Quasi-star family:

$$\left(n-1, \underbrace{\frac{k}{k-1}, \dots, \frac{k}{k-1}}_{\times (k-1)}, \underbrace{\frac{k-1}{\frac{(k-1)^2}{k} + \frac{1}{n-1}}, \dots, \frac{k-1}{\frac{(k-1)^2}{k} + \frac{1}{n-1}}}_{\times (n-k)} \right); \quad (4.29)$$

where $k \in [2, n-1]$.

3. Lower partial family:

$$\left(n-1, \frac{k(n-1)}{kn-2k-n+1}, \underbrace{\frac{k}{k-1}, \dots, \frac{k}{k-1}}_{n-2} \right), \quad (4.30)$$

where $k \in [3, n-1]$.

Global maximum We can now find all global maximisers for the relaxed maximisation problem. Via the Bauer maximum principle, the Defender must attain the maximum payoff by choosing one of the sequences that belong to families (4.21)–(4.24) or (4.28)–(4.30). We also know from Lemma 4.2 that the maximising degree sequence is conditional on the cost of a single edge, c . Thus, by comparing

the Defender's expected payoffs from games on sequences defined by extreme points of sets Υ^C and Υ^P for a given c , we can obtain all of the maximisers for the relaxed maximisation problem.

In Lemma 4.3, we demonstrate that the relaxed maximisation problem for the base case with $\delta = 1$ has three global maximisers: (1) a complete network if the cost of an edge is sufficiently low, $c < c_l$, (2) a sequence of an intermediate family with $q = 2$ if the cost of an edge attains intermediate values, $c_l < c < c_u$, and (3) a star network if the cost of an edge is sufficiently high, $c_u < c < 2$.

Lemma 4.3. *The relaxed maximisation problem has the following solutions: the optimal network for the Defender is a complete network if $c < c_l$, a sequence of the intermediate family with $q = 2$ if $c_l < c < c_u$, and a star network if $c_u < c < 2$, where $c_l = 2 - \frac{2}{n}$ and*

$$c_u = 2 - \frac{2n - 4}{n^2 - 2n + 2}.$$

The Defender is indifferent between a complete network and a sequence of the intermediate family with $q = 2$ if $c = c_l$, and between a sequence of the intermediate family with $q = 2$ and a star network if $c = c_u$.

Proof of Lemma 4.3. See Appendix C.1 ■

However, the solutions for the relaxed maximisation problem might not be feasible for the original problem since integrality (4.8) and Erdos-Gallai (4.10) constraints were relaxed. In the following subsection, we check whether the integrality and Erdos-Gallai constraints are satisfied by candidate solutions described in Lemma 4.3.

4.3.3 Feasibility of solutions

By Lemma 4.3 the relaxed maximisation problem yields three potential maximisers for the original problem: (1) a star network, (2) a complete network, or (3) an intermediate family sequence with two nodes of value $v^q = n - 1$. While a star network and a complete network are always graphic, it is not always the case for

the sequence of the intermediate family (4.23) with $q = 2$. For instance, consider the sequence of the intermediate family (4.23) with $n = 7$ and $q = 2$. This sequence has two nodes of value 6 and five nodes of value 3. Since the overall sum of degrees in this sequence is odd ($2 \times 6 + 5 \times 3 = 27$), the sequence is not graphic.

In this subsection, we demonstrate that when the intermediate sequence with $q = 2$ is not feasible, the Defender then chooses between star, complete, or a sequence of the intermediate family with the minimum possible q such that this sequence is graphic. In Subsection 4.3.4 we numerically verify that this result holds for networks of up to 20 nodes.

First, we rule out the possibility that when the sequence of the intermediate extreme point family with $q = 2$ is not feasible, the Defender chooses a sequence that belongs to one of the families (4.24) or (4.28)–(4.30). This is because the majority of extreme points that belong to families (4.24) and (4.28)–(4.30) yield sequences that are not graphic (or even integer). The only exception is a degree sequence that belongs to k -family (4.28) with $n = 4$, but it is always dominated by a star network. Therefore, it cannot be an equilibrium solution to the original maximisation problem.

Claim 4.10. *Sequences that belong to a lower family (4.24), quasi-star family (4.29), and lower partial family (4.30) are never graphic; the k -family of extreme points (4.28) yields one sequence that is graphic when $n = 4$ and $k = 2$, but a star network always dominates this sequence.*

Proof of Claim 4.10. See Appendix C.1. ■

However, if a sequence of the intermediate family with $q = 2$ is not graphic, it might still be possible for the Defender to choose another sequence that belongs to the intermediate family with $q \geq 2$ and graphic. Moreover, any of those sequences dominate star and complete networks for intermediate values of c .

Lemma 4.4. *The Defender prefers a feasible intermediate sequence to star and complete networks if $c_l < c < c_u^q$, where $c_l = 2 - \frac{2}{n}$ and*

$$c_u^q = 2 - \frac{(n-1)(6q-4) - 2n^2(q-1)}{(n^2 - 2n + 2)(n(q-1) - q)}.$$

Proof of Lemma 4.4. See Appendix C.1. ■

Moreover, the Defender is always better off choosing an intermediate sequence with the lowest possible amount of nodes of value $v^q = n - 1$. To see this, consider the following maximisation problem in which the Defender chooses the number of nodes of value v^q , optimising the expected value at the beginning of the game. Substituting $v^k = n - 1$ and $v^s = \frac{(q-1)(n-1)}{q}$ into the objective function (4.25) yields the following univariate maximisation problem:

$$\begin{aligned} \max_q \quad & U_I = -\frac{(n-1)(n(c-2)(q-1) + cq - 2)}{2q} \\ \text{s.t.} \quad & 2 \leq q \leq n. \end{aligned} \tag{4.31}$$

Choosing the minimum possible quantity of fully connected nodes when the game is played on a intermediate family sequence always yields the highest payoff for the Defender.

Lemma 4.5. *The Defender's expected payoff from a game on a intermediate family sequence (4.23) increases in quantity of completely connected nodes q if $c < 2 - \frac{2}{n}$, and decreases in q if $c > 2 - \frac{2}{n}$. The Defender is indifferent if $c = 2 - \frac{2}{n}$.*

Proof of Lemma 4.5. See Appendix C.1. ■

It follows from Lemmas 4.4 and 4.5 that whenever a sequence of the intermediate family is optimal for the Defender, she must choose the minimum possible number of fully connected nodes q .

We now analyse the conditions under which intermediate family sequences are graphic and feasible. We say that if some intermediate sequence is graphic, it yields a *maxi-core network*.

Definition 4.5 (Maxi-core Networks). *A q -maxi-core network is a network which has q vertices of value $v^q = n - 1$, which constitute the core and $n - q$ peripheral vertices of value $v^b = \frac{(n-1)(q-1)}{q} \in \mathbb{Z}$. We refer to a q -maxi-core network with minimum possible number of completely connected vertices $m = \min \left\{ q : \frac{(n-1)(q-1)}{q} \in \mathbb{Z} \right\}$ as an m -maxi-core network.*

The example of an m-maxi-core network on nine nodes is demonstrated in Figure 4.5. The illustrated network has two core nodes, $q = 2$, and seven peripheral nodes.

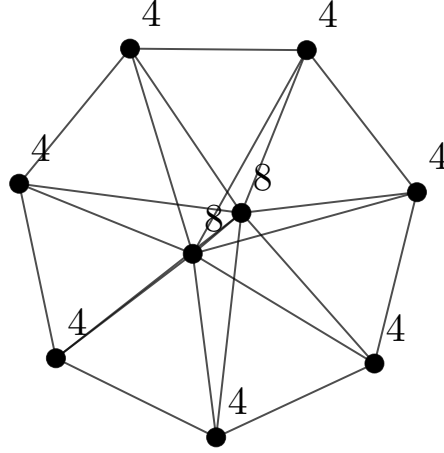


Figure 4.5: M-maxi-core network on nine vertices with two vertices in the core.

In Appendix C.4 we demonstrate that any sequence of the intermediate extreme point family always satisfies Erdos-Gallai constraints (4.10) if $q \in [2, n-2]$. However, q -maxi-core networks are feasible if and only if the Defender has an odd number of nodes¹². If the Defender has an even number of nodes, q -maxi-core networks are never feasible since for any $q \in [2, n-2]$, the value of a peripheral node is either not an integer, $v^b = \frac{(n-1)(q-1)}{q} \notin \mathbb{Z}$, or if it is an integer then the sum of all degrees is odd, which results in a non-graphic sequence.

It follows that if the Defender has an even number of nodes, the only two options that she has are a star or a complete network.

Proposition 4.5. *If the Defender has an even number of nodes and chooses among all feasible sequences which appear as extreme points of $\text{Conv}(\Upsilon^C)$ and $\text{Conv}(\Upsilon^P)$, a star network yields the largest payoff if $c > c_i$, and a complete network if $c < c_i$, where*

$$c_i = 2 - \frac{2(n^2 - 2 + 1)}{n(n^2 - 2n + 2)}$$

The Defender is indifferent between a star and complete network if $c = c_i$.

¹²For instance, it is always possible to create a network with $q = \frac{n-1}{2}$. In this case, any peripheral node has value $v^b = n-3$. Such a sequence is always is always graphic as it satisfies Erdos-Gallai constraints (4.10) and the sum of the degrees is even.

Proof of Proposition 4.5. See Appendix C.1. ■

If the Defender has an odd number of nodes, she also has an intermediate option—an *m-maxi-core network*. Utilising the results of Lemma 4.5 and Proposition 4.3, we can formalise the optimality conditions for all feasible networks of the Defender in case she has an odd number of nodes.

Proposition 4.6. *If the Defender has an odd number of nodes and chooses among all feasible sequences which appear as extreme points of $\text{Conv}(\Upsilon^C)$ and $\text{Conv}(\Upsilon^P)$, a star network yields the largest payoff if $c > c_u^m$, an m-maxi-core network if $c_u^m > c > c_l$, and a complete network if $c < c_l$, where $c_l = 2 - \frac{2}{n}$ and*

$$c_u^m = 2 - \frac{(6m - 4)(n - 1) - 2(m - 1)n^2}{(n^2 - 2n + 2)((m - 1)n - m)},$$

where $m = \min \left\{ q : \frac{(n-1)(q-1)}{q} \in \mathbb{Z} \right\}$.

The Defender is indifferent between a star network and an m-maxi-core network if $c = c_u^m$ and between an m-maxi-core network and a complete network if $c = c_l$.

Hence, if an edge is sufficiently cheap, it is optimal for the Defender to have a complete network, despite the potential high loss from an attack, as the loss can be balanced out by value induced by increased connectivity. On the contrary, when the cost of a link is too expensive, it is efficient for the Defender to choose the least connected network with the highest defence capabilities—a star network. If the cost of a connection is moderate, the complete network is still too expensive, but the Defender can increase the efficiency by choosing an m-maxi-core network.

Propositions 4.5 and 4.6 also demonstrate that contrary to the literature (e.g. Goyal & Vigier, 2010, 2014), a star network does not yield the highest payoff for the Defender in most scenarios and appears as the equilibrium only when the cost of connection is sufficiently high.

In the following subsection, we numerically verify that the optimal solution described in Propositions 4.5 and 4.6 are indeed the solutions to the original game.

4.3.4 Numerical verification

Equilibrium analysis in the previous subsection relies on the results obtained with the help of the convex relaxation method. To back up our findings, we offer a numerical verification of our results in this subsection. We perform an iterative maximisation procedure to find degree sequences that maximise the Defender's expected payoff at the beginning of the game for a given number of nodes, n , and a given cost of an edge, c . We describe the numerical optimisation procedure for a given number of nodes below.

1. We compose the set of degree sequences, D_G , of all known complete graphs based on the *Wolfram GraphData database* (Wolfram Research, 2007). Each entry in the data set is a degree sequence corresponding to some connected graph with n nodes. For $n \leq 10$ the database is exhaustive.
2. We assign each entry of the data set D_G to a certain class of networks based on the minimum size of the Attacker's support $|\text{supp } A| = k$ utilising condition (4.4).
3. We then create a set of corresponding expected utility functions, D_U . Each entry of D_U is the Defender's expected utility function of the form:

$$U(c) = \left(1 - \frac{c}{2}\right) \sum_{i=1}^n v_i - \frac{(k-1)}{\sum_{i=1}^k \frac{1}{v_i}}, \quad i \in [1, n].$$

where the size of the Attacker's support k corresponds to the class of the networks the entry belongs to.

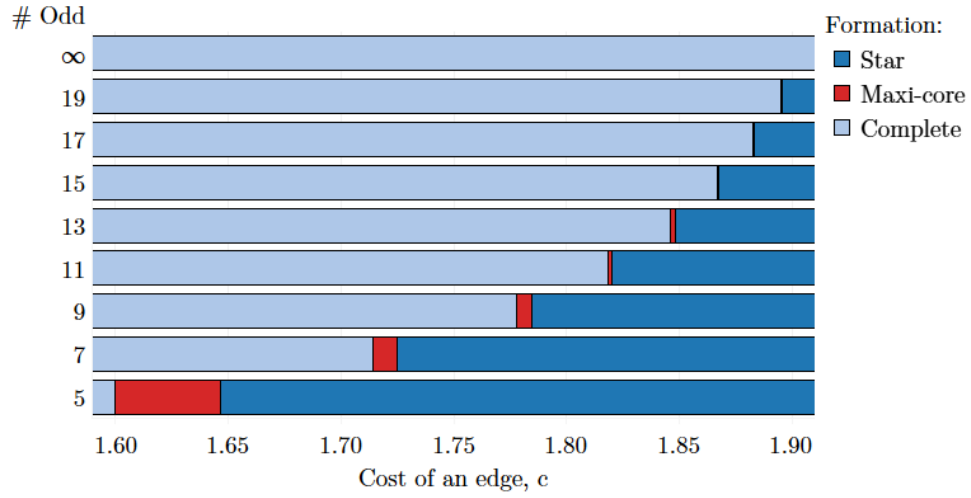
4. We then iterate the cost of an edge, c , with a stepsize 10^{-5} and find the entry in a set D_U which attains the maximum value for a given c , $\max \{D_U : c\}$.
5. The entry with a maximum value for a given c is then traced back to the original degree sequence in D_G .
6. Finally, we create a data set of all global maximisers for the corresponding ranges of cost value c .

The iterative numerical procedure allows us to efficiently identify global maximisers for the original integer non-linear maximisation problem (4.7) for networks with $n \leq 20$ nodes. The results of the numerical optimisation are illustrated in Figure 4.6. Figure 4.6a demonstrates that when the Defender has an odd number of nodes $n \leq 19$ and a single defensive resource ($\delta = 1$), three networks arise in the equilibrium: a complete network if the cost of an edge is sufficiently low, an m-maxi-core network if the cost of an edge is moderate, and a star network if the cost of an edge is sufficiently high. At the same time, if the Defender has an even number of nodes, her choice is limited to either a complete network, if the cost of an edge is sufficiently small, or a star network otherwise. Thus, the numerical analysis confirms that the optima identified in Propositions 4.5 and 4.6 are indeed the solutions to the game if $n \leq 20$.

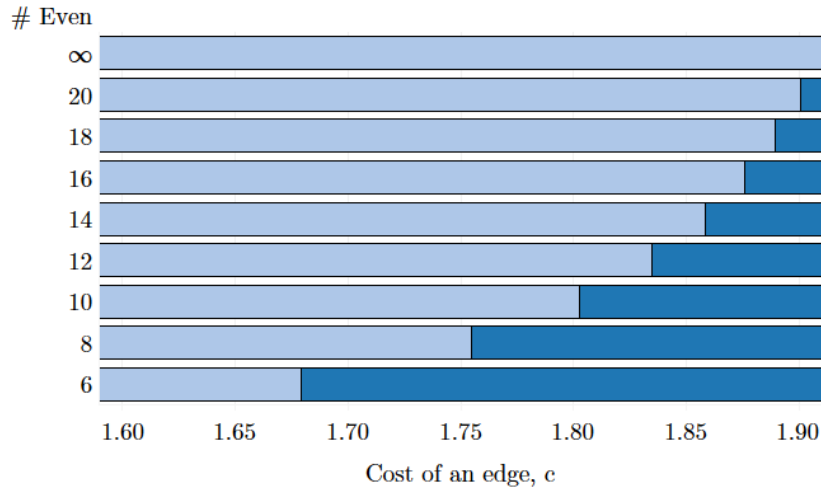
Observe also from Figure 4.6a that the interval of m-maxi-core network optimality diminishes in the number of nodes. For instance, whenever $n = 19$, maxi-core network is an optimal choice for the Defender only if $1.89474 < c < 1.89521$. Thus, while this network might appear in equilibrium when the number of nodes is small, it might not be an optimal choice for the Defender if she possesses a large number of nodes.

It is also evident that the cost range for complete network optimality increases in the number of nodes for networks with both even and odd quantities of vertices, which confirms the result of Proposition 4.4. Thus, with an increasing number of nodes, eventually, a complete network becomes the only optimal choice for the Defender.

To conclude, for the base case with a single defensive resource $\delta = 1$, the solutions to the original problem could be found utilising a convex relaxation technique and convex hull analysis. We observe that if some extreme point of the convex hull of the set of the original problem yields an integer and graphic sequence, this sequence is highly likely to be the solution for the original problem. This echoes the results of Geoffrion (1971), who stated that if a relaxation of some integer maximisation problem yields a solution that is feasible for the original



(a) Maximising degree sequences with an odd number of vertices



(b) Maximising degree sequences with an even number of vertices

Figure 4.6: Results of numerical optimisation. See Appendix C.5 for precise values of cost boundaries and more details on optimal maxi-core networks.

problem and the objective function in both original and relaxed problems are the same, then this feasible solution to the relaxed problem is highly likely to be the solution to the original MINLP problem.

However, the numerical verification must be taken with caution. Firstly, our database is complete only for networks with $n \leq 10$, while if $n \in [11, 20]$, it includes only classes of graphs discovered and extensively described in graph theory literature. Still, for a large number of nodes, the optimality range for the complete network covers almost the entire feasible interval for the cost of a single connection, and,

therefore, a complete network must be the only optimal choice. Secondly, this method might overlook some maximising degree sequence if $\delta \geq 2$; we discuss this limitation in Section 4.5.

4.4 Weighted version of the model

In our benchmark model, in case of a successful attack, the Defender loses the value of a compromised vertex. However, there are two more scenarios to consider: (1) when a successful attack results in a partial node loss rather than a complete loss, and (2) when the damage from a successful attack is much more considerable than the value of vertex alone. Both of those scenarios can be seen in military logistics planning¹³, civil transportation planning¹⁴, and cybersecurity.

In cybersecurity, scenarios in which some node is not lost entirely or is simply disrupted can be observed in cloud computing protection. For instance, a DDoS attack might disrupt a cloud server dealing significant damage, but the disruption is temporary, and the attack does not usually result in a complete loss of a cloud computing node (Deshmukh & Devadkar, 2015; Somani et al., 2016). However, cybersecurity incidents can often result in damage that extends way beyond the value of a damaged node. For instance, Spanos and Angelis (2016), following their systematic literature review on the impact of information breaches on the stock market, revealed that the majority of studies (25 studies out of 28) find a negative impact of data breaches on the victim’s stock price. Moreover, companies may face significant financial losses due to the reputational damage caused by a data breach. Schwartz and Janger (2007) argue that the contemporary legal regime is focused on incentivising business entities to address cybersecurity risks by imposing significant reputational sanctions in case of a breach (e.g. mandatory public disclosure of a breach). Makridis (2021) demonstrates that a data breach could lead to an

¹³For instance, Rogers et al. (2018) utilised simulations and “what if?” approach to analyse efficient military logistics planning. The authors considered a situation in which some nodes and vertices of military logistics networks are either completely destroyed or disrupted by an enemy.

¹⁴Ye and Kim (2019) analysed the vulnerability of heavy rail system networks and considered the influence of complete and partial node failure on the efficiency of the railway networks.

up to 9% decline in reputational intangible capital, which inevitably results in a significant loss of long-run profits.

We introduce the weighted version of the model to account for the possibility of augmentation or reduction of cyberattack damage. In this extension, we allow the expected value function of the Defender to be a weighted linear combination of the overall network value and the expected loss from an attack. By changing the linear coefficients, we can adjust the Defender's expected value function to situations in which an attack's damage extends beyond the value of a vertex or constitutes only a share of the vertex's value. The objective function of the modified Defender's problem then has the following form:

$$\max_{v_i} \quad \alpha \left(\sum_{i=1}^n v_i - \frac{c}{2} \sum_{i=1}^n v_i \right) - (1 - \alpha) \frac{(k - \delta)}{\sum_{i=1}^k \frac{1}{v_i}}, \quad (4.32)$$

where $\alpha \in [0, 1]$.

We now demonstrate that depending on the value of linear coefficient α , the cost boundaries between different networks can shift.

Lemma 4.6. *Any cost boundary, c_i , weakly increases in the weight of the overall network value, α .*

Proof of Lemma 4.6. See Appendix C.1 ■

Given that introduction of the weights does not affect the set of feasible degree sequences Γ , it must be the case that the set of potential maximising degree sequences for the Defender remains the same as in the unweighted case. However, Lemma 4.6 suggests that depending on the weight α , cost boundaries that separate different maximising degree sequences can shift. In practice, if the Defender prioritises the overall network value over the potential losses from an attack ($\alpha > 0.5$), denser networks will appear as maximisers for the broader range of c . For instance, it might happen if the Defender is confident that an attack will not cause any reputation losses and the compromised vertex can be easily replaced (e.g. a cyberattack on some NGO, which does not possess any financial data or trade secrets). On the

contrary, if an event of a cyberattack is capable of significant reputational damage or can even lead to the loss of a competitive edge, the Defender might put a higher weight on an expected loss. In this case, the cost range for which less dense networks are preferred would be extended (e.g. a cyberattack on a financial organisation, which heavily relies on clients' data).

Moreover, the cost range for some networks' optimality could be diminished or increased for the different values of α , as cost boundaries rates of change w.r.t. α might differ. For instance, in the base case with $\delta = 1$, if the Defender puts high weight on an expected loss from an attack ($\alpha < \frac{1}{2}$), the interval of m-maxi-core optimality might be significantly increased.

Claim 4.11. *If the Defender has one defensive resource $\delta = 1$, the optimality range for an m-maxi-core network, $\delta_{MC} = c_u^\alpha - c_l^\alpha$, strictly decreases in α , $\frac{\partial \delta_{MC}}{\partial \alpha} < 0$.*

Proof of Claim 4.11. See Appendix C.1. ■

Thus, it follows from Claim 4.11 that if $\alpha < \frac{1}{2}$ the optimality ranges for m-maxi-core networks can be larger than the ones demonstrated in Figure 4.6a.

Furthermore, observe that the boundaries can shift to the point where only a complete or star network can be optimal for the Defender. When $\alpha \rightarrow 1$, the benefits of building a well-connected network vastly outweigh the security risks, while when $\alpha \rightarrow 0$, the damage from an attack is so severe that the Defender's only concern is network security. This observation also implies that shifting the boundaries through changing the weight α can influence the size of the set of maximising degree sequences.

4.5 Limitations and future research

This section offers a concise discussion of limitations and future research.

4.5.1 Limitations

Section 4.3 of this chapter relies on the convex relaxation technique and characterises the solution for the relaxed Defender's maximisation problem. This technique has two limitations:

Non-convex regions of Γ . By convexifying set Γ and extending the set of feasible values to the convex hull, we potentially might lose some global maximisers that appear on the non-convex regions of the original set. While we numerically verified that the global maximisers for the original problem indeed coincide with the analytical solution to the relaxed problem when $\delta = 1$ and $n < 20$, it might not be true for a more general case. Moreover, any potential error from a convex relaxation approach is likely to be small, because a complete network is optimal for a wide range of c when $n \leq 20$ (e.g. for any $c < 1.9$ when $n = 20$) and will be optimal for an even bigger range of c as n increases further.

Relaxation of Erdos-Gallai graphicality conditions. During the integer relaxation performed as a part of a convex relaxation, we relaxed the integrality constraints (4.8) and Erdos-Gallai sufficient and necessary conditions for the sequence graphicality. While the former should not lead to a loss of solutions since the objective function (4.7) is convex, the latter might result in losing some of the extreme points defined by graphicality conditions. While we numerically verified for $\delta = 1$ and $n < 20$ that it is not the case, it might happen when the Defender has some ad-hoc number of defensive resources $\delta \geq 2$. This is because the expected gain from an attack strictly decreases in δ , implying that the Attacker's support conditions (4.11) are the tightest when $\delta = 1$. Thus, with additional defensive resources and loosening of the Attacker's support condition (4.11), the Erdos-Gallai conditions play a more prominent role. For instance, one can observe that the degree sequence (4, 2, 2, 1, 1) in Figures 4.2 and 4.3 arises as a global maximiser for some intermediate cost of a single edge $c_i < c < c_u$. However, this sequence does

not appear as an extreme point of the set defined by linear constraints (4.9) or Attacker's support constraints (4.11) and (4.12). Therefore, it must be the case that this degree sequence results from the intersection of Erdos-Gallai conditions with some other constraints of the set. It also implies that with additional defensive resources, the problem might have a larger set of maximising degree sequences.

To conclude, the presented approach might not be efficient in solving the Defender's optimisation when $\delta \geq 2$. We see two possibilities on how to approach a more general case:

1. **Numerical estimation.** The iterative numerical estimation demonstrated its efficiency for reasonably small networks. Once one possesses the complete data set of all available graphs on n nodes, there should be no computational limitations for obtaining optimal solutions for a given value of c . However, given the size of the data sets of all connected networks for a larger number of nodes, this method might not be efficient for larger n . In this case, the approximate solutions can be obtained utilising heuristic methods for non-convex and combinatorial optimisation e.g. Eichfelder et al., 2021; Xu et al., 2020. These methods are out of the scope of this chapter.
2. **Scope reduction.** It might also be possible to approach a general case with $\delta \geq 2$ by narrowing down the set of possible networks. For instance, reducing the scope to bipartite graphs¹⁵ or networks which only have two types of nodes will result in a much simpler model, which might have an analytical solution for the general case.

4.5.2 Future research

Large networks The presented model demonstrates that the only network which appears in the equilibrium of a game with a large number of nodes ($n \rightarrow \infty$) is

¹⁵A bipartite graph is a graph whose nodes can be separated into two disjoint sets T_1 and T_2 , i.e. is every edge connects a node in set T_1 to one in set T_2 .

a complete network. This is because a successful attack can compromise only an infinitesimal share of the network's overall value. We consider developing an alternative setting with a more sophisticated vertex value function (e.g. a network size inflated vertex value function $v_i = nd_i$, where d_i is the vertex's degree centrality and n is the number of vertices in the network), which might lead to new results about the formation of large networks.

Multiple offensive resources The baseline model considers the case where the Attacker has a single offensive resource. The framework is helpful for analysing the enterprise's cybersecurity capabilities against targeted information leakage attacks, as it is usually sufficient to compromise a single entry point to deal severe damage to the enterprise. Nevertheless, it does not cover the situation in which the network faces multiple attacks simultaneously (e.g. airstrikes on the military supply chain in several directions). The presented model can be extended to account for this possibility, but finding the equilibrium might be intractable in the general case. This is because if the Attacker has multiple offensive resources, the Attacker's support condition will no longer be monotonic, which makes it challenging to derive the corresponding constraints for the optimisation problem. We consider the development of an alternative attack mechanism in which the Attacker can choose some set of nodes for an attack following a defined rule (e.g. he can choose to attack all nodes with a specific value or choose a share of nodes in the network which will be attacked at random). This extension can provide insights into the network defence against various attack regimes.

Probabilistic defence function The deterministic defence function (when the Defender always wins the contest on the same vertex as the Attacker) might not correspond well to some real-life scenarios (e.g. air traffic security). Our model might be extended by imposing some probabilistic contest success function in the fashion of a rent-seeking contest (Gradstein & Konrad, 1999). We consider this extension for future research.

Edge attacks The baseline model assumes that the Attacker can only compromise nodes. Nevertheless, it is still possible to encounter situations where the attack targets connections between the nodes rather than the nodes themselves. We consider building an extended model which allows the Attacker to attack both edges and nodes to account for this opportunity.

Artificial intelligence Our game-theoretical formulation of the network design and defence problem can be used to create more sophisticated reward functions and serve as a benchmark for comparing reinforcement learning algorithms. We consider this possibility for future research.

4.6 Conclusion

In this chapter, we characterise efficient network formation in the presence of an intelligent attacker. In particular, the study investigates the trade-off between efficiency and security when the values of vertices are determined endogenously and are allowed to be heterogeneous. The model produces three main results. Firstly, it shows that a star network, contrary to the literature (e.g. Goyal & Vigier, 2010), does not generally yield the highest payoff for the Defender. A star network yields the lowest expected loss from an attack. Still, it is the least efficient connected network structure, which appears in equilibrium only if the cost of a single connection is sufficiently high. Secondly, we demonstrate that the density of an optimal network decreases with an increase in the cost of an edge. Finally, the study reveals a novel type of network, which appears in the equilibrium of games with scarce defensive resources—a maxi-core network with a completely connected core and sparse periphery. Moreover, we have presented a weighted model modification, which allows us to account for potential reputational losses from a cyberattack and partial vertex loss.

5

General Conclusion

This dissertation addresses the lack of socio-economic knowledge about industrial cyberespionage. We employ an economic approach to extend the understanding of industrial cyberespionage and its implications for organisational competitive and security incentives, contributing to the emerging field of economics of information security. This research can inform a better cybersecurity policy design, motivating higher-quality innovation, aid in optimising networked structures of enterprises for efficiency and security, and extend current cyberinsurance models to account for research-intensity.

Since industrial cyberespionage is a complex multidimensional phenomenon, it cannot be comprehensively covered in a single work. Instead, we focus on several socio-economic aspects of industrial cyberespionage. First, in Chapter 2, we use the game-theoretical approach to investigate industrial espionage in a competitive R&D setting to understand its motivation and influence on research companies' strategic incentives, payoffs, and social welfare. The theoretical investigation demonstrates that cyberespionage incentives might arise from the difference in relative positions of the competitors and that it has an ambiguous influence on the overall effort exerted in a race by research companies and social welfare. While industrial cyberespionage, in most cases, harms companies' payoffs and social welfare, we demonstrate that there might be an optimal espionage power that maximises companies' payoffs and

the quality of end products. The results of the theoretical investigation are then utilised to inform our empirical study in Chapter 3.

The statistical model presented in Chapter 3 studies the link between research-intensity and information leakage attack rate in a given industry. It reveals a strong positive association between information leakage attack rate and research-intensity indices (in a given industry in a given country), which can be considered as supportive evidence for the theoretical results obtained in Chapter 2. Companies in research-intensive industries are more likely to engage in R&D races, which, in turn, create a fruitful ground for industrial cyberespionage. Finally, in Chapter 4, we study how an organisation can protect itself from the potentially devastating effects of cyberespionage with the help of a new theoretical framework that focuses on the design and defence of organisational networks in the presence of an intelligent attacker. Individual contributions of each chapter are described in greater detail below.

Chapter 2 studies the implications of industrial cyberespionage for research and development races. We develop a dynamic version of Lazear and Rosen's rank-order tournament to study research incentives in R&D races with industrial cyberespionage. The model reveals that industrial cyberespionage has an ambiguous influence on firms' overall investments in the research and development of innovative products. It discourages early investments in cutting-edge technology, as the opponent can steal it. At the same time, it weakens the discouragement effect in a later competition that arises if one firm achieves a technological edge and gets ahead before the final stage of product development. Moreover, industrial cyberespionage has an ambiguous effect on the quality of the product delivered to the market by the winner of the R&D race. Thus, we conclude that aggressive enforcement of cybersecurity policies might have knock-off effects on competition that may be negative rather than positive. This model can inform better cybersecurity and intellectual property policies and laws.

Chapter 3 studies the relationship between an industry's research-intensity, data reliance, and information leakage attacks. Using multivariate fractional outcome regression analysis on Eurostat aggregated data on enterprises' innovative and

digital activity, we provide new evidence that research-intensive industries are particularly susceptible to information leakage attacks. The study also distinguishes two industry-specific associations: high-tech manufacturing industries are prone to experiencing sophisticated targeted attacks, while knowledge-intensive service industries are more likely to fall victim to opportunistic, financially-driven attacks.

Chapter 4 investigates efficient network formation in the presence of an intelligent attacker. We develop a theoretical framework in which two asymmetric players, the Defender and the Attacker, compete over the two stages in a network design and defence game with heterogeneous vertices' values. We assume that nodes are valued according to their degree centrality, allowing the Defender to manipulate the network value and expected losses from an attack by creating costly links. Moreover, we assume that players are not perfectly informed about the actions of each other. The framework focuses on the trade-off between network efficiency and security. The study produces three main results. Firstly, contrary to the literature, a centrally protected star network does not yield the maximum payoff for the defending side in most circumstances, even being the most secure network formation. Secondly, it reveals a new family of networks that often arises in an equilibrium of the network design and defence games with scarce defensive resources—the maxi-core network with a completely connected core and sparse periphery. Thirdly, the model demonstrates that the density of optimal network formations decreases with the cost of a connection. This framework can guide the design of an organisation's social and technical structures, aid the development of the algorithms for automated intrusion detection systems, or be used to develop optimal cybersecurity awareness training plans.

Appendices



Appendices for Chapter 2

Contents

A.1 Mathematical Appendix	162
A.1.1 Equilibrium derivation and analysis	162
A.1.2 Social welfare function derivation and analysis	170

A.1 Mathematical Appendix

A.1.1 Equilibrium derivation and analysis

Proof of Lemma 2.1. Assumption 2.1 and the fact that $c'(\cdot)$ is monotone guarantee the existence of a unique pure strategy equilibrium. It is easy to observe that in pure strategy equilibrium, it must be the case that $\hat{x}_i = \hat{x}_{-i}$ as FOCs (2.9) are symmetric for both players.

Given that $f(\hat{x}_i - \hat{x}_{-i})$ is a symmetric function with zero mean, both companies then choose:

$$\hat{x}_i^s = c'^{-1} [f(0)].$$

SOCs (2.10) are always satisfied around $\hat{x}_i^s$ as $f'(0) = 0$ and $c'(\hat{x}_i) > 0$ for all $\hat{x}_i > 0$ by Assumption 2.1.

Notice that corner solutions $\hat{x}_i = 0$ cannot be an equilibrium since the density functions are assumed to have infinite support (are unbounded) and $c'(0) = 0$. ■

Proof of Lemma 2.2. The proof follows the same pattern as the proof for the symmetric case of the stage. Assumption 2.1 and the monotonicity of a cost function guarantee a single intersect from below of marginal benefits and marginal costs, so the equilibrium is unique. It follows from the symmetry of the FOCs that both companies will choose the same levels of investment:

$$\hat{x}_i^a = c'^{-1} [f(h(1 - \alpha))].$$

Corner solutions $\hat{x}_i = 0$ cannot be an equilibrium since density functions are assumed to have an infinite support (are unbounded) and $c'(0) = 0$.

Assumption 2.1 also guarantees that the set of SOC's is always satisfied:

$$|f'(h(1 - \alpha))| < c''(\hat{x}_i^a). \quad (\text{A.1})$$

■

Proof of Proposition 2.2. The sufficient condition for the proposition to be true is:

$$\frac{\partial}{\partial x} \left(F(x) + c \left[c'^{-1} [f(x)] \right] \right) > 0, \quad (\text{A.2})$$

or:

$$-f(x) > \frac{f(x)f'(x)}{c'' \left[c'^{-1} [f(x)] \right]}.$$

Simplifying yields:

$$c'' \left[c'^{-1} [f(x)] \right] > f'(x), \quad (\text{A.3})$$

which is always true via Assumption 2.1. ■

Proof of Proposition 2.3. Observe that companies are identical at the beginning of the race. Moreover, they have symmetric marginal benefits and identical monotone cost functions. Thus, it must be the case that both companies choose identical equilibrium strategies.

Substituting $x_i = x_{-i}$, (2.12), and (2.20) into (2.25) and simplifying yields:

$$x_i^* = c'^{-1} \left[f(k) \left(2F(h(1-\alpha)) - 1 \right) \right].$$

In order to ensure the uniqueness of the equilibrium, it is sufficient to impose the following set of second-order conditions, which guarantees the single intersect of marginal benefits and marginal costs:

$$f'(x_i - x_i^* + k) \left(\Pi^S - \Pi_2^A \right) + f'(x_i - x_i^* - k) \left(\Pi_1^A - \Pi^S \right) < c''(x_i). \quad (\text{A.4})$$

Inequality (A.4) is satisfied if:

$$\max_{x_i} \underbrace{f'(x_i - x_i^* + k) \left(\Pi^S - \Pi_2^A \right)}_{t^1} + \underbrace{f'(x_i - x_i^* - k) \left(\Pi_1^A - \Pi^S \right)}_{t^2} < \min c''(x_i). \quad (\text{A.5})$$

Term t^1 of inequality (A.5) is always positive given that condition (A.2) of Proposition 2.2 always holds under Assumption 2.1:

$$t^1 = F(h(1-\alpha)) + c \left[c'^{-1} [f(h(1-\alpha))] \right] - \frac{1}{2} - c \left[c'^{-1} [f(0)] \right] > 0.$$

Term t^2 of inequality (A.5) is always positive given the monotonicity of the cost function and the fact that the random variable ζ has a single-peaked distribution with zero mean:

$$t^2 = F(h(1 - \alpha)) - c \left[c'^{-1} [f(h(1 - \alpha))] \right] - \frac{1}{2} + c \left[c'^{-1} [f(0)] \right] > 0.$$

Observe now that inequality (A.5) always holds if:¹

$$\max f'(\cdot) \underbrace{(\Pi_1^A - \Pi_2^A)}_{\rho} < \min c''(x_i), \quad (\text{A.6})$$

where $\rho = 2F(h(1 - \alpha)) - 1$.

It is easy to spot that the maximum value that ρ can attain is 1. It follows that (A.6) always holds whenever:

$$\max f'(\cdot) < \min c''(x_i), \quad (\text{A.7})$$

which is precisely the form of Assumption 2.1.

It immediately follows that SOC's are automatically satisfied under Assumption 2.1.

A corner solution ($x_i^* = 0$) cannot be an equilibrium since the density functions are assumed to have infinite support (are unbounded) and $c'(0) = 0$. It is easy to verify by substituting $x_i = x_{-i} = 0$ into the FOCs.

It is also trivial to demonstrate that a firm's investment in the research stage decrease in an espionage power α :

$$\frac{\partial x_i^*}{\partial \alpha} = -2hf(k)f(h(1 - \alpha))(c'^{-1})' \left[f(k) (2F(h(1 - \alpha)) - 1) \right] < 0.$$

■

¹At this step, we assumed that both $f'(x_i - x_i^* + k)$ and $f'(x_i - x_i^* - k)$ attain their maximum possible values – $\max f'(x_i - x_i^* + k) = \max f'(x_i - x_i^* - k) = f'(\cdot)$.

Proof of Proposition 2.4. The overall investments are defined as the sum of expected investments in the research and the development stages:

$$X \equiv 2x_i^* + 2P^S \hat{x}_i^s + 2P_1^A \hat{x}_i^a + 2P_2^A \hat{x}_i^a. \quad (\text{A.8})$$

Expanding the function and substituting $\Delta x = 0$ yields:

$$\begin{aligned} X = & \frac{2}{c} \left(f(k) \left(2F(h(1-\alpha)) + A - 1 \right) + \right. \\ & \left. + (F(k) - F(-k)) f(0) + 2F(-k) f(h(1-\alpha)) \right). \end{aligned} \quad (\text{A.9})$$

Taking the first derivative of (A.9) w.r.t. to α and setting it to zero yields:

$$\frac{\partial X}{\partial \alpha} = \frac{2h}{c} \left(\underbrace{-f(k)f(h(1-\alpha))}_{<0} \underbrace{-2F(-k)f'(h(1-\alpha))}_{>0} \right) = 0. \quad (\text{A.10})$$

The derivative has an ambiguous sign. It is caused by the fact that espionage power has an opposite effect on investment levels in different stages: an increase in α causes a decrease in optimal investment level during the research stage and an increase during the development stage.

Simplifying and rearranging (A.10) yields the following candidate solution:

$$\alpha^\iota = 1 - \frac{1}{h} \sqrt{\frac{2}{\pi}} \frac{\sigma}{\eta} \exp\left(-\frac{k^2}{2\sigma^2}\right), \quad (\text{A.11})$$

where $\eta = \text{erfc}\left(\frac{k}{\sqrt{2}s}\right)$. The candidate solution is bounded from above; however, it does not have a defined lower bond. Therefore, α^ι might take negative values, which contradicts the initial assumption (that $\alpha \in [0, 1]$). Thus, the interior solution indeed exists only for sufficiently large h :

$$h > \sqrt{\frac{2}{\pi}} \frac{\sigma}{\eta} \exp\left(-\frac{k^2}{2\sigma^2}\right) \equiv \bar{h}.$$

Otherwise, it is optimal to choose $a^\iota = 0$ to maximise the overall investment. Verifying the sign of the second-order condition around a^ι confirms that it is indeed a maximum point:

$$\frac{\partial^2 X}{\partial \alpha^2} = -\sqrt{\frac{2}{\pi}} \frac{h^2 \eta}{\sigma^3} \exp\left(-\frac{\exp(k^2/\sigma^2)}{\pi \eta^2}\right) < 0. \quad (\text{A.12})$$

The solution assuming normally distributed shocks can then be summarised as follows:

$$\begin{cases} \alpha^I = \alpha', & \text{if } h > \sqrt{\frac{1}{2\pi}} \frac{\sigma}{\Phi(-k)} \exp\left(-\frac{k^2}{2\sigma^2}\right), \\ \alpha^I = 0, & \text{otherwise,} \end{cases}$$

where Φ is the CDF of the normal distribution with standard deviation σ and zero mean. ■

Proof of Proposition 2.5. Candidate solution (A.11) increases in h , however the influence of innovation threshold k is not clear. Consider the following equation that was derived from (A.10):

$$\frac{f(k)}{F(-k)} = \frac{f'(h(1-\alpha))}{f(h(1-\alpha))}.$$

Taking the derivative of the left-hand side of the equation w.r.t. k yields:

$$\frac{\partial}{\partial k} \left(\frac{f(k)}{F(-k)} \right) = \frac{f'(k)F(-k) + f(k)^2}{F(-k)^2}. \quad (\text{A.13})$$

I now make use of two facts:

$$f'(k) = -kf(k),$$

and

$$\lambda(-k) = \frac{f(-k)}{1 - F(k)} \quad \forall k.$$

which could be immediately recognised as the inverse Mill's ratio. Simplifying (A.13) yields:

$$\lambda(-k) (\lambda(-k) - k) > 0, \quad (\text{A.14})$$

which is positive for $\forall k > 0$ given that normal distribution is symmetric around zero and the fact the inverse Mill's ratio is the expectation of the truncated normal:

$$E[\zeta | \zeta > k] = \frac{f(k)}{1 - F(k)} = \frac{f(-k)}{F(-k)} = \lambda(-k) > k.$$

Consider now the derivative of the right-hand side of equation (A.13) w.r.t. α :

$$\frac{\partial}{\partial \alpha} \left(\frac{f'(h(1-\alpha))}{f(h(1-\alpha))} \right) = (\alpha - 1)h < 0. \quad (\text{A.15})$$

Combining results (A.14) and (A.15) it is evident that there an inverse relationship between innovation threshold k and optimal attack power $\frac{\partial \alpha^I}{\partial k} < 0$. ■

Proof of Lemma 2.3. Consider Υ that is a simplified version of Π_i (2.28), which includes only members that are the functions of α :

$$\Upsilon \equiv -f(h(1-\alpha))^2 + \kappa F(h(1-\alpha)) F(-h(1-\alpha)), \quad (\text{A.16})$$

where $\kappa = \frac{2f(k)^2}{F(-k)}$. Taking the first derivative and setting it to zero yields:

$$\frac{\partial \Pi_i}{\partial \alpha} = \frac{\partial \Upsilon}{\partial \alpha} = hf(h(1-\alpha)) \left(2f'(h(1-\alpha)) + \kappa(1 - 2F(-h(1-\alpha))) \right) = 0.$$

It is easy to spot that $\alpha = 1$ is a candidate maximum point. It is now necessary to verify a second-order condition:

$$\begin{aligned} \frac{\partial^2 \Upsilon}{\partial \alpha^2} = & -h^2 f'(h(1-\alpha)) \left(2f'(h(1-\alpha)) + \kappa(1 - 2F(-h(1-\alpha))) \right) + \\ & + hf(h(1-\alpha)) \left(-2f''(h(1-\alpha)) - 2\kappa f(h(1-\alpha)) \right) < 0, \end{aligned}$$

which simplifies around $\alpha = 1$ to :

$$\left. \frac{\partial^2 \Upsilon}{\partial \alpha^2} \right|_{\alpha=1} = 2hf(0) \left(-f''(0) - \kappa f(0) \right) < 0. \quad (\text{A.17})$$

It must then be true that:

$$2f(k)^2 > -\frac{f''(0)}{f(0)} F(-k). \quad (\text{A.18})$$

Rearranging (A.18) yields:

$$2f(k)^2 - \frac{f''(0)}{f(0)} F(k) > -\frac{f''(0)}{f(0)}. \quad (\text{A.19})$$

Given that $-\frac{f''(0)}{f(0)} = \frac{1}{\sigma^2}$ inequality (A.19) becomes:

$$K \equiv 2\sigma^2 f(k)^2 - F(k) > 1. \quad (\text{A.20})$$

Consider now the first derivative of K w.r.t. k :

$$\frac{\partial K}{\partial k} = f(k) \left(4\sigma^2 f'(k) - 1 \right), \quad (\text{A.21})$$

which is always negative for any k in $\mathbb{R}^+$. As $k \geq 0$, it must be the case that K attains its maximum at $k = 0$. Evaluating (A.20) at $k = 0$ yields:

$$f(0)^2 > \frac{1}{4\sigma^2},$$

or

$$\frac{1}{2\pi\sigma^2} > \frac{1}{4\sigma^2},$$

which never holds. Therefore, $\alpha^* = 1$ is always a local minimum point. ■

Proof of Proposition 2.6. Consider Υ that is a simplified version of Π_i (2.28), which includes only terms that are the functions of α :

$$\Upsilon \equiv -f(h(1-\alpha))^2 + \kappa F(h(1-\alpha)) F(-h(1-a)), \quad (\text{A.22})$$

where $\kappa = \frac{2f(k)^2}{F(-k)}$. Substituting $H = h(1-\alpha)$ yields:

$$\Upsilon = -f(H)^2 + \kappa F(H)F(-H). \quad (\text{A.23})$$

The first-order condition is then:

$$\frac{\partial \Upsilon}{\partial H} = f(H) (\kappa F(-H) - \kappa F(H) - 2f'(H)) = 0, \quad (\text{A.24})$$

It is evident that $H = 0$ is a candidate solution for the FOC. However, it follows from Lemma 2.3 that $H = 0$ is a point of local minimum. Observe now that:

$$\left. \frac{\partial \Upsilon}{\partial H} \right|_{H \rightarrow \infty} = -\kappa f(H)F(H) < 0. \quad (\text{A.25})$$

Combining equation (A.25) and results from Lemma 2.3 it should be evident that the function $\Upsilon'(H)$ should intersect axis of abscissas at some point $H^* \in (0, \infty)$. Moreover, the dynamics of the function suggest that H^* should be indeed the local maximum point. Given that H is a linear combination of h and α it follows that:

$$\begin{cases} \alpha^* = \frac{h-H^*}{H^*}, & \text{if } h > H^*, \\ \alpha^* = 0, & \text{if } h \leq H^*. \end{cases}$$

It is now sufficient to show that H^* always exists. By rearranging (A.24), observe that the FOC is satisfied whenever:

$$\underbrace{\frac{f'(H)}{1 - 2F(H)}}_{q(H)} = \frac{\kappa}{2}. \quad (\text{A.26})$$

It is easy to spot that function $q(H)$ obtains its maximum at $H = 0$ as $\max f'(H) = f'(0) = 0$ and $1 - 2F(H)$ obtains its minimum positive value when $H = 0$. I now demonstrate that this condition is always satisfied as

$$\max_H q(H) = q(0) = \frac{1}{\sigma^2} > \kappa, \forall H \& \forall k,$$

which follows from Lemma 2.3. Consider also the sign of :

$$q'(H) = - \frac{\exp\left(-\frac{H^2}{\sigma^2}\right) \left(-2H\sigma - \exp\left(-\frac{H^2}{2\sigma^2}\right) (H^2 - \sigma^2) \operatorname{erf}\left(\frac{H}{\sqrt{2}\sigma}\right)\right)}{2\pi\sigma^5 \operatorname{erf}\left(\frac{H}{\sqrt{2}}\right)}. \quad (\text{A.27})$$

Simplifying (A.27) yields:

$$q'(H) \equiv -2H\sigma - \exp\left(-\frac{H^2}{2\sigma^2}\right) (H^2 - \sigma^2) \operatorname{erf}\left(\frac{H}{\sqrt{2}\sigma}\right) < 0, \quad (\text{A.28})$$

which is always negative given that:

$$\max_H q(H) = \lim_{H \rightarrow 0} \left(-2H\sigma - \exp\left(-\frac{H^2}{2\sigma^2}\right) (H^2 - \sigma^2) \operatorname{erf}\left(-\frac{H}{\sqrt{2}\sigma}\right)\right) = 0. \quad (\text{A.29})$$

Therefore, $q'(H) < 0$ and $q(H)$ is a monotone strictly decreasing function of H . Given that the right-hand side of (A.26) is a constant, there must be a unique value H^* that satisfies the equation, and H^* must be a global maximum. ■

Proof of Proposition 2.7. Rearranging (A.26) yields:

$$Q \equiv F(-k)q(H) - f(k)^2 = 0.$$

I now implicitly differentiate Q to obtain:

$$\frac{\partial H^*}{\partial k} = - \frac{\frac{\partial Q}{\partial k}}{q'(H)F(-k)}. \quad (\text{A.30})$$

It follows from Proposition 2.6 that the denominator of the equation is negative. Therefore, the sign of (A.30) depends solely on the sign of the numerator, which, however, is not obvious. We consider it separately:

$$\frac{\partial Q}{\partial k} = f(k) \left(-q(H) - 2f'(k) \right).$$

The sign of the equation above depends on the value of $q(H)$. Consider two possibilities:

1. if H^* is sufficiently small and $q(H^*) \in \left(\frac{\sqrt{2}}{\sqrt{e\pi}\sigma^2}, \frac{1}{2\sigma^2} \right)$ then $Q'(k) < 0$;
2. if H^* is sufficiently large and $q(H^*) \in \left(0, \frac{\sqrt{2}}{\sqrt{e\pi}\sigma^2} \right)$ then given the shape of $-2f'(k)$ it intersects $q(H^*)$ twice at k_1^* and k_2^* with $k_1^* < k_2^*$. It then gives three intervals: (1) $Q'(k) < 0$ for any $k \in (0, k_1^*)$, (2) $Q'(k) > 0$ for any $k \in (k_1^*, k_2^*)$, and (3) $Q'(k) < 0$ for any $k \in (k_2^*, \infty)$.

As mentioned before, the sign of $\frac{\partial H^*}{\partial k}$ follows the same pattern. It is now straightforward to translate the result into espionage power terms. Observe that $\frac{\partial H^*}{\partial k} = 0$ if $h < H^*$. Therefore, consider only the cases with $h > H^*$:

1. if H^* is sufficiently small then $\frac{\partial \alpha^*}{\partial k} > 0$;
2. if H^* is sufficiently large then (1) $\frac{\partial \alpha^*}{\partial k} < 0$ if $k \in (k_1^*, k_1^*)$, (2) $\frac{\partial \alpha^*}{\partial k} > 0$ for any $k \in (0, k_1^*) \cup (k_2^*, \infty)$.

■

A.1.2 Social welfare function derivation and analysis

Derivation of social welfare function

SWF is derived as an overall output exerted by the winner of the game minus the overall expected costs exerted by both competitors during the game. Due to the complexity of the function, the derivation is divided into five parts and then combined: (1) overall effort exerted by the winner of the game, (2) contribution of the symmetric case shocks to the overall expected shock, (3) contribution of

asymmetric case shocks when a leader wins, (4) contribution of asymmetric case shocks when an underdog wins, and (5) overall expected costs.

1. Overall effort Both players have similar expected investment levels at the beginning of the race. Moreover, the winner will choose the same level of investment as the loser. Expected investments of the winner are then expressed as:

$$X_i \equiv x_i^* + P^S \hat{x}_i^s + P_1^A \hat{x}_i^a + P_2^A \hat{x}_i^a.$$

or expanding:

$$X_i = f(-k) (2F(H) - 1) + (F(k) - F(-k)) f(0) + 2F(-k)f(H). \quad (\text{A.31})$$

2. Contribution of symmetric case shocks The contribution of wining firm shocks to the overall expected shock in case the winner is decided in the symmetric case of the development stage is:

$$\begin{aligned} E^S &= P^S (E(\epsilon_i | -k < \epsilon_i - \epsilon_{-i} < k) + E(\hat{\epsilon}_i | \hat{\epsilon}_i > \hat{\epsilon}_{-i})) = \\ &= P^S \left(\underbrace{\frac{\int_{-\infty}^{\infty} \int_{\epsilon_{-i}-k}^{\epsilon_{-i}+k} \epsilon_i f(\epsilon_i) f(\epsilon_{-i}) d\epsilon_i d\epsilon_{-i}}{F(k) - F(-k)}}_{=0 \text{ for symmetric distributions}} + 2 \int_{-\infty}^{\infty} \int_{\hat{\epsilon}_{-i}}^{\infty} \hat{\epsilon}_i f(\hat{\epsilon}_i) f(\hat{\epsilon}_{-i}) d\hat{\epsilon}_i d\hat{\epsilon}_{-i} \right). \end{aligned}$$

Substituting $P^S = F(k) - F(-k)$ into the equation yields:

$$E^S = 2 (F(k) - F(-k)) \int_{-\infty}^{\infty} \int_{\hat{\epsilon}_{-i}}^{\infty} \hat{\epsilon}_i f(\hat{\epsilon}_i) f(\hat{\epsilon}_{-i}) d\hat{\epsilon}_i d\hat{\epsilon}_{-i}. \quad (\text{A.32})$$

Asymmetric case shocks contributions The asymmetric development stage occurs with the probability $P^A = 2P_1^A = 2P_2^A$. There are two sub-cases to consider in the asymmetric development stage: (3) the leader wins, or (4) the underdog wins.

3. Leader wins Contribution of the asymmetric case shocks to the overall shock if an underdog wins:

$$\begin{aligned} E^L &= P^A \hat{P}_1^a \left(E(\epsilon_i | \epsilon_i > \epsilon_{-i} + k) + E(\hat{\epsilon}_1 | \hat{\epsilon}_1 > \hat{\epsilon}_2 - H) \right) = \\ &= P^A \hat{P}_1^a \left(\frac{\int_{-\infty}^{\infty} \int_{\epsilon_{-i}+k}^{\infty} \epsilon_i f(\epsilon_i) f(\epsilon_{-i}) d\epsilon_i d\epsilon_{-i}}{F(-k)} + \right. \\ &\quad \left. + \frac{\int_{-\infty}^{\infty} \int_{\hat{\epsilon}_2-H}^{\infty} \hat{\epsilon}_1 f(\hat{\epsilon}_1) f(\hat{\epsilon}_2) d\hat{\epsilon}_1 d\hat{\epsilon}_2}{F(H)} \right). \quad (\text{A.33}) \end{aligned}$$

Substituting $P^A = 2F(-k)$ and $\hat{P}_1^a = F(H)$ in (A.33) and simplifying yields:

$$\begin{aligned} E^L &= 2F(H) \int_{-\infty}^{\infty} \int_{\epsilon_{-i}+k}^{\infty} \epsilon_i f(\epsilon_i) f(\epsilon_{-i}) d\epsilon_i d\epsilon_{-i} + \\ &\quad + 2F(-k) \int_{-\infty}^{\infty} \int_{\hat{\epsilon}_2-H}^{\infty} \hat{\epsilon}_1 f(\hat{\epsilon}_1) f(\hat{\epsilon}_2) d\hat{\epsilon}_1 d\hat{\epsilon}_2. \quad (\text{A.34}) \end{aligned}$$

4. Underdog wins Similarly, the contribution of shocks of underdog-winner of the race to the overall expected shock is as follows:

$$\begin{aligned} E^U &= P^A \hat{P}_2^a \left(E(\epsilon_{-i} | \epsilon_{-i} < \epsilon_i - k) + E(\hat{\epsilon}_2 | \hat{\epsilon}_2 > \hat{\epsilon}_1 + H) \right) = \\ &= P^A \hat{P}_2^a \left(\frac{\int_{-\infty}^{\infty} \int_{-\infty}^{\epsilon_i-k} \epsilon_{-i} f(\epsilon_{-i}) f(\epsilon_i) d\epsilon_{-i} d\epsilon_i}{F(-k)} + \right. \\ &\quad \left. + \frac{\int_{-\infty}^{\infty} \int_{\hat{\epsilon}_2+H}^{\infty} \hat{\epsilon}_2 f(\hat{\epsilon}_2) f(\hat{\epsilon}_1) d\hat{\epsilon}_2 d\hat{\epsilon}_1}{F(-H)} \right), \quad (\text{A.35}) \end{aligned}$$

where $\hat{P}_2^a = 1 - \hat{P}_1^a = F(-H)$ is the probability an underdog wins in the asymmetric development stage: Expanding (A.35) and simplifying yields:

$$\begin{aligned} E^U &= 2F(-H) \int_{-\infty}^{\infty} \int_{-\infty}^{\epsilon_i-k} \epsilon_{-i} f(\epsilon_{-i}) f(\epsilon_i) d\epsilon_{-i} d\epsilon_i + \\ &\quad + 2F(-k) \int_{-\infty}^{\infty} \int_{\hat{\epsilon}_1+H}^{\infty} \hat{\epsilon}_2 f(\hat{\epsilon}_2) f(\hat{\epsilon}_1) d\hat{\epsilon}_2 d\hat{\epsilon}_1. \quad (\text{A.36}) \end{aligned}$$

5. Overall expected costs Overall expected costs of the R&D race can be expressed as:

$$C \equiv c(x_i^*) + P^S c(\hat{x}_i^s) + P_1^A c(\hat{x}_i^a) + P_2^A c(\hat{x}_i^a). \quad (\text{A.37})$$

Social Welfare Function

The expected quality social welfare function can be derived by combining (A.31), (A.32), (A.34), (A.36), and (A.37):

$$\begin{aligned}
W = & f(-k) (2F(H) - 1) - f(-k)^2 (2F(H) - 1)^2 + \\
& + (F(k) - F(-k)) (f(0) - f(0)^2) \\
& + 2F(-k) (f(H) - f(H)^2) \\
& + 2(F(k) - F(-k)) \int_{-\infty}^{\infty} \int_{\hat{\epsilon}_{-i}}^{\infty} \hat{\epsilon}_i f(\hat{\epsilon}_i) f(\hat{\epsilon}_{-i}) d\hat{\epsilon}_i d\hat{\epsilon}_{-i} + \\
& + 2F(H) \int_{-\infty}^{\infty} \int_{\epsilon_{-i}+k}^{\infty} \epsilon_i f(\epsilon_i) f(\epsilon_{-i}) d\epsilon_i d\epsilon_{-i} + \\
& + 2F(-k) \int_{-\infty}^{\infty} \int_{\hat{\epsilon}_2-H}^{\infty} \hat{\epsilon}_1 f(\hat{\epsilon}_1) f(\hat{\epsilon}_2) d\hat{\epsilon}_1 d\hat{\epsilon}_2 + \\
& + 2F(-H) \int_{-\infty}^{\infty} \int_{-\infty}^{\epsilon_i-k} \epsilon_{-i} f(\epsilon_{-i}) f(\epsilon_i) d\epsilon_{-i} d\epsilon_i + \\
& + 2F(-k) \int_{-\infty}^{\infty} \int_{\hat{\epsilon}_1+H}^{\infty} \hat{\epsilon}_2 f(\hat{\epsilon}_2) f(\hat{\epsilon}_1) d\hat{\epsilon}_2 d\hat{\epsilon}_1.
\end{aligned} \tag{A.38}$$

Proof of Proposition 2.8. It is sufficient to study the properties of the simplified function, which consist only of the members of SWF that are functions of attack power α :

$$\begin{aligned}
\psi = & F(H) (2g(k) + f(k) + F(-H) f(k)^2) + \\
& + F(-k) (2i(H) + f(H) - f(H)^2)
\end{aligned} \tag{A.39}$$

Consider the FOC of the function (A.39) w.r.t. H :

$$\begin{aligned}
\frac{\partial \psi}{\partial H} = & f(H) \left(2(g(k) - F(-k)) f'(H) \right) + f(-k)^2 (1 - 2F(H)) + f(-k) \Big) + \\
& + F(-k) (2i'(H) + f'(H)) = 0, \tag{A.40}
\end{aligned}$$

Studying the FOC around $H = 0$ yields:

$$\left. \frac{\partial \psi}{\partial H} \right|_{H=0} = (f(-k) + 2g(k)) f(0) > 0.$$

Therefore, the function is always positive around $H = 0$. Therefore, $H = 0$ is neither a maximum nor a minimum point. It is also evident that:

$$\left. \frac{\partial \psi}{\partial H} \right|_{H \rightarrow \infty} = 0.$$

Now I consider the sign of the $\frac{\partial \psi}{\partial H}$ as it approaches 0. Expanding FOC's terms and simplifying yields:

$$\begin{aligned} \frac{\partial \psi}{\partial H} = & \underbrace{\exp\left(-\frac{H^2 + k^2}{\sigma^2}\right) \exp\left(\frac{H^2}{2\sigma^2}\right) \left(\kappa^1 - 2\sqrt{2}\sigma \operatorname{erf}\left(\frac{H}{\sqrt{2}\sigma}\right)\right)}_{Q^1} + \\ & + \underbrace{H\sqrt{\pi} \operatorname{erfc}\left(\frac{k}{\sqrt{2}\sigma}\right) \exp\left(\frac{k^2}{\sigma^2}\right) \left(4 \exp\left(\frac{-H^2}{2\sigma^2}\right) - 2\sqrt{2\pi}\sigma - 2 \exp\left(\frac{H^2}{4\sigma^2}\right) \sqrt{\pi}\sigma^3\right)}_{Q^2}, \end{aligned} \quad (\text{A.41})$$

where:

$$\kappa^1 = 4 \exp\left(\frac{k^2}{2\sigma^2}\right) \sqrt{\pi}\sigma^2 \left(1 + \sqrt{2} \exp\left(\frac{k^2}{4\sigma^2}\right)\right).$$

The overall sign of the FOC is negative as $Q^1|_{H \rightarrow \infty} = 0^+$ and $Q^2|_{H \rightarrow \infty} = -\infty$:

$$\left. \frac{\partial \psi}{\partial H} \right|_{H \rightarrow \infty} < 0.$$

Therefore, given the function is smooth and continuous, there will always be a value of H^W between 0 and ∞ that maximises the social welfare function. Therefore, there exists a value of $a^W = \frac{h-H^W}{h}$ that maximises the social welfare function for any $h > H^W$. ■

B

Appendices for Chapter 3

Contents

B.1	List of industries' aggregates	177
B.2	Independent variables histograms	178
B.3	Multiple Imputation conditions	179
B.4	Assumptions for GLM class models	181
B.5	Eurostat indicators on high-tech industry and knowledge-intensive services	181
B.6	Collinearity in multivariate specifications	182
B.7	RESET	184
B.8	Alternative specifications test	185
B.9	Kernel density plots of residuals of imputation models	186
B.10	Complete case regression analysis results for γ^{per} . . .	187
B.11	Multiple imputation regression analysis results for γ^{per}	190
B.12	Complete case regression analysis results for γ^{res} . . .	193
B.13	Multiple imputation regression analysis results results for γ^{res}	195
B.14	Complete case regression analysis results for γ^{sup} . . .	198
B.15	Multiple imputation regression analysis results results for γ^{sup}	201
B.16	Complete case regression analysis results for models with research-intensity proxies and CRM variable . . .	204
B.17	Multiple imputation regression analysis results results for models with research-intensity proxies and CRM variable	206
B.18	Multiple imputation regression analysis results results for models with research-intensity proxies and CRM variable for manufacturing industries	208

B.19 Multiple imputation regression analysis results results for models with research-intensity proxies and CRM variable for service industries	210
B.20 Robustness checks	212

B.1 List of industries' aggregates

Table B.1: Industries according to NACE rev.2

NACE	Label
C13-C15	Manufacture of textile, leather and clothes
C16-C18	Manufacture of wood and products of wood and cork, except furniture; articles of straw and plaiting materials
C19-C23	Manufacture of coke, refined petroleum, chemical and basic pharmaceutical products, rubber and plastics, other non-metallic mineral products
C24-C25	Manufacture of basic metals and fabricated metal products excluding machines and equipment
C26	Manufacture of computer, electronic and optical products
C27-C28	Manufacture of electrical equipment, machinery and equipment
C29-C30	Manufacture of motor vehicles, trailers and semi-trailers, other transport equipment
C31-C33	Manufacture of furniture and other manufacturing; repair and installation of machinery and equipment
D35-E39	Electricity, gas, steam, air conditioning and water supply
F41-F43	Construction and civil engineering
G45-G47	Trade of motor vehicles and motorcycles
H49-H53	Transport
I55-I56	Accommodation and Food and beverage service activities
J58-J63	Publishing of books, software and media
K64	Monetary intermediation
K65	Insurance, reinsurance and pension funding, except compulsory social security
K66	Activities auxiliary to financial services and insurance activities
L68	Real estate activities
M69-M74	Professional, scientific and technical activities
N77-N82	Activities for rental and leasing, employment, security and investigation, services to buildings and landscape, office administrative
S951	Repair of computers and communication equipment

B.2 Independent variables histograms

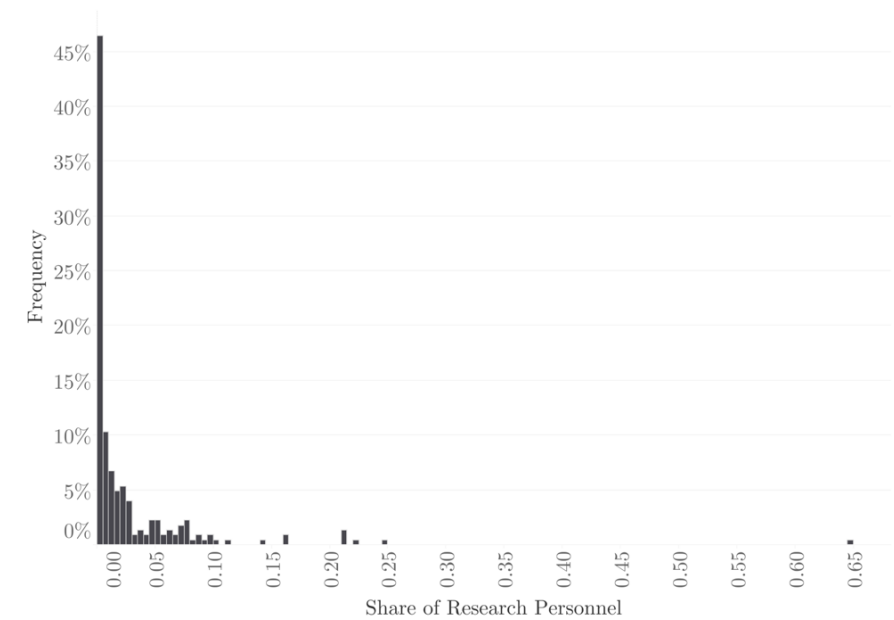


Figure B.1: Share of research personnel histogram

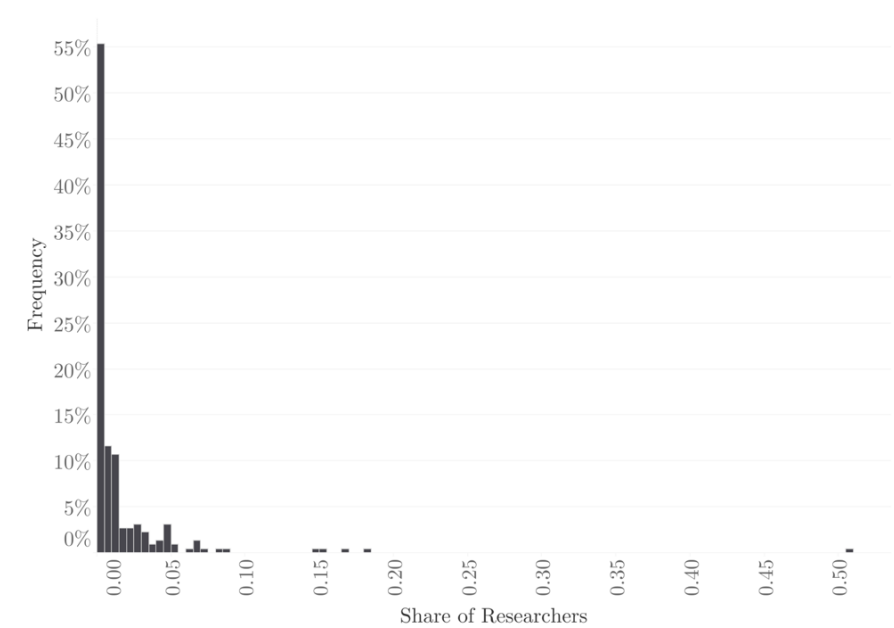


Figure B.2: Share of researchers histogram

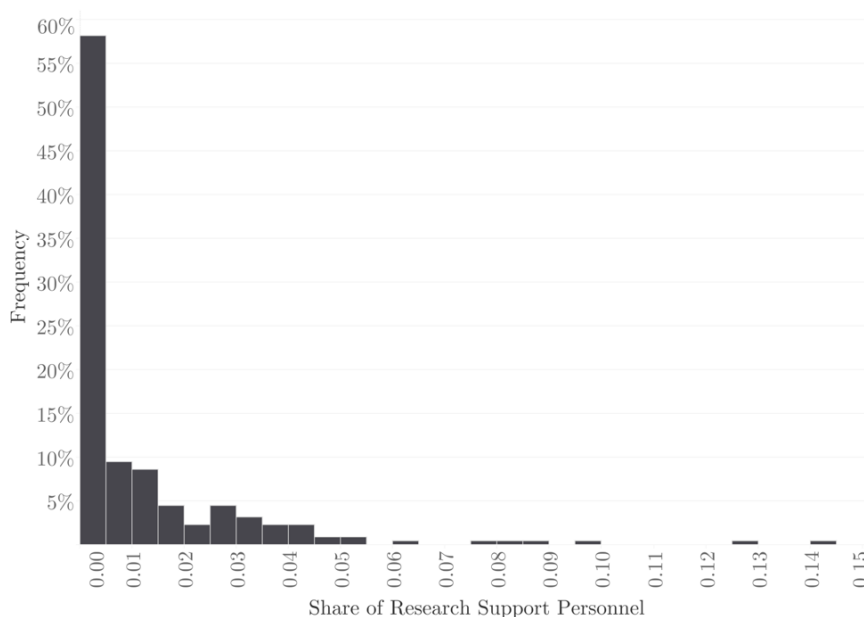


Figure B.3: Share of R&D support personnel histogram

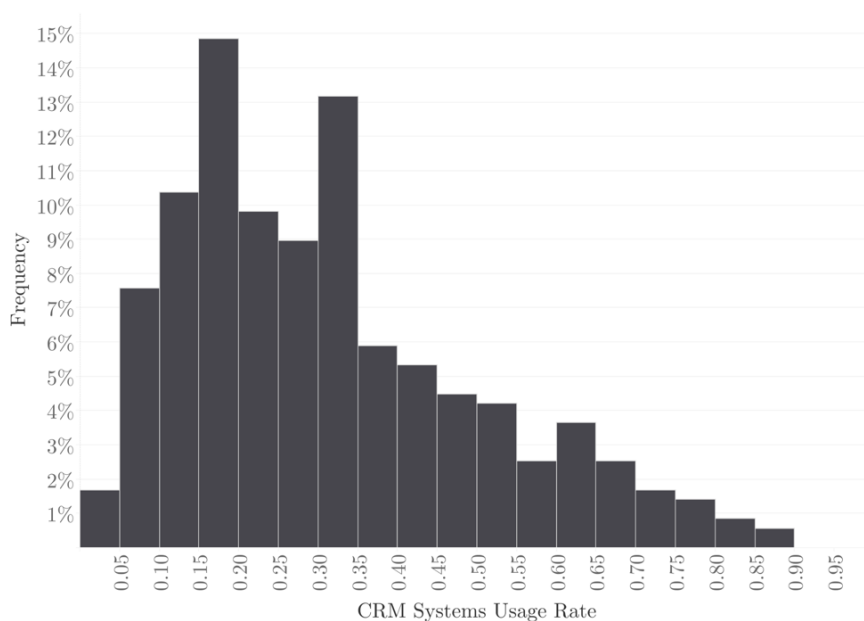


Figure B.4: CRM usage rate histogram

B.3 Multiple Imputation conditions

An appropriate imputation model is the model which preserves all relationships inside the data that are relevant to the analysis. It means that imputation models should include all variables (including the dependent variable in question) utilised in

the analytical models, which is ensured in this research by design. The appropriate imputation model is required for a proper MI procedure.

A *Proper MI procedure* is defined as the one that yields consistent asymptotically normal estimators and approximately unbiased estimator of its asymptomatic variance in common regular models (Rubin's rule) (Nielsen, 2007). In practice, it means that if the multiple imputations are proper, then complete data estimates can be obtained by a simple averaging of within imputation estimates. Their variances are then obtained by combining the within- and between-imputation variances.

Missing-data mechanism should be ignorable in order for the MI technique to produce unbiased results. That means that the pattern of missing observations must not depend on unobserved data. This condition is satisfied when: (1) data are missing completely at random (MCAR) and depend neither on observed nor unobserved data, or (2) data are missing at random (MAR), which allows the missing data pattern to be dependent on the observed data.

MCAR assumption is seldom satisfied in real life. Moreover, there is no formal statistical procedure to test the MAR assumption. Therefore, the ignorability of the missing pattern is generally assumed based on the nature of the data set and the available knowledge about it (Little & Rubin, 2002; X. Liu, 2016). Given that observations are likely missing due to industry-specific reasons, adopted control variables can be used to model the missingness mechanism. It then enables the research to assume that data is MAR (Bland, 2015).

Additionally, it was shown by Collins et al. (2001) and Schafer and Graham (2002) that model-based imputation methods (e.g. maximum likelihood or multiple imputation) perform well and do not create substantial bias even when the data are not MAR.

The sufficient quantity of completed data sets is another crucial point to determine to achieve stable unbiased results. Following the recommendations by Graham et al. (2007), the research chooses 40 imputations, which correspond to up to 50% of missing data. It is more than what was initially recommended by Rubin but should allow for more consistent point estimates and standard errors (von Hippel, 2020).

B.4 Assumptions for GLM class models

1. Linearity assumption is relaxed, meaning that an outcome variable does not need to have a linear relationship with regressors. The outcome variable is assumed to be generated from a particular distribution in an exponential family,
2. The fractional outcome model assumes a correct conditional mean specification and is agnostic about other moments of the outcome. It utilises robust estimators of a variance-covariance matrix to achieve consistent estimates of the unknown standard errors, which implies the that homoscedasticity assumption is relaxed (Pinzon, 2016; White, 1994),
3. Normality of error terms is not necessary.

B.5 Eurostat indicators on high-tech industry and knowledge-intensive services

Table B.2: Aggregations of manufacturing based on NACE Rev. 2

Manufacturing industries	NACE Rev. 2
High-technology	C21, C26
Medium-high-technology	C20, C27-C30
Medium-low-technology	C19, C22-C25, C33
Low-technology	C10-C18, C31-C32

Table B.3: Aggregations of services based on NACE Rev. 2

Knowledge based services	NACE Rev. 2
Knowledge-intensive services (KIS)	J58-I63, K64-K66, L68, M69-M74
Less knowledge-intensive services (LKIS)	D35-E39, F41-F43, G45-G47, H49-H53, I55-I56, N77-N82, S951

B.6 Collinearity in multivariate specifications

Tables B.4 and B.5 provide variance inflation factors for the main independent variables in multivariate specification with and without industry controls correspondingly. Sufficiently high VIF (> 4) is conventionally considered as a sign of collinearity. The removal of industry controls mitigates the problem. Table B.6 demonstrates that within industry variance of CRM usage is sufficiently small.

Table B.4: Variance Inflation Factor for the main independent variables in multivariate specification with industry controls

Variable	VIF	1/VIF
r_{ij}^L	2.15	0.464536
D_{ij}	5.20	0.192332

Table B.5: Variance Inflation Factor for the main independent variables in multivariate specification without industry controls

Variable	VIF	1/VIF
r_{ij}^L	1.39	0.719559
D_{ij}	1.91	0.522411

Table B.6: Within industry variance of CRM system usage

Industry NACE r2 code	Variance
C10-C12	0.003526
C13-C15	0.012441
C16-C18	0.01804
C19-C23	0.014848
C24-C25	0.010181
C26	0.058687
C27-C28	0.019917
C29-C30	0.008458
C31-C33	0.011255
D35-E39	0.014388
F41-F43	0.004358
G45-G47	0.009418
H49-H53	0.006454
I55-I56	0.006485
J58-J63	0.010973
K64	0.03126
K65	0.021616
K66	0.054944
M69-M74	0.013961
N77-N82	0.011742
S951	0.050869

B.7 RESET

H0: all non-linear combinations coefficients are zero—the model is correctly specified

1. RESET results for univariate specification with probit link function

Maximum polynomial power: 2

Statistic 1.072

P-value 0.3005 > 0.05

Result: fail to reject H0

2. RESET results for univariate specification with logit link function

Maximum polynomial power: 2

Statistic 1.353

P-value 0.2448 > 0.05

Result: fail to reject H0

3. RESET results for base multivariate specification with probit link function

Maximum polynomial power: 2

Statistic 1.315

P-value 0.2516 > 0.05

Result: fail to reject H0

4. RESET results for base multivariate specification with logit link function

Maximum polynomial power: 2

Statistic 0.724

P-value 0.3948 > 0.05

Result: fail to reject H0

B.8 Alternative specifications test

Univariate fractional probit vs. logit models

H0: Univariate fractional probit model

H1: Univariate fractional logit model

Version: t-test

Statistic: -0.190

P-value: 0.8496

Result: fail to reject H0

Multivariate fractional probit vs. logit models

H0: Multivariate fractional probit model

H1: Multivariate fractional logit model

Version: t-test

Statistic: 0.915

P-value: 0.3614

Result: fail to reject H0

B.9 Kernel density plots of residuals of imputation models

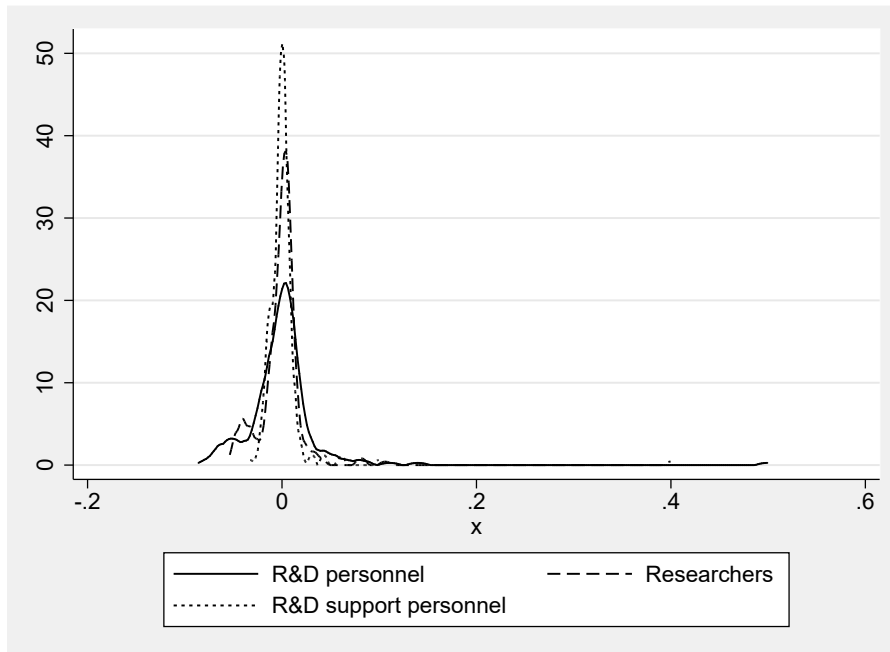


Figure B.5: Residuals distributions for univariate imputation models

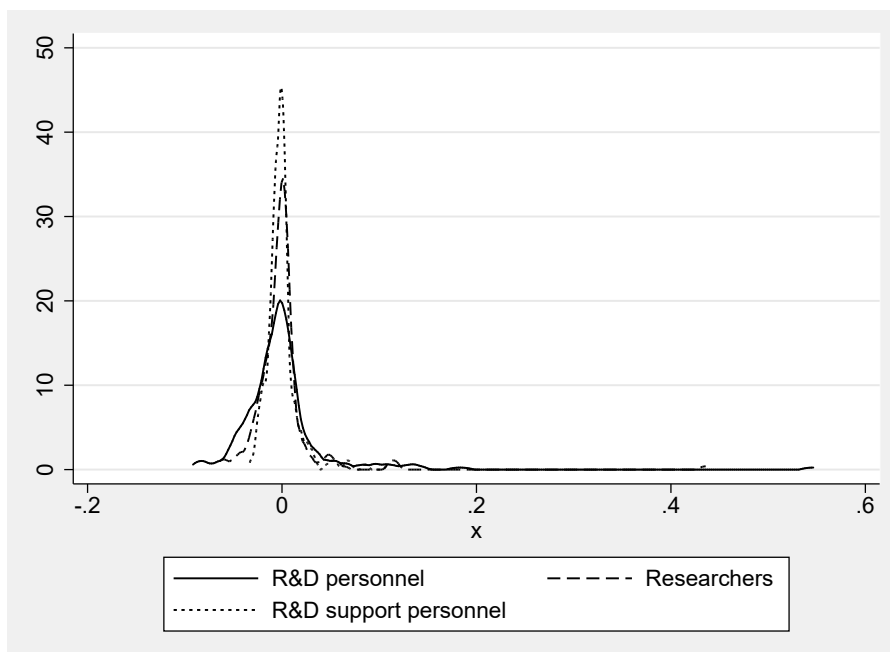


Figure B.6: Residuals distributions for multivariate imputation models

B.10 Complete case regression analysis results for r^{per}

Table B.7: Regression results: fractional probit regressions on complete data with share of research personnel as an independent variable

	I Base model	II Country controls	III Industry controls	IV All controls
r^{per}	0.965**	0.612**	1.217**	1.036***
Complete	(2.08)	(1.96)	(2.04)	(3.11)
BG		-0.730*** (-3.01)		-0.760*** (-2.96)
CY		-0.820*** (-3.81)		-0.847*** (-3.90)
CZ		-0.257* (-1.69)		-0.295** (-2.01)
DK		-0.341** (-2.17)		-0.417** (-2.57)
ES		-0.209 (-1.43)		-0.240 (-1.50)
FI		-0.373** (-2.04)		-0.441** (-2.37)
FR		-0.111 (-0.62)		-0.182 (-1.15)
HR		-0.454** (-2.07)		-0.486** (-2.52)
HU		-0.865*** (-2.73)		-0.822*** (-2.78)
IE		-0.414** (-2.34)		-0.430** (-2.42)
IT		-0.196 (-1.17)		-0.207 (-1.25)
LT		0.182 (1.20)		0.168 (1.11)
LV		-0.199 (-1.06)		-0.198 (-1.11)
MT		-0.0241 (-0.12)		-0.0703 (-0.39)
NL		0.268* (1.81)		0.244 (1.61)
NO		-0.452** (-2.44)		-0.502*** (-2.70)

Table B.7: Regression results: fractional probit regressions on complete data with share of research personnel as an independent variable (continued)

PT	-0.269*	-0.338**
	(-1.89)	(-1.98)
RO	0.0236	0.0289
	(0.15)	(0.18)
SI	-1.193***	-1.213***
	(-3.84)	(-3.87)
SK	0.325**	0.285*
	(2.01)	(1.82)
C13-C15	-0.240	-0.112
	(-1.31)	(-1.03)
C16-C18	-0.165	-0.0419
	(-1.10)	(-0.41)
C19-C23	-0.177	-0.0981
	(-0.88)	(-0.62)
C24-C25	0.0384	0.140
	(0.27)	(1.39)
C26	-0.251	-0.187
	(-1.30)	(-1.60)
C27-C28	-0.124	-0.0568
	(-0.71)	(-0.48)
C29-C30	-0.0138	0.0656
	(-0.09)	(0.64)
C31-C33	-0.106	-0.0426
	(-0.59)	(-0.41)
D35-E39	-0.121	0.0563
	(-0.81)	(0.66)
F41-F43	-0.0576	-0.00865
	(-0.46)	(-0.11)
G45-G47	-0.0333	0.0909
	(-0.29)	(1.47)
H49-H53	-0.0699	0.0448
	(-0.52)	(0.57)
I55-I56	-0.294**	0.00450
	(-2.49)	(0.03)
J58-J63	0.106	0.247***
	(0.72)	(2.97)
L68	0.0217	0.157
	(0.12)	(1.36)
M69-M74	-0.190	-0.0591
	(-1.11)	(-0.59)
N77-N82	-0.00764	0.0610
	(-0.06)	(0.90)

Table B.7: Regression results: fractional probit regressions on complete data with share of research personnel as an independent variable (continued)

S951			-0.518 (-1.64)	-0.110 (-0.42)
Constant	-2.214*** (-71.13)	-2.080*** (-14.90)	-2.139*** (-23.27)	-2.093*** (-13.96)
Observations	224	224	224	224

z statistics in parentheses; Significance levels: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$

C10–C12 and BE are chosen as base levels

B.11 Multiple imputation regression analysis results for r^{per}

Table B.8: Regression results: fractional probit regressions on imputed data with share of research personnel as an independent variable

	I Base model	II Country controls	III Industry controls	IV All controls
r^{per}	1.377** (2.23)	1.011* (1.95)	1.820** (2.44)	1.285** (2.43)
BG		-0.651*** (-3.01)		-0.716*** (-3.11)
CY		-0.660*** (-3.13)		-0.749*** (-3.82)
CZ		-0.243 (-1.58)		-0.340** (-2.15)
DE		-3.371*** (-14.00)		-3.568*** (-11.93)
DK		-0.341** (-2.16)		-0.445*** (-2.69)
EL		-0.621* (-1.81)		-0.648* (-1.92)
ES		-0.131 (-0.83)		-0.279* (-1.77)
FI		-0.409** (-2.17)		-0.484** (-2.48)
FR		-0.0809 (-0.46)		-0.200 (-1.27)
HR		-0.493** (-2.18)		-0.554*** (-2.71)
HU		-0.750*** (-3.04)		-0.839*** (-3.68)
IE		-0.404** (-2.35)		-0.468*** (-2.59)
IS		-0.564** (-2.14)		-0.658*** (-2.85)
IT		-0.116 (-0.68)		-0.193 (-1.18)
LT		0.262 (1.52)		0.201 (1.21)
LU		-0.509** (-2.25)		-0.595*** (-2.86)

Table B.8: Regression results: fractional probit regressions on imputed data with share of research personnel as an independent variable (continued)

LV	-0.157 (-0.84)	-0.221 (-1.27)
MK	-3.379*** (-21.20)	-3.471*** (-19.04)
MT	0.0411 (0.23)	-0.0583 (-0.34)
NL	0.292* (1.96)	0.228 (1.47)
NO	-0.350* (-1.74)	-0.451** (-2.33)
PL	-0.241 (-1.55)	-0.693*** (-3.68)
PT	-0.265* (-1.83)	-0.304* (-1.78)
RO	-0.0278 (-0.17)	-0.0566 (-0.33)
SE	-0.491*** (-2.72)	-0.570*** (-3.02)
SI	-1.198*** (-3.85)	-1.267*** (-3.95)
SK	0.395** (2.39)	0.311** (2.00)
TR	-0.0219 (-0.14)	-0.0801 (-0.49)
UK	-3.354*** (-14.01)	-3.545*** (-12.12)
C13-C15		-0.0825 (-0.54)
C16-C18		-0.0953 (-0.68)
C19-C23		-0.127 (-0.97)
C24-C25		0.136 (0.46)
C26		-0.210* (-1.60)
C27-C28		0.0470 (0.03)
C29-C30		0.0384 (0.33)
C31-C33		-0.0755 (-0.61)

Table B.8: Regression results: fractional probit regressions on imputed data with share of research personnel as an independent variable (continued)

D35-E39			-0.0621 (-0.46)	0.0105 (0.13)
F41-F43			0.0210 (0.18)	0.0340 (0.42)
G45-G47			0.0552 (0.49)	0.108* (1.74)
H49-H53			0.0633 (0.50)	0.0713 (0.90)
I55-I56			-0.0986 (-0.77)	0.0119 (0.13)
J58-J63			0.166 (1.24)	0.284*** (3.69)
K			0.350** (2.32)	0.460*** (4.42)
L68			0.0239 (0.14)	0.0981 (0.85)
M69-M74			-0.186 (-1.01)	-0.0305 (-0.26)
N77-N82			0.0295 (0.23)	0.0746 (1.21)
S951			0.105 (0.44)	0.197 (1.12)
Constant	-2.233*** (-67.86)	-2.100*** (-14.56)	-2.252*** (-24.76)	-2.110*** (-13.69)
Observations	351	351	351	351

z statistics in parentheses; Significance levels: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$

C10–C12 and BE are chosen as base levels

B.12 Complete case regression analysis results for r^{res}

Table B.9: Regression results: fractional probit regressions on complete data with share of researchers as an independent variable

	I Base model	II Country controls	III Industry controls	IV All controls
r^{res}	0.990*	0.669	1.106*	1.013***
Complete	(1.89)	(1.50)	(1.82)	(2.66)
BG		-0.676*** (-3.17)		-0.706*** (-3.21)
CY		-0.827*** (-3.81)		-0.848*** (-3.88)
CZ		-0.258* (-1.68)		-0.287* (-1.94)
DK		-0.339** (-2.14)		-0.416** (-2.57)
ES		-0.212 (-1.43)		-0.238 (-1.48)
FI		-0.373** (-2.03)		-0.443** (-2.40)
FR		-0.110 (-0.61)		-0.182 (-1.14)
HR		-0.458** (-2.07)		-0.495** (-2.52)
HU		-0.870*** (-2.74)		-0.842*** (-2.82)
IE		-0.420** (-2.35)		-0.443** (-2.47)
IT		-0.194 (-1.14)		-0.205 (-1.21)
LT		0.177 (1.16)		0.167 (1.09)
LV		-0.290 (-1.50)		-0.279 (-1.48)
MT		-0.0259 (-0.13)		-0.0708 (-0.39)
NL		0.272* (1.82)		0.261* (1.69)
NO		-0.443** (-2.40)		-0.481*** (-2.62)

Table B.9: Regression results: fractional probit regressions on complete data with share of researchers as an independent variable (continued)

PT	-0.264*	-0.330*
	(-1.84)	(-1.92)
RO	0.0167	0.00790
	(0.11)	(0.05)
SI	-1.194***	-1.206***
	(-3.84)	(-3.84)
SK	0.335**	0.300*
	(2.06)	(1.91)
C13-C15	-0.236	-0.102
	(-1.29)	(-0.96)
C16-C18	-0.167	-0.0328
	(-1.12)	(-0.31)
C19-C23	-0.148	-0.0695
	(-0.73)	(-0.43)
C24-C25	0.0433	0.150
	(0.30)	(1.47)
C26	-0.192	-0.135
	(-0.96)	(-1.05)
C27-C28	-0.0897	-0.0152
	(-0.50)	(-0.13)
C29-C30	-0.0219	0.0792
	(-0.14)	(0.77)
C31-C33	-0.101	-0.0330
	(-0.55)	(-0.33)
D35-E39	-0.243**	0.0864
	(-2.01)	(0.80)
F41-F43	-0.0613	-0.00823
	(-0.49)	(-0.10)
G45-G47	-0.0382	0.0895
	(-0.33)	(1.45)
H49-H53	-0.0751	0.0477
	(-0.55)	(0.60)
I55-I56	-0.300**	0.00182
	(-2.53)	(0.01)
J58-J63	0.131	0.268***
	(0.88)	(3.16)
L68	-0.0836	0.0246
	(-0.39)	(0.28)
M69-M74	-0.137	-0.0195
	(-0.86)	(-0.21)
N77-N82	-0.0125	0.0576
	(-0.09)	(0.86)

Table B.9: Regression results: fractional probit regressions on complete data with share of researchers as an independent variable (continued)

S951			-0.521*	-0.115
			(-1.65)	(-0.44)
Constant	-2.209***	-2.073***	-2.133***	-2.092***
	(-71.27)	(-14.59)	(-23.12)	(-13.79)
Observations	224	224	224	224

z statistics in parentheses; Significance levels: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$

C10–C12 and BE are chosen as base levels

B.13 Multiple imputation regression analysis results results for r^{res}

Table B.10: Regression results: fractional probit regressions on imputed data with share of researchers as an independent variable

	I Base model	II Country controls	III Industry controls	IV All controls
r^{res}	1.366* (1.68)	1.052 (1.53)	1.747* (1.96)	1.315** (2.07)
BG		-0.664*** (-3.07)		-0.736*** (-3.17)
CY		-0.670*** (-3.15)		-0.759*** (-3.81)
CZ		-0.248 (-1.59)		-0.347** (-2.16)
DE		-3.365*** (-14.01)		-3.579*** (-12.03)
DK		-0.338** (-2.11)		-0.447*** (-2.69)
EL		-0.610* (-1.76)		-0.648* (-1.90)
ES		-0.136 (-0.85)		-0.284* (-1.77)
FI		-0.409** (-2.17)		-0.480** (-2.44)
FR		-0.0773 (-0.44)		-0.201 (-1.27)
HR		-0.501** (-2.20)		-0.571*** (-2.73)

Table B.10: Regression results: fractional probit regressions on imputed data with share of researchers as an independent variable (continued)

HU	-0.765*** (-3.09)	-0.871*** (-3.80)
IE	-0.417** (-2.41)	-0.487*** (-2.65)
IS	-0.566** (-2.15)	-0.663*** (-2.88)
IT	-0.115 (-0.66)	-0.191 (-1.14)
LT	0.256 (1.45)	0.193 (1.14)
LU	-0.506** (-2.22)	-0.591*** (-2.83)
LV	-0.166 (-0.89)	-0.234 (-1.32)
MK	-3.369*** (-21.23)	-3.479*** (-18.35)
MT	0.0409 (0.22)	-0.0543 (-0.31)
NL	0.303** (2.01)	0.238 (1.51)
NO	-0.330* (-1.67)	-0.428** (-2.22)
PL	-0.242 (-1.60)	-0.708*** (-3.91)
PT	-0.258* (-1.77)	-0.302* (-1.74)
RO	-0.0428 (-0.26)	-0.0831 (-0.48)
SE	-0.489*** (-2.71)	-0.572*** (-3.00)
SI	-1.200*** (-3.85)	-1.274*** (-3.94)
SK	0.411** (2.47)	0.332** (2.12)
TR	-0.0144 (-0.09)	-0.0729 (-0.45)
UK	-3.342*** (-13.86)	-3.548*** (-11.96)
C13-C15	-0.0704 (-0.46)	-0.0807 (-0.87)
C16-C18	-0.0951 (-0.69)	-0.0441 (-0.43)

Table B.10: Regression results: fractional probit regressions on imputed data with share of researchers as an independent variable (continued)

C19-C23			-0.154 (-0.79)	-0.106 (-0.69)
C24-C25			0.0699 (0.51)	0.140 (1.46)
C26			-0.232 (-1.19)	-0.160 (-1.24)
C27-C28			0.0461 (0.30)	0.0875 (0.76)
C29-C30			-0.0285 (-0.18)	0.0545 (0.54)
C31-C33			-0.102 (-0.58)	-0.0756 (-0.77)
D35-E39			-0.0747 (-0.56)	-0.00105 (-0.01)
F41-F43			0.0163 (0.14)	0.0267 (0.32)
G45-G47			0.0460 (0.41)	0.0984 (1.56)
H49-H53			0.0507 (0.41)	0.0606 (0.78)
I55-I56			-0.113 (-0.88)	0.00463 (0.05)
J58-J63			0.203 (1.52)	0.308*** (4.00)
K			0.370** (2.42)	0.480*** (4.64)
L68			0.0183 (0.11)	0.0915 (0.79)
M69-M74			-0.107 (-0.64)	0.0100 (0.09)
N77-N82			0.0212 (0.17)	0.0671 (1.09)
S951			0.127 (0.53)	0.220 (1.23)
Constant	-2.212*** (-70.84)	-2.086*** (-14.27)	-2.239*** (-24.96)	-2.101*** (-13.38)
Observations	351	351	351	351

z statistics in parentheses; Significance levels: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$

C10–C12 and BE are chosen as base levels

B.14 Complete case regression analysis results for r^{sup}

Table B.11: Regression results: fractional probit regressions on complete data with share of R&D support personnel as an independent variable

	I Base model	II Country controls	III Industry controls	IV All controls
r^{sup}	4.584***	2.324***	6.548***	4.302***
Complete	(3.71)	(2.60)	(3.52)	(3.51)
BG		-0.726*** (-3.02)		-0.722*** (-2.92)
CY		-0.815*** (-3.85)		-0.827*** (-3.93)
CZ		-0.262* (-1.77)		-0.291** (-2.06)
DK		-0.349** (-2.28)		-0.415*** (-2.63)
ES		-0.213 (-1.50)		-0.250 (-1.61)
FI		-0.380** (-2.10)		-0.449** (-2.44)
FR		-0.110 (-0.62)		-0.169 (-1.08)
HR		-0.453** (-2.11)		-0.480*** (-2.59)
HU		-0.860*** (-2.74)		-0.786*** (-2.69)
IE		-0.410** (-2.38)		-0.407** (-2.39)
IT		-0.214 (-1.33)		-0.234 (-1.49)
LT		0.186 (1.28)		0.193 (1.34)
LV		-0.289 (-1.52)		-0.261 (-1.45)
MT		-0.0287 (-0.15)		-0.0790 (-0.45)
NL		0.237 (1.61)		0.211 (1.43)
NO		-0.447** (-2.53)		-0.482*** (-2.82)

Table B.11: Regression results: fractional probit regressions on complete data with share of R&D support personnel as an independent variable (continued)

PT	-0.289** (-2.06)	-0.353** (-2.14)
RO	0.0310 (0.21)	0.0646 (0.42)
SI	-1.202*** (-3.87)	-1.213*** (-3.91)
SK	0.300* (1.90)	0.249* (1.65)
C13-C15	-0.260 (-1.42)	-0.126 (-1.16)
C16-C18	-0.160 (-1.06)	-0.0417 (-0.42)
C19-C23	-0.271 (-1.37)	-0.142 (-0.96)
C24-C25	0.0207 (0.15)	0.132 (1.34)
C26	-0.419** (-2.21)	-0.270** (-2.43)
C27-C28	-0.246 (-1.44)	-0.138 (-1.10)
C29-C30	-0.0797 (-0.50)	0.0344 (0.34)
C31-C33	-0.124 (-0.72)	-0.0515 (-0.52)
D35-E39	-0.235* (-1.89)	0.0386 (0.34)
F41-F43	-0.0463 (-0.37)	0.00867 (0.11)
G45-G47	-0.0172 (-0.15)	0.105* (1.74)
H49-H53	-0.0537 (-0.40)	0.0587 (0.73)
I55-I56	-0.275** (-2.33)	0.0131 (0.08)
J58-J63	0.0358 (0.26)	0.205** (2.48)
L68	-0.0601 (-0.28)	0.0415 (0.57)
M69-M74	-0.267 (-1.48)	-0.0868 (-0.82)
N77-N82	0.00775 (0.06)	0.0680 (1.01)

Table B.11: Regression results: fractional probit regressions on complete data with share of R&D support personnel as an independent variable (continued)

S951			-0.508 (-1.63)	-0.113 (-0.43)
Constant	-2.254*** (-70.11)	-2.087*** (-15.60)	-2.158*** (-23.55)	-2.100*** (-14.72)
Observations	222	222	222	222

z statistics in parentheses; Significance levels: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$

C10–C12 and BE are chosen as base levels

B.15 Multiple imputation regression analysis results results for r^{sup}

Table B.12: Regression results: fractional probit regressions on imputed data with share of R&D support personnel as an independent variable

	I Base model	II Country controls	III Industry controls	IV All controls
r^{sup}	6.225*** (3.81)	3.838*** (2.80)	9.106*** (4.38)	5.521*** (3.67)
BG		-0.642*** (-3.05)		-0.684*** (-3.06)
CY		-0.650*** (-3.18)		-0.727*** (-3.81)
CZ		-0.256* (-1.77)		-0.353** (-2.32)
DE		-3.311*** (-13.78)		-3.481*** (-12.05)
DK		-0.355** (-2.37)		-0.446*** (-2.77)
EL		-0.629* (-1.89)		-0.633* (-1.89)
ES		-0.142 (-0.96)		-0.294* (-1.93)
FI		-0.420** (-2.29)		-0.498*** (-2.59)
FR		-0.0781 (-0.47)		-0.187 (-1.22)
HR		-0.497** (-2.27)		-0.542*** (-2.79)
HU		-0.737*** (-3.06)		-0.788*** (-3.53)
IE		-0.386** (-2.37)		-0.429** (-2.49)
IS		-0.566** (-2.20)		-0.653*** (-2.88)
IT		-0.148 (-0.93)		-0.229 (-1.48)
LT		0.268* (1.65)		0.220 (1.40)
LU		-0.494** (-2.23)		-0.573*** (-2.77)

Table B.12: Regression results: fractional probit regressions on imputed data with share of R&D support personnel as an independent variable (continued)

LV	-0.157 (-0.88)	-0.208 (-1.24)
MK	-3.370*** (-22.27)	-3.444*** (-19.04)
MT	0.0312 (0.18)	-0.0685 (-0.41)
NL	0.244* (1.69)	0.176 (1.18)
NO	-0.338* (-1.84)	-0.434** (-2.44)
PL	-0.246* (-1.68)	-0.673*** (-3.64)
PT	-0.299** (-2.19)	-0.323** (-1.96)
RO	-0.0139 (-0.09)	-0.0135 (-0.08)
SE	-0.477*** (-2.72)	-0.540*** (-2.92)
SI	-1.216*** (-3.91)	-1.274*** (-4.02)
SK	0.358** (2.30)	0.258* (1.70)
TR	-0.0195 (-0.13)	-0.0673 (-0.43)
UK	-3.318*** (-13.74)	-3.492*** (-12.00)
C13-C15	-0.112 (-0.71)	-0.106 (-1.10)
C16-C18	-0.0898 (-0.62)	-0.0458 (-0.47)
C19-C23	-0.307 (-1.57)	-0.190 (-1.32)
C24-C25	0.0496 (0.35)	0.122 (1.29)
C26	-0.531*** (-2.66)	-0.324*** (-2.71)
C27-C28	-0.174 (-1.05)	-0.0587 (-0.44)
C29-C30	-0.130 (-0.77)	-0.00396 (-0.04)
C31-C33	-0.127 (-0.73)	-0.0917 (-0.91)

Table B.12: Regression results: fractional probit regressions on imputed data with share of R&D support personnel as an independent variable (continued)

D35-E39			-0.0467 (-0.34)	0.0148 (0.17)
F41-F43			0.0544 (0.44)	0.0552 (0.67)
G45-G47			0.0804 (0.68)	0.122* (1.93)
H49-H53			0.0821 (0.62)	0.0843 (1.01)
I55-I56			-0.0645 (-0.46)	0.0279 (0.28)
J58-J63			0.0846 (0.63)	0.234*** (2.96)
K			0.285** (1.97)	0.423*** (4.15)
L68			0.0341 (0.20)	0.103 (0.89)
M69-M74			-0.270 (-1.40)	-0.0700 (-0.59)
N77-N82			0.0593 (0.45)	0.0871 (1.34)
S951			0.102 (0.44)	0.184 (1.13)
Constant	-2.273*** (-68.83)	-2.110*** (-15.78)	-2.279*** (-23.31)	-2.118*** (-14.35)
Observations	351	351	351	351

z statistics in parentheses; Significance levels: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$

C10–C12 and BE are chosen as base levels

B.16 Complete case regression analysis results for models with research-intensity proxies and CRM variable

Table B.13: Regression results: fractional probit regressions on complete data with research proxies and CRM variable

	I R&D personnel model	II Researchers model	III R&D support model
γ^{per}	0.332 (0.94)		
γ^{res}		0.225 (0.37)	
γ^{sup}			1.606* (1.71)
d	0.487** (2.29)	0.524** (2.48)	0.440** (1.98)
BG	-0.576** (-2.25)	-0.522** (-2.32)	-0.582** (-2.29)
CY	-0.696*** (-3.10)	-0.697*** (-3.10)	-0.700*** (-3.14)
CZ	-0.144 (-0.83)	-0.139 (-0.81)	-0.158 (-0.92)
DK	-0.285* (-1.65)	-0.281 (-1.63)	-0.296* (-1.74)
ES	-0.139 (-0.87)	-0.140 (-0.87)	-0.147 (-0.93)
FI	-0.354* (-1.85)	-0.354* (-1.86)	-0.360* (-1.89)
FR	-0.0176 (-0.09)	-0.00982 (-0.05)	-0.0260 (-0.14)
HR	-0.346 (-1.55)	-0.344 (-1.55)	-0.354 (-1.60)
HU	-0.692** (-2.12)	-0.684** (-2.10)	-0.704** (-2.16)
IE	-0.368* (-1.95)	-0.372** (-1.97)	-0.367** (-1.98)
IT	-0.0950 (-0.53)	-0.0905 (-0.50)	-0.116 (-0.65)
LT	0.315* (1.86)	0.319* (1.89)	0.305* (1.82)

Table B.13: Regression results: fractional probit regressions on complete data with research proxies and CRM variable (continued)

LV	-0.141 (-0.68)	-0.136 (-0.66)	-0.155 (-0.75)
MT	0.0759 (0.38)	0.0787 (0.39)	0.0649 (0.33)
NL	0.372** (2.30)	0.381** (2.36)	0.340** (2.07)
NO	-0.427** (-2.15)	-0.416** (-2.11)	-0.432** (-2.24)
RO	0.128 (0.76)	0.128 (0.75)	0.126 (0.76)
SI	-1.023*** (-3.19)	-1.016*** (-3.18)	-1.043*** (-3.24)
SK	0.404** (2.31)	0.415** (2.39)	0.378** (2.18)
Constant	-2.293*** (-13.69)	-2.301*** (-13.80)	-2.280*** (-13.65)
Observations	214	215	213

z statistics in parentheses; Significance levels: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$

C10–C12 is chosen as a base level

B.17 Multiple imputation regression analysis results results for models with research-intensity proxies and CRM variable

Table B.14: Regression results: fractional probit regressions on imputed data with research poxies and CRM variable

	I R&D personnel model	II Researchers model	III R&D support model
γ^{per}	0.562 (1.22)		
γ^{res}		0.410 (0.53)	
γ^{sup}			3.041** (2.35)
d	0.715*** (3.63)	0.762*** (3.78)	0.641*** (3.48)
BG	-0.485** (-2.15)	-0.485** (-2.16)	-0.487** (-2.20)
CY	-0.545*** (-2.70)	-0.550*** (-2.72)	-0.543*** (-2.74)
CZ	-0.104 (-0.63)	-0.103 (-0.62)	-0.122 (-0.77)
DE	-3.436*** (-13.59)	-3.450*** (-13.70)	-3.417*** (-13.06)
DK	-0.265 (-1.55)	-0.258 (-1.52)	-0.285* (-1.74)
EL	-0.615* (-1.66)	-0.625* (-1.67)	-0.604* (-1.69)
ES	-0.0930 (-0.61)	-0.0939 (-0.61)	-0.0949 (-0.64)
FI	-0.380** (-2.02)	-0.381** (-2.04)	-0.391** (-2.11)
FR	0.0259 (0.15)	0.0346 (0.20)	0.0159 (0.10)
HR	-0.372 (-1.60)	-0.368 (-1.60)	-0.381* (-1.68)
HU	-0.528** (-2.14)	-0.527** (-2.14)	-0.532** (-2.20)
IE	-0.333* (-1.85)	-0.342* (-1.90)	-0.319* (-1.85)

Table B.14: Regression results: fractional probit regressions on imputed data with research proxies and CRM variable (continued)

IS	-0.519** (-1.99)	-0.523** (-2.01)	-0.521** (-2.03)
IT	-0.00988 (-0.06)	-0.00298 (-0.02)	-0.0410 (-0.25)
LT	0.453** (2.42)	0.456** (2.43)	0.442** (2.51)
LU	-0.426* (-1.91)	-0.428* (-1.92)	-0.423* (-1.93)
LV	0.0182 (0.09)	0.0198 (0.10)	0.00470 (0.03)
MK	-3.218*** (-18.38)	-3.216*** (-18.34)	-3.229*** (-19.55)
MT	0.170 (0.95)	0.173 (0.96)	0.157 (0.90)
NL	0.431*** (2.77)	0.442*** (2.87)	0.379** (2.49)
NO	-0.315 (-1.58)	-0.300 (-1.52)	-0.327* (-1.74)
PL	-0.305* (-1.95)	-0.313** (-2.00)	-0.272* (-1.80)
PT	-0.274* (-1.69)	-0.272* (-1.66)	-0.285* (-1.83)
RO	0.120 (0.69)	0.115 (0.66)	0.129 (0.78)
SE	-0.457** (-2.37)	-0.461** (-2.39)	-0.451** (-2.40)
SI	-1.038*** (-3.21)	-1.033*** (-3.19)	-1.064*** (-3.31)
SK	0.524*** (3.06)	0.538*** (3.17)	0.467*** (2.87)
TR	0.0441 (0.27)	0.0433 (0.27)	0.0375 (0.24)
UK	-3.444*** (-13.99)	-3.459*** (-14.10)	-3.411*** (-14.18)
Constant	-2.409*** (-14.62)	-2.416*** (-14.64)	-2.395*** (-15.37)
Observations	351	351	351

z statistics in parentheses; Significance levels: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$

C10–C12 is chosen as a base level

B.18 Multiple imputation regression analysis results results for models with research-intensity proxies and CRM variable for manufacturing industries

Table B.15: Regression results: fractional probit regressions on imputed data with research poxies and CRM variable for manufacturing industries

	I All manufacturing	II Low-tech manufacturing	III High-tech manufacturing
r^{sup}	4.488*** (2.77)	3.200 (0.66)	7.206*** (3.13)
d	-0.776 (-1.41)	-0.651 (-0.87)	-1.337 (-1.29)
BG	-1.222*** (-3.50)	-0.794** (-2.29)	-4.560*** (-10.74)
CY	-4.216*** (-16.02)	-4.037*** (-13.17)	-4.612*** (-9.70)
CZ	-0.682*** (-2.71)	-0.408 (-1.49)	-1.115*** (-2.77)
DK	-0.643*** (-3.25)	-0.468*** (-2.78)	-0.927*** (-3.03)
EL	-0.821** (-2.07)	-3.930*** (-16.98)	-0.893* (-1.78)
ES	-0.650*** (-3.36)		-0.870*** (-2.83)
FI	-4.134*** (-18.11)	-3.900*** (-19.28)	-4.573*** (-13.66)
FR	-0.643*** (-2.82)	-0.557** (-2.46)	-0.875** (-2.49)
HR	-4.119*** (-18.30)	-3.992*** (-13.51)	-4.426*** (-13.67)
HU	-1.287*** (-3.38)	-4.039*** (-13.27)	-1.418*** (-2.60)
IE	-0.771*** (-3.31)	-0.535** (-2.20)	-1.168*** (-3.02)
IS	-4.146*** (-17.01)	-4.013*** (-13.60)	-4.376*** (-13.05)
IT	-0.607*** (-2.58)	-0.624*** (-3.25)	-0.758** (-1.98)

Table B.15: Regression results: fractional probit regressions on imputed data with research proxies and CRM variable for manufacturing industries (continued)

LT	-0.253 (-1.00)	-0.0903 (-0.33)	-0.550 (-1.23)
LU	-4.079*** (-16.99)	-3.962*** (-13.54)	-4.238*** (-14.04)
LV	-0.880*** (-2.81)	-0.480 (-1.39)	-4.656*** (-10.96)
MK	-4.188*** (-16.57)	-3.993*** (-12.98)	-4.621*** (-10.64)
MT	-0.515 (-1.46)	-0.223 (-0.83)	-4.995*** (-8.86)
NL	-0.277 (-1.17)	-0.0948 (-0.49)	-0.681* (-1.67)
NO	-0.994*** (-2.85)	-0.559* (-1.86)	-4.406*** (-15.56)
PT	-0.524*** (-2.76)		-0.655** (-2.41)
RO	-0.319 (-1.42)	-0.171 (-0.79)	-0.532 (-1.42)
SE	-1.203*** (-3.55)	-0.902*** (-2.77)	-4.348*** (-15.02)
SI	-4.274*** (-16.59)	-4.027*** (-15.59)	-4.851*** (-11.11)
SK	-0.188 (-0.86)	-0.0713 (-0.47)	-0.420 (-1.14)
TR	-0.292 (-1.50)	-0.205 (-1.40)	-0.392 (-1.47)
Constant	-1.549*** (-5.43)	-1.705*** (-4.98)	-1.223** (-2.24)
Observations	152	81	71

z statistics in parentheses; Significance levels: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$

C10–C12 is chosen as a base level

B.19 Multiple imputation regression analysis results results for models with research-intensity proxies and CRM variable for service industries

Table B.16: Regression results: fractional probit regressions on imputed data with research poxies and CRM variable for service industries

	I All knowledge- based services	II Less knowledge- intensive services	III Knowledge- intensive services
r^{sup}	3.698* (1.82)	10.70* (1.89)	1.862 (1.01)
d	0.729*** (3.48)	0.238 (0.43)	0.985*** (4.55)
BG	-0.237 (-1.09)	-0.401 (-1.21)	-0.182 (-0.56)
CY	-0.272* (-1.77)	-0.230 (-1.19)	-0.412 (-1.54)
CZ	0.0182 (0.16)	-0.119 (-0.55)	0.0828 (0.75)
DE	-3.395*** (-10.45)	-3.463*** (-9.16)	
DK	-0.216 (-1.49)	-0.285 (-0.93)	-0.201 (-1.48)
EL	-3.547*** (-11.77)	-3.621*** (-9.13)	-3.343*** (-14.28)
ES	0.0543 (0.60)	0.0000345 (0.00)	0.0168 (0.17)
FI	-0.170 (-1.26)	-0.325 (-1.41)	-0.109 (-0.64)
FR	0.206* (1.66)	-0.0498 (-0.32)	0.295** (2.40)
HR	-0.0212 (-0.15)	-0.175 (-0.78)	0.0735 (0.36)
HU	-0.331 (-1.22)	-3.228*** (-9.32)	-0.0782 (-0.36)
IE	-0.157 (-0.97)	-0.180 (-0.95)	-0.163 (-0.58)
IS	-0.128 (-0.64)	-0.566* (-1.68)	0.0304 (0.15)

Table B.16: Regression results: fractional probit regressions on imputed data with research proxies and CRM variable for service industries (continued)

IT	0.136 (1.35)	-0.0439 (-0.28)	0.195** (2.01)
LT	0.605*** (3.67)	0.549** (2.03)	0.429*** (3.20)
LU	-0.179 (-1.06)	-0.279 (-0.91)	-0.101 (-0.53)
LV	0.298** (2.00)	0.0346 (0.14)	0.527*** (4.26)
MK	-3.232*** (-13.50)	-3.326*** (-10.27)	-3.442*** (-18.72)
MT	0.339*** (3.14)	0.225 (1.36)	0.410*** (3.13)
NL	0.607*** (6.77)	0.470** (2.46)	0.640*** (7.98)
NO	-0.0540 (-0.35)	-0.0128 (-0.06)	-0.0806 (-0.68)
PL	-0.155 (-1.51)		-0.146 (-1.59)
RO	0.144 (0.92)	-0.312 (-0.98)	0.343** (2.37)
SE	-0.150 (-0.97)	-0.0449 (-0.19)	-0.315* (-1.77)
SI	-0.761** (-2.39)	-0.629* (-1.84)	-3.431*** (-21.46)
SK	0.701*** (5.90)	0.499** (2.59)	0.811*** (6.04)
TR	0.181 (1.51)	0.0171 (0.09)	0.274* (1.70)
UK	-3.401*** (-12.05)	-3.449*** (-10.40)	
Constant	-2.573*** (-23.41)	-2.345*** (-8.96)	-2.698*** (-27.33)
Observations	199	105	94

z statistics in parentheses; Significance levels: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$

C10–C12 is chosen as a base level

B.20 Robustness checks

Table B.17: Regression results: robustness checks

	I A.Surface controls	II Turnover controls	III E-comm controls	IV Sysadmin controls	V All controls
r^{sup}	3.277*** (2.68)	3.235*** (2.72)	3.783*** (3.25)	2.999** (2.01)	3.981** (2.53)
d	0.585*** (3.02)	0.581*** (2.96)	0.285 (1.34)	0.540*** (2.88)	0.310 (1.45)
A.Surface	-18E-7 (-0.17)				703E-9 (0.07)
Turnover		0.000120 (0.32)			0.0000710 (0.19)
E-comm			0.00768** (2.22)		0.00787** (2.30)
Sysadmin				0.0822 (0.33)	-0.0711 (-0.29)
BG	-0.504** (-2.30)	-0.502** (-2.29)	-0.391* (-1.69)	-0.500** (-2.26)	-0.390* (-1.70)
CY	-0.553*** (-2.79)	-0.552*** (-2.79)	-0.465** (-2.29)	-0.549*** (-2.76)	-0.466** (-2.31)
CZ	-0.136 (-0.85)	-0.134 (-0.85)	-0.182 (-1.11)	-0.146 (-0.90)	-0.173 (-1.05)
DE	-3.365*** (-13.52)	-3.367*** (-13.60)	-3.449*** (-13.85)	-3.388*** (-13.15)	-3.429*** (-13.17)
DK	-0.290* (-1.78)	-0.291* (-1.80)	-0.318* (-1.91)	-0.286* (-1.74)	-0.320* (-1.91)
EL	-0.639* (-1.87)	-0.639* (-1.87)	-0.459 (-1.31)	-0.637* (-1.87)	-0.456 (-1.31)
ES	-0.106 (-0.70)	-0.110 (-0.75)	-0.0715 (-0.47)	-0.104 (-0.71)	-0.0755 (-0.49)
FI	-0.393** (-2.13)	-0.392** (-2.13)	-0.347* (-1.77)	-0.392** (-2.12)	-0.346* (-1.77)
FR	0.00912 (0.05)	-0.00145 (-0.01)	0.0251 (0.15)	0.00698 (0.04)	0.0110 (0.06)
HR	-0.397* (-1.74)	-0.395* (-1.73)	-0.446* (-1.94)	-0.399* (-1.73)	-0.445* (-1.95)
HU	-0.558** (-2.33)	-0.558** (-2.33)	-0.536** (-2.28)	-0.577** (-2.39)	-0.520** (-2.18)
IE	-0.325* (-1.90)	-0.322* (-1.89)	-0.361** (-2.01)	-0.330* (-1.89)	-0.355** (-1.99)
IS	-0.520**	-0.522**	-0.515**	-0.522**	-0.515**

Table B.17: Regression results: robustness checks (continued)

	(-2.09)	(-2.10)	(-2.06)	(-2.10)	(-2.08)
IT	-0.0501	-0.0583	0.0666	-0.0489	0.0570
	(-0.30)	(-0.36)	(0.38)	(-0.30)	(0.32)
LT	0.421**	0.424**	0.352**	0.417**	0.355**
	(2.38)	(2.36)	(2.08)	(2.41)	(2.15)
LU	-0.449**	-0.451**	-0.378*	-0.451**	-0.378*
	(-2.04)	(-2.06)	(-1.70)	(-2.04)	(-1.70)
LV	-0.0259	-0.0243	0.0495	-0.0312	0.0552
	(-0.14)	(-0.13)	(0.26)	(-0.17)	(0.29)
MK	-3.280***	-3.280***	-3.162***	-3.288***	-3.154***
	(-18.69)	(-18.92)	(-16.59)	(-18.39)	(-16.55)
MT	0.133	0.134	0.151	0.137	0.149
	(0.76)	(0.77)	(0.85)	(0.78)	(0.84)
NL	0.363**	0.361**	0.314**	0.354**	0.317**
	(2.37)	(2.37)	(1.99)	(2.32)	(2.01)
NO	-0.337*	-0.334*	-0.445**	-0.338*	-0.441**
	(-1.82)	(-1.81)	(-2.34)	(-1.81)	(-2.35)
PL	-0.313**	-0.314**	-0.347**	-0.325**	-0.340**
	(-2.00)	(-2.03)	(-2.09)	(-2.03)	(-1.99)
PT	-0.302*	-0.301*	-0.423***	-0.293*	-0.429***
	(-1.92)	(-1.91)	(-2.63)	(-1.88)	(-2.69)
RO	0.113	0.115	0.221	0.129	0.211
	(0.69)	(0.70)	(1.24)	(0.75)	(1.17)
SE	-0.432**	-0.432**	-0.442**	-0.430**	-0.443**
	(-2.31)	(-2.32)	(-2.35)	(-2.29)	(-2.37)
SI	-1.090***	-1.088***	-1.082***	-1.091***	-1.081***
	(-3.37)	(-3.37)	(-3.35)	(-3.37)	(-3.35)
SK	0.446***	0.445***	0.547***	0.456***	0.540***
	(2.80)	(2.74)	(3.23)	(2.80)	(3.15)
TR	-0.00319	-0.00292	0.113	-0.00303	0.118
	(-0.02)	(-0.02)	(0.66)	(-0.02)	(0.70)
UK	-3.419***	-3.415***	-3.337***	-3.432***	-3.322***
	(-14.03)	(-13.99)	(-13.27)	(-13.54)	(-12.71)
Constant	-2.370***	-2.371***	-2.464***	-2.380***	-2.458***
	(-15.16)	(-15.08)	(-14.51)	(-14.61)	(-14.41)
Observations	351	351	351	351	351

z statistics in parentheses; Significance levels: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$

C10–C12 is chosen as a base level; See Appendix B.19 for the full table

C

Appendices for Chapter 4

Contents

C.1 Mathematical appendix	215
C.2 Supplementary material	234
C.3 Extreme points characteristics and the convex hull . .	235
C.3.1 Full support	235
C.3.2 Convex hull construction for the full support case	245
C.3.3 Partial support	249
C.3.4 Convex hull construction for the partial support case . .	261
C.3.5 Extreme points of convex regions	263
C.4 Feasibility of intermediate family sequences	268
C.5 Supplementary material for Subsection 4.3.4	273

C.1 Mathematical appendix

Proof of Claim 4.4. Suppose that vertex $m < l$ is not attacked in some equilibrium, in which l is attacked. Then m is not defended by Claim 4.3. The expected payoff from attacking m is v_m , which is greater than the upper bound of the expected payoff from attacking l , v_l . Therefore, there is a profitable deviation: attacking m with probability 1. If the inequality is weak, i.e. $v_l = v_m$, we can relabel the nodes and preserve the monotonicity of the payoffs in the assumption. ■

Proof of Claim 4.5. From (4.2) we have the following condition for the top k vertices to be attacked:

$$\frac{v_k}{v_1} + \frac{v_k}{v_2} + \cdots + \frac{v_k}{v_{k-1}} + \frac{v_k}{v_k} \geq k - \delta.$$

Rearranging and simplifying yields:

$$\frac{v_k}{v_1} + \frac{v_k}{v_2} + \cdots + \frac{v_k}{v_{k-1}} \geq k - 1 - \delta,$$

which, because $v_{k-1} \geq v_k$, implies condition for the top $k - 1$ vertices. ■

Proof of Claim 4.6. Suppose the condition (4.4) holds strictly for some node k and does hold for node $k + 1$.

The Attacker's gain from an attack on top k nodes is then:

$$A(\mathbf{v}_k) = \frac{(k - \delta)}{\sum_{i=1}^k \frac{1}{v_i}}.$$

If the Attacker chooses to mix over $k - 1$ nodes, then by Lemma 4.4, the Defender protects only $k - 1$ nodes with positive probability. Then, the Attacker's gain is

$$A(\mathbf{v}_{k-1}) = \frac{(k - \delta - 1)}{\sum_{i=1}^{k-1} \frac{1}{v_i}}.$$

We now demonstrate that:

$$A(\mathbf{v}_k) > A(\mathbf{v}_{k-1}). \tag{C.1}$$

Expanding (C.1) yields:

$$\begin{aligned} \frac{(k - \delta)}{\sum_{i=1}^k \frac{1}{v_i}} &> \frac{(k - \delta - 1)}{\sum_{i=1}^{k-1} \frac{1}{v_i}}, \\ (k - \delta) \sum_{i=1}^{k-1} \frac{1}{v_i} &> (k - \delta - 1) \sum_{i=1}^k \frac{1}{v_i}, \\ (k - \delta) \sum_{i=1}^{k-1} \frac{1}{v_i} - (k - \delta - 1) \sum_{i=1}^k \frac{1}{v_i} &> 0, \end{aligned}$$

and finally:

$$v_k > \frac{(k - \delta - 1)}{\sum_{i=1}^{k-1} \frac{1}{v_i}},$$

which is always satisfied by the assumption that that condition (4.4) holds strictly for node k .

Thus, if the Attacker deviates to smaller support, he receives a strictly smaller payoff. ■

Proof of Claim 4.7. Suppose that for some node y condition (4.4) holds strictly, while for nodes $\nu \in (y, y + z]$, where $z \in [1, n - y]$ and $z \in \mathbb{Z}$, it holds with equality.

The expected gain of the Attacker from an attack on top k nodes is:

$$A(\vec{v}_y) = \frac{(y - \delta)}{\sum_{i=1}^y \frac{1}{v_i}}.$$

Then it follows from condition (4.4) must be the case that each node indexed $\nu \in (y, y + z]$ must be $v_\nu = A(\vec{v}_y) = -D(\vec{v}_y)$.

Then the expected payoff from an attack on top $y + z$ nodes is:

$$\begin{aligned} A(\vec{v}_{y+z}) &= \frac{(y + z - \delta)}{\sum_{i=1}^y \frac{1}{v_i} + \frac{z}{\sum_{i=1}^y \frac{1}{v_i}}} = \\ &= \frac{(y + z - \delta)}{\sum_{i=1}^y \frac{y+z-\delta}{(y-\delta)v_i}} = \frac{(y - \delta)}{\sum_{i=1}^y \frac{1}{v_i}} = A(\vec{v}_y). \end{aligned} \tag{C.2}$$

Since equality (C.2) is satisfied regardless of the quantity of nodes that pass condition (4.4), it must be the case that $A(\vec{v}_{y+z}) = -D(\vec{v}_{y+z})$ for any $|\text{supp}_A| \in [y, y + z]$. ■

Proof of Lemma 4.1. Let $(y_1, \dots, y_k) \in \mathbb{R}_+^k$ and $(z_1, \dots, z_k) \in \mathbb{R}_+^k$.

Consider Defender's expected loss function (4.5) rewritten as:

$$A(y) = \frac{k - \delta}{g(y)},$$

where $g(y) = \frac{1}{y_1} + \dots + \frac{1}{y_k}$.

Consider now function $g(y)$. Via the Cauchy-Schwarz inequality:

$$\begin{aligned} \frac{1}{y_1} + \dots + \frac{1}{y_k} &= \left(\frac{\sqrt{z_1}}{y_1}, \dots, \frac{\sqrt{z_k}}{y_k} \right) \cdot \left(\frac{1}{\sqrt{z_1}}, \dots, \frac{1}{\sqrt{z_k}} \right) \\ &\leq \left(\frac{z_1}{y_1^2} + \dots + \frac{z_k}{y_k^2} \right)^{\frac{1}{2}} \left(\frac{1}{z_1} + \dots + \frac{1}{z_k} \right)^{\frac{1}{2}}. \end{aligned} \tag{C.3}$$

Now we rewrite the function as:

$$\left(\frac{1}{z_1} + \dots + \frac{1}{z_k}\right)^{-1} \leq \left(\frac{z_1}{y_1^2} + \dots + \frac{z_k}{y_k^2}\right) \left(\frac{1}{y_1} + \dots + \frac{1}{y_k}\right)^{-2}. \quad (\text{C.4})$$

Observe that the right-hand side of (C.4) is exactly the tangent plane of the left-hand side at some point $(y_1, \dots, y_k)$. Since the inequality holds for *any* such point $(y_1, \dots, y_n)$, the left-hand side is a concave function. It immediately follows that $A(\vec{v}_k)$ is concave, while $D(\vec{v}_k)$ is convex. ■

Proof of Proposition 4.1. Consider the expected attack on k nodes in general form:

$$A(\mathbf{v}_k) = \frac{k - \delta}{\sum_{i \in K} \frac{1}{v_i}},$$

where $K = \text{supp}_A$.

Consider now the first derivative of $A(\mathbf{v}_k)$ w.r.t. the value of some node $j \in K$:

$$\frac{\partial A(\mathbf{v}_k)}{\partial v_j} = \frac{k - \delta}{\left(\sum_{i \in K} \frac{1}{v_i}\right)^2 v_j^2} > 0, \quad (\text{C.5})$$

which is always positive.

It follows that the expected gain from an attack on k nodes attains its maximum value when all nodes attain value $v^q = n - 1$. Substituting $v_i = v^q = n - 1$, $\forall i \in K$ yields:

$$A^{\max}(\mathbf{v}_k) = \frac{(k - \delta)(n - 1)}{k}. \quad (\text{C.6})$$

Now consider the first derivative of $A^{\max}(\mathbf{v}_k)$ w.r.t. k :

$$\frac{\partial A(\mathbf{v}_k)}{\partial k} = \frac{\delta(n - 1)}{k^2} > 0.$$

It implies that the maximum gain from an attack is attained whenever $k = n$ and the game is played on a complete network. In this case, the Attacker's payoff is:

$$A^c = \frac{(n - 1)(n - \delta)}{n}.$$

We demonstrate that the Attacker obtains the minimum payoff when the game is played on a star network. Consider two cases: (1) full support case $k = n$, and (2) partial support case with $k < n$.

Case 1. $k = n$. By Lemma 4.1, the expected gain from an attack on k nodes is concave. Thus, by the Bauer minimum principle, it must attain the minimum value at some extreme point of the set of feasible values. In the full support case each $v_i \in \{1, n-1\}$. Therefore, each extreme point can be represented as the following ordered set:

$$\mathbf{v}_n^e = \left(\underbrace{n-1, \dots, n-1}_q, \underbrace{1, \dots, 1}_{n-q} \right). \quad (\text{C.7})$$

The expected attack on the network with degree sequence (C.7) is¹:

$$A^q(\mathbf{v}_n^e) = \frac{(n-\delta)}{\frac{q}{n-1} + n - q}.$$

Since we demonstrated that $\frac{\partial A^q(\mathbf{v}_n)}{\partial v_i} > 0$, $\forall i \in N$, it must be the case that $A^q(\mathbf{v}_n^e)$ is minimised when the Defender chooses the minimum possible number of completely connected nodes. As we assume that any network must be connected, $q^{\min} = 1$ ². Therefore, it implies that the expected gain from an attack attains its minimum when the game is played on a star network. It can be written down as follows:

$$A^s = \frac{(n-\delta)(n-1)}{(n-1)^2 + 1} < 1. \quad (\text{C.8})$$

Observe that (C.8) is strictly smaller than 1.

Case 2. $k < n$. Observe that if k nodes are included in the Attacker's support, it implies that the expected attack on the top k nodes is strictly larger than 1. If it is smaller than 1, then any node of value $v_i \geq 1$ must also be included in the Attacker's support, which yields a full support case and contradicts the partial

¹We ignore the fact that any node of value $v^u = 1$ would not be included in the Attacker's support for now

²Observe that choosing $q = 0$ yields a disconnected network with degree sequence $(1, \dots, 1)$.

support assumption.

Therefore, the Attacker's payoff is minimised when the game is played on a star network. ■

Proof of Proposition 4.2. Observe that expected value function of the Defender (4.6) is linear in c . Therefore, we can represent the utility function for any possible network formation as a straight down-sloping line defined over $c \in [0, 2]$. Consider two cases:

Case 1. $c = 2$. In this case, the Defender does not attain any value from network's connectivity and her payoff is equal to the expected loss in the encounter stage.

Now note that line equations associated with star and path networks have the lowest possible slope among all the networks, since both of those networks have the lowest possible overall nodes' value, which is equal to $S = -2(n - 1)$. However, we know from Claim 4.1 that the Defender minimises her losses whenever she chooses a star, and, therefore, an equation associated with star has a larger intercept. An intercept of a line associated with a star network is $I = 2(n - 1) - \frac{(n-1)(n-\delta)}{(n-1)^2+1}$. It implies that a star is certainly preferred to a path network around $c = 2$. Therefore, there must exist some point $c_u < 2$ in the neighbourhood of 2, where the line spawned by star formation is intersected by some line spawned by the network, which is optimal for the Defender if $c = c_u - \epsilon$ (where ϵ is some small positive real number). We call c_u an upper-cost boundary.

Case 2. $c = 0$. If the edges are free for the Defender, the expected value function becomes:

$$U_G(\vec{v}_n) = \sum_{i=1}^n v_i + D(\vec{v}_k), \quad (\text{C.9})$$

which is also convex and, therefore, must attain its maximum at some extreme point of the set it is defined over.

Firstly, observe that (C.9) is maximised whenever the Defender chooses a complete network. Consider the first derivative of the function w.r.t. to the value of some vertex inside the Attacker's support:

$$\frac{\partial U_G(\vec{v}_n)}{\partial v_i} = 1 + \underbrace{\frac{(k - \delta) \prod_{j=1}^k v_j \sum_{j \in K \setminus \{i\}} \prod_{i \neq j} v_i}{\left(\sum_{j=1}^k \prod_{i \neq j} v_i \right)^2}}_{m_1} - \underbrace{\frac{(k - \delta) \prod_{j=1}^k v_j}{v_k \sum_{j=1}^k \prod_{i \neq j} v_i}}_{m_2} > 0, \quad (\text{C.10})$$

where K is the set of indices of all vertices in the Attacker's support.

Observe that term m_1 is positive and term $m_2 \leq 1$ as it is exactly the Attacker's support condition (4.4). It follows that the function (C.9) always increases in the value of vertex inside the Attacker's support. If some vertex j is outside the Attacker's support the first derivative of $U_G(\vec{v}_n)$ w.r.t. to v_j is constant and equal to one, $\frac{\partial U_G(\vec{v}_n)}{\partial v_i} = 1$.

The vertices outside the Attacker's support are also bounded from above by the Attacker's support condition. Given that the Defender chooses $v_i = v^q = n - 1$, we can write down the upper boundary for the vertices outside the Defender's support as $v_j \leq \frac{(k - \delta)(n - 1)}{k}$. It is now left to demonstrate that if $c = 0$, the Defender is strictly better off to choose a complete network than any other network which does not have full support:

$$n(n - 1) - \frac{(n - \delta)(n - 1)}{n} > k(n - 1) + (n - k - 1) \frac{(k - \delta)(n - 1)}{k},$$

which is true for any $n > k > 2$.

We deduce that it must be the case that at extreme point $c = 0$, the Defender receives a maximum possible payoff whenever she chooses a complete network formation. It also follows that the line spawned by complete network formation must have the highest intercept point with the ordinate axis. Given that the line also must have the steepest slope, we conclude that there must exist some dot $c_l > 0$ in the neighbourhood of 0 where it is intersected by some other line spawned by the network formation, which is optimal if $c = c_l + \epsilon$. We call c_l a lower cost-boundary. ■

Proof of Proposition 4.3. Observe that the value of the network (4.1) strictly increases in v_i for any $c < 2$. We also know from Lemma 4.1 that the expected damage from an attack weakly increases in the value of a single vertex v_i .

Now observe the indifference condition for the Defender between some dense structure G_D and some sparse structure G_S :

$$V_D - \frac{c}{2}V_D + D_D = V_S - \frac{c}{2}V_S + D_S,$$

or:

$$c = 2 + \frac{D_D - D_S}{V_D - V_S} = c_i. \quad (\text{C.11})$$

Now observe that the dense graph is preferred to the sparse graph if $c < c_i$, and vice versa if $c > c_i$. As this fact applies to any pair of network structures, it must be the case that less dense maximising degree sequences are associated with a higher cost of a single edge, while higher density maximising degree sequences are associated with a lower edge cost. ■

Proof of Proposition 4.4. If the Defender possesses $n \rightarrow \infty$ vertices and $\delta \geq 1$ defensive resources, there might be ∞ potential maximising degree sequences and ∞ cost boundaries as a result. The only fact we know about all other non-extreme potential maximising degree sequences is that they must be of lower density than the complete network. Consider some cost boundary c_i^f , which is derived by comparing the complete network and some other network with Attacker's support $k \in [\delta + 1, n]$:

$$\begin{aligned} c_i^f &= 2 - \frac{\frac{(n-\delta)(n-1)}{n} - \frac{n-\delta}{\sum_{i=1}^k \frac{1}{(n-1)e_i}}}{n(n-1) - \sum_{i=1}^n (n-1)e_i} = \\ &= 2 - \frac{(n-d) \left(\frac{1}{n} - \frac{1}{\sum_{i=1}^k \frac{1}{e_i}} \right)}{n - \sum_{i=1}^n e_i}, \end{aligned} \quad (\text{C.12})$$

where $e_i = \frac{v_i}{n-1}$ and v_i is the value of some node i of some other network. Since $v_i \in [1, n-1]$, it follows that $e_i \in \left[\frac{1}{n-1}, 1 \right]$.

Thus, e_i is a coefficient calculated as a portion of the number of connections incident upon a vertex and the maximum possible connections that the node might have $(n - 1)$. Given that we assume that nodes' values of each network are weakly ordered it follows that coefficients must be order weakly ordered too: $e_1 \geq e_2 \geq \dots \geq e_n$. Note that any incomplete network must have $1 \leq q < n$ completely connected vertices, for which $e_i^C = 1$, and $n - q$ vertices with lower degree centralities, for which $e_i^{NC} \in [\frac{1}{n-1}, 1)$.

Now, let

$$HM(e_1, \dots, e_k) = \sum_{i=1}^k \frac{1}{e_i}$$

be the Harmonic Mean of first k coefficients and

$$AM(e_1, \dots, e_n) = \frac{1}{n} \sum_{i=1}^n e_i$$

be the Arithmetic Mean of all n coefficients.

Then expression (C.12) can be rewritten as:

$$c_i^f = 2 - \left(\frac{n - \delta}{n^2} \right) \underbrace{\left(\frac{1 - HM(e_1, \dots, e_k)}{1 - AM(e_1, \dots, e_n)} \right)}_{\tau}. \quad (\text{C.13})$$

Now we consider the upper-bound for the numerator of τ and the lower bound of the denominator of τ .

The upper-bound of the numerator of τ can be expressed as (Sýkora, 2009):

$$1 - HM(e_1, \dots, e_k) \leq 1 - \min(e_1, \dots, e_k) = 1 - e_k < 1.$$

Consider now the denominator of τ . Since e_i 's are nonincreasing and $e_i < 1$ for all $i > q$, there must exist $\iota \in \mathbb{N}$ such that for all $n \geq \iota$:

$$AM(e_1, \dots, e_n) \leq AM(e_1, \dots, e_\iota) < 1.$$

Then the lower bound of the denominator of τ can be expressed as:

$$1 - AM(e_1, \dots, e_n) \geq 1 - AM(e_1, \dots, e_\iota) = \tau_L > 0,$$

which is a positive constant for any fixed number $\iota \in \mathbb{N}$ and $\iota > q$.

Thus, we can now write the upper and lower bounds of expression (C.13):

$$2 - \left(\frac{n - \delta}{n^2} \right) \left(\frac{1}{\tau_L} \right) \leq c_i^f \leq 2,$$

where the upper boundary is achieved when the Attacker has only completely connected nodes in his support, $k \in [1, q]$.

Finally, consider now the limit of the lower boundary of c_i^f when $n \rightarrow \infty$:

$$\lim_{n \rightarrow \infty} 2 - \left(\frac{n - \delta}{n^2} \right) \left(\frac{1}{\tau_L} \right) = 2.$$

Then, by Squeeze Theorem, $c_i^f = 2$ as $n \rightarrow \infty$ (Weisstein, n.d.). ■

Proof of Lemma 4.2. Consider the first derivative of expected utility function of the Defender (4.25) w.r.t. v_j^s :

$$\frac{\partial U^P}{\partial v_j^s} = 1 - \frac{c}{2} \geq 0.$$

Therefore, the Defender must always choose the maximum possible value for the nodes outside the Attacker's support. Observe that the upper limit of a value of a node outside that Attacker's support is given by constraint (4.18) and is equal to the expected gain from an Attack on top k nodes. ■

Proof of Lemma 4.3. From Proposition 4.2, we know that a complete and a star network must be maximisers at the extreme values of c . Therefore, we sequentially compare all the sequences (4.23)-(4.30) against a star network and a complete network and determine which sequences can act as maximisers for intermediate values of c .

First, observe the Defender's expected payoffs from a game on a star network:

$$U_* = 2 \left(1 - \frac{c}{2} \right) (n - 1) - \frac{(n - 1)^2}{(n - 1)^2 + 1},$$

and complete network:

$$U_C = n \left(1 - \frac{c}{2} \right) (n - 1) - \frac{(n - 1)^2}{n}.$$

Intermediate sequence An expected Defender's payoff from a game on an intermediate sequence (4.23) can be stated as follow:

$$U_I = \left(1 - \frac{c}{2}\right) \left(\frac{(n-1)(q-1)(n-q)}{q} + (n-1)q\right) - \frac{(n-1)(q-1)}{q},$$

where $q \in [2, n-1]$.

The Defender prefers an intermediate sequence to a star network when:

$$U_I > U_*,$$

which is satisfied IFF:

$$c < 2 - \frac{2n^2(q-1) - n(6q-4) + 6q-4}{(n^2-2n+2)(n(q-1)-q)} = c_u. \quad (\text{C.14})$$

Similarly, the Defender prefers an intermediate sequence to a complete network when:

$$U_I > U_c,$$

which is satisfied when:

$$c > 2 - \frac{2}{n} = c_l.$$

Now observe that $c_u > c_l$:

$$2 - \frac{2n^2(q-1) - n(6q-4) + 6q-4}{(n^2-2n+2)(n(q-1)-q)} > 2 - \frac{2}{n}. \quad (\text{C.15})$$

Rearranging and simplifying (C.15) yields:

$$(n-2)n(n^2-2n+2)q(n(q-1)-q) > 0,$$

which is always satisfied for any $n > q \geq 2$.

Therefore, if $c_l < c < c_u$, the optimal choice for the Defender is an intermediate sequence.

Now observe, that if $c_l < c < c_u$, the optimal intermediate sequence has $q = 2$.

To see that observe the first derivative of U_I w.r.t. to q :

$$\frac{\partial U_I}{\partial q} = -\frac{(n-1)((c-2)n+2)}{2q^2},$$

which is always positive if $c > 2 - \frac{2}{n} = c_l$ and negative otherwise.

Since when $c > c_l$ it is always optimal for the Defender to choose a complete network, it must be the case that in the range of intermediate sequence optimality, the Defender must choose an intermediate sequence with $q = 2$. Then the condition (C.14) becomes:

$$2 - \frac{2n - 4}{n^2 - 2n + 2}.$$

Lower family The Defender's expected payoff in a game in which she chooses a lower family sequence (4.24) is:

$$U_L = \left(1 - \frac{c}{2}\right) \left(\frac{n-1}{n-2} + 2n - 3\right) - 1.$$

The Defender prefers a lower family sequence to the star network when:

$$U_L > U_*,$$

which is satisfied when:

$$c < 2 - \frac{2n - 3}{n^2 - 2n + 2} = c_L^1.$$

The Defender prefers a lower family sequence to an intermediate sequence with $q = 2$ when:

$$U_L > U_I,$$

which is satisfied when:

$$c > 2 - \frac{2n - 3}{n^2 - 2n + 2} = c_L^2.$$

Since $c_L^1 = c_L^2$, the Defender never prefers a lower family sequence to a star or an intermediate sequence. However, she is indifferent between star, lower family sequence, and intermediate sequence if $c = c_L^1 = c_L^2$.

k-family An expected payoff in a game in which the Defender chooses a k-family sequence (4.28) is:

$$U_K = \left(1 - \frac{c}{2}\right) \left(\frac{k^2}{k-1} - k + n\right) - 1.$$

Observe that the Defender prefers a k-family sequence to a star sequence if:

$$U_K > U_*,$$

which satisfied IFF:

$$c > 2 + \frac{2(k-1)}{((n-2)n+2)(k(n-3)-n+2)} > 2,$$

which is always larger than 2 for any $n \geq 4$.

It implies that any sequence of the k-family cannot be a maximiser for the relaxed optimisation problem.

Quasi-star partial family The Defender's expected payoff from a game when she chooses a family (4.29) is:

$$U_{QS} = \left(1 - \frac{c}{2}\right) \left(\frac{(k-1)(n-k)}{\frac{(k-1)^2}{k} + \frac{1}{n-1}} + k + n - 1\right) - \frac{k-1}{\frac{(k-1)^2}{k} + \frac{1}{n-1}}.$$

First, observe that $U_{QS} > U_*$ IFF:

$$c < 2 - \frac{2(n-1) \left(k^2 + k(n-3)n + k - (n-1)^2\right)}{((n-2)n+2) \left(k^2 + k(n-2)(n-1) - (n-1)^2\right)} = c_{qs}^1.$$

Now observe that $U_{QS} > U_C$ IFF:

$$c > 2 - \frac{2(n-1) \left((k-1)^2 n^2 - 3(k-2)kn + (k-3)k - 2n + 1\right)}{n \left(k^2(n((n-4)n+5) - 3) - 2k(n-2)(n-1)^2 + (n-1)^3\right)} = c_{qs}^2.$$

We now demonstrate that $c_{qs}^2 > c_{qs}^1$:

$$\begin{aligned} 2 - \frac{2(n-1) \left((k-1)^2 n^2 - 3(k-2)kn + (k-3)k - 2n + 1\right)}{n \left(k^2(n((n-4)n+5) - 3) - 2k(n-2)(n-1)^2 + (n-1)^3\right)} &> \\ 2 - \frac{2(n-1) \left(k^2 + k(n-3)n + k - (n-1)^2\right)}{((n-2)n+2) \left(k^2 + k(n-2)(n-1) - (n-1)^2\right)}, & \end{aligned} \quad (C.16)$$

Rearranging and simplifying (C.16) yields:

$$\begin{aligned}
 & \underbrace{\left(k^2(n((n-4)n+5)-3)-2k(n-2)(n-1)^2+(n-1)^3\right)}_{\tau_1} \\
 & \underbrace{\left(k^2+k(n-2)(n-1)-(n-1)^2\right)}_{\tau_2} \underbrace{\left(k(k(n-1)-2n+3)+n-1\right)}_{\tau_3} \quad (C.17) \\
 & \underbrace{\left(k^2+k((n-4)n+2)-(n-1)^2\right)}_{\tau_4} > 0
 \end{aligned}$$

We now consider first derivatives of τ_1 , τ_2 , τ_3 , and τ_4 w.r.t. k .

τ_1 :

$$\frac{\partial \tau_1}{\partial k} = 2k(n((n-4)n+5)-3)-2(n-2)(n-1)^2. \quad (C.18)$$

To analyse the sign of (C.18) we now consider the second derivative of τ_1 w.r.t. k :

$$\frac{\partial^2 \tau_1}{\partial k^2} = 2(n((n-4)n+5)-3) > 0, \quad (C.19)$$

which is always positive for any $n > 3$. Now consider the sign of the first derivative (C.18) when $k = 2$:

$$\frac{\partial \tau_1}{\partial k} = 2n((n-4)n+5)-8 > 0, \quad (C.20)$$

which is always positive for any $n > 3$.

τ_2 :

$$\frac{\partial \tau_2}{\partial k} = 2k + (n-2)(n-1) > 0,$$

which is always positive for any $n > k \geq 2$.

τ_3 :

$$\frac{\partial \tau_3}{\partial k} = 2k(n-1)-2n+3 > 0.$$

which is always positive for any $n > k \geq 2$.

τ_4 :

$$\frac{\partial \tau_4}{\partial k} = 2k + (n-4)n + 2 > 0.$$

which is always positive for any $n > k \geq 2$.

Since $\frac{\partial \tau_j}{\partial k} > 0$, $\forall j \in \{1, 2, 3, 4\}$, it is sufficient to verify condition (C.17) when $k = 2$. Substituting $k = 2$ and simplifying (C.17) yields:

$$(n-2)(n-1)n(n+1)((n-6)n+7)((n-4)n+7)((n-2)n+2)(n((n-3)n+3)-5) > 0,$$

which is always positive for any $n \geq 5$.

It implies that sequence (4.29) cannot be a maximiser for the relaxed maximisation problem.

Lower partial family The Defender's expected payoff from a game when she chooses a sequence (4.30):

$$U_{LP} = \left(1 - \frac{c}{2}\right) \left(\frac{k(n-2)}{k-1} + \frac{k(n-1)}{k(n-2)-n+1} + n-1\right) - \frac{k}{k-1}.$$

Observe that $U_{LP} > U_*$ IFF:

$$c < 2 - \frac{2(k+(n-1)^2)(k(n-2)-n+1)}{((n-2)n+2)(k^2+k(n-2)(n-1)-(n-1)^2)} = c_{LP}^1.$$

Similarly, $U_{LP} > U_I$ IFF:

$$c > 2 - \frac{2kn-4k-2n+2}{k((n-2)n+2)-n^2+n} = c_{LP}^2.$$

Now observe that $c_{LP}^2 > c_{LP}^1$:

$$2 - \frac{2(k+(n-1)^2)(k(n-2)-n+1)}{((n-2)n+2)(k^2+k(n-2)(n-1)-(n-1)^2)} < 2 - \frac{2kn-4k-2n+2}{k((n-2)n+2)-n^2+n}.$$

Rearranging and simplifying yields:

$$\lambda \left(k^2 + k(n-2)(n-1) - (n-1)^2\right) (k(n-2)-n+1) \left(k((n-2)n+2) - n^2 + n\right) > 0, \quad (\text{C.21})$$

where $\lambda = (k-1)(n-2)(n-1)^2((n-2)n+2)(k(n-2)-n+1) > 0$.

Dividing both sides of (C.21) by λ yields:

$$\underbrace{\left(k^2 + k(n-2)(n-1) - (n-1)^2\right)}_{\iota_1} \underbrace{\left(k((n-2)n+2) - n^2 + n\right)}_{\iota_2} > 0.$$

Now consider first derivatives of ι_1 and ι_2 w.r.t. k :

$$\frac{\partial \iota_1}{\partial k} = 2k + (n-2)(n-1) > 0, \quad \frac{\partial \iota_2}{\partial k} = (n-2)n + 2 > 0,$$

which are always positive.

Thus, it is sufficient to verify that the condition is satisfied when $q = 2$:

$$(n^2 - 4n + 7)(n^2 - 3n + 4) > 0,$$

which is always positive.

It follows that sequence (4.30) cannot be a maximiser either.

■

Proof of Claim 4.10. First consider family (4.24). The second largest element of this family always has the following form:

$$v_2 = \frac{n-1}{n-2} = 1 + \frac{1}{n-1},$$

which is not integer.

Now consider family (4.28). The first k elements of the k -family have the following form:

$$\frac{k}{k-1} = 1 + \frac{1}{k-1},$$

which is integer only when $k = 2$.

Thus, the only sequence which belongs to k -family and graphic is a degree sequence with $n = 4$ and $k = 2^3$. However, this sequence is always dominated by a

³Any sequence with $n > 4$ does not pass the Erdos-Gallai constraint (4.10) and, therefore, is not graphic.

star network. To see then consider the Defender's expected utility functions from games on a star network (U_*) and on a k -family network (U_K), when $n = 4$ and $q = 2$:

$$U_* = 2\left(1 - \frac{c}{2}\right)(n-1) - \frac{(n-1)^2}{(n-1)^2 + 1} = 51/10 - 3c,$$

$$U_K = \left(1 - \frac{c}{2}\right) \left(\frac{2k}{k-1} - k + n\right) - 1 = 5 - 3c.$$

Thus, $U_* > U_K$.

Now consider a quasi-star family (4.29). As was stated above any element which has value $\frac{k}{k-1}$ is integer IFF $k = 2$. Then elements indexed $j \in (k, n]$ have value:

$$v_k = \frac{k-1}{\frac{(k-1)^2}{k} + \frac{1}{n-1}} = \frac{2(n-1)}{n+1} = 2 - \frac{4}{n+1},$$

which is integer IFF $n = 3$.

However, eqrefquasi-star family sequence with $n = 3$ elements and $k = 2$ is not graphic $(2, 2, 1)$.

The last sequence to consider is *lower partial family* sequence. However, this family exists only when $k \geq 3$, and therefore any element of value $\frac{k}{k-1}$ is not integer. ■

Proof of Lemma 4.4. The Defender's expected payoff at the beginning of the game on a star network:

$$U_* = 2\left(1 - \frac{c}{2}\right)(n-1) - \frac{(n-1)^2}{(n-1)^2 + 1},$$

and on a complete network:

$$U_C = n\left(1 - \frac{c}{2}\right)(n-1) - \frac{(n-1)^2}{n}.$$

An expected Defender's payoff from a game on a sequence of the intermediate family (4.23) is:

$$U_I = \left(1 - \frac{c}{2}\right) \left(\frac{(n-1)(q-1)(n-q)}{q} + (n-1)q\right) - \frac{(n-1)(q-1)}{q},$$

where $q \in [2, n-1]$.

An intermediate sequence is preferred to a star network if:

$$U_I > U_*,$$

which is satisfied IFF:

$$c < 2 - \frac{2n^2(q-1) - n(6q-4) + 6q-4}{(n^2-2n+2)(n(q-1)-q)} = c_u^q. \quad (\text{C.22})$$

Similarly, the Defender prefers an intermediate sequence to a complete network when:

$$U_I > U_C,$$

which is satisfied when:

$$c > 2 - \frac{2}{n} = c_l.$$

Now observe that $c_u^q > c_l$:

$$2 - \frac{2n^2(q-1) - n(6q-4) + 6q-4}{(n^2-2n+2)(n(q-1)-q)} > 2 - \frac{2}{n}. \quad (\text{C.23})$$

Rearranging and simplifying (C.23) yields:

$$(n-2)n(n^2-2n+2)q(n(q-1)-q) > 0,$$

which is always satisfied for any $n > q \geq 2$.

Therefore, if $c_l < c < c_u^q$, the optimal choice for the Defender is an intermediate family sequence. ■

Proof of Lemma 4.5. Consider the first derivative of the objective function of maximisation problem (4.5) w.r.t. q :

$$\frac{\partial U_I}{\partial q} = -\frac{(n-1)(n(c-2)+2)}{2q^2}. \quad (\text{C.24})$$

Observe that the sign of (C.24) does not depend on q itself, but rather on the cost of connection, c , relative to the overall number of vertices. Partial derivative $\frac{\partial U_I}{\partial q}$ is positive whenever $c < 2 - \frac{2}{n}$ and negative whenever $c > 2 - \frac{2}{n}$. The Defender is indifferent between any quantity of completely connected nodes if $c = 2 - \frac{2}{n}$. ■

Proof of Proposition 4.5. Consider expected payoffs of the Defender at the beginning of the game on a star network:

$$U_* = 2 \left(1 - \frac{c}{2}\right) (n-1) - \frac{(n-1)^2}{(n-1)^2 + 1},$$

and on a complete network:

$$U_C = n \left(1 - \frac{c}{2}\right) (n-1) - \frac{(n-1)^2}{n}.$$

A star network is preferred to a complete network if $U_* > U_C$, which is true IFF:

$$c > 2 - \frac{2(n^2 - 2 + 1)}{n(n^2 - 2n + 2)} = c_i$$

■

Proof of Lemma 4.6. Suppose the Defender must choose between two network structures G_1 and G_2 . Assume that the overall value of degree sequence G_1 is larger than the overall value of G_2 : $V_1 > V_2$. Then it must be the case that the expected payoff of the Defender in the encounter stage from the game on G_1 is weakly smaller than an expected payoff on the game G_2 : $D_1 \leq D_2$. Now observe that the Defender is indifferent between the two structures IFF:

$$c_{12} = 2 - \frac{(1 - \alpha)(D_2 - D_1)}{\alpha(V_1 - V_2)}. \quad (\text{C.25})$$

Now consider the first derivative of the boundary (C.25) with respect to α :

$$\frac{\partial c_{12}}{\partial \alpha} = \frac{D_2 - D_1}{\alpha^2 (V_1 - V_2)} \geq 0.$$

Therefore, any cost boundary weakly increases in α .

■

Proof of Claim 4.11. Consider weighted expected utility functions of the Defender from games on a star, complete, and m-maxi-core networks:

$$U_*^\alpha = 2\alpha \left(1 - \frac{c}{2}\right) (n-1) - \frac{(1-\alpha)(n-1)^2}{(n-1)^2 + 1}, \quad (\text{C.26})$$

$$U_C^\alpha = \alpha \left(1 - \frac{c}{2}\right) (n-1)n - \frac{(1-\alpha)(n-1)^2}{n}, \quad (\text{C.27})$$

$$U_{MC}^\alpha = \alpha \left(1 - \frac{c}{2}\right) \left(\frac{(n-1)(q-1)(n-q)}{q} + (n-1)q \right) - \frac{(1-\alpha)(n-1)(q-1)}{q}. \quad (\text{C.28})$$

An m-maxi-core network is preferred to a star and complete network IFF:

$$2 + \underbrace{\frac{2(\alpha-1)}{\alpha n}}_{=c_l^\alpha} < c < 2 + \underbrace{\frac{2(\alpha-1)(n(n(q-1)-3q+2)+3q-2)}{\alpha((n-2)n+2)(n(q-1)-q)}}_{=c_u^\alpha}.$$

Now consider the size of optimality interval $\delta_{MC} = c_u^\alpha - c_l^\alpha$:

$$\delta_{MC} = \frac{2(\alpha-1)(n-2)q}{\alpha n((n-2)n+2)(n(-q)+n+q)}.$$

Now consider the first derivative of δ_{MC} w.r.t α :

$$\frac{\partial \delta_{MC}}{\partial \alpha} = \frac{2(n-2)q}{\alpha^2 n((n-2)n+2)(n+q-nq)} < 0.$$

Since $(n+q-nq) < 0$, it follows that $\frac{\partial \delta_{MC}}{\partial \alpha}$ is always negative. ■

C.2 Supplementary material

Erdos-Gallai theorem is a known results from graph theory proposed by Erdős and Gallai (1960). The theorem solves the graph realisation problem, i.e. provides a set of necessary and sufficient conditions for a finite sequence of positive integer numbers to be the degree sequence of a graph. The sequence that satisfies the conditions is said to be “graphic”. The original theorem demonstrated that a sequence of positive integers $d_1 \geq \dots \geq d_n$ is graphic IFF $\sum_{i=1}^n d_i$ is even and:

$$\sum_{i=1}^k d_i \leq k(k-1) + \sum_{i=k+1}^n \min(d_i, k)$$

holds for $\forall k \in [1, n]$.

Therefore, Erdos-Gallai conditions ensure that any given node does not have more links than can be supported by the degrees of all other nodes in the network.

C.3 Extreme points characteristics and the convex hull

The non-convexity of set Υ is caused by non-convex constraints (4.17). Given that non-convex constraints are defined by smooth concave functions (by Lemma 4.1), we can construct a convex hull of set Υ , $\text{Conv}(\Upsilon)$ by replacing them with linear constraints defined by hyperplanes spanned by extreme intersection points of non-convex constraints (4.17) and the surface of a hypercube defined by linear constraints (C.29) (Boyd & Vandenberghe, 2004). We denote extreme points defined by intersections of non-convex constraints (4.17) and linear constraints (4.16) as *extreme points of non-convex regions*.

This process is often referred to as set “convexification” (Li et al., 2001). However, the convexification process for complete and partial support cases is different. While in the partial support case, there are $n - k$ convex constraints (4.18), there are no convex constraints in the full support case. Thus, the set of extreme points of non-convex regions can differ. We consider full and partial support cases separately.

C.3.1 Full support

Consider set Υ^C defined by the following inequality constraints:

$$1 \leq v_i \leq n - 1 \quad \forall i \in [1, n], \quad (\text{C.29})$$

$$\frac{n - 2}{\sum_{j \in N \setminus \{i\}} \frac{1}{v_j}} - v_i \leq 0 \quad \forall i \in [1, n], \quad (\text{C.30})$$

To construct $\text{Conv}(\Upsilon^C)$, we first characterise all of the extreme points of non-convex regions of the original set Υ^C . We denote the set of extreme points of Υ^C as $E(\Upsilon^C)$. There might be two types of extreme points:

Type 1 extreme points defined by the intersection of a single non-convex constraint (C.30) and the surface of a hypercube (C.29);

Type 2 extreme points defined by the intersection of $t \geq 2$ non-convex constraints (C.30) and the surface of a hypercube (C.29).

Type 1 extreme points. Consider a non-convex constraint for the value of some node $x \in N$:

$$v_x \geq \frac{n-2}{\sum_{j \in N \setminus \{x\}} \frac{1}{v_j}} = A_x, \quad (\text{C.31})$$

where A_x is an expected gain from an attack on nodes indexed $j \in N \setminus \{x\}$. Non-convex constraint intersects the surface of a hyperplane (C.29) in two cases:

Case 1a. Node x takes values $v_x = A_x \in \{1, n-1\}$;

Case 2a. Nodes indexed $j \in N \setminus \{x\}$ take extreme values $v_j \in \{1, n-1\}$ and $v_x = A_x$.

Also observe that expected gain from an attack on nodes indexed $j \in N \setminus \{x\}$ does not depend on the order of values in set $N \setminus \{x\}$. It follows that for a given set of values that satisfy (C.31) there are $(n-1)!$ ways to write coordinates of some point in which coordinate v_x has a fixed value, while all other coordinates permute. We denote the set of all $(v_{\sigma(1)}, \dots, v_{\sigma(n)})$ as σ ranges over permutations of $\{1, \dots, n\}$ that fix coordinate x , as a *family of points*. To simplify the notation, we represent each family of points with a partially ordered set in which first $n-1$ elements v_j indexed $j \in N \setminus \{x\}$ are ordered from the largest to the smallest and the last element is v_x —we denote this partially ordered set as the *composition of coordinates*.

Case 1a: $v_x = A_x \in \{1, n-1\}$. First observe that any point which has coordinate $v_x = n-1$ cannot be an extreme point of non-convex region. This is because the maximum possible expected value from an attack on $n-1$ nodes is $\max_{v_j} A = \frac{(n-1)(n-2)}{n-1} = n-2 < n-1$, where $i \in N \setminus \{x\}$.

Claim C.1. Any point which has coordinate $v_x = A_x = n-1$ is not an extreme point of non-convex region of Υ^C .

Now consider some point which has coordinate $v_x = 1$. First observe that when $v_x = 1$, the nodes indexed $j \in N \setminus \{x\}$ cannot all simultaneously attain extreme values $v_j \in \{1, n-1\}$, i.e. $A_x = 1$.

Claim C.2. *If $v_j \in \{1, n-1\}$, $\forall j \in N \setminus \{x\}$, then $A_x \neq 1$.*

Proof of Claim C.2. Consider some point in which q coordinates have value $v^q = n-1$, and $n-q-1$ coordinates have value $v^u = 1$ and $v_x = A_x = 1$. Substituting the sequence with this composition of coordinates in (C.31) yields:

$$A_x = 1 = \frac{n-2}{\frac{q}{n-1} + n-q-1}. \quad (\text{C.32})$$

Solving (C.32) for q yields:

$$q = \frac{n-1}{n-2},$$

which is not feasible since q is the number of coordinates that attain value $v^q = n-1$, $q \in \mathbb{Z}$. ■

However, we can still find a point which has composition of coordinates, i.e. $v_x = 1$, $v_j \in \{1, n-1\}$, $\forall j \in N \setminus \{x, y\}$, and some coordinate $v_y \in [1, n-1]$. The composition of coordinates of the family mentioned above can be written as follows:

$$\mathbf{v}_b^q = \left(\underbrace{n-1, \dots, n-1}_{\times q}, v_y, \underbrace{1, \dots, 1}_{\times (n-q-2)}, \underbrace{1}_{=A_x} \right), \quad (\text{C.33})$$

where q is a number of nodes which attain value $n-1$.

We now demonstrate that the only case in which $\mathbf{v}_b^q \in E(\Upsilon^C)$ is when $q = 1$.

Claim C.3. $\mathbf{v}_b^q \in E(\Upsilon^C)$ IFF: $q = 1$. In this case, $v_y = \frac{n-1}{n-2}$.

Proof of Claim C.3. First, consider points which have composition of coordinates $\mathbf{v}_b^0$ in which $q = 0$:

$$\mathbf{v}_b^0 = \left(v_y, \underbrace{1, \dots, 1}_{\times (n-2)}, \underbrace{1}_{=A_x} \right). \quad (\text{C.34})$$

Substituting (C.34) into (C.31) yields:

$$A_x^0 = 1 = \frac{n-2}{(n-2) + \frac{1}{v_y}},$$

which does not have a solution for $v_y \in \mathbb{R}_+$.

Now consider some points which have composition of coordinates $\mathbf{v}_b^{q \geq 1}$. Substituting (C.33) in (C.31) yields:

$$A_x^{q \geq 1} = 1 = \frac{n-2}{\frac{q}{n-1} + (n-2-q) + \frac{1}{v_y}}. \quad (\text{C.35})$$

Solving equation (C.35) for v_y yields:

$$v_y = \frac{n-1}{(n-2)q}. \quad (\text{C.36})$$

Observe that the RHS of equation (C.36) is smaller than 1 for any $q \geq 2$. Therefore, any point with composition of coordinates $v_b^{\geq 2}$ lies beyond a hypercube defined by linear constraints (C.29) and $\mathbf{v}_b^{q \geq 2} \notin E(\Upsilon^C)$.

Thus, the only possibility in which $v_b^{q \geq 1} \in E(\Upsilon^C)$ is when $q = 1$. In this case:

$$v_y = \frac{n-1}{n-2} > 1.$$

■

It follows from Claims C.1 and C.3 that the only family of extreme points that *Case 1a* yields must have the following composition of coordinates:

$$\mathbf{v}_b = \left(n-1, \frac{n-2}{n-1}, \underbrace{1, \dots, 1}_{\times(n-3)}, \underbrace{1}_{=A_x} \right). \quad (\text{C.37})$$

As in sequence (C.37), coordinate v_x attains the lowest possible value $v_x = 1$, we denote $\mathbf{v}_b$ as the *lower family of extreme points*.

Case 2a. $v_j \in \{1, n-1\}$, $\forall j \in N \setminus \{x\}$. In this case, any point has the following composition of the coordinates:

$$\mathbf{v}_t^q = \left(\underbrace{n-1, \dots, n-1}_q, \underbrace{1, \dots, 1}_{n-q-1}, A_x \right), \quad (\text{C.38})$$

where q is the number of nodes that attain value $n - 1$.

The value A_x given the composition (C.38) can be expressed as:

$$A_x^q = \frac{(n-2)(n-1)}{(n-1)(n-q-1) + q}. \quad (\text{C.39})$$

First, we demonstrate that families of extreme points (C.38) with $q = \{0, 1\}$ are not extreme points of set Υ^C .

Claim C.4. $\mathbf{v}_t^q \notin E(\Upsilon^C)$ if $q \in \{0, 1\}$.

Proof of Claim C.4. Substituting composition (C.38) in which $q = 0$, $\mathbf{v}_t^0$ into (C.39) yields:

$$A_x^0 = \frac{(n-2)(n-1)}{(n-1)(n-1)} = \frac{n-2}{n-1} < 1.$$

Therefore, points of family $\mathbf{v}_t^0 \notin E(\Upsilon^C)$.

Substituting composition (C.38) in which $q = 1$, $\mathbf{v}_t^1$ into (C.39) yields:

$$A_x^1 = \frac{(n-2)(n-1)}{(n-1)(n-2) + 1} < 1.$$

Thus, points of family $\mathbf{v}_t^1 \notin E(\Upsilon^C)$ as well. ■

We now demonstrate that the only family of extreme points that is included in $E(\Upsilon^C)$ must have the composition of coordinates (C.38) with $q = n - 1$.

Claim C.5. $\mathbf{v}_t^q \in E(\Upsilon^C)$ IFF: $q = n - 1$.

Proof of Claim C.5. Consider families of points $\mathbf{v}_t^q$ with $q \in [2, n - 2]$. Any of those families has at least one node of value $v^u = 1$. Now consider non-convex constraint (C.30) written for some node $l \in N \setminus \{x\}$ that attains value v^u :

$$v_l \geq \underbrace{\frac{(n-2)^2}{n^2 - n(q+3) + 2q + 3}}_{\iota}. \quad (\text{C.40})$$

Consider the first derivative of ι w.r.t. q :

$$\frac{\partial \iota}{\partial q} = \frac{(n-2)^3}{(n^2 - n(q+3) + 2q + 3)^2} > 0.$$

Therefore, it is sufficient to rule-out (C.40) when $q = 2$:

$$v_l \geq \frac{(n-2)^2}{(n-5)n+7}. \quad (\text{C.41})$$

Substituting $v_l = 1$ in (C.41) and simplifying yields:

$$n \leq 3.$$

Therefore, constraint (C.40) is never satisfied for any $n > 3$. It follows that $\mathbf{v}_t^q \notin E(\Upsilon^C)$, $\forall q \in [2, n-2]$.

The only case left to verify is when $q = n-1$. Substituting $\mathbf{v}_t^{n-1}$ into (C.39) yields:

$$A_x^{n-1} = \frac{(n-2)(n-1)}{(n-1)(n-n+1-1) + n-1} = n-2.$$

It follows that $\mathbf{v}_t^{n-1} \in E(\Upsilon^C)$ as all nodes always satisfy all constraints. $\blacksquare$

Therefore the only family of extreme points which *Case 2a* yields must have the following composition of coordinates:

$$\mathbf{v}_t = \left(\underbrace{n-1, \dots, n-1}_{\times(n-1)}, \underbrace{n-2}_{=v_x} \right). \quad (\text{C.42})$$

As in composition (C.42), coordinate $v_x = n-2$, we denote $\mathbf{v}_t$ as the *upper family of extreme points*.

Type 2 extreme points. This type of extreme points results from intersection of a surface of a hypercube (C.29) with $t \in [2, n]$ non-convex constraints (C.30).

We first characterise the function of intersection of $t \leq n$ non-convex constraints (C.30). We denote the set of nodes whose constraints intersect as $T \subseteq N$, $|T| = t$. Without loss of generality we relabel the nodes which belong to set T and write their values as a set $\mathbf{v}_z = (v_{z_1}, \dots, v_{z_t})$. Set $\mathbf{v}_z$ is not necessarily ordered. The function of intersection must satisfy the following system of equations:

$$\begin{cases} v_{z_1} = \frac{n-2}{\sum_{j \in N \setminus \{z_1\}} \frac{1}{v_j}}, \\ \vdots \\ v_{z_t} = \frac{n-2}{\sum_{j \in N \setminus \{z_t\}} \frac{1}{v_j}}. \end{cases} \quad (\text{C.43})$$

We now demonstrate that system (C.43) is satisfied IFF $v_{z_i} = \frac{n-1-t}{\sum_{j \in N \setminus T} \frac{1}{v_j}}$, $\forall z_i \in T$.

Claim C.6. *System of equations (C.43) is satisfied IFF*

$$v_{z_i} = \frac{n-1-t}{\sum_{j \in N \setminus T} \frac{1}{v_j}}, \quad \forall z_i \in T. \quad (\text{C.44})$$

Proof of Claim C.6. We prove by induction on the number of constraints that intersect, t .

Base case, $t = 2$. Consider the base case in which constraints for some nodes z_1 and z_2 intersect. The corresponding system of equations is then:

$$\begin{cases} v_{z_1} = \frac{n-2}{\frac{1}{z_2} + \sum_{j \in N \setminus \{z_1, z_2\}} \frac{1}{v_j}}, \\ v_{z_2} = \frac{n-2}{\frac{1}{z_1} + \sum_{j \in N \setminus \{z_1, z_2\}} \frac{1}{v_j}}. \end{cases} \quad (\text{C.45})$$

$$(\text{C.46})$$

Substituting (C.46) into (C.45) yields:

$$v_{z_1} = \frac{n-2}{\frac{\frac{1}{z_1} + \sum_{j \in N \setminus \{z_1, z_2\}} \frac{1}{v_j}}{n-2} + \sum_{j \in N \setminus \{z_1, z_2\}} \frac{1}{v_j}}.$$

Rearranging and simplifying yields:

$$v_{z_1} = \frac{n-3}{\sum_{j \in N \setminus \{z_1, z_2\}} \frac{1}{v_j}}. \quad (\text{C.47})$$

Substituting (C.47) into (C.46), rearranging, and simplifying yields:

$$v_{z_2} = \frac{n-3}{\sum_{j \in N \setminus \{z_1, z_2\}} \frac{1}{v_j}}. \quad (\text{C.48})$$

Thus, the system is satisfied whenever $v_{z_1} = v_{z_2} = \frac{n-3}{\sum_{j \in N \setminus \{z_1, z_2\}} \frac{1}{v_j}}$ and the claim holds.

Induction step. Assume that the statement holds for some $2 < w < n$. Now consider the case in which $w+1$ constraints intersect:

$$\begin{cases} v_{z_1} = \frac{n-2}{\sum_{j \in N \setminus \{z_1\}} \frac{1}{v_j}}, \\ \vdots \\ v_{z_{w+1}} = \frac{n-2}{\sum_{j \in N \setminus \{z_{w+1}\}} \frac{1}{v_j}}. \end{cases} \quad (\text{C.49})$$

As we assumed that statement holds for w it must be the case that solving first w equations of the system (C.49) for v_{z_i} , $i \in [1, w]$ yields:

$$v_{z_1} = v_{z_2} = \cdots = v_{z_w} = \frac{n - w - 1}{\sum_{j \in N \setminus W} \frac{1}{v_j}}, \quad (\text{C.50})$$

where $W \in \{z_1, \dots, z_w\}$.

Substituting (C.50) into equation $w + 1$ of system (C.49) yields:

$$v_{z_{w+1}} = \frac{n - 2}{\frac{\frac{1}{v_{z_{w+1}}} + \sum_{j \in N \setminus W^1} \frac{1}{v_j}}{n - 1 - w} + \sum_{j \in N \setminus W^1} \frac{1}{v_j}}, \quad (\text{C.51})$$

where $W^1 \in \{z_1, \dots, z_{w+1}\}$. Solving (C.51) for $v_{z_{w+1}}$ yields:

$$v_{z_{w+1}} = \frac{n - w - 2}{\sum_{j \in N \setminus W^1} \frac{1}{v_j}}. \quad (\text{C.52})$$

We now substitute (C.52) into (C.50):

$$v_{z_i} = \frac{n - w - 1}{\frac{\sum_{j \in N \setminus W^1} \frac{1}{v_j}}{n - w - 2} + \sum_{j \in N \setminus W^1} \frac{1}{v_j}}. \quad (\text{C.53})$$

Simplifying (C.54) yields:

$$v_{z_i} = \frac{n - w - 2}{\sum_{j \in N \setminus W^1} \frac{1}{v_j}} = v_{z_{w+1}}, \quad \forall z_i \in W. \quad (\text{C.54})$$

Thus, $v_{z_1} = \cdots = v_{z_{w+1}} = \frac{n - w - 2}{\sum_{j \in N \setminus W^1} \frac{1}{v_j}}$. ■

It immediately follows from Claim C.6 that the maximum number of non-convex constraints that can intersect must be $t = n - 2$, since otherwise system (C.43) does not have solutions in $\mathbb{R}_+^t$.

Claim C.7. *The maximum possible number of constraints that can intersect is $t = n - 2$.*

Proof of Claim C.7. Observe that intersection function (C.44) is positive IFF:

$$n - 1 - t \geq 1,$$

which can be rearranged as:

$$t \leq n - 2.$$

■

We can now characterise the extreme points that result from intersection of linear constraints (C.29) and t non-convex constraints (C.30). To simplify the notation we denote the set of nodes whose non-convex constraints are not intersected as $M = N \setminus T$. Without loss of generality we relabel nodes in set M and write the values of those nodes as a set $\mathbf{v}_x = (v_{x_1}, \dots, v_{x_{n-t}})$. Set $\mathbf{v}_x$ is not necessarily ordered.

The composition of coordinates of each extreme point that results from intersection of t non-convex constraints and linear constraints then can be expressed as:

$$\mathbf{v}_i = (v_{x_1}, \dots, v_{x_{n-t}}, v_{z_1}, \dots, v_{z_t}), \quad (\text{C.55})$$

where $x_i \in M$ and $z_i \in T$, and $v_{z_i} = \frac{n-1-t}{\sum_{j \in M} \frac{1}{v_j}}$

We now consider two cases:

Case 1b. Each node indexed z_i attains value $v_{z_i} \in \{n-1, 1\}$;

Case 2b. Each node indexed x_i attains value $v_{x_i} \in \{n-1, 1\}$.

Case 1b. $v_{z_i} \in \{1, n-1\}$. First observe that none of nodes indexed $z_i \in T$ can attain value $v_{z_i} = n-1$, as the maximum expected gain from an attack is equal to $A^c = \frac{(n-1)^2}{n} < n-1$.

Second, if nodes $z_i \in T$ each attain unit values $v_{z_i} = 1$, it yields the extreme point described by composition (C.37). Therefore, this case does not yield any new extreme points.

Case 2b. $v_{x_i} = \{1, n-1\}$. First, observe that nodes indexed $x_i \in M$ cannot all simultaneously attain unit values $v_{x_i} = v^u = 1$, as in this case $v_{z_i} < 1$, meaning that points like this lie beyond the hypercube (C.29).

Claim C.8. $\mathbf{v}_i \notin E(\Upsilon^C)$ if $v_{x_i} = 1, \forall x_i \in M$.

Proof of Claim C.8. Substituting $v_{x_i} = 1, \forall x_i \in M$ into (C.44) yields: $v_{z_i} = \frac{n-1-t}{n-t} < 1$ ■

Moreover, $\mathbf{v}_i \in E(\Upsilon^C)$ only when all nodes indexed $x_i \in M$ all attain maximum possible value $v_{x_i} = n-1$.

Claim C.9. $\mathbf{v}_i \in E(\Upsilon^C)$ IFF $v_{x_i} = n-1, \forall x_i \in M$.

Proof of Claim C.9. First, consider the case in which some node $x_b \in M$ attains value $v_b = n-1$, while the rest nodes $x_j, \forall j \in M \setminus b$ attain a unit value, $v_{x_j} = v^u = 1$. Then by Proposition 4.1, the expected attack on nodes $x_i \in M$ must be less than 1. Therefore, these points lie beyond the hypercube defined by (C.29).

Now consider the case in which $q \geq 2$ nodes indexed $x_i \in Q$ where $Q \subset M$ attain value $v_{x_i} = v^q = n-1$ while the rest $n-q-t$ nodes indexed $x_j \in M \setminus Q$ attain a unit value, $v_{x_j} = v^u = 1$. The corresponding composition of coordinates is then:

$$\mathbf{v}_i^{qu} = \left(\underbrace{n-1, \dots, n-1}_{\times q}, \underbrace{1, \dots, 1}_{\times (n-q-t)}, v_{z_1}, \dots, v_{z_t} \right),$$

where $z_i \in M$, and $v_{z_i} = \frac{n-1-t}{\sum_{j \in M} \frac{1}{v_j}}$.

However, composition $\mathbf{v}_i^{qu}$ violates the full support assumption as expected attack on $q \geq 2$ nodes of value $v^q = n-1$, $A^q = \frac{(n-1)(q-1)}{q} > 1$ for any $n > q \geq 2$.

Thus, $\mathbf{v}_i \in E(\Upsilon^C)$ IFF all nodes indexed $x_i \in M$ attain value $v_{x_i} = v^q = n-1$. ■

Therefore, coordinates of each extreme point defined by the intersection of t non-convex constraints and a hypercube defined by linear constraints (C.29) can be written as follows:

$$\mathbf{v}_i^q = \left(\underbrace{n-1, \dots, n-1}_{\times q}, \underbrace{\frac{(n-1)(q-1)}{q} \dots \frac{(n-1)(q-1)}{q}}_{\times (n-q)} \right). \quad (\text{C.56})$$

By Claim C.7, any set Υ^C must have $n-2$ families of extreme points that have composition of coordinates (C.56). Each family differs in the number of coordinates $q \in [2, n-2]$. We denote the families of extreme points which have composition of coordinates (C.56) as *intermediate families of extreme points*.

Observe that intermediate families of extreme points (C.56) and the upper family of extreme points (C.42) can be described by a composition (C.56) allowing $q \in [2, n-1]$. Thus, we add the upper family of extreme points to the intermediate extreme points.

Therefore, set Υ^C has two classes of families of extreme points:

1. lower family of extreme points:

$$\mathbf{v}_b = \left(n-1, \frac{n-2}{n-1}, \underbrace{1, \dots, 1}_{\times (n-2)} \right). \quad (\text{C.57})$$

2. intermediate families of extreme points:

$$\mathbf{v}_i^q = \left(\underbrace{n-1, \dots, n-1}_{\times q}, \underbrace{\frac{(n-1)(q-1)}{q} \dots \frac{(n-1)(q-1)}{q}}_{\times (n-q)} \right), \quad (\text{C.58})$$

where $q \in [2, n-1]$.

C.3.2 Convex hull construction for the full support case

In order to construct $\text{Conv}(\Upsilon)$, it is sufficient to find a series of hyperplanes which span through all extreme points of non-convex regions and derive the corresponding constraints.

We first derive the form of the equation of a hyperplane that spans through all of extreme points of some family.

Claim C.10. *An equation of a hyperplane that spans through all of the extreme points of some family can be expressed as:*

$$\sum_{j \in N \setminus \{x\}} v_j + C v_x + D = 0,$$

where C and D are constants.

Proof of Claim C.10. Every extreme point from the same family has one fixed coordinate v_x , while all other coordinates $v_i \in N \setminus \{x\}$ permute. Let $P = N \setminus \{x\}$. Without loss of generality we relabel all coordinates $v_i \in P$, i.e. $i \in \{p_1, \dots, p_{n-1}\}$. Then let S_P denote the set of all permutations of elements of P , i.e. each element $\sigma_j \in S_P$, $\sigma_j = (\sigma_j(v_1), \dots, \sigma_j(v_{n-1}))$, where $j \in [1, (n-1)!]$.

If there exists a hyperplane which spans through all of the points of some family then the following system of equations must be satisfied for any $y \in [1, n-1]$ and $-y \in [1, n-1]$:

$$\left\{ \begin{array}{l} \sum_{i=1}^{n-1} a_i \sigma_y(v_{p_i}) + a_x v_x + a_0 = 0, \end{array} \right. \quad (\text{C.59})$$

$$\left\{ \begin{array}{l} \sum_{i=1}^{n-1} a_i \sigma_{-y}(v_{p_i}) + a_x v_x + a_0 = 0. \end{array} \right. \quad (\text{C.60})$$

Now assume that $a_1 = a_2 = \dots = a_{n-1} = a$. Then, subtracting equation (C.60) from equation (C.59) and rearranging yields:

$$a \sum_{i=1}^{n-1} \sigma_y(v_{p_i}) + a_x v_x + a_0 = a \sum_{i=1}^{n-1} \sigma_{-y}(v_{p_i}) + a_x v_x + a_0. \quad (\text{C.61})$$

By simplifying (C.61) we obtain:

$$\sum_{i=1}^{n-1} \sigma_y(v_i) = \sum_{i=1}^{n-1} \sigma_{-y}(v_i), \quad (\text{C.62})$$

which always holds as the sum of elements of any permutation is always constant.

Therefore, the equation of a hyperplane that spans through points of one family can be expressed as follows:

$$a \sum_{j \in N \setminus \{x\}} v_j + a_x v_x + a_0 = 0. \quad (\text{C.63})$$

Dividing both sides of (C.63) by a yields:

$$\sum_{j \in N \setminus \{x\}} v_j + \frac{a_x}{a} v_x + \frac{a_0}{a} = 0, \quad (\text{C.64})$$

or

$$\sum_{j \in N \setminus \{x\}} v_j + C v_x + D = 0, \quad (\text{C.65})$$

where $C = \frac{a_x}{a}$ and $D = \frac{a_0}{a}$. ■

We now demonstrate that extreme points of all intermediate families (C.58) lie on the same hyperplane.

Claim C.11. *There exists a hyperplane that spans through all extreme points of intermediate families. The corresponding hyperplane equation is:*

$$\sum_{j \in N \setminus \{x\}} v_j - (n-1)v_x - (n-1) = 0. \quad (\text{C.66})$$

Proof of Claim C.11. Assume that there exists a hyperplane that spans through extreme points of two intermediate families with exactly q and $q+1$ coordinates attaining value $v^q = n-1$. Then the coefficients of corresponding hyperplane equation must satisfy the following system of equations:

$$\begin{cases} q(n-1) + (n-q-1)\frac{(n-1)(q-1)}{q} + \frac{(n-1)(q-1)}{q}C + D = 0, \\ (q+1)(n-1) + (n-q-2)\frac{(n-1)q}{q+1} + \frac{(n-1)q}{q+1}C + D = 0. \end{cases} \quad (\text{C.67})$$

Solving system (C.67) for C and D yields:

$$C = -\frac{n^2 - 2n + 1}{n-1} = -(n-1), \quad D = -(n-1).$$

Since neither C nor D depend on q , it is possible to build a hyperplane that spans through extreme points of all intermediate families simultaneously. The corresponding hyperplane equation is then:

$$\sum_{j \in N \setminus \{x\}} v_j - (n-1)v_x - (n-1) = 0.$$

■

However, extreme points of the lower family do not lie on the hyperplane (C.66). Thus, to cover all of the non-convex regions it is required to span one additional hyperplane. In the next claim we find a hyperplane which spans through lower family of nodes and the closest family of intermediate nodes⁴.

Claim C.12. *The hyperplane which spans through extreme points of the lower family and the closest intermediate family of extreme points (with $q = 2$) can be described by the following hyperplane equation:*

$$\sum_{j \in N \setminus \{i\}} v_j - \frac{n^2 - 3n + 4}{n - 2} v_i - \frac{(n - 3)(n - 1)}{n - 2} = 0. \quad (\text{C.68})$$

Proof of Claim C.12. To derive a corresponding hyperplane the following system of equations must be solved:

$$\begin{cases} 2(n - 1) + (n - 3)\frac{(n-1)}{2} + \frac{(n-1)}{2}C + D = 0, \\ (n - 1) + \frac{n-1}{n-2} + (n - 3) + C + D = 0. \end{cases} \quad (\text{C.69})$$

Solving system (C.69) for C and D yields:

$$C = -\frac{n^2 - 3n + 4}{n - 1} \quad D = -\frac{(n - 3)(n - 1)}{n - 2}.$$

The corresponding hyperplane equation is then:

$$\sum_{j \in N \setminus \{i\}} v_j - \frac{n^2 - 3n + 4}{n - 2} v_i - \frac{(n - 3)(n - 1)}{n - 2} = 0. \quad (\text{C.70})$$

■

Therefore, in order to obtain $\text{Conv}(\Upsilon^C)$, each non-convex constraint for some node $x \in [1, n]$ must be replaced with two linear constraints of the form:

$$\begin{cases} \sum_{j \in N \setminus x} v_j - (n - 1)v_x - (n - 1) \leq 0, \\ \sum_{j \in N \setminus x} v_j - \frac{n^2 - 3n + 4}{n - 1} v_x - \frac{(n - 3)(n - 1)}{n - 2} \leq 0. \end{cases} \quad (\text{C.71})$$

Thus, $\text{Conv}(\Upsilon^C)$ is defined by the following set of constraints:

⁴The difference between coordinates of a ‘lower’ family extreme points and ‘intermediate’ extreme points is minimal when $q = 2$.

$$1 \leq v_i \leq n - 1 \quad \forall i \in [1, n], \quad (\text{C.72})$$

$$\sum_{j \in N \setminus \{i\}} v_j - (n - 1)v_i - (n - 1) \leq 0 \quad \forall i \in [1, n], \quad (\text{C.73})$$

$$\sum_{j \in N \setminus \{i\}} v_j - \frac{n^2 - 3n + 4}{n - 2}v_i - \frac{(n - 3)(n - 1)}{n - 2} \leq 0 \quad \forall i \in [1, n]. \quad (\text{C.74})$$

We demonstrate the construction of a convex hull for the full support case with the example in which $n = 4$.

Example, $n = 4$. Consider the original set Υ^4 for the full support maximisation problem with $n = 4$:

$$\begin{aligned} 1 \leq v_i \leq 3 & \quad \forall i \in [1, 4], \\ \frac{2}{\sum_{j \in N \setminus \{i\}} \frac{1}{v_j}} - v_i \leq 0 & \quad \forall i \in [1, 4]. \end{aligned}$$

Then, $\text{Conv}(\Upsilon^4)$ is defined by the following constraints:

$$\begin{aligned} 1 \leq v_i \leq 3 & \quad \forall i \in [1, 4], \\ \sum_{j \in N \setminus \{i\}} v_j - 3v_i - 3 \leq 0 & \quad \forall i \in [1, 4], \\ \sum_{j \in N \setminus \{i\}} v_j - 4v_i - \frac{3}{2} \leq 0 & \quad \forall i \in [1, 4]. \end{aligned}$$

Figure 4.4 in the main text illustrates the set of feasible values for v_i , where $i \in [1, 3]$ along the edge v_4 .

△

C.3.3 Partial support

In this subsection we characterise the set of all extreme points of non-convex regions for the partial support case.

Let Υ^P be the set of feasible values for nodes $i \in N$ for the partial support case. Set Υ^P is defined by the following constraints:

$$1 \leq v_i \leq n - 1 \quad \forall i \in [1, k], \quad (\text{C.75})$$

$$\frac{k - 2}{\sum_{j \in K \setminus \{i\}} \frac{1}{v_j}} - v_i \leq 0 \quad \forall i \in [1, k], \quad (\text{C.76})$$

$$-\frac{(k - 1)}{\sum_{i=1}^k \frac{1}{v_i}} - v_j \leq 0 \quad \forall j \in (k, n]. \quad (\text{C.77})$$

This set of constraints can be significantly simplified. First, by Lemma 4.2 constraints (C.77) can be written with equality sign:

$$-\frac{(k - 1)}{\sum_{i=1}^k \frac{1}{v_i}} = v_j \quad \forall j \in (k, n].$$

Thus, every node outside set K attains a value that equals exactly the expected gain from an attack on nodes in set K .

We can now reduce the problem's dimensionality by ensuring that the gain from an attack on top k nodes is always larger than 1 (since any node outside the Attacker's support cannot attain a value smaller than 1). Moreover, nodes inside the Attacker's support cannot attain a value of 1 as well, as, in this case, all the nodes must then included in the Attacker's support violating the partial support assumption. To guarantee that all the nodes inside the Attacker's support are larger than 1 and all the nodes outside the Attacker's support are not smaller than 1, the lower boundary of constraint (C.75) must be increased. Since the expected attack on top k nodes attains the smallest value when all of the nodes attain the minimum possible value v_{min} , in order to find a new lower boundary for (C.75), it is sufficient to solve the following equation:

$$\frac{k - 1}{\frac{k}{v_{min}}} = 1. \quad (\text{C.78})$$

Solving equation (C.78) for v_{min} yields:

$$v_{min} = \frac{k}{k - 1}.$$

Therefore, if $\frac{k}{k-1} \geq v_i, \forall i \in K$, all of the constraints for nodes $v_j, j \in N \setminus K$ are always satisfied. The constraints that define set Υ^P then can be restated as follows:

$$\frac{k}{k-1} \leq v_i \leq n-1 \quad \forall i \in [1, k], \quad (\text{C.79})$$

$$\frac{k-2}{\sum_{j \in K \setminus \{i\}} \frac{1}{v_j}} - v_i \leq 0 \quad \forall i \in [1, k], \quad (\text{C.80})$$

$$-\frac{(k-1)}{\sum_{i=1}^k \frac{1}{v_i}} = v_j \quad \forall j \in (k, n]. \quad (\text{C.81})$$

It follows that similarly to the full support case, all non-convex regions of set Υ^P have extreme points defined by the intersection of one or several non-convex constraints (C.80) and the surface of hypercube defined by (C.79). We denote the set of extreme points of Υ^P as $E(\Upsilon^P)$. The set of nodes which pass the Attacker's support conditions with equality is denoted as $\Theta, \Theta = N \setminus K$.

Before approaching the search of extreme points for the general partial support case, we first consider a special case in which $k = 2$.

Special case, k=2. Consider non-convex constraint (C.80) for nodes $i \in K = \{1, 2\}$:

$$v_i \geq \frac{1}{\frac{1}{v_i} + \frac{1}{v_{-i}}}.$$

Rearranging and simplifying yields:

$$v_i \leq v_i + v_{-i},$$

which is always satisfied.

Therefore, if $k = 2$, non-convex constraints do not influence the set of feasible values $\Upsilon_{k=2}^P$ and this set is always convex. It follows that all extreme points in this case are defined by linear constraints (C.79). In the next claim we derive the general form of all extreme points for the special case $k = 2$.

Claim C.13. *If $k = 2$, the set of feasible values for $i \in K = \{1, 2\}$ has three families of extreme points:*

$$\mathbf{w}_{k=2}^1 = \left(2, 2, \underbrace{1, \dots, 1}_{\times \theta} \right), \quad (\text{C.82})$$

$$\mathbf{w}_{k=2}^2 = \left(n-1, n-1, \underbrace{\frac{n-1}{2}, \dots, \frac{n-1}{2}}_{\times \theta} \right), \quad (\text{C.83})$$

$$\mathbf{w}_{k=2}^3 = \left(n-1, 2, \underbrace{\frac{2(n-1)}{n+1}, \dots, \frac{2(n-1)}{n+1}}_{\times \theta} \right). \quad (\text{C.84})$$

where $\theta = |\Theta|$.

Proof of Claim C.13. First, observe that if $k = 2$, the minimum value that any node $i \in K$ can attain is $v_{\min} = \frac{k}{k-1} = \frac{2}{2-1} = 2$.

We now consider three possibilities in which $v_i \in \{2, n-1\}$

1. $v_i = v_{-i} = 2, i \in \{1, 2\}$. Substituting $k = 2, v_i = v_{-i} = 2$ into (C.81) yields:

$$v_j = \frac{1}{\frac{1}{2} + \frac{1}{2}} = 2 \quad \forall j \in \Theta.$$

Therefore, this case yields the extreme point with composition of coordinates (C.82).

2. $v_i = v_{-i} = n-1, i \in \{1, 2\}$. Substituting $k = 2, v_i = v_{-i} = n-1$ into (C.81) yields:

$$v_j = \frac{n-1}{2} \quad \forall j \in \Theta.$$

Therefore, this case yields the extreme point with composition of coordinates (C.83).

3. $v_i = n-1, v_{-i} = 2, i \in \{1, 2\}$. Substituting $k = 2, v_i = n-1$ and $v_{-i} = n-1$ into (C.81) yields:

$$v_j = \frac{2(n-1)}{n+1} \quad \forall j \in \Theta.$$

Therefore, this case yields the extreme point with composition of coordinates (C.84). ■

From this point we focus on a more general case in which $k \geq 3$. As in the full support case, we distinguish two types of extreme points of non-convex regions: *Type*

1 extreme points which result from intersection of a single non-convex constraint and linear constraints; and *Type 2 extreme points* which result from intersection of several non-convex constraints and linear constraints.

Type 1 extreme points. We demonstrate that similarly to full support case, there are two families of *Type 1 extreme points* that belong to set $E(\Upsilon^P)$: lower family and upper family.

First, consider a non convex constraint (C.80) for the value of some node $x \in K$:

$$v_x \geq \frac{k-2}{\sum_{j \in K \setminus \{x\}} \frac{1}{v_j}} = A_x. \quad (\text{C.85})$$

There are two cases to consider:

Case 1c. Intersections in which node x takes extreme values $v_x \in \{\frac{k}{k-1}, n-1\}$,

Case 2c. Intersection in which nodes indexed $i \in K \setminus x$ take values $v_i \in \{1, n-1\}$.

Case 1c. $v_x = A_x \in \{1, n-1\}$. We have already demonstrated in Claim C.1 that coordinate v_x cannot attain value $n-1$ in any extreme point of a non-convex region.

We now search for an extreme point in which coordinate $v_x = A_x = \frac{k}{k-1}$, some coordinate $v_y \in [\frac{k}{k-1}, n-1]$ and the rest coordinates $v_i \in \{\frac{k}{k-1}, n-1\}$, $\forall i \in K \setminus \{x, y\}$. The corresponding composition of coordinates can be written in the following form:

$$\mathbf{w}_b^q = \left(\underbrace{n-1}_{\times q}, \underbrace{\frac{k}{k-1}}_{\times k-q-2}, v_y, \underbrace{\frac{k}{k-1}}_{=A_x}, \underbrace{A^k}_{\times \theta} \right), \quad (\text{C.86})$$

where q is a number of nodes that attain value $v^q = n-1$, A^k is an expected attack on top k nodes, and $\theta = |\Theta|$. Also note that in any sequence (C.86) it must be the case that $q \in [1, k-2]$.

We now demonstrate that the only case in which $\mathbf{w}_b^q \in E(\Upsilon^P)$ is when $q = 1$ and $v_y = \frac{k(n-1)}{kn-2k-n+1}$.

Claim C.14. $\mathbf{w}_b^q \in E(\Upsilon^P)$ IFF $q = 1$ and $v_y = \frac{k(n-1)}{kn-2k-n+1}$

Proof of Claim C.14. First, we demonstrate that if $q = 0$, $\mathbf{w}_b^0 \notin E(\Upsilon^P)$. Substituting sequence (C.86) with $q = 0$ into (C.85) yields:

$$A_x = \frac{k}{k-1} = \frac{k-2}{\frac{k-2}{k-1} + \frac{1}{v_y}}. \quad (\text{C.87})$$

By rearranging and simplifying equation (C.87) we obtain:

$$\frac{k}{(k-1)(k^2v_y - 3kv_y + k + 2vy)} = 0,$$

which does not have a solution for $v_y \in \mathbb{R}$.

Now consider the case in which $q \geq 1$. Substituting composition (C.86) in (C.85):

$$A_x = \frac{k}{k-1} = \frac{k-2}{\frac{q}{n-1} + \frac{\frac{k-q-2}{k-1}}{\frac{k}{k-1}} + \frac{1}{v_y}}. \quad (\text{C.88})$$

Solving (C.88) for v_y yields:

$$v_y = \frac{k(n-1)}{q(kn - 2k - n + 1)}. \quad (\text{C.89})$$

We now demonstrate that in (C.89) $v_y < \frac{k}{k-1}$ for any $q \geq 2$.

Consider the first derivative of the RHS of equation (C.89) w.r.t. q :

$$\frac{\partial}{\partial q} \frac{k(n-1)}{q(kn - 2k - n + 1)} = -\frac{k(n-1)}{q^2(kn - 2k - n + 1)} < 0,$$

which is always negative as $n(k-1) > 2k$. Therefore, it is sufficient to consider (C.89) when $q = 2$. Substituting $q = 2$ in (C.89) yields:

$$v_y = \frac{k(n-1)}{2k(n-2) - 2n + 2}. \quad (\text{C.90})$$

We now show that:

$$\frac{k(n-1)}{2k(n-2) - 2n + 2} < \frac{k}{k-2}. \quad (\text{C.91})$$

Rearranging and simplifying (C.91) yields:

$$(k-1)k(k(n-3) - n + 1)(k(n-2) - n + 1) > 0,$$

or:

$$\underbrace{(n(k-1) - 3k + 1)}_{\iota_1} (n(k-1) - 2k + 1) > 0.$$

Observe that $\frac{\partial \iota}{\partial k} = n - 3 > 0$ for any $n > 3$. Thus, consider ι when $k = 4$ (the minimum value that k can attain in this composition is 4, as the composition has two nodes of value v^q , one node of value v^q and one node of value $\frac{k}{k-1}$):

$$\iota = 3n - 12 + 1 > 0,$$

which is always positive since $n > k = 4$.

It follows that inequality (C.91) is always satisfied and the intersection happens beyond the hypercube defined by linear constraints (C.79) for any $q \geq 2$. Thus, any family of points with composition $\mathbf{w}_b^{q \geq 2} \notin E(\Upsilon^P)$.

The only possibility that is left to verify is when $q = 1$. Substituting sequence (C.86) with $q = 1$ into (C.89) yields:

$$v_y = \frac{k(n-1)}{kn - 2k - n + 1}. \quad (\text{C.92})$$

We now demonstrate that:

$$\frac{k(n-1)}{kn - 2k - n + 1} > \frac{k}{k-1}. \quad (\text{C.93})$$

Rearranging and simplifying (C.93) yields:

$$(k-1)k^2(k(n-2) - n + 1) > 0,$$

or

$$k(n-2) - n + 1 > 0,$$

which is always satisfied for any $n > k \geq 3$. ■

Therefore, *Case 1c* yields only one family of extreme points with the following composition of coordinates:

$$\mathbf{w}_b^q = \left(n-1, \frac{k(n-1)}{kn - 2k - n + 1}, \underbrace{\frac{k}{k-1}, \dots, \frac{k}{k-1}}_{\times n-2} \right). \quad (\text{C.94})$$

We denote the family of extreme points with composition (C.94) as the *the lower family of extreme points of the partial support case*.

Case 2c. $v_i \in \{\frac{k}{k-1}, n-1\}$, $\forall i \in K \setminus \{x\}$. The general composition of coordinates for this case can be written as:

$$\mathbf{w}_u^q = \left(\underbrace{n-1, \dots, n-1}_{\times q}, \underbrace{\frac{k}{k-1}, \dots, \frac{k}{k-1}}_{\times k-q-2}, A_x, \underbrace{A^k, \dots, A^k}_{\theta} \right), \quad (\text{C.95})$$

We demonstrate that *Case 2c* yields only one extreme point family with composition of coordinates $\mathbf{w}_u^q$ in which $q = k-1$.

Claim C.15. $\mathbf{w}_u^q \in E(\Upsilon^P)$ IFF $q = k-1$.

Proof of Claim C.15. We first demonstrate that $\mathbf{w}_u^q \notin E(\Upsilon^P)$ if $q \in \{0, 1\}$.

Substituting composition (C.95) with $q = 0$ in (C.85) yields:

$$A_x = \frac{k-2}{\frac{(k-1)^2}{k}} = \frac{(k-2)k}{(k-1)^2} = 1 - \frac{1}{(k-1)^2} < \frac{k}{k-1}.$$

It follows that this intersection happens beyond the hypercube defined by (C.79) and $\mathbf{w}_u^0 \notin E(\Upsilon)$.

Now consider the case in which $q = 1$. Substituting composition (C.95) with $q = 1$ in (C.85) yields:

$$A_x = \frac{k-2}{\frac{(k-2)(k-1)}{k} + \frac{1}{n-1}}.$$

We now demonstrate that:

$$\frac{k-2}{\frac{(k-2)(k-1)}{k} + \frac{1}{n-1}} < \frac{k}{k-1}. \quad (\text{C.96})$$

Rearranging and simplifying (C.97) yields:

$$(k-1)k^2 \left(k^2(n-1) + k(4-3n) + 2(n-1) \right) > 0,$$

or:

$$\underbrace{k^2(n-1) + k(4-3n) + 2(n-1)}_{\tau} > 0. \quad (\text{C.97})$$

Consider the first derivative of τ w.r.t to k :

$$\frac{\partial \tau}{\partial k} = 4 - 2k - 3n + 2kn > 0,$$

which is always positive since $2n(k-3) \geq 2k$ for any $k \geq 3$. Thus, it is sufficient to verify inequality (C.97) at $k = 3$. Substituting $k = 3$ into (C.97) yields:

$$2n + 1 > 0,$$

which always holds.

It follows that $\mathbf{w}_u^1 \notin E(\Upsilon^P)$.

Now we demonstrate that any family with composition (C.95) with $q \in [2, k-2]$ and $k \geq 4$ cannot belong to $E(\Upsilon^P)$. To see that, consider the gain from an expected attack on $q \geq 2$ nodes:

$$A^q = \frac{(n-1)(q-1)}{q}.$$

Since $\frac{\partial A^q}{\partial q} = \frac{n-1}{q^2}$, it follows that the minimum value of A^q is achieved whenever $q = 2$:

$$A^2 = \frac{n-1}{2}.$$

Now observe that:

$$\frac{n-1}{2} > \frac{k}{k-1},$$

for any $n > k \geq 4$.

It follows that in any composition (C.95) with $q \in [2, k-2]$ nodes of value $v^q = n-1$, and $k-q-1$ nodes of value $v^k = \frac{k}{k-1}$ with $n > k \geq 4$, any node of value v^k is not included in the Attacker support. Thus, extreme points of this composition do not belong to Υ^P .

The only possibility left to verify is when $q = k-1$. In this case, there is no nodes of value v^k and every node of value $n-1$ must be included in the Attacker's support. Substituting composition (C.95) in (C.85) with $q = n-1$ yields:

$$A_x = \frac{(n-1)(q-1)}{q} > \frac{k}{k-1}.$$

■

It immediately follows from Claim (C.15) that the only family of extreme points that *Case 2c* yields is:

$$\mathbf{w}_u = \left(\underbrace{n-1, \dots, n-1}_{\times k-1}, \underbrace{\frac{(n-1)(q-1)}{q}, \dots, \frac{(n-1)(q-1)}{q}}_{\times n+1-k} \right). \quad (\text{C.98})$$

We denote the family which have composition of coordinates (C.98) as the *upper family of extreme points of the partial support case*.

Type 2 extreme points. This type of extreme points results from the intersection of a surface of a hypercube defined by (C.79) and $t \geq 2$ non-convex constraints (C.80).

Utilising the results of Claim (C.6) we can write down the general form of intersection of t constraints straight away. Let T denote the set of nodes whose constraints intersect and $M = K \setminus T$ denote the set of nodes whose constraints do not intersect. Without loss of generality we relabel nodes in set T as $\mathbf{v}_z = (v_{z_1}, \dots, v_{z_t})$ and nodes in set M as $\mathbf{v}_z = (v_{x_1}, \dots, v_{x_{k-t}})$. Then any *Type 2* extreme point family can be written as:

$$\mathbf{w}_i = \left(v_{x_1}, \dots, v_{x_{k-t}}, v_{z_1}, \dots, v_{z_t}, \underbrace{A^k}_{\times \theta} \right), \quad (\text{C.99})$$

where $v_{z_i} = \frac{k-1-t}{\sum_{j \in M} \frac{1}{v_j}}$.

As previously, there are two cases to consider:

Case 1d. Each node indexed $i \in T$ attains value $v_i \in \left\{ \frac{k}{k-1}, n-1 \right\}$;

Case 2d. Each node indexed $j \in M$ attains value $v_j \in \left\{ \frac{k}{k-1}, n-1 \right\}$.

Case 1d. $v_i \in \left\{ \frac{k}{k-1}, n-1 \right\}$, $\forall i \in T$. First, note that similarly to the full support case, none of nodes indexed $i \in T$ can attain value $v_{z_i} = v^q = n-1$, since the maximum gain from an attack is $\frac{(n-1)^2}{n} < n-1$.

Second, if nodes $v_i = \frac{k}{k-1}$, $\forall i \in T$ attain value $v_i = v^k = \frac{k}{k-1}$, then it yields intersection point which has composition (C.94). Therefore, this case does not yield any new extreme points.

Case 2d. $v_j \in \left\{\frac{k}{k-1}, n-1\right\}, \forall j \in M$. We demonstrate that $\mathbf{w}_i \in E(\Upsilon^P)$ only if all nodes indexed $j \in M$ take value $v_j = v^q = n-1$.

Claim C.16. $\mathbf{w}_i \in E(\Upsilon^P)$ IFF $v_j = v^q = n-1, \forall j \in M$.

Proof of Claim C.16. First, consider the case in which all the nodes $j \in M$ take value $v_j = v^k = \frac{k}{k-1}$. In this case any node indexed $i \in M$ must attain value:

$$v_{z_i} = \frac{k-1-t}{\frac{k-t}{\frac{k}{k-1}}} = \frac{k(k-t-1)}{(k-1)(k-t)}.$$

Now observe that:

$$\frac{k-1-t}{\frac{k-t}{\frac{k}{k-1}}} = \frac{k(k-t-1)}{(k-1)(k-t)} < \frac{k}{k-1}. \quad (\text{C.100})$$

Simplifying (C.100) yields:

$$\frac{k-t-1}{k-t} < 1,$$

which is always satisfied for any $t \geq 2$.

Second, consider the case in which $q = 1$. In this case, any node indexed $i \in M$ must attain value:

$$v_{z_i} = \frac{k-t-1}{\frac{(k-1)(k-t-1)}{k} + \frac{1}{n-1}}.$$

Similarly observe that:

$$\frac{k-t-1}{\frac{(k-1)(k-t-1)}{k} + \frac{1}{n-1}} < \frac{k}{k-1},$$

which can be rearranged and simplified as

$$\underbrace{k^2(n-1) + k(-n(t+2) + t+3) + (n-1)(t+1)}_{\psi} > 0. \quad (\text{C.101})$$

Consider the first derivative of ψ w.r.t to t :

$$\frac{\partial \psi}{\partial t} = k-1 - (k-1)n < 0,$$

which is always negative.

Now consider condition (C.101) when t attains its maximum value $t = k - 2$:

$$(k - 1)n + 1 > 0,$$

which is always satisfied.

Therefore, the case in which $q = 1$ does not yield any extreme points.

Third, consider the case in which $q \in [2, k - t)$ and $k > q$. In this case, any node that attains value $v^k = \frac{k}{k-1}$ cannot be included in the Attacker's support as $\frac{(n-1)(q-1)}{q} > \frac{k}{k-1}$ for any $k \geq 4$.

Therefore, the only feasible extreme point has the composition of coordinates in which $v_j = v^q = n - 1$, $\forall j \in M$ since, in this case, all the nodes $i \in M$ are included in the Attacker's support and:

$$v_{z_i} = \frac{(n-1)(q-1)}{q} > \frac{k}{k-1},$$

for any $k \geq 4$. ■

It follows from Claim C.16 that the only family of extreme points that *Case 2d* yields must have the following composition of coordinates:

$$\mathbf{w}_i^q = \left(\underbrace{n-1, \dots, n-1}_{\times q}, \underbrace{\frac{(n-1)(q-1)}{q}, \dots, \frac{(n-1)(q-1)}{q}}_{\times (n-q)} \right), \quad (\text{C.102})$$

where $q \in [2, k - 2]$.

We denote extreme points which have composition of coordinates (C.102) as the *intermediate family of extreme points of the partial support case*. Also observe that by allowing $q \in [2, k - 1]$ and $k \in [2, n - 1]$, the composition (C.102) also covers the upper family of extreme points of the partial support case (C.98) and the family (C.83). We combine those three types of families.

Therefore set Υ^P with $k \geq 3$ has two families of extreme points:

1. lower family of extreme points of the partial support case:

$$\mathbf{w}_b^q = \left(n-1, \frac{k(n-1)}{kn-2k-n+1}, \underbrace{\frac{k}{k-1}, \dots, \frac{k}{k-1}}_{\times(n-2)} \right). \quad (\text{C.103})$$

2. intermediate family of extreme points of the partial support case:

$$\mathbf{w}_i^q = \left(\underbrace{n-1, \dots, n-1}_{\times q}, \underbrace{\frac{(n-1)(q-1)}{q}, \dots, \frac{(n-1)(q-1)}{q}}_{\times(n-q)} \right), \quad (\text{C.104})$$

where $q \in [2, k-1]$.

Note that any point of family (C.104) can be characterised by the composition (C.58) found in the full support case.

C.3.4 Convex hull construction for the partial support case

As in the full support case, in order to construct a convex hull of Υ^P , it is sufficient to build a series of hyperplanes spanned by all the extreme points of non-convex regions and derive the corresponding constraints.

From Claim C.11 we know that hyperplane that spans through all extreme points of some family can be expressed in the following form:

$$\sum_{j \in K \setminus x} v_j + C v_x + D = 0,$$

where C and D are coefficients.

We now demonstrate that all the extreme points of the intermediate families of extreme points lie on the same hyperplane.

Claim C.17. *There exists a hyperplane that spans through all extreme points of all intermediate families. The corresponding hyperplane equation is:*

$$\sum_{j \in K \setminus x} v_j - (k-1)v_x - (n-1) = 0.$$

Proof of Claim C.17. Assume that there exists a hyperplane which spans through extreme points of two intermediate families which have exactly q and $q + 1$ nodes of value $v^q = n - 1$ correspondingly. Then the corresponding hyperplane equation must satisfy the following system of equations:

$$\begin{cases} q(n - 1) + (k - q - 1)\frac{(n-1)(q-1)}{q} + \frac{(n-1)(q-1)}{q}C + D = 0, \\ (q + 1)(n - 1) + (k - q - 2)\frac{(n-1)q}{q+1} + \frac{(n-1)q}{q+1}C + D = 0. \end{cases} \quad (\text{C.105})$$

Solving system (C.105) for C and D yields:

$$C = -(k - 1), \quad D = -(n - 1).$$

Since neither C nor D depend on q , it is possible to build a hyperplane that spans through extreme points of all intermediate families simultaneously. The corresponding hyperplane equation is then:

$$\sum_{j \in N \setminus x} v_j - (k - 1)v_x - (n - 1) = 0.$$

■

We now characterise a hyperplane which spans through nodes of the lower family of extreme points and the closest to that family extreme points of the intermediate family (which have the composition of coordinates $\mathbf{w}_i^2$).

Claim C.18. *The hyperplane which spans through extreme points of the lower family and the closest intermediate family of extreme points ($\mathbf{w}_i^2$) can be described by the following hyperplane equation:*

$$\sum_{j \in K \setminus i} v_j - \frac{k^2(n - 2) + k(5 - 2n) + n - 1}{k(n - 2) - n + 1}v_i - \frac{(n - 1)(kn - 3k - n + 1)}{kn - 2k - n + 1} = 0. \quad (\text{C.106})$$

Proof of Claim C.18. To derive a corresponding hyperplane the following system of equations must be solved:

$$\begin{cases} 2(n-1) + (k-3)\frac{(n-1)}{2} + \frac{(n-1)}{2}C + D = 0, \\ (n-1) + \frac{k(n-1)}{kn-2k-n+1} + (k-3)\frac{k}{k-1} + \frac{k}{k-1}C + D = 0. \end{cases} \quad (\text{C.107})$$

Solving system (C.107) for C and D yields:

$$C = -\frac{k^2(n-2) + k(5-2n) + n-1}{k(n-2) - n+1}, \quad D = -\frac{(n-1)(kn-3k-n+1)}{kn-2k-n+1}.$$

The corresponding hyperplane equation is then:

$$\sum_{j \in K \setminus i} v_j - \frac{k^2(n-2) + k(5-2n) + n-1}{k(n-2) - n+1} v_i - \frac{(n-1)(kn-3k-n+1)}{kn-2k-n+1} = 0. \quad (\text{C.108})$$

■

Therefore, in order to obtain $\text{Conv}(\Upsilon^P)$ each non-convex constraint $x \in [1, k]$ must be replaced with two linear constraints of the form:

$$\begin{cases} \sum_{j \in K \setminus \{x\}} v_j - (k-1)v_x - (n-1) \leq 0, \\ \sum_{j \in K \setminus \{x\}} v_j - \frac{k^2(n-2) + k(5-2n) + n-1}{k(n-2) - n+1} v_x - \frac{(n-1)(kn-3k-n+1)}{kn-2k-n+1} \leq 0. \end{cases}$$

Thus, $\text{Conv}(\Upsilon^P)$ is defined by the following set of constraints:

$$\begin{aligned} \frac{k}{k-1} &\leq v_i \leq n-1 && \forall i \in [1, k], \\ \sum_{j \in K \setminus \{i\}} v_j - (k-1)v_i - (n-1) &\leq 0 && \forall i \in [1, k], \\ \sum_{j \in K \setminus \{i\}} v_j + Cv_i + D &\leq 0 && \forall i \in [1, k], \\ -\frac{(k-1)}{\sum_{i=1}^k \frac{1}{v_i}} &= v_j && \forall j \in (k, n], \end{aligned}$$

where $C = -\frac{k^2(n-2) + k(5-2n) + n-1}{k(n-2) - n+1}$ and $D = -\frac{(n-1)(kn-3k-n+1)}{kn-2k-n+1}$.

C.3.5 Extreme points of convex regions

Apart from extreme points of non-convex regions, sets Υ^C and Υ^P have extreme points defined solely by linear constraints (C.29) and (C.79). In this subsection, we characterise these points. Full and partial support cases are considered separately.

Full support Any extreme points of a hypercube defined by linear constraints (C.29) can have coordinates of value either $v^q = n - 1$ or $v^u = 1$. Any family of those points must have the following composition:

$$\mathbf{v}_f^q = \left(\underbrace{n - 1, \dots, n - 1}_q, \underbrace{1, \dots, 1}_{n-q} \right), \quad (\text{C.109})$$

where $q \in [0, n]$.

We now demonstrate that only families in which $q = 0$ (sequence of unit values), $q = 1$ (star network sequence), and $q = n$ (complete network sequence) are included in $E(\Upsilon^C)$.

Claim C.19. $\mathbf{v}_f^q \in E(\Upsilon^C)$ IFF $q \in \{0, 1, n\}$.

Proof of Claim C.19. First consider composition (C.109) with $q = 0$. Each coordinate in this composition takes value $v^u = 1$. Since all the coordinates in this sequence attain the same value, then by Claim (4.4) they must all be included in the Attacker's support. It follows that $\mathbf{v}_f^0 \in E(\Upsilon^C)$.

By the same argument, it must be the case that any family with a composition of coordinates (C.109) with $q = n$ must be included in the set of extreme points too, $\mathbf{v}_f^0 \in E(\Upsilon^C)$.

We now analyse families of nodes with composition (C.109) in which $q \geq 1$. We consider two cases:

Case 1. $q = 1$. This composition yields a star network degree sequence, which must be always included in $E(\Upsilon^C)$.

Case 2. $q \geq 2$. In this case, sequence (C.109) must have q coordinates of value $v^q = n - 1$ and $n - q$ coordinates of unit value $v^u = 1$. However, any node which attains a unit value is not included in the Attacker's support if

$q \geq 2$. To see that observe that the expected gain from an attack on q nodes of values $n - 1$ is:

$$A^q = \frac{(n-1)(q-1)}{q} > 1,$$

which is larger than 1 for any $n > q \geq 2$.

It follows that $\mathbf{v}_f^{q \geq 2} \notin E(\Upsilon^C)$. ■

Therefore, three extreme point families are defined purely by linear constraints in the full support case: a star network, a complete network, and a sequence of unit values. However, we are not considering the latter as an eligible solution for the maximisation problem. Sequence $\mathbf{v}_f^0$ appears as a maximiser for the relaxed maximisation problem at extreme values of $c \rightarrow 2$, outperforming a star network. However, the sequence is clearly not-graphic as the minimum possible number of connections that a connected network can have is $(n - 1)$, and this sequence has only $\frac{1}{2}n$. Moreover, by Lemma 4.2 we know that a star network must always be an optimal choice for the Defender for the sufficiently high cost of an edge, $c > c_u$.

Partial support Similarly, we now determine extreme points of a hypercube defined by linear constraints (C.79). As case with $k = 2$ was already considered in Subsection C.3.3, we consider a case in which $k \geq 3$. Any family of extreme points defined by linear constraints (C.79) have coordinates either of value $v^q = n - 1$ or $v^k = \frac{k}{k-1}$. Thus, they can be represented by the following composition of coordinates:

$$\mathbf{w}_p^q = \left(\underbrace{n-1, \dots, n-1}_q, \underbrace{\frac{k}{k-1}, \dots, \frac{k}{k-1}}_{k-q}, \underbrace{A_p^q}_{\times n-k} \right), \quad (\text{C.110})$$

where $q \geq 2$, and

$$A_p^q = \frac{k-1}{\frac{(k-1)(k-q)}{k} + \frac{q}{n-1}}.$$

We now demonstrate that $\mathbf{w}_p^q \in E(\Upsilon^P)$ IFF $q \in \{0, 1, k\}$.

Claim C.20. $\mathbf{w}_f^q \in E(\Upsilon^P)$ IFF $q \in \{0, 1, k\}$.

Proof of Claim C.20. First, consider the case in which $q = 0$. The composition of coordinates then has k coordinates which attain value $v^k = \frac{k}{k-1}$ and $n - k$ coordinates of value $v_u = 1$. Since every of the top k nodes in the composition attains the same value, then by Claim 4.4 they must be included in the Attacker's support. It follows that $\mathbf{w}_0^q \in E(\Upsilon^C)$.

By the same argument, the composition in which $q = k$ is included in the set of extreme values of the set Υ^C , $\mathbf{w}_0^k \in E(\Upsilon^C)$.

Now consider sequence (C.20) in which $q \geq 1$.

Case 1. $q = 1$. In this case, the composition has one coordinate of value $v^q = n - 1$, $k - 1$ coordinates of value $v^k = \frac{k}{k-1}$ and $n - k$ coordinates that attain the value equal to the expected gain from an attack on the top k nodes. Now consider constraint (C.76) for some node y of value v^k .

$$v_y \geq \frac{k-2}{\frac{(k-2)(k-1)}{k} + \frac{1}{n-1}}. \quad (\text{C.111})$$

Substituting $v_y = \frac{k}{k-1}$ into (C.111):

$$\frac{k}{k-1} \geq \frac{k-2}{\frac{(k-2)(k-1)}{k} + \frac{1}{n-1}}. \quad (\text{C.112})$$

Rearranging and simplifying (C.112) yields:

$$(k-1)k^2 \left(k^2(n-1) + k(4-3n) + 2(n-1) \right) \geq 0,$$

Which is satisfied whenever:

$$\underbrace{k^2(n-1) + k(4-3n) + 2(n-1)}_{\iota} \geq 0. \quad (\text{C.113})$$

Consider the first derivative of ι w.r.t. n :

$$\frac{\partial \iota}{\partial n} = k^2 - 3k + 2 > 0,$$

which is larger than zero for any $n \geq 3$.

Thus, we verify inequality (C.113) when $n = 4$ (since it is the minimum possible number of nodes that network can have if $n > k \geq 3$):

$$3k^2 - 8k + 6 > 0,$$

which is always satisfied.

It follows that $\mathbf{w}_f^1 \in E(\Upsilon^P)$.

Case 2. $q \geq 2$. We demonstrate that any node which attains value v^k is not included in the Attacker's support in this case. To see that observe that an expected gain from an attack on q nodes of value v^q is always larger than v^k :

$$\frac{(n-1)(q-1)}{q} > \frac{k}{k-1},$$

which is always satisfied for any $k > q \geq 2$ as the minimum value that the LHS of inequality can attain is 2 (when $n = 4$ and $q = 2$), and the RHS of inequality is strictly smaller than 2 for any $k \geq 3$. Thus, $\mathbf{w}_f^{q \geq 2} \notin E(\Upsilon^P)$. ■

To conclude, partial support case yields three families of extreme points defined purely by linear constraints (C.79):

(a) k -family:

$$\mathbf{w}_p^0 = \left(\underbrace{\frac{k}{k-1}, \dots, \frac{k}{k-1}}_{\times k}, \underbrace{1, \dots, 1}_{\times (n-k)} \right),$$

where $k \in [2, n-1]$. Note that when, $k = 2$ it yields family (C.82) found in the convex case of partial support extreme points analysis;

(b) Quasi-star family:

$$\left(n-1, \underbrace{\frac{k}{k-1}, \dots, \frac{k}{k-1}}_{\times (k-1)}, \underbrace{\frac{k-1}{\frac{(k-1)^2}{k} + \frac{1}{n-1}}, \dots, \frac{k-1}{\frac{(k-1)^2}{k} + \frac{1}{n-1}}}_{\times (n-k)} \right),$$

where $k \in [2, n-1]$. Note that when $k = 2$, it yields family (C.84) found in the convex case of partial support extreme points analysis;

(c) Intermediate family:

$$\mathbf{w}_p^k = \left(\underbrace{n-1, \dots, n-1}_{\times k}, \underbrace{\frac{(n-1)(k-1)}{k}, \dots, \frac{(n-1)(k-1)}{k}}_{\times (n-k)} \right)$$

which is equivalent to intermediate families of extreme points (C.58) and (C.104) that were discovered in both full and partial support cases.

C.4 Feasibility of intermediate family sequences

In order for an intermediate family sequence to yield a maxi-core network, three conditions must be satisfied:

- (a) $v^b = \frac{(n-1)(q-1)}{q}$ must be integer, $v^b \in \mathbb{Z}$;
- (b) the sequence must pass Erdos-Gallai graphicality conditions (4.10);
- (c) the sum of all degrees in the sequence must be even.

We now consider each of those conditions separately.

Integrity of the sequence The value of peripheral nodes indexed $j > q$ can be restated as follows:

$$v^b = \frac{(n-1)(q-1)}{q} = n-1 - \underbrace{\frac{n-1}{q}}_{\tau}.$$

Thus, peripheral nodes' value is integer IFF $\tau \in \mathbb{Z}$.

Claim C.21. *Peripheral nodes' value $v^b \in \mathbb{Z}$ IFF $\frac{n-1}{q} \in \mathbb{Z}$.*

Sequence graphicality We now demonstrate that any sequence of an intermediate family with $q \in [2, n-2]$ always satisfies graphicality constraints.

Claim C.22. *Erdos-Gallai graphicality constraints (4.10) are always satisfied for any sequence of intermediate family (4.23) with $q \in [2, n-2]$.*

Proof of Claim C.22. As intermediate sequence (4.23) has only two types of nodes, it is sufficient to check Erdos-Gallai graphicality conditions for nodes indexed q and n (Tripathi & Vijay, 2003).

First, observe that:

$$q < \frac{(n-1)(q-1)}{q},$$

if $n > \frac{q^2+q-1}{q-1} = q + \frac{1}{q-1} + 2$ or, since $n \in \mathbb{Z}$, when $q < n - 2$.

Consider two cases:

Case 1. $q < n - 2$. In this case, Erdos-Gallai conditions for node indexed q can be written as follows:

$$q(n-1) \leq q(q-1) + (n-q) \underbrace{\min(q, \frac{(n-1)(q-1)}{q})}_{=q},$$

which is always satisfied as:

$$q(n-1) - q(q-1) - (n-q)q \leq 0,$$

$$0 \leq 0.$$

We first verify the condition for node indexed n :

$$q(n-1) + (n-q) \frac{(n-1)(q-1)}{q} \leq n(n-1),$$

which can be rearranged and simplified as:

$$\frac{(n-1)(n-q)}{q} \geq 0.$$

Thus, in *Case 1*, Erdos-Gallai conditions are always satisfied.

Case 1. $q \geq n - 2$. Erdos-Gallai conditions for node indexed q are:

$$q(n-1) \leq q(q-1) + (n-q) \underbrace{\min(q, \frac{(n-1)(q-1)}{q})}_{=\frac{(n-1)(q-1)}{q}},$$

$$q(n-1) \leq q(q-1) + (n-q) \frac{(n-1)(q-1)}{q},$$

which can be rearranged as:

$$q(n-q) \leq \frac{(n-1)(q-1)(n-q)}{q},$$

or:

$$q \leq \frac{(n-1)(q-1)}{q},$$

which is never satisfied for any $q \in \{n-2, n-1\}$. ■

Sum of degrees parity The last thing left to verify is the parity of the sum of all degrees in the intermediate families that satisfy Erdos-Gallai and integrality constraints. First, observe that a q -maxi-core network is always feasible if the Defender has an odd number of nodes. For instance, it is always possible to create a network with $q = \frac{n-1}{2}$ and $v_b = n-3$, which satisfies both Erdos-Gallai and integrality constraints. If the Defender has an even number of vertices, q -maxi-core networks are not feasible since the value of peripheral nodes is either $v^b = \frac{(q-1)(n-1)}{q} \notin \mathbb{Z}$ for any $q \in [2, n-2]$ or the sum of all degrees is not even.

Claim C.23. *If the Defender has an odd number of vertices, maxi-core networks are always feasible. If the Defender has an even number of nodes, maxi-core networks are never feasible.*

Proof of Claim C.23.

The Defender has an odd number of nodes. Suppose that Claim C.21 holds and n is odd. Then peripheral nodes of value v^b in the intermediate family sequence can be either odd or even depending on the parity of q . We consider two cases: (1) q is even, and (2) q is odd.

Case 1. q is even. In this case, the value of the peripheral node can be stated as follows:

$$v^b = \underbrace{n-1}_{\text{even}} - \underbrace{\frac{n-1}{q}}_{\text{even/odd}}.$$

Since $\frac{n-1}{q}$ can be either even or odd, we consider two subcases.

(1.1) $\frac{n-1}{q}$ is even.

Suppose $\frac{n-1}{q}$ is even then: $v^b = \underbrace{n-1}_{\text{even}} - \underbrace{\frac{n-1}{q}}_{\text{even}} = \text{even}$. Then the overall sum of degrees must be:

$$\underbrace{q(n-1)}_{\text{even}} + \underbrace{(n-q)\frac{(n-1)(q-1)}{q}}_{\text{even}} = \text{even}.$$

Therefore, if both q and $\frac{n-1}{q}$ are even, then the resulting degree sequence is graphic.

(1.2) $\frac{n-1}{q}$ is odd.

Suppose $\frac{n-1}{q}$ is odd then: $v^b = \underbrace{n-1}_{\text{even}} - \underbrace{\frac{n-1}{q}}_{\text{odd}} = \text{odd}$.

The overall degree sum is then:

$$\underbrace{q(n-1)}_{\text{even}} + \underbrace{(n-q)\frac{(n-1)(q-1)}{q}}_{\text{odd}} = \text{odd}.$$

Thus, if the sequence with an even number of vertices q of value v^q , but $\frac{n-1}{q}$ is odd, such a sequence is not graphic.

Case 2. q is odd. Suppose $\frac{n-1}{q}$ is odd then:

$$v^b = \underbrace{n-1}_{\text{even}} - \underbrace{\frac{n-1}{q}}_{\text{even}} = \text{even}.$$

The sum of degrees is then:

$$\underbrace{q(n-1)}_{\text{even}} + \underbrace{(n-q) \frac{(n-1)(q-1)}{q}}_{\text{even}} = \text{even}.$$

Thus, if q is odd and Claim C.21 holds, the resulting sequence is always graphic.

The Defender has an even number of nodes. Suppose that n is even.

Then in order for $v^b = n-1 - \frac{n-1}{q}$ to be an integer, number of nodes of value $v^q = n-1$ must be odd. It then follows that v^b is always even:

$$v^b = \underbrace{n-1}_{\text{odd}} - \underbrace{\frac{n-1}{q}}_{\text{odd}} = \text{even}.$$

Then the overall sum of degrees must be:

$$\underbrace{q}_{\text{odd}} \cdot \underbrace{(n-1)}_{\text{odd}} + \underbrace{(n-q)}_{\text{odd}} \cdot \underbrace{\frac{(q-1)(n-1)}{q}}_{\text{even}} = \text{odd}.$$

It immediately follows that a sequence like that is not graphic.

Therefore, maxi-core networks are not feasible if the Defender has an even number of nodes. ■

Therefore, if the Defender has an even number of vertices, the relaxed maximisation problem provides only two feasible solutions: a star and complete networks.

C.5 Supplementary material for Subsection 4.3.4

Table C.1: Estimations of edge cost threshold between complete and star formations for networks with 3 and 4 nodes

#	Complete/star
3	1.62353
4	1.55

Table C.2: Estimations of edge cost thresholds between complete, maxi-core (MC), and star formations for networks with odd number of nodes and $n \leq 19$, where q stands for the core size, v^s is a value of a peripheral node in an optimal maxi-core network

#	Complete/MC	MC/Star	q	v^s
5	1.6	1.64706	2	2
7	1.71429	1.72482	3	4
9	1.77778	1.78462	2	4
11	1.81818	1.82026	5	8
13	1.84615	1.84828	2	6
15	1.86667	1.86741	7	12
17	1.88235	1.88327	2	8
19	1.89474	1.89521	3	12
∞	2	2	n/a	n/a

Table C.3: Estimations of edge cost threshold between the star and complete formations for networks with an even number of nodes, $n \leq 20$

#	Complete/star
6	1.67949
8	1.755
10	1.80244
12	1.8347
14	1.85798
16	1.87555
18	1.88927
20	1.90028
∞	2

References

- Acemoglu, D., Malekian, A., & Ozdaglar, A. (2016). Network security and contagion. *Journal of Economic Theory*, 166, 536–585.
- Ahmad, A., Webb, J., Desouza, K. C., & Boorman, J. (2019). Strategically-motivated advanced persistent threat: Definition, process, tactics and a disinformation model of counterattack. *Computers and Security*, 86, 402–418.
- Ahmed, M., Kambam, H. R., Liu, Y., & Uddin, M. N. (2020). Impact of Human Factors in Cloud Data Breach. *Proceedings of the 4th International Conference on Intelligent, Interactive Systems and Applications (IISA2019)*, 568–577.
- Anderson, R. (2001). *Security Engineering: A Guide to Building Dependable Distributed Systems*. Wiley.
- Anderson, R., Barton, C., Böhme, R., Clayton, R., van Eeten, M. J., Levi, M., Moore, T., & Savage, S. (2013). Measuring the Cost of Cybercrime. In R. Böhme (Ed.), *The economics of information security and privacy* (pp. 265–300). Springer.
- Anderson, R., & Moore, T. (2006). The economics of information security. *Science*, 314(5799), 610–613.
- Androulidakis, I., & Kioupakis, F.-E. (2016). *Industrial Espionage and Technical Surveillance Counter Measurers*. Springer.
- Angst, C. M., Block, E. S., D’Arcy, J., & Kelley, K. (2017). When do it security investments matter? Accounting for the influence of institutional factors in the context of healthcare data breaches. *MIS Quarterly: Management Information Systems*, 41(3), 893–916.
- Baker, K. (2022). What is Cyber Espionage? Retrieved June 7, 2022, from <https://www.crowdstrike.com/cybersecurity-101/cyberattacks/cyber-espionage/>
- Baker, N. R., Green, S. G., & Bean, A. S. (1986). Why R&D projects succeed or fail. *Research management*, 29(6).
- Bardopoulos, J. (2020). *Cyber-insurance pricing models* [Doctoral dissertation, University of Cape Town].
- Barracchini, C., & Addessi, M. E. (2014). Cyber Risk and Insurance Coverage: An Actuarial Multistate Approach. *Review of Economics & Finance*, 4, 57–69.
- Barrachina, A., Tauman, Y., & Urbano, A. (2014). Entry and espionage with noisy signals. *Games and Economic Behavior*, 83, 127–146.
- Bauer, H. (1958). Minimalstellen von Funktionen und Extrempunkte. *Archiv der Mathematik* 1957 9:4, 9(4), 389–393.
- Berentsen, A., & Lengwiler, Y. (2005). *Fraudulent Accounting and Other Doping Games (Zurich IEER Working Paper No. 175)*.

- Bertrand, J. (1989). Review of Walras's *Théorie mathématique de la richesse sociale* and Cournot's *Recherches sur les principes mathématiques de la théorie des richesses*. In A. F. Daughety (Ed.), *Cournot oligopoly* (pp. 73–81). Cambridge University Press.
- Biener, C., Eling, M., & Wirfs, J. H. (2015). Insurability of cyber risk: An empirical analysis. *Geneva Papers on Risk and Insurance: Issues and Practice*, 40, 131–158.
- Billand, P., Bravard, C., Chakrabarti, S., & Sarangi, S. (2012). *Spying in Multi-Market Oligopolies (FEEM Working Paper No. 117.2010)*.
- Bland, M. (2015). *An introduction to medical statistics* (4th ed.). Oxford University Press.
- Böhme, R. (2005). Cyber-Insurance Revisited. *Proceedings of the 4th Workshop on the Economics of Information Security (WEIS 2005)*.
- Böhme, R., & Kataria, G. (2006). On the limits of cyber-insurance. In S. Fischer-Hübner, S. Furnell, & C. Lambrinoudakis (Eds.), *Lecture notes in computer science* (pp. 31–40). Springer, Berlin, Heidelberg.
- Boyd, S. S., & Vandenberghe, L. (2004). *Convex optimization*. Cambridge University Press.
- Busso, M., & Galiani, S. (2014). The Causal Effect of Competition on Prices and Quality: Evidence from a Field Experiment (NBER Working Paper 20054).
- Callegati, F., Cerroni, W., & Ramilli, M. (2009). Man-in-the-Middle Attack to the HTTPS Protocol. *IEEE Security & Privacy Magazine*, 7(1), 78–81.
- Cerdeiro, D., Dziubiński, M., & Goyal, S. (2015). *Contagion Risk and Network Design (FEEM Working Paper No. 056.2015)*.
- Chen, P.-L. (2016). Cross-Country Economic Espionage and Investment in Research and Development. *International Journal of Economics and Finance*, 8(4).
- Cheng, L., Liu, F., & Yao, D. D. (2017). Enterprise data breach: causes, challenges, prevention, and future directions. *Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery*, 7(5).
- Choi, B. C., Kim, S. S., & Jiang, Z. (2016). Influence of Firm's Recovery Endeavors upon Privacy Breach on Online Customer Behavior. *Journal of Management Information Systems*, 33(3), 904–933.
- Cisco. (2019). *Data Privacy Benchmark Study: Maximizing the value of your data privacy investments* (tech. rep.).
https://www.cisco.com/c/dam/en_us/about/doing_business/trust-center/docs/dpbs-2019.pdf
- Coe, D. T., & Helpman, E. (1995). International R&D spillovers. *European Economic Review*, 39(5), 859–887.
- Collins, L. M., Schafer, J. L., & Kam, C. M. (2001). A comparison of inclusive and restrictive strategies in modern missing data procedures. *Psychological Methods*, 6(3), 330–351.
- Comanor, W. S. (1965). Research and Technical Change in the Pharmaceutical Industry. *The Review of Economics and Statistics*, 47(2), 182.
- Crane, A. (2005). In the company of spies: When competitive intelligence gathering becomes industrial espionage. *Business Horizons*, 48(3), 233–240.
- Crete-Nishihata, M., Dalek, J., Maynier, E., & Scott-Railton, J. (2018). *Spying on a budget. Inside a Phishing Operation with Targets in the Tibetan*

- Community Suggested Citation Copyright* (tech. rep.).
<https://citizenlab.ca/2018/01/spying->
- Curry, P. A., & Mongrain, S. (2000). Deterrence in Rank-Order Tournaments. *Review of Law & Economics*, 5(1).
- Darvari, V.-A., Hailes, S., & Musolesi, M. (2021). Goal-directed graph construction using reinforcement learning. *Proceedings of the Royal Society A: Mathematical, Physical and Engineering Sciences*, 477(2254).
- Davidson, R., & MacKinnon, J. G. (1981). Several Tests for Model Specification in the Presence of Alternative Hypotheses. *Econometrica*, 49(3), 781–793.
- DCMS. (2018). *The Network and Information Systems Regulations 2018* (tech. rep.). Queen’s Printer of Acts of Parliament.
<https://www.gov.uk/government/collections/nis-directive-and-nis-regulations-2018>
- DCMS. (2022). *Cyber Security Breaches Survey 2022* (tech. rep. No. 4).
<https://www.gov.uk/government/publications/cyber-security-breaches-survey-2020/cyber-security-breaches-survey-2020>
- Denter, P., & Sisak, D. (2015). Do polls create momentum in political competition? *Journal of Public Economics*, 130, 1–14.
- Denter, P., & Sisak, D. (2016). Head starts in dynamic tournaments? *Economics Letters*, 149, 94–97.
- Deshmukh, R. V., & Devadkar, K. K. (2015). Understanding DDoS Attack & its Effect in Cloud Environment. *Procedia Computer Science*, 49, 202–210.
- Dhillon, G., Smith, K., & Dissanayaka, I. (2021). Information systems security research agenda: Exploring the gap between research and practice. *The Journal of Strategic Information Systems*, 30(4).
- Donaldson, S., Navin Shah, J., Pedley, D., Crozier, D., & Furnell, S. (2020). *UK Cyber Security Sectoral Analysis 2020* (tech. rep.).
<https://www.ncsc.gov.uk/news/annual-review-2019>
- Dresher, M. (1961). *Games of Strategy: Theory and Applications*. Prentice Hall.
- Dziubiński, M., & Goyal, S. (2013). Network design and defence. *Games and Economic Behavior*, 79(1), 30–43.
- Dziubiński, M., & Goyal, S. (2017). How do you defend a network? *Theoretical Economics*, 12(1), 331–376.
- Eckhart, M., Meixner, K., Winkler, D., & Ekelhart, A. (2019). Securing the testing process for industrial automation software. *Computers & Security*, 85, 156–180.
- Eichfelder, G., Klamroth, K., & Niebling, J. (2021). Nonconvex constrained optimization by a filtering branch and bound. *Journal of Global Optimization*, 80(1), 31–61.
- Eling, M., & Schnell, W. (2016). What do we know about cyber risk and cyber risk insurance? *The Journal of Risk Finance*, 17(5), 474–491.
- Erdős, P., & Gallai, T. (1960). Gráfok előírt fokszámú pontokkal. *Matematikai Lapok*, 11, 264–274.
- European Commission. (2008). *NACE Rev. 2 – Statistical classification of economic activities in the European Community* (tech. rep.). Office for Official Publications of the European Communities.
<https://ec.europa.eu/eurostat/documents/3859598/5902521/KS-RA-07-015-EN.PDF>

- Eurostat. (2010a). Digital Economy and Society [Data set]. <https://ec.europa.eu/eurostat/web/digital-economy-and-society/data/database>
- Eurostat. (2010b). Science, Technology and Innovation [Data set]. <https://ec.europa.eu/eurostat/web/science-technology-innovation/data/database>
- Eurostat. (2010c). Structural Business Statistics & Global Business Activities [Data set]. <https://ec.europa.eu/eurostat/web/structural-business-statistics/data/database>
- Eurostat. (2016). *High-tech industry and knowledge-intensive services (HTEC)* (tech. rep.). https://ec.europa.eu/eurostat/cache/metadata/en/htec_esms.htm
- Gandal, N., Moore, T., Riordan, M., & Barnir, N. (2022). Empirically Evaluating the Effect of Cybersecurity Precautions on Incidents in Israeli Enterprises. *Proceedings of the 21st Workshop on the Economics of Information Security (WEIS 2022)*.
- Gelman, A., Carlin, J. B., Stern, H. S., Dunson, D. B., Vehtari, A., & Rubin, D. B. (2013). *Bayesian Data Analysis* (3rd ed.). Chapman & Hall/CRC.
- Geoffrion, A. M. (1971). Duality in Nonlinear Programming: A Simplified Applications-Oriented Development. *SIAM Review*, 13(1), 1–37.
- Gilpatric, S. (2011). Cheating in contests. *Economic Inquiry*, 49(4), 1042–1053.
- Goyal, S. (2007). *Connections: An introduction to the economics of networks*. Princeton University Press.
- Goyal, S., Jabbari, S., Kearns, M., Khanna, S., & Morgenstern, J. (2016). Strategic Network Formation with Attack and Immunization. In E. Bertino, W. Gao, B. Steffen, & M. Yung (Eds.), *Lecture notes in computer science* (pp. 429–443, Vol. 10123). Springer.
- Goyal, S., & Vigier, A. (2010). *Robust Networks (STICERD Working Paper)*. <https://sticerd.lse.ac.uk/seminarpapers/et11032010.pdf>
- Goyal, S., & Vigier, A. (2014). Attack, Defence, and Contagion in Networks. *The Review of Economic Studies*, 81(4), 1518–1542.
- Gradstein, M., & Konrad, K. A. (1999). Orchestrating Rent Seeking Contests. *The Economic Journal*, 109(458), 536–545.
- Graham, J. W., Olchowski, A. E., & Gilreath, T. D. (2007). How Many Imputations are Really Needed? Some Practical Clarifications of Multiple Imputation Theory. *Prevention Science*, 8(3), 206–213.
- Gruber, W., Mehta, D., & Vernon, R. (1967). The R&D Factor in International Trade and International Investment of United States Industries. *Journal of Political Economy*, 75(1), 20–37.
- Gueye, A., Walrand, J. C., & Anantharam, V. (2011). How to Choose Communication Links in an Adversarial Environment? In R. Jain & R. Kannan (Eds.), *Game theory for networks* (pp. 233–248). Springer.
- Haislip, J., Lim, J.-H., & Pinsker, R. (2021). The Impact of Executives’ IT Expertise on Reported Data Security Breaches. *Information Systems Research*, 32(2), 318–334.
- Harris, J. R. (1985). The Rolt Memorial Lecture, 1984 Industrial Espionage in the Eighteenth Century. *Industrial Archaeology Review*, 7(2), 127–138.

- Harris, J. R. (2017). French Industrial Espionage in Eighteenth-Century Britain and Technology Transfer: an Evaluation. In *Industrial espionage and technology transfer* (pp. 544–565). Routledge.
- Hart, S. (2008). Discrete Colonel Blotto and General Lotto games. *International Journal of Game Theory*, 36(3-4), 441–460.
- Heitzenrater, C. D. (2016). *Software security investment modelling for decision-support* [Doctoral dissertation, University of Oxford]. Oxford.
<https://ora.ox.ac.uk/objects/uuid:64ddd45e-87ab-4c92-a085-df2d0d4e22e0>
- Henningsen, S., Dietzel, S., & Scheuermann, B. (2018). Misbehavior Detection in Industrial Wireless Networks: Challenges and Directions. *Mobile Networks and Applications*, 23(5), 1330–1336.
- Herath, H., & Herath, T. (2007). Cyber-Insurance: Copula Pricing Framework and Implication for Risk Management. *Proceedings of the 6th Workshop on the Economics of Information Security (WEIS 2007)*.
- Herath, H. S., & Herath, T. C. (2011). Copula-based actuarial model for pricing cyber-insurance policies. *Insurance Markets and Companies: Analysis and Actuarial Computation*, 2(1).
- Higgs, J. L., Pinsker, R. E., Smith, T. J., & Young, G. R. (2016). The Relationship between Board-Level Technology Committees and Reported Security Breaches. *Journal of Information Systems*, 30(3), 79–98.
- Hoo, K. J. S. (2000). *How Much Is Enough? A Risk-Management Approach to Computer Security*, University of Stanford.
https://cisac.fsi.stanford.edu/publications/how_much_is_enough_a_riskmanagement_approach_to_computer_security
- HOSAC. (2018). *Understanding the costs of cyber crime A report of key findings from the Costs of Cyber Crime Working Group* (tech. rep.). Home Office Science Advisory Council.
https://assets.publishing.service.gov.uk/government/uploads/system/uploads/attachment_data/file/674046/understanding-costs-of-cyber-crime-horr96.pdf
- Hou, T., & Wang, V. (2020). Industrial espionage – A systematic literature review (SLR). *Computers & Security*, 98.
- Hovav, A., Han, J., & Kim, J. (2017). Market Reaction to Security Breach Announcements. *ACM SIGMIS Database: the DATABASE for Advances in Information Systems*, 48(1), 11–52.
- Hoyer, B., & Jaegher, K. D. (2016). Strategic Network Disruption and Defense. *Journal of Public Economic Theory*, 18(5).
- Hsu, C., & Wang, T. (2021). Too Busy to Monitor? Board Busyness and the Occurrence of Reported Information Security Incidents. *Proceedings of the Annual Hawaii International Conference on System Sciences*.
- IBM Security. (2020). *Cost of a Data Breach Study 2020* (tech. rep.). International Business Machines.
<https://www.ibm.com/account/reg/uk-en/signup?formid=urx-46542>
- IBM Security. (2021). *Cost of a Data Breach Report 2021* (tech. rep.). International Business Machines.
<https://www.ibm.com/account/reg/uk-en/signup?formid=urx-50915>

- Innerhofer-Oberperfler, F., & Breu, R. (2010). Potential Rating Indicators for Cyberinsurance: An Exploratory Qualitative Study. *Economics of Information Security and Privacy*, 249–278.
- IRM. (2014). *Cyber Risk Executive Summary* (tech. rep.). Institute of Risk Management.
<https://www.cgma.org/resources/reports/downloadabledocuments/irm-cyber-risk-report-executive-summary.pdf>
- Jackson, M. O. (2010). *Social and Economic Networks*. Princeton University Press.
- Jackson, M. O., & Rogers, B. W. (2005). The Economics of Small Worlds. *Journal of the European Economic Association*, 3(2-3), 617–627.
- Jackson, M. O., & Wolinsky, A. (1996). A Strategic Model of Social and Economic Networks. *Journal of Economic Theory*, 71(1), 44–74.
- Jones, C. I. (2003). Human Capital, Ideas and Economic Growth. In *Finance, research, education and growth* (pp. 51–74). Palgrave Macmillan UK.
- Juniper Research. (2019). *Business Losses to Cybercrime Data Breaches to Exceed \$5 trillion* (tech. rep.). Juniper Research.
<https://www.juniperresearch.com/press/press-releases/business-losses-cybercrime-data-breaches>
- Khan, F., Kim, J. H., Mathiassen, L., & Moore, R. (2021). Data Breach Management: An Integrated Risk Model. *Information & Management*, 58(1).
- Khraisat, A., Gondal, I., Vamplew, P., & Kamruzzaman, J. (2019). Survey of intrusion detection systems: techniques, datasets and challenges. *Cybersecurity*, 2(1), 20.
- Kianpour, M., Kowalski, S. J., & Øverby, H. (2021). Systematically Understanding Cybersecurity Economics: A Survey. *Sustainability*, 13(24).
- Kidmose, E., Stevanovic, M., & Pedersen, J. M. (2018). Detection of Malicious domains through lexical analysis. *2018 International Conference on Cyber Security and Protection of Digital Services*, 1–5.
- Kim, S. H., & Kwon, J. (2019). How Do EHRs and a Meaningful Use Initiative Affect Breaches of Patient Information? *Information Systems Research*, 30(4), 1184–1202.
- Kolevski, D., Michael, K., Abbas, R., & Freeman, M. (2021). Cloud Data Breach Disclosures: the Consumer and their Personally Identifiable Information (PII)? *2021 IEEE Conference on Norbert Wiener in the 21st Century (21CW)*, 1–9.
- Konrad, K. A. (2009). *Strategy and Dynamics in Contests*. Oxford University Press.
- Kovenock, D., & Roberson, B. (2008). Electoral Poaching and Party Identification. *Journal of Theoretical Politics*, 20(3), 275–302.
- Kovenock, D., & Roberson, B. (2012). Conflicts with Multiple Battlefields. In M. R. Garfinkel & S. Skaperdas (Eds.), *The oxford handbook of the economics of peace and conflict* (pp. 503–531). Oxford University Press.
- Kovenock, D., & Roberson, B. (2021). Generalizations of the General Lotto and Colonel Blotto games. *Economic Theory*, 71(3), 997–1032.
- Kräkel, M. (2007). Doping and cheating in contest-like situations. *European Journal of Political Economy*, 23(4), 988–1006.

- Kwon, J., Ulmer, J. R., & Wang, T. (2013). The Association between Top Management Involvement and Compensation and Information Security Breaches. *Journal of Information Systems*, 27(1), 219–236.
- Laszka, A., Johnson, B., Schöttle, P., Grossklags, J., & Böhme, R. (2014). Secure Team Composition to Thwart Insider Threats and Cyber-Espionage. *ACM Transactions on Internet Technology*, 14(2-3), 1–22.
- Lazear, E. P., & Rosen, S. (1981). Rank-Order Tournaments as Optimum Labor Contracts. *Journal of Political Economy*, 89(5), 841–864.
- Leita, C., Pham, V., Thonnard, O., Ramirez-Silva, E., Pouget, F., Kirda, E., & Dacier, M. (2008). The Leurre.com Project: Collecting Internet Threats Information Using a Worldwide Distributed Honeynet. *2008 WOMBAT Workshop on Information Security Threats Data Collection and Sharing*, 40–57.
- Lemley, M. A. (2011). The Myth of the Sole Inventor. *SSRN Electronic Journal*, 110(5), 709–760.
- Lewis, J. A., & Baker, S. (2013). *The Economic Impact of Cybercrime and Cyber Espionage* (tech. rep.). Center for Strategic & International Studies. <https://www.csis.org/analysis/economic-impact-cybercrime-and-cyber-espionage>
- Lewis, P. A. (2017). Technician Roles, Skills, and Training in Industrial Biotechnology: An Analysis. *SSRN Electronic Journal*.
- Li, D., Sun, X. L., Biswal, M. P., & Gao, F. (2001). Convexification, Concavification and Monotonization in Global Optimization. *Annals of Operations Research*, 105, 213–226.
- Liginlal, D., Sim, I., & Khansa, L. (2009). How significant is human error as a cause of privacy breaches? An empirical study and a framework for error management. *Computers & Security*, 28(3-4), 215–228.
- Little, R. J. A., & Rubin, D. B. (2002). *Statistical Analysis with Missing Data* (D. J. Balding, N. A. C. Cressie, G. M. Fitzmaurice, G. H. Givens, H. Goldstein, G. Molenberghs, D. W. Scott, A. F. M. Smith, & R. S. Tsay, Eds.; 2nd). Wiley.
- Liu, C.-W., Huang, P., & Lucas, H. C. (2020). Centralized IT Decision Making and Cybersecurity Breaches: Evidence from U.S. Higher Education Institutions. *Journal of Management Information Systems*, 37(3), 758–787.
- Liu, W., Tanaka, H., & Matsuura, K. (2007). Empirical-Analysis Methodology for Information-Security Investment and Its Application to Reliable Survey of Japanese Firms. *IPSJ Digital Courier*.
- Liu, W., & Xin, J. (2014). *Modeling Fractional Outcomes with SAS*. <http://support.sas.com/resources/papers/proceedings14/1304-2014.pdf>
- Liu, X. (2016). Methods for handling missing data. In *Methods and applications of longitudinal data analysis* (pp. 441–473). Elsevier.
- Lynch, S. (2018). U.S. indicts Chinese, Taiwan firms for targeting Micron trade secrets [Reuters]. <https://uk.reuters.com/article/us-usa-justice-china-espionage/u-s-indicts-chinese-taiwan-firms-for-targeting-micron-trade-secrets-idUKKCN1N65R2>
- Madley-Dowd, P., Hughes, R., Tilling, K., & Heron, J. (2019). The proportion of missing data should not be used to guide decisions on multiple imputation. *Journal of Clinical Epidemiology*, 110, 63–73.

- Makridis, C. A. (2021). Do data breaches damage reputation? Evidence from 45 companies between 2002 and 2018. *Journal of Cybersecurity*, 7(1).
- Malliouris, D. D., & Simpson, A. (2020). Underlying and Consequential Costs of Cyber Security Breaches: Changes in Systematic Risk. *Proceedings of the 19th Workshop on the Economics of Information Security (WEIS 2020)*.
- Matsui, A. (1989). Information leakage forces cooperation. *Games and Economic Behavior*, 1(1), 94–115.
- McAfee. (2020). *The Hidden Costs of Cybercrime* (tech. rep.).
<https://www.mcafee.com/enterprise/en-us/assets/reports/rp-hidden-costs-of-cybercrime.pdf>
- McDonald, J. F., & Moffitt, R. A. (1980). The Uses of Tobit Analysis. *The Review of Economics and Statistics*, 62(2), 318.
- McLeod, A., & Dolezel, D. (2018). Cyber-analytics: Modeling factors associated with healthcare data breaches. *Decision Support Systems*, 108, 57–68.
- McMullen, C. (1987). Families of Rational Maps and Iterative Root-Finding Algorithms. *The Annals of Mathematics*, 125(3), 467.
- Moore, T., & Anderson, R. (2011). *Economics and Internet Security: A Survey of Recent Analytical, Empirical and Behavioral Research* (tech. rep.). Harvard Computer Science Group.
- Moore, T., & Clayton, R. (2011). The Impact of Public Information on Phishing Attack and Defense. *Communications & Strategies*, (81), 45–68.
- Morgan, S. (2020). *2019 Official Annual Cybercrime Report* (tech. rep.). Cybersecurity Ventures.
<https://www.herjavecgroup.com/wp-content/uploads/2018/12/CV-HG-2019-Official-Annual-Cybercrime-Report.pdf>
- Mukkamala, S., Sung, A., & Abraham, A. (2005). Cyber security challenges: Designing efficient intrusion detection systems and antivirus tools. In V. V Rao (Ed.), *Enhancing computer security with smart technology*. (pp. 125–163). Auerbach Publications. <https://usermanual.wiki/Document/EnhancingComputerSecuritywithSmartTechnology.2134762856.pdf>
- Nahari, S. (2018). Malware, Mistakes and Meaningful Measures to Protect Critical Infrastructure. Retrieved July 2, 2022, from
<https://www.cyberark.com/resources/blog/malware-mistakes-and-meaningful-measures-to-protect-critical-infrastructure>
- Nalebuff, B. J., & Stiglitz, J. E. (1983). Prizes and Incentives: Towards a General Theory of Compensation and Competition. *The Bell Journal of Economics*, 14(1), 21–43.
- National Cyber Security Centre. (2015). How cyber attacks work. Retrieved January 9, 2021, from
<https://www.ncsc.gov.uk/information/how-cyber-attacks-work>
- Naumowicz, A. (2014). A Note on the Seven Bridges of Königsberg Problem. *Formalized Mathematics*, 22(2), 177–178.
- Nielsen, S. F. (2007). Proper and Improper Multiple Imputation. *International Statistical Review*, 71(3), 593–607.
- OECD. (2015). *Frascati Manual 2015: Guidelines for Collecting and Reporting Data on Research, Experimental Development*. Organisation for Economic Co-operation and Development Publishing, Paris.
- Osborne, M. J. (2003). *An Introduction to Game Theory*. Oxford University Press.

- Oxford Economics. (2014). *Cyber-attacks: Effects on UK Companies* (tech. rep.). Oxford Economics.
<https://www.oxfordeconomics.com/resource/cyber-attacks-effects-on-uk/>
- Parrend, P., Navarro, J., Guigou, F., Deruyver, A., & Collet, P. (2018). Foundations and applications of artificial Intelligence for zero-day and multi-step attack detection. *EURASIP Journal on Information Security*, 2018(1), 4.
- Paul Markillie. (2012). A third industrial revolution [The Economist].
<https://www.economist.com/leaders/2012/04/21/the-third-industrial-revolution>
- Pinzon, E. (2016). Two faces of misspecification in maximum likelihood: Heteroskedasticity and robust standard errors. Retrieved November 23, 2020, from
<https://blog.stata.com/2016/08/30/two-faces-of-misspecification-in-maximum-likelihood-heteroskedasticity-and-robust-standard-errors/>
- Png, I. P. L., Wang, C.-Y., & Wang, Q.-H. (2008). The Deterrent and Displacement Effects of Information Security Enforcement: International Evidence. *Journal of Management Information Systems*, 25(2), 125–144.
- Porter, M., & Stern, S. (2000). *Measuring the "Ideas" Production Function: Evidence from International Patent Output*, National Bureau of Economic Research.
- Pulleyblank, W. (1989). Chapter V Polyhedral combinatorics. In *Handbooks in operations research and management science* (pp. 371–446).
- PwC. (2020). *23rd Annual Global CEO Survey. Navigating the rising tide of uncertainty* (tech. rep.). PricewaterhouseCoopers.
<https://www.pwc.com/gx/en/ceo-survey/2020/reports/pwc-23rd-global-ceo-survey.pdf>
- Ramsey, J. B. (1969). Tests for Specification Errors in Classical Linear Least-Squares Regression Analysis. *Journal of the Royal Statistical Society: Series B (Methodological)*, 31(2), 350–371.
- Ransbotham, S. (2010). An Empirical Analysis of Exploitation Attempts Based on Vulnerabilities in Open Source Software. *Proceedings of the 9th Workshop on the Economics of Information Security (WEIS 2010)*.
- RMS. (2016). *Cyber Accumulation Risk Management* (tech. rep.). Cambridge Centre for Risk Studies, Risk Management Solutions, Inc.
http://forms2.rms.com/rs/729-DJX-565/images/Cyber_datasheet.pdf
- Roberson, B. (2006). The Colonel Blotto game. *Economic Theory*, 29(1), 1–24.
- Rodwell, L., Lee, K. J., Romaniuk, H., & Carlin, J. B. (2014). Comparison of methods for imputing limited-range variables: a simulation study. *BMC Medical Research Methodology*, 14(1).
- Rogers, M. B., McConnell, B. M., Hodgson, T. J., Kay, M. G., King, R. E., Parlier, G., & Thoney-Barletta, K. (2018). A military logistics network planning system. *Military Operations Research*, 23(4), 5–25.
- Romanosky, S. (2016). Examining the costs and causes of cyber incidents. *Journal of Cybersecurity*, 2(2), 121–135.
- Romanosky, S., Ablon, L., Kuehn, A., & Jones, T. (2017). Content Analysis of Cyber Insurance Policies: How Do Carriers Write Policies and Price Cyber Risk? *SSRN Electronic Journal*.

- Rubin, D. B. (1987). *Multiple Imputation for Nonresponse in Surveys* (D. B. Rubin, Ed.). Wiley.
- Sadok, M., Welch, C., & Bednar, P. (2020). A socio-technical perspective to counter cyber-enabled industrial espionage. *Security Journal*, 33(1), 27–42.
- Sahinidis, N. V. (2019). Mixed-integer nonlinear programming 2018. *Optimization and Engineering*, 20(2), 301–306.
- Sawa, Y., Bhakta, R., Harris, I. G., & Hadnagy, C. (2016). Detection of Social Engineering Attacks Through Natural Language Processing of Conversations. *2016 IEEE Tenth International Conference on Semantic Computing (ICSC)*, 262–265.
- Say, G., & Vasudeva, G. (2020). Learning from Digital Failures? The Effectiveness of Firms' Divestiture and Management Turnover Responses to Data Breaches. *Strategy Science*, 5(2), 117–142.
- Schafer, J. L., & Graham, J. W. (2002). Missing data: Our view of the state of the art. *Psychological Methods*, 7(2), 147–177.
- Schlackl, F., Link, N., & Hoehle, H. (2022). Antecedents and consequences of data breaches: A systematic review. *Information & Management*, 59(4).
- Schwartz, P. M., & Janger, E. J. (2007). Notification of data security breaches. *Michigan Law Review*, 105(5).
- Scotchmer, S. (2004). *Innovation and incentives*. MIT Press.
- Sen, R., & Borle, S. (2015). Estimating the Contextual Risk of Data Breach: An Empirical Approach. *Journal of Management Information Systems*, 32(2), 314–341.
- Send, J. (2022). Contest Copycats: Adversarial Duplication of Effort in Contests. *Defence and Peace Economics*.
- Shrout, R. (2018). Intel and AMD both dive into many-core CPU race [tech.pinions]. Retrieved May 31, 2020, from <https://techpinions.com/intel-and-amd-both-dive-into-many-core-cpu-race/53040>
- Shubber, K. (2018). US charges Chinese group with theft of Micron trade secrets [Financial Times]. <https://www.ft.com/content/d34d9b58-ddff-11e8-8f50-cbae5495d92b>
- Siddiqi, M. A., Pak, W., & Siddiqi, M. A. (2022). A Study on the Psychology of Social Engineering-Based Cyberattacks and Existing Countermeasures. *Applied Sciences*, 12(12).
- Siegel, R. (2014). Asymmetric Contests with Head Starts and Nonmonotonic Costs. *American Economic Journal: Microeconomics*, 6(3), 59–105.
- Slikker, M., & van den Nouweland, A. (2000). Network formation models with costs for establishing links. *Review of Economic Design*, 5(3), 333–362.
- Sobers, R. (2018). GDPR's Impact So Far: Must-Know Stats and Takeaways. Retrieved June 25, 2020, from <https://www.varonis.com/blog/gdpr-effect-review/>
- Solan, E., & Yariv, L. (2004). Games with espionage. *Games and Economic Behavior*, 47(1), 172–199.
- Somani, G., Gaur, M. S., Sanghi, D., & Conti, M. (2016). DDoS attacks in cloud computing: Collateral damage to non-targets. *Computer Networks*, 109, 157–171.

- Spanos, G., & Angelis, L. (2016). The impact of information security events to the stock market: A systematic literature review. *Computers & Security*, 58, 216–229.
- Stata Corporation. (2013). Stata Multilevel Mixed-Effects Reference Manual. Retrieved November 23, 2020, from <https://www.stata.com/manuals13/me.pdf>
- Steger, M. B., & James, P. (2011). Three Dimensions of Subjective Globalization. *ProtoSociology*, 27(3), 53–70.
- Sterne, J. A., White, I. R., Carlin, J. B., Spratt, M., Royston, P., Kenward, M. G., Wood, A. M., & Carpenter, J. R. (2009). Multiple imputation for missing data in epidemiological and clinical research: Potential and pitfalls. *BMJ*, 339(7713), 157–160.
- Stigler, G. J. (1961). The Economics of Information. *Journal of Political Economy*, 69(3), 213–225.
- Stone, A. (2007). Natural-Language Processing for Intrusion Detection. *Computer*, 40(12), 103–105.
- Stowe, C. J., & Gilpatric, S. M. (2010). Cheating and Enforcement in Asymmetric Rank-Order Tournaments. *Southern Economic Journal*, 77(1), 1–14.
- Stupak, O. (2022a). Industrial cyberespionage in research and development races. *Research and Innovation Policies in Europe: Evolution, Scope and Perspectives*.
- Stupak, O. (2022b). No research, no risk? An empirical study of determinants of information leakage attacks. *Proceedings of the 21st Workshop on the Economics of Information Security (WEIS 2022)*.
- Stupak, O. (2022c). *Secure and efficient networks*, University of Cambridge.
- Sweeney, J. A. (2015). Infectious Connectivity: Affect and the Internet in Postnormal Times. In *Public administration and information technology* (pp. 107–123, Vol. 17). Springer International Publishing Switzerland.
- Syed, R. (2019). Enterprise reputation threats on social media: A case of data breach framing. *The Journal of Strategic Information Systems*, 28(3), 257–274.
- Sýkora, S. (2009). Mathematical means and averages: basic properties.
- Tanriverdi, H., Roumani, Y., & Nwankpa, J. K. (2019). Structural complexity and data breach risk. *40th International Conference on Information Systems, ICIS 2019*.
- Terleckyj, N. (1960). *Sources of productivity advance : A pilot study of manufacturing industries, 1899-1953* [Doctoral dissertation, Columbia University].
- The Commission on the Theft of American Intellectual Property. (2017). *The IP commission report* (tech. rep.). Center for Innovation, Trade, and Strategy. <https://www.nbr.org/program/commission-on-the-theft-of-intellectual-property/>
- Tripathi, A., & Vijay, S. (2003). A note on a theorem of Erdős & Gallai. *Discrete Mathematics*, 265(1-3), 417–420.
- Tullock, G. (2001). Efficient Rent Seeking. In G. Tullock & A. A. Lockard (Eds.), *Efficient rent-seeking* (pp. 3–16). Springer US.

- Tuy, H., & Van Thuong, N. (1988). On the global minimization of a convex function under general nonconvex constraints. *Applied Mathematics and Optimization*, 18(1), 119–142.
- Van Alsenoy, B. (2019). General Data Protection Regulation. In *Data protection law in the eu: Roles, responsibilities and liability* (pp. 279–324). Intersentia.
- van Buuren, S., Boshuizen, H. C., & Knook, D. L. (1999). Multiple imputation of missing blood pressure covariates in survival analysis. *Statistics in Medicine*, 18(6), 681–694.
- Verizon Business. (2011). *2011 Data Breach Investigations Report (DBIR)* (tech. rep.). Verizon.
https://www.wired.com/images_blogs/threatlevel/2011/04/Verizon-2011-DBIR_04-13-11.pdf
- Verizon Business. (2020). 2020 The Veris Community Database (VCDB) [Data set].
<http://veriscommunity.net/vcdb.html>
- Villadsen, A. R., & Wulff, J. N. (2021). Are you 110% sure? Modeling of fractions and proportions in strategy and management research. *Strategic Organization*, 19(2), 312–337.
- von Hippel, P. T. (2013). Should a Normal Imputation Model be Modified to Impute Skewed Variables? *Sociological Methods & Research*, 42(1), 105–138.
- von Hippel, P. T. (2020). How Many Imputations Do You Need? A Two-stage Calculation Using a Quadratic Rule. *Sociological Methods & Research*, 49(3), 699–718.
- von Neumann, J. (1928). Zur Theorie der Gesellschaftsspiele. *Mathematische Annalen*, 100(1), 295–320.
- Wang, Q. H., & Kim, S. H. (2009). Cyber attacks: Does physical boundary matter? *Proceedings of 30th Annual International Conference on Information Systems*.
- Wang, T., & Hsu, C. (2010). The impact of board structure on information security breaches. *PACIS 2010 - 14th Pacific Asia Conference on Information Systems*.
- Wang, Y., Sinha, A., Ch-Wang, S., & Wellman, M. P. (2021). Building Action Sets in a Deep Reinforcement Learner. *2021 20th IEEE International Conference on Machine Learning and Applications (ICMLA)*, 484–489.
- Wangen, G. (2015). The Role of Malware in Reported Cyber Espionage: A Review of the Impact and Mechanism. *Information*, 6(2), 183–211.
- Warrell, H., & Manson, K. (2020). State-backed hackers using virus to increase spying, UK and US warn [Financial Times].
<https://www.ft.com/content/37149106-eb16-4b4e-879b-2913b99da84f>
- Washburn, A. (2013). OR Forum—Blotto Politics. *Operations Research*, 61(3), 532–543.
- Wedderburn, R. (1974). Quasi-likelihood functions, generalized linear models, and the Gauss—Newton method. *Biometrika*, 61(3), 439–447.
- WEF. (2015). *Global Risks 2015 10th Edition Insight Report* (tech. rep.). World Economic Forum. https://www3.weforum.org/docs/WEF_Global_Risks_2015_Report15.pdf
- Weisstein, E. W. (n.d.). Squeezing Theorem. Retrieved May 9, 2022, from <https://mathworld.wolfram.com/SqueezingTheorem.html>

- White, H. (1994). *Estimation, Inference and Specification Analysis*. Cambridge University Press.
- Whitney, M. E., & Gaisford, J. D. (1996). Economic Espionage as Strategic Trade Policy. *The Canadian Journal of Economics*, 29(2), 627–632.
- Whitney, M. E., & Gaisford, J. D. (1999). An Inquiry Into the Rationale for Economic Espionage. *International Economic Journal*, 13(2), 103–123.
- Williams, R. (2019). Analyzing Proportions: Fractional Response and Zero One Inflated Beta Models. Retrieved June 9, 2022, from <https://www3.nd.edu/~rwilliam/stats3/fractionalresponsemodels.pdf>
- Wolfram Research. (2007). GraphData [Data set]. <https://reference.wolfram.com/language/ref/GraphData.html>
- Wooldridge, J. M. (2010). *Econometric Analysis of Cross Section and Panel Data*. MIT Press.
- Xu, P., Roosta, F., & Mahoney, M. W. (2020). Newton-type methods for non-convex optimization under inexact Hessian information. *Mathematical Programming*, 184(1-2), 35–70.
- Ye, Q., & Kim, H. (2019). Assessing network vulnerability of heavy rail systems with the impact of partial node failures. *Transportation*.
- Yuan, S., & Wu, X. (2021). Deep learning for insider threat detection: Review, challenges and opportunities. *Computers & Security*, 104.