



LUND UNIVERSITY
School of Economics and Management
Department of Informatics

AI-policy formulation and practise

**A study of AI-policy and organisations AI practises
emergence and co-existence**

Master thesis 15 HEC, course INFM12, Master Thesis in Information Systems

Authors: Astrid Andersson
Fanny Hedin

Supervisor: Paul Pierce

Grading Teachers: Stig Strandbæk Nyman
Osama Mansour

AI-policy formulation and practise: A study of AI-policy and organisations AI practises emergence and co-existence

AUTHORS: Astrid Andersson and Fanny Hedin

PUBLISHER: Department of Informatics, Lund School of Economics and Management,
Lund University

PRESENTED: May, 2026

DOCUMENT TYPE: Master Thesis

FORMAL EXAMINER: Osama Mansour, Associate professor

NUMBER OF PAGES: 74

KEY WORDS: AI-policy, policy formulation, AI governance, Institutional theory, isomorphism, decoupling, Responsible AI

ABSTRACT (MAX. 200 WORDS):

The usage of AI in organisations is increasing at a fast pace and is creating a growing demand for AI governance and AI-policies. Despite this, the knowledge of how policies are formulated and translated into practise is at large limited. This study aims to qualitatively explore how Nordic organisations navigate both the process of formulating and translating their AI-policies into practise. This study is based on seven qualitative semi-structured interviews with people in management in medium to large organisations operating in different industries.

This study applies Institutional theory as the theoretical and analytical framework and identifies three isomorphic mechanisms in policy formulation: coercive, mimetic and normative isomorphism. Two additional views of considering the AI-policy's function and purpose also arose, namely seeing it as an external communication tool for stakeholders or seeing it as an internal governance tool.

The results show that organizations are actively working to achieve AI-policy understanding rather than simply policy awareness. This is shown by the inclusion of training efforts,

executive engagement as well as short and concise policy documents. This study also identifies risks of policy-practise decoupling, especially as rapid technological development makes it difficult to keep policies relevant over time.

Content

1	Introduction.....	9
1.1	Background	9
1.2	Problem	10
1.3	Purpose and objective.....	11
1.4	Research Question.....	11
1.5	Delimitation.....	11
2	Theoretical background.....	13
2.1	AI in organisations	13
2.1.1	AI adoption.....	13
2.1.2	Responsible AI	14
2.1.3	Summary	14
2.2	Policy formulation.....	15
2.2.1	AI governance	15
2.2.2	Policy drivers and motivations	16
2.2.3	Summary	17
2.3	Policy in practise	18
2.3.1	Policy-practise decoupling	18
2.3.2	Means-ends decoupling.....	19
2.3.3	Handling of future AI practises	19
2.3.4	Summary	20
2.4	Institutional Theory	20
2.4.1	Strong and weak institutional pressure.....	21
2.4.2	Isomorphism.....	22
2.4.3	Decoupling	23
2.5	Literature matrix.....	24
3	Methodology.....	26
3.1	Research philosophy	26
3.2	Research approach.....	27
3.3	Literature collection and review.....	27
3.4	Data collection.....	30
3.4.1	Semi-structured interviews.....	30
3.4.2	Interview Guide.....	31
3.4.3	Participant selection	34
3.5	Data analysis and preparation process	37

3.5.1	Transcription	37
3.5.2	Data storing	37
3.5.3	Transcript validation	38
3.5.4	Data analysis	38
3.6	Research Quality	42
3.7	Ethics	43
4	Findings.....	46
4.1	AI in organisations	46
4.1.1	AI in use	46
4.2	Policy formulation.....	47
4.2.1	Purpose of policy.....	47
4.2.2	Policy formulation process.....	50
4.3	Policy in practise	52
4.3.1	Communication	52
4.3.2	Compliance.....	54
4.3.3	Future handling	57
5	Discussion	59
5.1	AI in organisations	59
5.1.1	AI in use	59
5.2	Policy formulation.....	60
5.2.1	Purpose of policy.....	60
5.2.2	Policy formulation process.....	61
5.3	Policy in practise	64
5.3.1	Communication	64
5.3.2	Compliance.....	65
5.3.3	Future handling	67
6	Conclusion	69
6.1	Future research	70
	Appendix: AI contribution statement	71
	References	72

Figures:

Figure 2.1: Strong and weak institutional pressures (Minkkinen & Mäntymäki, 2023).....	22
Figure 2.2: Institutional logics and pressures to comply (Berente et al., 2019).....	23
Figure 3.1: Example of initial pitch sent by email and via LinkedIn	36
Figure 3.2: Participants rights and consent form.....	44

Tables:

Table 2.1: Literature matrix	24
Table 3.1: Synthesized research model	28
Table 3.2: Source Mapping of Interview Guide and Literature review	32
Table 3.3: Table of informants	35
Table 3.4: Coding	39
Table 3.5: Frequency of themes coded.....	41

1 Introduction

In this chapter the research purpose, aim and question will be presented, the chapter will commence in an introductory background to AI-policy formulation and practise.

1.1 Background

As organisations are increasingly adopting Artificial Intelligence (AI), the challenge of governing AI has become an increasing concern (Rivera et al., 2025). Despite the growing regulatory pressures and widespread adoption, organisations still struggle at formulating AI-policies and translating these into practise (Horneber, 2025) and therefore risk negatively affecting both established organisational structures and processes (Van Giffen & Ludwig, 2023).

AI is often described as a system's ability to replicate human cognitive abilities, such as reasoning and logical thinking (Janiesch et al., 2021; Lee, 2020). Contemporary AI applications span from robotics, autonomous vehicles, facial recognition, natural language processing, and virtual agents, that transform operations across numerous business sectors (Van Giffen & Ludwig, 2023). Within the Information Systems (IS) discipline, AI is not viewed as a static technological artefact but rather as an evolving domain of computational innovations that emulate human intelligence to enhance complex decision-making (Berente et al., 2021). Given the field's traditional focus on the relationship between technology and human behavior, the rise of AI marks a potential paradigm shift in the collaboration between humans and digital technologies (Conboy et al., 2022). Through its diverse abilities, it has become a strategic enabler and can be used as a competitive edge for organisations that redefine their competing dynamics and business models (Van Giffen & Ludwig, 2023). However, despite the strategic opportunities, organisations still struggle with navigating AI governance (Mäntymäki et al., 2023).

The Artificial Intelligence Index Report 2025 indicates that AI adoption accelerated rapidly in 2024, transitioning to becoming a central driver in organisations. Organizational adoption jumped from 55% to 78% in a single year (Maslej et al., 2025). AI now represents the forefront of computational systems, with forecasts suggesting a global economic impact of \$15.7 trillion by 2030 (Minkkinen & Mäntymäki, 2023). Yet, despite this rapid rise in use, there remains a persistent gap between the potential and realized value of AI. Ransbotham et al. (2019) found that although 90% of organizations have invested in AI, only 40% report measurable business benefits even though they make substantial investments in AI.

This gap illustrates a broader governing challenge. As adoption expands, governance and control challenges intensify (Berente et al., 2021; Horneber, 2025). Berente et al. (2021) further highlight that managing AI differs materially from traditional IT management because of three interrelated traits of AI. Firstly, autonomy which means the ability to independently operate without human input. Secondly, learning which refers to the model's ability to self-improvement through data exposure. Thirdly, inscrutability which is the nature of advanced models to be “black box” and opaque. Consequently, interest in Responsible AI (RAI) has surged, emphasizing ethical development, deployment, and governance that aligns with legal, moral, and societal standards to maximize benefits and minimize risks (Horneber, 2025).

Despite this the mechanisms through which AI-policies are formulated and enacted remains poorly understood (Horneber, 2025; Rivera et al., 2025; Williams et al., 2024).

1.2 Problem

Researchers in the IS field have extensively examined AI governance frameworks, RAI principles, and institutional pressures shaping AI adoption (Berente et al., 2021; Horneber, 2025; Rivera et al., 2025; Conboy et al., 2022). Despite this, limited empirical research has explored how internal AI-policies are formulated within organizations and how these policies actually influence the development and use of AI systems (Horneber, 2025; Rivera et al., 2025; Williams et al., 2024). In particular, the socio-technical translation between policy formulation and practical technological execution still remains underexplored (Holmström, 2024; Horneber, 2025; Rivera et al., 2025; Conboy et al., 2022). While this foundation in the IS field provides a thorough understanding of what AI should be and the guidelines it should follow, a gap still remains in how these abstract ideals are translated into practise by organisations who are facing multiple real-world pressures (Rivera et al., 2025).

Horneber (2025) argues that the pressures impacting AI implementation comes down to the tension between external and internal drivers. Some of the external drivers he highlights are the legal compliance and reputational concerns. This is also mentioned by Rakova et al. (2021) who press on the reputational risk as driver for RAI, while also highlighting the complexities of legal compliance as a significant challenge. As for internal drivers, Horneber (2025) lifts the competing demand for functional AI and Rakova et al. (2021) also raise concern of the simultaneous lack of formal accountability mechanisms. Horneber (2025) argues that this tension often leads to policy-practise decoupling, which is where an organization adopts a policy merely symbolically to gain legitimacy but in reality, fails to implement those policies into actual practises.

The topic of researching decoupling is both timely and highly relevant, as AI continues to reshape how individuals and organizations interact with digital technologies (Holmström, 2024; Berente et al., 2021). Organizations are currently moving from experimenting with AI adoption to institutionalizing and regulating it instead (Van Giffen & Ludwig, 2023; Grøder et al., 2025; Minkkinen & Mäntymäki, 2023). Furthermore, with the introduction of regulations such as the EU AI Act, expectations concerning accountability and risk control are being brought up, and thus incentivizing organizations to create more structured AI governance preventative measures (Van Giffen & Ludwig, 2023; Minkkinen & Mäntymäki, 2023; Horneber, 2025). This offers a timely opportunity to explore how AI-policies are developed and acted upon before they become too embedded and unquestioned organizational routines (Grøder et al., 2025). Further investigation is therefore beneficial, and also called for by Holmström (2024), in order to further explore how AI and human practises co-evolve within organizational contexts.

This topic is further highly relevant to the IS field since it addresses both the social and technical aspects of organisations management of AI (Berente et al., 2021). Furthermore Ågerfalk (2020), defines the IS field as:

“... the design and use of information and communication technologies in practical contexts and critically uncovering the practises that emerge as digital technology is appropriated.” (Ågerfalk, 2020 p.1)

He suggests that with the social aspect on one side and the technological aspect on the other side, IS researchers are at the perfect place to explore what practises emerge when the two meet in practise. Therefore, by positioning this study in the IS field, the aim of studying how internal AI-policies are formulated and how these policies shape the AI use within organisations is well aligned with the purpose of the IS research field.

As a response to the gaps presented by Horneber's (2025) in understanding the motivations and strategies AI practitioners use to navigate conflicting organizational pressures that in turn shapes the structure of policies. As well as the gap identified by Rivera et al. (2025) as to how these policies are translated into the realities of organisations. This study will focus on these two perspectives of the AI-policy, both in formulation and in practise, and the interplay between these two perspectives.

Concluding, while AI is increasingly relevant and used by organizations of all sizes, the mechanisms for how this evolution is controlled and how effective AI-policies are created still remain poorly understood (Holmström, 2024). This study is motivated by the critical need to bridge the gap between high-level governance ideals and the technical reality of AI systems in practise (Grøder et al., 2025; Horneber, 2025).

1.3 Purpose and objective

Existing literature highlights AI governance as an increasingly complex challenge for organisations, as the rapid diffusion of AI technologies disrupts established organisational structures and processes (Van Giffen & Ludwig, 2023). Despite growing pressure to govern AI responsibly, driven by regulatory developments such as the EU AI Act and shifting stakeholder expectations, limited empirical research has explored how organisations actually formulate their AI-policies and translate them into practise (Horneber, 2025; Rivera et al., 2025; Williams et al., 2024). This study aims to qualitatively explore how medium-to-large sized organisations operating in the Nordic region navigate the formulation of internal AI-policies and how these policies are enacted in practise, in order to deepen the understanding of how organisations manage the complexities of governing AI in an era of rapid technological change and uncertainty.

1.4 Research Question

How do organisations navigate AI-policy formulation and AI-policy in practise?

1.5 Delimitation

The purpose of this research is to evaluate the AI-policy formulation processes and its practical implementation within organizations. The choice of organizations will be limited to

medium-to-large sized organisations active in the Nordic Region utilising AI. The research is by design not limited to a specific sector, and while the research aim is to not compare different sectors, the objective is instead focused on identifying patterns spanning over multiple different industries, as to not be constrained by industry specific regulations. The scope will not extend to comparison or mapping of the content of policies but rather gaining an understanding of the factors impacting policy formulation and policy in practise.

As this study seeks to understand the processes that affect AI-policy formulation and enactment, rather than the content of policies themselves, triangulation with AI-policy documents has not been performed. Similarly, the study is not limited to a specific AI technology, as the aim is not to examine how organisations regulate particular tools, but rather to identify patterns in how AI governance processes unfold more broadly.

The research timeline is constrained to 2-3 months which limits the scope of the amount of semi-structured interviews and could potentially reduce data depth. This time restriction, combined with potential challenges in accessing domain experts within target sectors, may also impact the sample size and participant diversity, thereby affecting the study's generalizability. The data collection is also limited by design to informants on management level and this study will therefore not focus on the employee perspective. The choice of the type of informants for the data collection being on a managerial level was made due to time constraints. The exclusion of an employee perspective was made since it was considered to be a complementary perspective but not the primary focus, the primary focus of this study is to collect data from the roles that has insights into the actual process of the AI-policy formulation and compliance work.

2 Theoretical background

In this chapter the relevant literature and previous research surrounding AI and AI-Policies will be presented. After reviewing the current literature, the different related aspects have been divided across three main themes: 1. AI in organisations, which will describe AI and its role in organisations. 2. Policy formulation, which will present the drivers and processes of AI-policy formulation. 3. AI-policy in practise, where the practical implementation of AI-policy will be presented. Further section 2.4 of this chapter will present this study's theoretical framework Institutional Theory.

2.1 AI in organisations

This chapter introduces the definition of AI and further discusses AI adoption and Responsible AI.

Creating an universal definition of AI is a battle that may never be completely solved, as there are many different ways of perceiving it (Dignum, 2019). One way of perceiving AI is through the idea that AI is a frontier instead of a specific technology or phenomenon (Berente et al., 2021). Berente et al. (2021) argue that AI can be perceived as a process, further defining it as "*the frontier of computational advancements that references human intelligence in addressing ever more complex decision-making problems*". Berente et al. (2021) pushes the idea that, ultimately, AI is a concept representing the ever-moving front of computational advancements rather than the actual computer programs or algorithms themselves.

AI as a scientific discipline has its start in the 1950s but has just over recent years become widely recognised. In earlier years it was defined as computers' abilities to imitate and understand natural intelligence, in other words human intelligence. Today, it most commonly refers to computational abilities that interpret large volumes of data to derive decisions. (Dignum, 2019) Further, present AI systems move at an increasing rate from passive and human-controlled tools toward being autonomous digital agents capable of making independent decisions with real-world effects (Conboy et al., 2022). As this shift spans a broad range of technologies and applications (Van Giffen & Ludwig, 2023), this study therefore does not focus on any particular AI technology but rather discusses AI from a broader organisational perspective.

2.1.1 AI adoption

Over the last years the increased computational power alongside the enhanced abilities to utilize large data has enabled the rise in AI and particular machine learning abilities (Ayling & Chapman 2022). Today, AI is used to enhance quality, decisions, and processes within organisations spanning multiple fields both in the public and private sector (Mäntymäki et al., 2023).

As a result of rapid AI diffusion, society and organizations face new and complex challenges in mitigating harm (Rivera et al., 2025). Ayling and Chapman (2022) describe *epistemic* concerns surrounding the black-box nature of AI algorithms, which leads to an inherent lack of scrutability, but also the risk connected to AI systems being largely probabilistic in nature, as well as unreliability in the data used in training and execution of the systems. They also

describe risks that are *normative* in nature, these risks include concerns about maintaining fairness, privacy and AIs effect on surveillance and profiling. Furthermore, they highlight that when using AI systems, it often becomes unclear who is responsible for system outcomes, which leads to difficulties in assigning accountabilities (Ayling & Chapman 2022). To meet these challenges, more pressure is put on creation of ethical frameworks both in public policy, corporations and in society as a whole (Ayling & Chapman 2022; Rivera et al., 2025).

2.1.2 *Responsible AI*

The continuous advancements in AI, along with its rapid and widespread diffusion across a wide range of organisational activities and sectors, have placed increasing pressure on organisations ethical and responsible use and effective governance of AI systems (de Laat, 2021; Gianni et al., 2022; Papagiannidis et al., 2025). AI is effectively transforming both every day and organisational life and has the ability to promote substantial societal and organisational value; however, it can also give rise to severe risks (Vassilakopoulou et al., 2022). AI-related incidents have highlighted the potential risk AI poses to both individuals and society as a whole. As a result, ethical issues arising from AI use are widely debated in law, research, and within corporations (Papagiannidis et al., 2025; Mäntymäki et al., 2023; Horneber, 2025; Ayling & Chapman, 2022). Ethical risks stemming from AI use include concerns to human rights, such as unemployment, bias and discrimination as well as environmental and sustainability risks, and concerns regarding privacy and data governance (Fioravante, 2024).

As a result, organisations face increasing public and regulatory pressures to mitigate ethical risks when deploying and utilising AI (Horneber, 2025). In response, the term Responsible AI (RAI) is increasingly being discussed as an approach to ensuring alignment with ethical, legal and societal norms (Horneber, 2025; Rivera et al., 2025; Dignum 2019). RAI is defined as:

“Rather, it is the development of intelligent systems according to fundamental human principles and values. Responsibility is about ensuring that results are beneficial for many instead of a source of revenue for a few.” (Dignum, 2019 p.6)

RAI is thus concerned with the ethical consequences and impacts of AI and seeks to provide direction for a behavioural code both for AI systems and humans (Dignum 2019). However, it is important to note that defining what constitutes RAI remains a challenging task, as governmental bodies, practitioners and scholars are simultaneously developing frameworks and policies to ensure responsible and ethical use (Papagiannidis et al., 2025; Rivera et al., 2025). Furthermore, Ayling and Chapman (2022) criticises the current AI ethics debate as too focused on defining what is ethical, rather than how these ethics are applied in practise, resulting in organisations struggling to balance profitability with ethical considerations.

2.1.3 *Summary*

As AI is increasingly being adopted across organisations in both the private and public sector (Mäntymäki et al., 2023), debate has emerged regarding how organisations and society should approach this adoption, particularly with respect to the potential risks involved (Ayling & Chapman, 2022). In response, greater pressure is placed on ethical and responsible AI use (de Laat, 2021; Gianni et al., 2022; Papagiannidis et al., 2025). As a result, Responsible AI (RAI)

is increasingly being discussed as a means of addressing these challenges (Horneber, 2025; Rivera et al., 2025; Dignum, 2019).

Although research on RAI has advanced considerably, significant work remains in balancing the collective maximisation of societal value with the prioritisation of harm mitigation (Rivera et al., 2025). A critical insight gained from the literature is that although RAI and the ethics of AI is being discussed, there still a lack of perspective on how to apply these values in practise, leading to organisations struggling to balance values with profitability (Ayling and Chapman, 2022).

2.2 Policy formulation

This chapter introduces policy, AI-policy and AI governance, before delving into the key drivers and motivations behind them.

According to Aphale et al. (2013) policies are created to steer and control the actions of participants within a system. They further define policies as the representations of ideal behaviour. Schneider et al. (2023), on the other hand, extends this definition beyond rules and guidelines to also encompass the main objectives, responsibilities, roles and accountability.

A policy is not a fixed object that emerges spontaneously, but rather the result of considerable efforts invested in its creation (Flowerday & Tuyikeze, 2016). Policy creation is, moreover, only one stage of the policy lifecycle, according to Flowerday and Tuyikeze (2016). They argue that the full cycle also includes, among others, “Policy Construction, Policy Implementation, Policy Compliance and Policy Monitoring” (Flowerday & Tuyikeze, 2016, p. 173). Since policy creation is where the key drivers and motivations are most evident, this study will primarily focus on the creation and maintenance stages. While policies govern overall desired behaviour, AI-policies delves more specifically into the behaviors and ethical considerations that must be addressed when managing the use of AI (Mäntymäki et al., 2023).

The inclusion of policy in IS research is important, since information technology plays a significant role in many organisations and simultaneously a large proportion of those systems are shaped by policies (King & Kraemer, 2019). This chapter therefore examines AI-policy and AI governance along with their key drivers.

2.2.1 AI governance

The adoption of AI has the potential to become a transformative force in organisations and is reshaping the organisational landscape across many industries. However, it demands considerable efforts and capital investments, often under conditions of uncertainty, while also generating compliance and reputational risks (Van Giffen & Ludwig, 2023).

In light of these challenges, AI governance is becoming increasingly important. AI governance refers to sociotechnical phenomena encompassing the rules, processes and tools an organisation employs to ensure that its AI use is in line with its organisation's broader strategic goals and values. AI governance should thus ensure that AI use aligns with the ethical, legal and strategic frameworks the organisation operates within. In other words, it should mitigate

the risks while maximising the value derived from AI initiatives (Mäntymäki et al., 2022; Schneider et al., 2023)

The current state of AI governance faces challenges that potentially hinder the translation of ethical principles and objectives into actionable praxis. This is described as a translation problem, where principles remain as representations of abstract ethical considerations rather than being translated into actual implementation (Minkkinen & Mäntymäki, 2023). Butcher and Beridze (2019) further argue that the global AI governance is a disorganised field; Mäntymäki et al. (2022) build upon this and contend that such disorganisation also extends to the governance of AI within individual organisations. In earlier literature, AI governance was viewed as an isolated phenomenon, however, Mäntymäki et al. (2022) argue that it exists within a broader and more complex landscape of organisational governance, encompassing areas such as “IT governance”, “corporate governance” and “data governance”. By treating AI governance as an isolated phenomenon, research overlooks the broader organisational factors that shape it (Mäntymäki et al., 2022). Consistent with this, in 2019, 75% of managers acknowledged the importance of implementing AI effectively to avoid becoming obsolete, yet by 2023, a majority of executives were still struggling to effectively govern AI and clarifying their own role in ensuring responsible and efficient use (Van Giffen & Ludwig, 2023).

A further issue in AI governance is the regulatory vacuum that arises when national and international policies fail to keep the pace with technological evolution. This results in an expanded grey area in which the organisations themselves are responsible for defining and developing ethical AI. This creates situations in which organisations, both those who develop AI and those that utilize it, are confronted with the value-alignment problem and must balance profitability with ethical considerations in their AI solutions. Further exacerbating this are competitive pressures that may lead to the prioritisation of speed and resource efficiency, overriding ethical concerns. (Fioravante, 2024)

2.2.2 Policy drivers and motivations

As previously discussed, AI-policies are shaped by a variety of factors. To unpack the different logics behind this, the pressures on organisation can be categorised into three layers: the *environmental*, *organisational* and *systems layer* (Mäntymäki et al., 2023).

Environmental:

According to Mäntymäki et al. (2023) the environmental layer contains the contextual environment surrounding the organisation, including regulations or pressures from actors beyond the organisation's control. Mäntymäki et al. (2023) names the three aspects of the environmental layer as normative regulation, for example hard law such as GDPR, self-regulation, for example principles and guidelines, and lastly stakeholder pressure. By adopting responsible AI governance, organisations can enhance their legitimacy amongst society while also improving productivity, and employee loyalty (Papagiannidis et al., 2025). Papagiannidis et al. (2025) also argue that developing and using responsible AI practices can improve an organisation's reputation.

Organisational:

The organisational layer is defined by Mäntymäki et al. (2023) as the alignment of AI governance with both the strategic and the ethical values of the organisations. This layer is therefore divided into two components: *Strategic Alignment* and *Value Alignment*. Strategic alignment refers to AI implementations that align with the broader strategic goals of the organisation. To this end, it is often necessary to establish clear expectations regarding AI use. This includes defining the direction of AI use, that is defining the *what*, the *why* and the *how*, while also ensuring the necessary resource allocation, internal capabilities, and processes are in place (Mäntymäki et al., 2023). In this regard, Daugherty and Wilson (2018) highlight the need for training and retraining staff to ensure that the necessary skills are present within organisations. Furthermore, Shneiderman (2020) emphasises the importance of leadership commitment, both in ensuring that such training initiatives are in place and in managing AI-related incidents. *Value alignment* refers to ensuring that AI use aligns with the organisation's ethical values and principles. It requires the organisation to establish a clear understanding of which risks, including regulatory and reputational risk, are deemed acceptable. Mäntymäki et al. (2023) state that it is especially important to ensure clarity regarding which risks are acceptable, as ethical considerations are rarely clear-cut and often require organisations to navigate complex trade-off between strategic and ethical values (Mäntymäki et al., 2023).

Flowerday and Tuyikeze (2016) also argue that, in order to achieve a functional balance between value alignment and strategic alignment, all the stages of the policy lifecycle have to be taken into account. A particularly important stage is policy creation, during which the focus of the policy, that is what the policy should aim to address, is established. Flowerday and Tuyikeze (2016) emphasise that this process should involve a broad range of stakeholders including both employees and senior management, in order to balance strategic and value alignment while also fostering a sense of inclusion and ownership. Such inclusion not only provides a more comprehensive overview of the objectives but also ensures the subsequent development and implementation phases are supported and accepted. (Flowerday & Tuyikeze, 2016) In particular, employee involvement is essential as the employees interact with the system and IT landscape on a daily basis and are often the group most directly affected by the policy (Flowerday & Tuyikeze, 2016).

Systems:

According to Mäntymäki et al. (2023) the Systems layer is defined as the layer that combines the internal plans from the organisational layer and the external demands from the environmental layer, translating them into practise. Mäntymäki et al. (2023) highlight that, since governance follows the guidelines set by the environmental layer, it must remain adaptable in order to stay relevant to the organisation's needs.

2.2.3 Summary

AI-policies are created in order to steer and guide behavior within an organisation, this is supported by Aphale et al. (2013) who suggests that policies are made to serve this very function, and should also serve as guidelines for understanding the goals, roles and responsibilities involved in directing the organisation (Schneider et al. 2023). AI-Policies are, in turn, a component of AI governance, which ensures that organisations use AI in a manner that aligns with the ethical, legal and strategic frameworks (Mäntymäki et al., 2022; Schneider et al., 2023). In light of this, Mäntymäki et al. (2023) identified three layers, environmental, organisational and systems layers, that impact the AI-governance within the organisations. However, a key

challenge remains in translating and aligning these values into actionable practise (Minkkinen & Mäntymäki, 2023).

2.3 Policy in practise

This chapter examines policy in practise, exploring how policy impacts the organisations and why it may differ from the intended outcome.

How policy is translated into practise is shaped by a range of organisational and environmental aspects and can give cause to a phenomenon called decoupling (Bromley & Powell, 2012; Jabbouri et al., 2019; Horneber, 2025). Thus Flowerday & Tuyikeze (2016) highlights the need for policy creation to take into account these different aspects to ensure the policy is not redundant, irrelevant nor constructed in a manner that inhibits it from supporting and being supported by the users.

As previously mentioned, Flowerday and Tuyikeze (2016) highlight the need for stakeholder engagement. Employee and management support are therefore critical for successful policy implementation. This includes involving employees in the policy creation and enabling them to take ownership throughout the full policy lifecycle rather than solely during the policy commitment stage. This fosters a sense of inclusion among employees and, in turn, makes the transition easier to implement, as a sense of ownership is established (Flowerday & Tuyikeze, 2016). Flowerday and Tuyikeze (2016) also identify management support as crucial for ensuring policy awareness and policy education.

“The committed participation of management will motivate employees to comply with the new policy requirements and will also help to increase employees’ acceptance and commitment to the organisation’s security policy.” (Flowerday & Tuyikeze, 2016 p. 176)

Understanding the context of a policy is increasingly important as more organizations face dilemmas in which they must navigate trade-off between ethics and profitability. In light of this, *ethics washing* is increasingly being discussed both in society and in academia, referring to situations where organisations use ethics as a marketing tool, projecting an image of responsibility without grounding it in reality (de Laat, 2021; Fioravante, 2024). Within the AI field, this is particularly challenging as its principles remain relatively new and lack general consensus on norms and best practise (de Laat, 2021). Furthermore, there is limited formal responsibility and accountability, as well as lack of adequate tools to enable ethical AI (de Laat, 2021). These factors ultimately lead to unclarity, making it difficult to translate ethical AI principles into meaningful impact (de Laat, 2021).

2.3.1 Policy-practise decoupling

When examining decoupling as a phenomenon, it is valuable to distinguish between it from two distinct levels: *Policy-practise Decoupling* and *Means-Ends Decoupling* (Bromley & Powell, 2012). Bromley and Powell (2012) argue that policy-practise decoupling seeks to explain why organisations’ implementations of formal policy and rules fail. Policy-practise decoupling occurs when formal rules are violated or remain unimplemented, and when mechanisms for control and evaluation of governance remain largely symbolic (Bromley & Powell,

2012; Horneber, 2025; Jabbouri et al., 2019). This type of decoupling may occur when policies are adopted solely to comply with public expectations and governmental pressures, yet never translate into actual organisational processes, rendering the policy merely symbolic (Acquisti & Steed, 2025; Jabbouri et al., 2019).

The use of policy as a symbol may result from organisations adopting ethical or responsible AI-policies as a means of gaining competitive advantage or enhancing their public image, without any genuine intention or mechanism for actual implementation (de Laat, 2021; Horneber, 2025). One risk factor is limited capacity to implement the policy, as well as a lack of policy reinforcement (Horneber, 2025; Bromley & Powell, 2012; Jabbouri et al., 2019). Regarding formal accountability measures, organisations are subject to hard laws such as the GDPR. However, Horneber (2025) argues that a lack of regulatory accountability principles persists. A potential reason for this is persistent gap between law and practise, where high-level legal principles are difficult to implement, and disagreements between researchers and policymakers remain on fundamental questions such as what constitutes AI (Rakova et al., 2021). Bromley and Powell (2012) also argue that weak legal reinforcement creates a risk of policy-practise decoupling.

2.3.2 *Means-ends decoupling*

Means-ends decoupling occurs when a policy is implemented successfully but fails to fulfil its intended purpose (Bromley & Powell, 2012; Horneber, 2025; Jabbouri et al., 2019). This stems from a weak relationship between the policies implemented and the organisation's core processes. Consequently, the rules are enacted, but their outcomes remain uncertain (Bromley & Powell, 2012). Jabbouri et al. (2019) argue that, despite significant efforts to ensure policy implementation and avoid policy-practise decoupling, organisations still struggle to achieve their intended goals through policy, highlighting the need to examine this challenge through the lens of means-ends decoupling. They further argue that, as control and monitoring have increased within audit culture, the environment that surrounds and shapes organisations is fundamentally changing, therefore affecting the way policy and practise is coupled.

Horneber (2025) highlights that means-ends decoupling occurs when organisations face multiple conflicting institutional pressures, and to comply with the pressures, policies are implemented symbolically but without their intended effect. He further states that this is especially prevalent in complex fields such as in AI, where regulatory pressures are highly intricate and uncertain, and where it is difficult to measure and compare actual outcomes against desired ones.

2.3.3 *Handling of future AI practises*

Looking ahead, given that AI has already transformed the way organisations operate, as well as the individuals within them, it is unsurprising that Holmström (2024) anticipates further change and considers organisations capacity to adapt to it increasingly crucial.

Managing this change will therefore play an important role in shaping future organisational practises. Van Giffen and Ludwig (2023) emphasise that high-level AI governance should not be viewed as a one-time project, but rather as an ongoing commitment. Van Giffen and Ludwig (2023) further emphasise that AI governance requires continuous learning and adaptation.

Furthermore, deeper insights are needed into how AI technologies should be monitored and controlled. Encompassing both how AI usage should be structured and who should be held accountable when AI is used in a way that is regarded as inappropriate (Rivera et al., 2025). Given the inherent complexity and opacity of AI models, balancing and measuring all of these aspects will remain a significant challenge (Rivera et al., 2025).

2.3.4 Summary

As a result of the complex organisational and external environment, policies face the risk of decoupling, meaning that the actual outcomes of policy implementation diverge from the intended outcome. Decoupling can occur on two levels: policy-practise decoupling and means-ends decoupling. Policy-practise decoupling occurs when policies are implemented symbolically and fail to affect daily processes, whereas means-ends decoupling occurs when a policy is successfully implemented and adopted but fails to achieve its intended goal. (Bromley & Powell, 2012; Jabbouri et al., 2019; Horneber, 2025) This highlights the persistent problem caused by the gap between abstract ideals and actual implementation in the face of complex internal and external pressures (Rivera et al., 2025).

In light of the pressures organisations face when developing and enacting AI-policy, understanding the forces that drive these processes becomes increasingly important. One way of approaching this issue is through an institutional perspective. Through the lens of institutional theory, policies are a means by which organisations navigate external pressures, including cultural norms, cognitive expectations and broader societal pressures (Powell & DiMaggio, 1991). Given that AI governance, and thus AI-policy, is shaped by a mix of regulatory demands, norms and organisational pressures (Mäntymäki et al. 2023), institutional theory provides a fitting lens to examine these. Specifically, this study, draws on the concepts of isomorphism, institutional pressure, and decoupling which will be further explained in Chapter 2.4 Institutional theory. The study will therefore explore how institutional pressures shape AI-policy formulation and practise. This framework will serve as an analytical and an interpretive lens to guide the discussion rather than as a theoretical framework to be empirically tested.

2.4 Institutional Theory

Classical institutionalization concerns changes in practises that are influenced by values beyond the specifications of the current tasks (Powell & Bromley, 2015). Neo-Institutional theory, on the other hand, has evolved and built upon the concepts from the classical institutional theory (Powell & Bromley, 2015). Defined by Powell and DiMaggio (1991, p. 8) as:

“The new institutionalism in organization theory and sociology comprises a rejection of rational-actor models, an interest in institutions as independent variables, a turn toward cognitive and cultural explanations, and an interest in properties of supraindividual units of analysis that cannot be reduced to aggregations or direct consequences of individuals’ attributes or motives”

As reflected in the Powell and DiMaggio (1991) definition, Neo-institutional theory shifts focus toward cognitive and cultural societal expectations and pressures, rather than concentrating solely on historically grounded, organizational-specific power dynamics. This is further

supported by Powell and Bromley (2015) who argue that Neo-institutional theory examines how external social pressures affect an organisation's structure and policies, and how this in turn impacts the organisation's perceived legitimacy. They further argue that Neo-institutional theory emphasises two key concepts: *isomorphism*, which concerns the process of organisations becoming homogenous since they are responding to similar pressures (DiMaggio & Powell, 1983); and *decoupling*, which refers to the misalignment between an organisation's technical and institutional level (Orton & Weick, 1990). These two concepts will be elaborated in sections 2.4.2 *Isomorphism* and 2.4.3 *Decoupling* below.

Horneber (2025) further highlights that, according to Neo-institutional theory, regulative, normative or cultural aspects of society can have an effect on organisations according to the Neo-institutional theory. He suggests that, through the lens of Neo-institutional theory, organisations must take these pressures into consideration in order to thrive and ultimately survive. Although complying with these pressures can earn an organisation legitimacy, he notes doing so risks conflicting with the organisation's efficiency goals. At the same time, he notes that the institutional pressures can themselves even be inconsistent and misaligned.

The regulatory aspect, according to Horneber (2025), typically concerns policies and government legislation, while the normative aspect concerns societal values and norms, each carrying formal and informal consequences if breached. The cultural aspect, according to Horneber (2025), concerns the shared beliefs and societal reality that shapes institutional pressures and allows organisations to gain legitimacy without altering their actual work processes, for example, by formulating formal policies as a response to the institutional pressures.

As demonstrated in earlier studies, policy-practise decoupling is likely to occur when a policy conflicts with the organisation's internal goals and values, or when the organisation does not invest sufficient time or resources for the actual implementation (Bromley & Powell, 2012). According to Powell and Bromley (2015), organisations may also adopt a policy to protect the organisation's technical core, while still maintaining their decoupled practises, thereby satisfying external pressures without altering internal processes.

2.4.1 Strong and weak institutional pressure

When applying an institutional perspective to organisations, two overarching categories of pressure can be identified: strong institutional pressures and weak institutional pressures (Berente et al., 2019; Minkkinen & Mäntymäki, 2023). According to Minkkinen and Mäntymäki (2023), strong institutional pressure can be defined as involuntary compliance pressures that influence an organisation through, for example, regulations that carry substantial consequences if not followed. Weak institutional pressures, on the other hand, are according to Minkkinen and Mäntymäki (2023) defined as voluntary adoption pressures, where failing to comply does not produce consequences that create a tangible impact. Berente et al. (2019) extends this definition by further noting that weak institutional pressure affords stakeholders the freedom to adopt if they choose to do so, without incurring significant negative consequences should they choose not to.

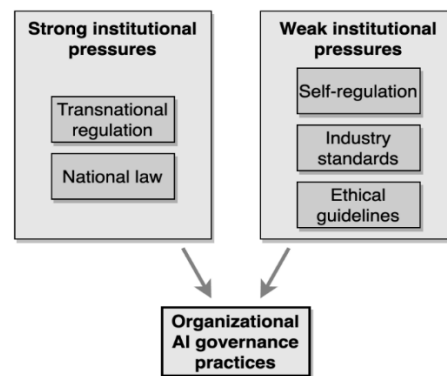


Figure 2.1: Strong and weak institutional pressures (Minkkinen & Mäntymäki, 2023)

2.4.2 Isomorphism

The Neo-institutional theory places a considerable emphasis on Isomorphism (Powell & Bromley, 2015). Isomorphism is the process in which organisations become increasingly homogenous as they respond to pressures from the same environmental context (DiMaggio & Powell, 1983). Isomorphism can, in turn, be divided into three subcategories: *coercive*, *mimetic* and *normative* isomorphism (DiMaggio & Powell, 1983). Each subcategory concerns a distinct type of influence and produces different effects (DiMaggio & Powell, 1983).

The first category is coercive isomorphism, which is defined as: “*both formal and informal pressures exerted on organizations by other organizations upon which they are dependent and by cultural expectations in the society within which organizations function.*” (DiMaggio & Powell, 1983, p. 150) According to DiMaggio and Powell (1983), these forces can take the form in various external pressures, such as invitations to organisational agreements or as a consequence of government regulations. Even though new regulations pressures organisations to change and those changes are initially ceremonial, DiMaggio and Powell (1983) point out that they can make a meaningful difference in the long term.

The second category, mimetic isomorphism, is defined by DiMaggio and Powell (1983) as the actions taken by organisation when faced with uncertainty. They further define mimetic actions of organisations as a process of modelling. In which, when organisations are faced with uncertainty and are seeking greater legitimacy, they look to other similar organisations that they consider superior or more legitimate in order to emulate their practises and advance themselves. One notable advantage they highlight is that modelling another organisation is relatively cost-effective, given the comparatively sound solutions it yields. Although modelling does not always occur intentionally, DiMaggio and Powell (1983) argue that it can take place through consultancy firms or employees who bring their ways of working from a previous organisation into a new one.

The third category, normative isomorphism, is defined by the notion that change stems from professionalization (DiMaggio & Powell, 1983). Professionalization is defined by DiMaggio and Powell (1983) as:

“the collective struggle of members of an occupation to define the conditions and methods of their work ... and to establish a cognitive base and legitimation for their occupational autonomy” DiMaggio and Powell (1983 p. 152).

They further argue that professionalization can be divided into two subcategories: education and legitimation, and professional network. The first subcategory emphasises the impact of formal education, while the second focuses on the professional networks that span across organizations through which modelling can be transferred (DiMaggio & Powell, 1983).

2.4.3 Decoupling

Loose coupling is a component of decoupling and refers to the degree of connectedness between elements within an organization (Orton & Weick, 1990). According to Orton and Weick (1990), it reflects the degree of alignment between the technical level, where external pressures are least prominent, and on the institutional level, where they are most prominent. They argue that tight coupling at the technical level can enable stability, while loose coupling on an institutional level can enable flexibility. One example of loose coupling can be found in a study conducted by Berente et al. (2019). He found that nurses experienced loose coupling during their daily work with the health record system, prioritising the care they deemed necessary for patients and only afterwards fulfilling the requirement of the information system. He emphasised that, since the nurses prioritised the patient care quality, they worked around the system and only used it retrospectively, not because they found it helpful but merely to satisfy its formal requirements.

As discussed in chapter 2.4.1 *Strong and weak institutional pressure*, strong institutional pressure arises from external factors such as regulations, while weak institutional pressures take the form of industry standards or guidelines (Minkinen & Mäntymäki, 2023).

		Institutional Logics	
		Congruent	Incongruent
Pressure to Comply	Strong	Faithful Appropriation	Loose Coupling
	Weak	Co-optation	Resistance

Figure 2.2: Institutional logics and pressures to comply (Berente et al., 2019)

Other than this, as seen in Figure 2.2 Institutional logics and pressures to comply, Institutional logics also play a role within Neo-institutional theory (Berente et al., 2019). According to Berente et al. (2019), institutional logics can be divided into two main categories of institutional logics: congruent and incongruent institutional logics. Congruent institutional logics are according to Berente et al. (2019), characterised by the logics of change, for example the introduction of a new system, being aligned with the values and logics of the organisation and the employees involved. He argues that the presence of congruent institutional logics increases the likelihood that new systems or operations will be used as intended. With incongruent institutional logics, on the other hand, he argues that, the misalignment between the new system and the organisations inherent logics will risk causing loose coupling, where the new system will not be used as intended or force a fundamental change in the nature of the organisations operations. Loose coupling is therefore a consequence of incongruent institutional logics and strong institutional pressure, as illustrated in Figure 2.2 Institutional logics and pressures to comply (Berente et al., 2019).

2.5 Literature matrix

The aim of this literature matrix is to provide a clear overview of the most important insights relevant to this study. Table 2.1 Literature Matrix provides an overview of the key themes this study follows, a summary of their respective main insights and the literature chapters and sources connected to those insights. This matrix is particularly relevant to this study as it reflects the intent and structure of the thesis established by the authors. These themes collectively form the frame of reference drawn upon in the empirical analysis.

Table 2.1: Literature matrix

Theme	Summary	Chapter	Sources
AI in Organisations	AI use is rapidly diffusing in organisations across sectors and is actively redefining the organisational landscape. This diffusion has sparked discussion within organisations, the legal field, and research regarding the risks associated with AI and Responsible AI.	<i>2.1 AI in organisations</i>	(Berente et al., 2021) (Rivera et al., 2025) (Dignum, 2019) (Horneber, 2025) (Conboy et al., 2022) (Ayling & Chapman 2022) (Mäntymäki et al., 2023) (de Laat, 2021) (Gianni et al., 2022) (Papagiannidis et al., 2025) (Vassilakopoulou et al., 2022) (Fioravante, 2024) (Van Giffen & Ludwig, 2023)
Policy Formulation	Policies are tools to steer behavior within organisations and act as guidelines for employees. There are several external and internal factors that determine how organisations formulate their policies. The pressures organisations experience can be split into three layers: environmental, organisational and systems layer.	<i>2.2 Policy Formulation</i>	(Van Giffen & Ludwig, 2023) (Berente et al., 2019) (Minkinen & Mäntymäki, 2023) (Mäntymäki et al., 2022) (Mäntymäki et al., 2023) (Papagiannidis et al., 2025) (Horneber, 2025) (V. Flowerday & Tuyikeze, 2016)

			(Aphale et al., 2013) (Schneider et al., 2023) (King & Kraemer, 2019) (Butcher & Beridze, 2019) (Daugherty & Wilson, 2018) (Shneiderman, 2020) (Fioravante, 2024)
Policy in practise	The practical implementation of policy is dependent on several environmental and organisational aspects. These aspects could impact the policy and create situations of policy-practise or means-ends decoupling, where the policy does not serve its intended purpose.	<i>2.3 Policy in practise</i>	(Bromley & Powell, 2012) (Jabbouri et al., 2019) (Horneber, 2025) (Flowerday & Tuyikeze, 2016) (de Laat, 2021) (Fioravante, 2024) (Acquisti & Steed, 2025) (Horneber, 2025) (Rakova et al., 2021) (Holmström, 2024) (Van Giffen & Ludwig, 2023) (Rivera et al., 2025)
Institutional theory	The chosen framework, The Institutional Theory, consists of three types of isomorphisms: coercive, mimetic and normative isomorphism. Institutional theory further highlights how institutional logics and pressure affects the coupling of policies.	<i>2.4 Institutional Theory</i>	(Powell & Bromley, 2015) (Powell & DiMaggio, 1991) (Bromley & Powell, 2012) (DiMaggio & Powell, 1983) (Horneber, 2025) (Orton & Weick, 1990) (Minkkinen & Mäntymäki, 2023) (Berente et al., 2019)

3 Methodology

In this chapter the research methodology will be presented, this chapter will present the research philosophy and approach as well as an explanation and motivation of the methodological selections. This documentation is included to enhance traceability and trustworthiness (Bryman, 2018) and allow for scrutiny of potential biases unconsciously included by the authors.

3.1 Research philosophy

The research philosophy was selected to fit the research purpose, Saunders et al. (2009) state that the philosophy selection shapes how knowledge is created for the study. This study adopted an interpretivist research philosophy, since this study's research question explores subjective stakeholder perspectives, motivations, and socio-technical dynamics surrounding AI-policy formulation and impact. Interpretivism emerged as a critique of positivism, rejecting the notion of a single, objective social reality (Bhattacharjee, 2019). Instead, it views reality as subjective and being actively shaped by individuals' experiences and their social contexts (Bhattacharjee, 2019). Saunders et al. (2009) explain that interpretivism is essential because humans cannot be studied like physical phenomena can. They further argue that social reality is shaped by cultural contexts and the unique meanings individuals create, making objective, natural-science approaches inadequate for understanding these complexities.

This perspective resonates strongly with this study's objectives. This research aims to interpret a specific socio-technical phenomenon, AI-policy, that hinges on each organization's unique social realities and cultural contexts. Rather than seeking universal laws about “how an AI-policy should be”, this study prioritizes the understanding of how policies emerge and play out in practise. This study seeks to understand the dynamics of AI-policy formulation and its real-world implementation. According to Saunders et al. (2009), the interpretivist paradigm is particularly valuable here, as it examines not only rational strategies but also irrationalities and unintended consequences, such as policy-practise decoupling or organizational dynamics that diverge from the stated goals. Goldkuhl (2012) further mentions that in the IS field, interpretivism is particularly valuable as it often deals with complex socio-technical phenomena that might be difficult to quantify.

As Kaplan & Maxwell (2005) further emphasizes, it is important to understand that technology does not exist in a vacuum and is embedded in the social and organizational context in which they exist. This means that they view the context as an inherent part of the technological phenomenon and technology can therefore not be studied without understanding the context. For this study, this means that AI and AI-policy cannot be studied without understanding the organisational context. This makes interpretivism philosophy valuable for understanding the social, cultural, and technological phenomena. This will be particularly valuable as this study strives to understand the greater context in which AI fits within an organisation, and more specifically how AI-policy is formulated and shapes the use of the AI technologies. The process for AI-policy is often subjective and highly dependent on the subjective and contextual processes. This interview-based study aims to capture these nuances through the lens of interpretivism, in order to avoid reducing the complexity through quantified generalizations.

This study has adopted a qualitative research methodology, which aligns with interpretivist knowledge creation by enabling the analysis of complex, rich data. This approach is commonly used in interpretivist studies (Goldkuhl, 2012; Saunders et al., 2009). Its specific implementation and further motivation are elaborated in Section 3.4 *Data Collection*.

3.2 Research approach

To address the research question and examine how internal AI-policies are formulated and how they translate into practise, this study has adopted a qualitative research approach. The goal of the qualitative approach is described as understanding complex situations or issues through the investigation of people's behaviours and perspectives within the specific context (Kaplan & Maxwell, 2005). This makes the approach suitable for gaining an understanding of how organisations understand their realities and unique motivators when creating and maintaining their AI-policy. The qualitative approach is also valuable in gaining an understanding of the different perceptions and experiences of stakeholders (Kaplan & Maxwell, 2005). A major strength of these methods lies in the ability to understand the broader social, cultural, organizational, and political context surrounding information technology (Kaplan & Maxwell, 2005). This approach is particularly suitable since the research focus is on understanding how different organisations perceive, interpret, and enact AI-policy in their everyday work practises. Rather than measuring predefined variables, the objective is to capture complex meanings, tensions, and experiences related to the AI governance in their respective contexts.

Patton (2015) emphasizes that a core strength of qualitative inquiry lies in examining the consequences of systems within their contexts, particularly how intended goals often diverge from its actual implementation. This insight is highly relevant for understanding AI-policy, where formal policies may articulate an ambition around AI, yet the everyday realities can lead to very different outcomes. A qualitative approach was therefore well suited to exploring AI-policy enactment across organizational stakeholders and to understanding how such policies are reflected, negotiated, and realized within organizations.

3.3 Literature collection and review

When searching for literature, library sources such as Finn, Springer, AIS eLibrary, ResearchGate and sources included in "Basket of 11" were prioritized. Furthermore, when possible, sources written after the year 2020 were prioritized in order to access the most current and relevant information possible. Despite this ambition, there are exceptions, especially concerning the sources used for the theoretical framework, where the original sources were preferred in order to access the original definitions, regardless of how old they were. Examples of keywords or phrases used during the literature synthesizing are:

- "Policy"
- "AI"
- "AI-policy"
- "Decoupling"
- "Policy practise decoupling"
- "Loose coupling"

- “AI in organisations”
- “Policy impact factors”
- “AI governance”
- “Institutional theory”
- “Responsible AI”
- “Policy regulations”
- “Policy Analysis”

During the literature search a technique called backwards snowballing was used. Backwards snowballing is the act of using an article for its information while also “digging deeper” into the sources the article has used and using those as well (Alvehus et al., 2025).

Below is the synthesized research model, compiling the themes and its relevant literature. The function of this matrix is to show the synthesis preparation for the literature collection and connection to the subsequent subthemes used in the rest of the study. Furthermore, this was used as a base for the interview guide, in order to anchor it in the literature. A more detailed connection between the interview guide and literature review can be found in Table 3.1 Synthesized research model.

Table 3.1: Synthesized research model

Main theme	Sub themes	Literature
<i>AI in organisation</i>	Overview of AI in organisations	Dignum (2019) Berente et al. (2021) Conboy et al., (2022)
	AI adoption	Ayling & Chapman (2022) Mäntymäki et al. (2023)
	Responsible AI	de Laat (2021) Gianni et al. (2022) Papagiannidis et al. (2025) Vasilakopoulou et al. (2022) Mäntymäki et al. (2023) Horneber, (2025) Ayling & Chapman (2022) Fioravante (2024) Rivera et al. (2025)
<i>Policy Formulation</i>	Overview of Policy formulation	Aphale et al. (2013) Schneider et al. (2023) Flowerday & Tuyikeze (2016) Mäntymäki et al. (2023) King & Kraemer (2019).
	AI governance	Van Giffen & Ludwig (2023) Mäntymäki et al. (2022)

		Schneider et al. (2023) Minkkinen & Mäntymäki (2023) Butcher & Beridze (2019) Mäntymäki et al. (2022) Fioravante (2024)
	Policy Drivers and motivations	Mäntymäki et al. (2023) Papagiannidis et al. (2025) Daugherty & Wilson (2018) Shneiderman (2020) Flowerday & Tuyikeze (2016)
Policy in Practise	Overview of Policy in Practise	Bromley & Powell (2012) Jabbouri et al. (2019) Horneber (2025) Flowerday & Tuyikeze (2016) de Laat (2021) Fioravante (2024)
	Policy-practise decoupling	Bromley & Powell (2012) Horneber (2025) Acquisti & Steed (2025) Jabbouri et al. (2019) de Laat (2021) Rakova et al. (2021)
	Means-ends decoupling	Bromley & Powell (2012) Horneber (2025) Jabbouri et al. (2019)
	Handling of Future Practises	Holmström (2024) Van Giffen & Ludwig (2023) Rivera et al. (2025)
Institutional Theory	Overview of Institutional Theory	Powell & Bromley (2015) Powell & DiMaggio (1991) Orton & Weick (1990) Horneber (2025) Bromley & Powell (2012)
	Strong and weak Institutional pressure	Berente et al. (2019) Minkkinen & Mäntymäki (2023)

	Isomorphism	Powell & Bromley (2015) DiMaggio & Powell (1983)
	Decoupling	Orton & Weick (1990) Berente et al. (2019) Minkkinen & Mäntymäki (2023)

3.4 Data collection

In this chapter the data collection methods will be presented and motivated.

3.4.1 Semi-structured interviews

As this study is relying on conducting qualitative research, conducting semi-structured interviews was assessed as the most appropriate approach. The goal of the interviews in qualitative research is to create rich data, which is characterized as detailed and comprehensive data that have the ability to explain a complex phenomenon or situation (Kaplan & Maxwell, 2005; Schultze & Avital, 2011). Furthermore, semi-structured interviews were used for this study, this entails, according to Bryman (2018), that interviews are based on a general interview guide with a few questions and themes. For this study, semi-structured interviews were chosen since it gives the interviewer more flexibility in adjusting the flow of the questions as well as offering the opportunity for follow-up questions (Bryman, 2018; Schultze & Avital, 2011). According to Kaplan & Maxwell (2005) this interview method combines the consistency of a structured approach with the flexibility to expand on details when needed. It is targeted and insightful, while also allowing researchers to clarify confusion in real-time. Furthermore, according to Kaplan and Maxwell (2005) it ensures coverage of core themes across interviews while still allowing for exploration of context-specific depth and nuance. Therefore, this approach was deemed well adapted for exploring AI-policy formulation and enactment.

Interviews have been conducted with key stakeholders who are involved with the organization's AI and AI-policy initiatives. The interviews were conducted through video calls, both through Microsoft Teams and Google Meet. To ensure comfort, the informants were allowed to choose which video-call service they preferred. Video interviews were conducted instead of in-person interviews as the time and resources were limited, and the informants were located in a wide range of locations within the Nordic Region. Bryman (2018) notes that remote interviews are less time-consuming than in-person meetings though they may limit the researcher's ability to capture non-verbal cues and facial expressions that according to Bhattacharjee (2019) enriches the data collection. However, Bryman (2018) further emphasizes that video interviews largely preserve this non-verbal contact. Additionally, conducting interviews through a video format allows participants to select their own comfortable environment, which according to Bryman (2018) enhances their sense of comfort during the conversation.

Conducting interviews presents several methodological challenges that must be carefully managed (Recker, 2021). One such risk is reflexivity, where the informants experience

pressure to respond in a manner that the informant feels the interviewer wants, this is particularly common in cases where subjectivity and bias is shown through poorly constructed interview questions (Recker, 2021), in the following section 3.4.2 *Interview Guide* the mitigation of this risk is discussed. Further ethical considerations surrounding the interview, including informed consent, anonymity, and data protection are elaborated further in section 3.7 *Ethics*. Despite the challenges surrounding conducting interviews, the method choice still seems appropriate because of the aforementioned potential for gaining rich and comprehensive insights (Kaplan & Maxwell, 2005; Schultze & Avital, 2011).

The interviews took about 30-50 minutes, and each interview ended with an open question where the informant had the opportunity to supplement and add information if there was something they had reflected on during the interview. In one interview the contacted informant, I6a, wanted to have their Head of AI, I6b, along in the interview to be able to give relevant information.

3.4.2 Interview Guide

The interview guide was divided into three main sections modelled after this study's literature review: 1. *AI in organisations*, 2. *Policy Formulation* and 3. *Policy in practise*, along with space for contextual opening and closing questions. Alvehus et al. (2025) states that by dividing the interview guide around themes, with questions on each, it enables flexibility in the interview as well as allowing the informants to discuss things not directly related to the question but the theme as a whole. Both Alvehus et al. (2025) and Jacobsen (2017) further stress the importance of having opening questions related to the interviewee as a chance for both understanding context but also for the informants to get comfortable with the interview. Therefore, each interview began with small talk, introducing both the researchers as students and the study as a whole before starting to record the interview.

The interview guide contained a total of 28 questions and could be described as a highly structured model with fully structured questions, this was to ensure that all informants had access to the questions before the interview. This highly structured interview guide could suggest it diverges from the openness ideals of qualitative research, however, Jacobsen (2017) states not structuring the data could result in complexities leading to the data analysis being difficult, resource intensive, and sometimes impossible to conduct. Furthermore, Jacobsen (2017) also states that there will always be some amount of bias and subconscious structuring and by stating this clearly these biases will be highlighted both to the authors and the readers. The interviews were still semi-structured meaning they left room for follow-up questions, asking for examples and clarifications when needed (Bryman, 2018; Schultze & Avital, 2011).

A primary concern when creating the interview guide was wording bias, where leading questions may pressure participants to respond in specific directions (Bryman, 2018; Recker, 2021). To mitigate this, interview questions were designed to avoid imposing researcher assumptions. While the literature review generated the general thematic areas, the questions remained open-ended and neutral to allow participants to steer the conversation themselves to reveal unanticipated perspectives. To further ensure this, the thesis supervisor reviewed the interview guide beforehand to help identify hidden wording bias.

The interview guide was chosen to be modelled on the theoretical background, as according to Bryman (2018) the goal of the interview is to seek the knowledge lacking in the current

literature. By having a theoretical background when creating the interview guide, it ensured that the questions were relevant and sought to fill the identified research gap. Jacobsen (2017) further emphasised that the interview guide should ensure that all relevant themes are highlighted but still leaving room for dialog. To ensure all themes were discussed a source mapping between each question and the relevant chapter in the literature was created.

Table 3.2: Source Mapping of Interview Guide and Literature review

Theme	Question	Literature Review Section
Opening Questions	Can you describe your professional role?	Contextual Question
	Which areas of responsibility do you have, connected to AI and policy	
AI in the organisation	What types of AI do you primarily use in your organisation?	<i>Sec. 2.1 AI in organisations</i>
	What is the purpose of this AI use?	<i>Sec. 2.1.1: AI adoption</i>
	Are there any specific tools you prefer or avoid in your organisation?	<i>Sec. 2.2.2: Responsible AI</i>
	Why do you choose to use the tools you are using?	<i>Sec. 2.2.2 Responsible AI</i> <i>Sec. 2.2.1: AI governance</i>
Policy Formulation	What is the primary goal of your AI-policy?	<i>Sec. 2.2.2: Policy Drivers and motivations</i>
	What principles guide your AI-policy?	<i>Sec. 2.2.2: Policy Drivers and motivations</i>
	What initiated the creation of your AI-policy?	<i>Sec. 2.2.2: Policy Drivers and motivations</i>
	When, in relation to the organisation's actual AI use, was the policy developed?	<i>Sec. 2.2.2: Policy Drivers and motivations</i> <i>Sec. 2.3.1 Policy-Practise Decoupling</i> <i>Sec. 2.3.2 means-ends decoupling</i>
	How did the AI-policy emerge and evolve?	<i>Sec. 2.2.2: Policy Drivers and motivations</i>
	Who was involved in the formulation of the AI-policy?	<i>Sec. 2.2.2: Policy Drivers and motivations</i>

		<i>Sec. 2.3.1 Policy-Practise Decoupling</i>
	How do you identify what should be included in the AI-policy?	<i>Sec. 2.2.2: Policy Drivers and motivations Sec. 2.1.2 Responsible AI</i>
	How do regulations and industry requirements shape your AI-policy development?	<i>Sec. 2.2.2 Policy Drivers and Motivation</i>
Policy in practise	How was the policy communicated when it was first implemented?	<i>Sec. 2.3 Policy in Practise</i>
	Is there a continuous communication of policy in the organisation?	<i>Sec. 2.3 Policy in Practise</i>
	Where is the policy? Is it easy to find and understand?	<i>Sec. 2.3.1 Policy-Practise Decoupling</i>
	Is there a separate policy for different departments, with specialized language and specifications?	<i>Sec. 2.3.1 Policy-Practise Decoupling</i>
	Are there any support tools to help employees how to comply with the policy?	<i>Sec. 2.3.1 Policy-Practise Decoupling</i>
	Who is responsible for policy compliance?	<i>Sec. 2.3.1 Policy-Practise Decoupling</i>
	What are the consequences of non-compliance with AI-policy? What does the follow-up process look like?	<i>Sec. 2.3.1 Policy-Practise Decoupling</i>
	In what ways does the AI-policy impact your daily work?	<i>Sec. 2.3.1 Policy-Practise Decoupling</i>
	Could you discuss challenges regarding the AI-policy?	<i>Sec. 2.3.1 Policy-Practise Decoupling Sec. 2.3.2 Means-Ends decoupling</i>
	Do you experience that the AI-policy achieves or does	<i>Sec. 2.3.1 Policy-Practise Decoupling</i>

	not achieve its purpose? And why?	<i>Sec. 2.3.2 Means-Ends decoupling</i>
Closing Questions	How does your organisation work with updating and maintaining relevance in the policy?	<i>2.3.3 Handling of future AI practises</i>
	When you look at your AI-policy in 1-2 years, do you expect it to change, and in that case how?	<i>2.3.3 Handling of future AI practises</i>
	Is there something you would like to add that we haven't touched on yet regarding this subject?	Contextual Question

3.4.3 Participant selection

For this study a purposive data collection method has been chosen. It was chosen since it allows the data collection to be done through gathering interviewees through the perceived information they might give, rather than through their specific title or department (Jacobsen, 2017). The purposive collection method was also applied to the choices made of the companies to contact. By purposely selecting companies to contact for interviews, the companies were not selected with the goal to find representative organizations for their industry, but to find companies that would have interesting and relevant information on this study's subject to offer (Jacobsen, 2017; Bryman, 2018). This ties into one goal of this study to capture a broad view spanning over multiple industries in order to not be limited by sector specific regulations.

The ideal interview participant for this study would be a person who is actively involved with the formulation of the AI-policy. Ideally the interviewee would be someone with the title "CAIO". Since AI is a relatively new area, companies may not have had the time to adopt that role, it was chosen to include AI-management and CIOs, CTOs and CSOs. Since there is also a restricted amount of time and resources available for this study, the scope of the working titles was broadened to more easily include people who were also deemed relevant to the study.

In order to find and contact the people to be interviewed, LinkedIn was used. Since this study uses a purposive sampling method it meant that people who had matching AI and policy management responsibilities present in their LinkedIn profile were contacted and checked that they were a fit for the participant goal profile. Search terms used on LinkedIn included:

- "CIO CTO AI-policy Norden"
- "AI-policy Nordics"

Some of the potential informants that were contacted turned out not to be the right fit for the subject and they were therefore asked if they would refer us to someone in their company that

would be willing to participate and would be a good fit, knowing the information that was sought after. Some informants were found on LinkedIn by searching for their title and were thereafter searched for online, or on their company website so that an email could be sent to them instead of a LinkedIn message. Emails were also sent directly to companies' career-email addresses and were thereafter directed internally by them to a person who would be knowledgeable in the subject of this study. After the participants agreed to be interviewed the interview guide along with the participants rights and consent form was sent out a few days ahead of the interview.

The number of interviews were established by the perceived amount of information gained, in order to achieve an appropriate amount of saturation. This is supported by Recker (2021) who states that saturation is gained by the amount of valuable information that can be obtained and not by the quantity of interviews being made. The amount of saturation was balanced against the time and resource limitations of the study as Jacobsen (2017) states that these perspectives must be taken into account when conducting interviews.

The informants were all made anonymous in order to preserve their privacy. The informants were also from seven different companies and sectors within the Nordic Region. Choosing interviewees from different industries was a conscious choice made in order to capture information across different sectors and not be constrained by branch specific regulations. Furthermore, the study was partly focused on how rules and regulations impact organisations and therefore it was crucial to not only investigate one industry, but rather many different industries, since they can have different regulatory constraints.

Table 3.3: Table of informants

Informant	Role	Interview length	Interview Date	Appendix
I1	Project Manager	36 minutes	10/4-2026	A
I2	Senior IT Manager	55 minutes	16/4-2026	B
I3	Operating Chief	31 minutes	16/4-2026	C
I4	CIO/CSO	47 minutes	17/4-2026	D
I5	Compliance Manager	51 minutes	17/4-2026	E
I6 a&b	CIO & Head of AI	44 minutes	21/4-2026	F
I7	CTO	32 minutes	22/4-2026	G

- *Informant 1*, is a Project Manager responsible for the organisation's national AI implementation. They work with the legal team, the DPO and the IT Manager concerning

the coordination of the AI and AI-policy. They were involved in the creation of the AI-policy and are currently also responsible for anchoring the AI-policy within the organisation.

- *Informant 2*, is a Senior IT Manager, and their role includes responsibility over the following departments: security, automation, IT infrastructure, IT support and IT compliance.
- *Informant 3*, is an Operating Chief, and their role includes a multitude of activities and responsibilities including organising the internal AI efforts. Alongside the IT-team they have created the organisation's AI-policy.
- *Informant 4*, is an CIO/CISO and is responsible for the information security within their organisation, in this security work the creation of an AI-policy is included.
- *Informant 5*, is a Compliance manager, and their role includes a share of the responsibility for their organisations policy, policy for responsible use of AI, and associated procedures and guidelines. They, together with other management, share the responsibility of the AI-policy.
- *Informant 6 a&b*, *I6a* is CIO and responsible of the IT-department and *I6b* is Head of AI which entails responsibility for the organisation AI strategy including the coordination and management of AI initiatives. They together share the responsibility for the AI-policy.
- *Informant 7*, is a CTO and is responsible for the technical leadership and the technology strategy. *I7* was responsible for formulating the first draft of the AI-policy.

When making the initial contact with the potential participants, both the private email and the school provided email addresses were used. The following email content was used, as seen in Figure 3.1. Example of initial pitch sent by email and via LinkedIn.

Hi!

Our names are Astrid and Fanny and we are studying Information Systems at Lund University. We are currently writing our master's thesis on how companies implement AI policies and how these are reflected within the organization.

We want to explore how companies manage AI, particularly how AI policies are developed and how they shape the use of AI in day-to-day operations. AI-policies and guidelines are a rapidly evolving area, and we aim to build an understanding of how companies navigate this. We would therefore love to hear about your experiences of putting AI-policy into practice.

Participation is of course entirely anonymous and voluntary, and the results will be used solely for academic purposes. You will also have the opportunity to review the finished study before we publish it.

Kinds regards,

Astrid Andersson & Fanny Hedin

Figure 3.1: Example of initial pitch sent by email and via LinkedIn

3.5 Data analysis and preparation process

In this chapter the method for the data analysis and preparation process will be described and motivated.

3.5.1 Transcription

All interviews were, with explicit consent, recorded using smartphones. Following Jacobsen's (2017) recommendations, interviews will be recorded rather than taking notes during the interviews, as he states that this improves conversational flow and allows for better eye contact. He describes that by using a tape recorder or a smartphone, the conversation can be thoroughly transcribed after, as this will help in the following data analysis. Alvehus (2023) further describes the transcription as the first step of the data analysis. Alvehus (2023) also mentions that recording the interview can make the participants feel that they are being accurately represented, as it will avoid the risk of misunderstanding and misrepresentation.

It was chosen to transcribe the recordings for the interview because transcription enables for easier analysis than only using audio-files and it is also a way to ensure transparency in the study, where others can review and evaluate the interpretations of the raw data (Jacobsen, 2017). The transcription was done with the tool Whisper. All the transcription through the tool was done locally and not in the cloud. After the transcription was done, the audio files were deleted. The interview recordings were recorded using an iPhone voice-memo app.

Whisper is a tool that analyses audio files, detects speech and produces a transcript based on the detected speech in the file. Whispers produced transcriptions aren't always accurate, so the transcripts were manually read and corrected afterwards, by listening to the actual audio files again. By transcribing this way it could be ensured that the transcriptions correctly captured the words and meaning of the interviewees' answers.

For the transcription *Intelligent Verbatim* was performed meaning the transcriptions are made to capture the essence of the spoken words but without including small utterances and filler words such as "um" or "mm" (McMullin, 2023). The intelligent verbatim method is used when such verbal cues are of lesser importance, and it creates opportunities for the written transcription to be more easily understood (McMullin, 2023). In the same manner for interviews carried out in Swedish and where quotes needed to be translated for section 4. *Findings* these quotes were translated to keep the intended meaning of the quotes rather than translating word for word.

3.5.2 Data storing

The data stored in the process of this study is audio-files, transcripts and emails collected before and after the interviews. These files and information were stored locally and subsequently deleted after they were not in use, in compliance with the "*Guidelines for the processing of personal data in student projects at Lund University School of Economics and Management*" (LUSEM, 2023).

3.5.3 Transcript validation

After the transcripts were completed, all of the interviewees received their transcripts through email and were therefore given the opportunity to correct or extract any possible information they did not feel represented their opinions or wanted to be used in the study. This was especially prioritised since the transcripts had been modified, in an *Intelligent Verbatim* manner, so that the participants could feel confident they had been correctly represented.

3.5.4 Data analysis

According to Kaplan & Maxwell (2005) coding can be done through deducing themes from the existing theory or deduced from the data itself. They highlight that the data is the most important aspect to be taken into account when deciding on what codes to utilize. Jacobsen (2017) highlights coding as a highly efficient method for working with large datasets, as it facilitates subsequent analysis by allowing the researcher to organize the material around themes rather than relying on extensive blocks of text. Using thematic coding is therefore well suited for this study because of the nature of semi-structured interviews producing large amounts of data (Patton, 2015). Since semi-structured interviews contain open ended questions, they present an opportunity for discovering new and unexpected patterns in the data (Kaplan & Maxwell, 2005). The goal of coding the data in themes was to simplify and make the analysis more efficient, while also improving the data by bringing it to a higher analytic level (Olsson, 2008).

Deductive coding, concerns testing existing theory by deriving logical consequences from general premises or known theories to evaluate hypotheses against new empirical data (Saunders et al., 2009). This study's coding has been done from a deductive standpoint, meaning it has been based on the theoretical themes based on the literature review (Guest et al., 2012). Deductive basis for coding is in this case the most appropriate analysis methodology since it allows for a comparable and solid theoretical backbone for the analysis (Bhattacharjee, 2019).

During the process of choosing the codes for the analysis, the three headings used for the literature chapter were considered. Later on, after considering the answers given by the interviewees a higher coding theme granularity was agreed upon. These original three themes emerged from the initial synthesized research model, which was the base for the literature synthesis. As the interview guide was based on the three main themes presented in the literature review, the precise connections between the literature chapter and interview question can be found above in Table 3.2: Source Mapping of Interview Guide and Literature review. Thereafter, it was deemed that also including the sub-themes for each of the three previously considered main themes, would ease the categorization and analyzation of information while still remaining sufficiently saturated. After conducting the coding, Table 3.5: Frequency of themes coded, was constructed to ensure the themes were adequately represented, in order to ensure that the level of saturation is sufficient in the findings.

These sub-themes were developed by both authors initially doing an individual mapping of potential sub-themes based on the literature. Afterwards, the potential sub-themes were discussed among the authors and the themes where both authors reached consensus were selected. Therefore, the codes used for this study's analysis are the following six:

- AI in use

- Purpose of Policy
- Policy Formulation Process
- Communication
- Compliance
- Future Handling

Below in Table 3.4 Coding, is an account of all the coding categories with their respective content and what chapter in the literature chapter they are connected to.

Table 3.4: Coding

Code Name	Abbreviation	Definition	Example from Data	Main Theme from literature	Literature Review Section
AI in use	AiU	How and what AI is used in the organisation	<i>“And for the rest of the organization, we try to support them with AI for things like increased collaboration and higher productivity.” (Appendix F #6)</i>	AI in organisation	<i>2.1 AI in organisation</i>
Purpose of Policy	PoP	Aim and purpose of the AI-policy	<i>“Yes, we see it both as a way to protect the company when it comes to ownership around the data and the source code. At the same time it is actually somewhat of a communication tool to encourage AI use as well.”(Appendix G #34)</i>	Policy Formulation	<i>2.2 Policy Formulation</i>
Policy Formulation Process	PFP	Who, when and what impacted AI-policy	<i>“So that was the starting point. To do a sort of, what is</i>	Policy Formulation	<i>2.2 Policy Formulation</i>

			<i>sometimes called a materiality analysis. A stakeholder analysis. And think about those who care about us. What questions do they have around AI? And how would we answer those questions?"</i> (Appendix C #23)		
Communication	COM	How is the policy communicated and implemented throughout the org.	<i>"We presented it at an information meeting and then sent it out by email."</i> (Appendix C #29)	Policy in practise	<i>2.3 Policy in practise</i>
Compliance	CMP	Responsibility of compliance and daily AI-policy impact	<i>"It's only one and a half A4 pages long. Because we were like, nobody reads these damn policy documents. So it has to be short."</i> (Appendix A #12)	Policy in practise	<i>2.3 Policy in practise</i>
Future Handling	FHa	Future view of AI and AI-policy and keeping it relevant	<i>"But the challenge remains to ensure that we have an AI-policy that reflects the technological development and how we want to</i>	Policy in practise	<i>2.3 Policy in practise</i>

			<i>work with the AI-policy” (Appendix C #47)</i>		
--	--	--	--	--	--

As seen below in Table 3.5: Frequency of themes coded the frequencies of the themes that were used as codes were calculated, where some of the themes were more frequently used than others. This can be explained by some codes being related to more questions than others.

Table 3.5: Frequency of themes coded

Theme	Count	Relative Frequency (%)
AI in use	30	~11
Purpose of Policy	60	~23
Policy Formulation Process	49	~19
Communication	41	~16
Compliance	57	~22
Future Handling	27	~10

The analytical method chosen for this study was thematic analysis, as this is a method of identifying and analysing different patterns within data (Braun & Clarke, 2006) this was deemed well suited to the nature of qualitative semi-structured interview data as it according to Patton (2015) tend to entail large data set and according to Kaplan & Maxwell, (2005) give opportunity for patterns to emerge. Within this method, thematic coding was used as a technique to systematically organise and categorise the data, serving as the foundation upon which the subsequent analysis and interpretation were later built (Braun & Clarke, 2006).

The first step of the research process, when the interviews were completed, was to transcribe and analyse the collected data (Kaplan & Maxwell, 2005). By conducting interviews, there had been a vast amount of generated data to have to be analysed (Patton, 2015). To be able to make sense of the collected data it was important to have scrutinized and coded it. Coding is a very useful technique of analyzing data in order to gain insights in the form of key ideas or concepts (Recker, 2021; Bhattacharjee, 2019).

The analysis culminated in interpreting the findings from the empirical data collection. This was done by systematically comparing the coded data across informants within each theme, identifying patterns and divergences (Braun & Clarke, 2006). These themes were then related back to the theoretical framework and literature review in order to draw analytical conclusions.

The data is presented in a thematic manner in section 4. *Findings* based on the three themes of the thesis 1. AI in organisations 2. Policy formulation 3. Policy in practise. The decision was made to present the data in themes rather than in individual interviewee sections as a way to synthesise and compare the data in the findings section. The quotes presented in 4. *Findings* were carefully translated in a manner that would preserve the meaning of the statement rather than simply directly translating the words. Each theme has one or multiple codes associated with it and these codes have served as the sub-headings for the findings section. The connection to the literary review section can be found in Table 3.4: Coding.

3.6 Research Quality

In qualitative research, reliability and validity concern the trustworthiness and credibility of the researcher's interpretation and methodological procedures, rather than statistical measures (Kirk & Miller, 2011). They further describe reliability as the degree of consistency in the research process, or as the extent of which similar findings can be generated under similar conditions. They describe validity, contrastingly, as the measurement of whether the study gives an accurate and credible representation of the topic under investigation.

Reliability is enforced by thoroughly documenting the analytical procedures and therefore enhancing transparency and traceability throughout the research process (Bryman, 2018). There are two perspectives of reliability often discussed, internal reliability concerns the replicability of the interviewee data, while external reliability concerns the replicability of the interviews themselves (Bryman, 2018). External reliability concerns the extent to which a study can be replicated, and in qualitative research, this is naturally challenging since social contexts and interactions are dynamic and cannot be reproduced under identical conditions (Bryman, 2018). Contrastingly, internal reliability means, for example, having replicability when it comes to whether members of the research team agree in how they interpret and code the data (Bryman, 2018). The coding of the content was done by both authors where each code was actively discussed to ensure agreement concerning the coding.

As for replicability, it is weakened by the fact that the number of interviews conducted is relatively low compared to quantitative interview studies (Kaplan & Maxwell, 2005). As a result of this, the study's internal reliability will be the primarily enforced of the two.

This study reinforced reliability through the use of a semi-structured interview guide, in order to ensure that all the interviews systematically address the same comparable themes. This ensures that the analysis of the interview data remains consistent and facilitates comparison between informants, while also enabling the informants the flexibility to expand on their answers if the situation arose.

In addition to reliability, validity must also be considered when assessing research quality. According to Bhattacharjee (2019), validity concerns how sound the conclusions drawn from a study are and whether the research truly measures or captures what it intends to explore. Bhattacharjee (2019) further differentiates between internal and external validity. Internal validity refers to the credibility of the causal relationships or interpretations developed within the study (Bhattacharjee, 2019). External validity, on the other hand, relates to the extent to which the findings of a study can be generalized beyond the specific research context (Bhattacharjee, 2019).

In this study, internal validity is strengthened through the systematic alignment between the research questions, data collection, and analytical procedures (Recker, 2021). The use of a semi-structured interview guide, along with a transparent coding process, ensures that the interpretations are grounded in the empirical material (Recker, 2021; Kaplan & Maxwell, 2005; Bhattacharjee, 2019). As this study does not have a focus on a specific organization, sector, or single country, the external validity is strengthened, however, it is still limited due to the fact that the focus have been put on the Nordic Region and this research only included an analysis of a few select organisations and sectors (Recker, 2021; Bhattacharjee, 2019). This means that while the findings may offer analytical insights applicable to similar settings, they are not intended to be statistically generalizable (Bhattacharjee, 2019).

3.7 Ethics

This study's data collection will ensure respondents' rights. Jacobsen (2002) describes three basic requirements for an ethical study as the right to informed consent, the right to privacy, and the right to be accurately represented.

Informed consent means participants should not be pressured into participating, and must be aware of the study's purpose, as well as any potential positive or negative effects from their participation (Jacobsen, 2002; Patton, 2015; Recker, 2021). Participants should also be able to revoke their consent at any time during the research process (Recker, 2021; Oates et al., 2022; Jacobsen, 2002). To ensure informed consent, the participants have approved a consent form (Recker, 2021). This consent form, as presented in figure 3.2 Participants rights and consent form below, was sent about a week before each interview where the participants via email could provide written consent and where the opportunity for questions was opened up. The aim of the consent form was to ensure participants are aware of their rights, and how their data will be used. The consent form was supplemented by an email disclosing the study's purpose once again. Following Patton's (2015) recommendation, this information was also disclosed in short when starting the interviews to further ensure informed consent from the participants and give them a chance to raise questions or concerns and give a final verbal consent to the consent form.

Participants rights and Consent form

I am voluntarily participating in this interview, and understand that I have the right to withdraw my consent or refuse to answer a question at any time.

I understand that in the report of the results of this research my identity will remain anonymous. This will be done by changing my name and disguising any details of my interview that may reveal my identity or the identity of the people I speak about.

I agree to having the interview recorded, transcribed and analyzed. In the final report the full transcripts will not be published, however, certain quotes from transcripts can be included.

I understand that the transcription will be done through a local AI model in order to preserve privacy. All transcriptions will also be reviewed by the authors.

I understand that I have the right to request to review the transcripts prior to publication.

I understand the nature and purpose of the study and I have the opportunity to ask questions to further clarify this.

I understand that I can withdraw permission to use data from my participation within two weeks after the participation, in which the material will be deleted.

I understand that audio recordings will be stored with the researchers for a period of time in accordance with research data management regulations and will be deleted after the study is done.

Figure 3.2: Participants rights and consent form

In regard to the right to privacy, all participants have had the right to be represented anonymously (Bhattacharjee, 2019; Oates et al., 2022). Since the interviews were conducted through a video call the study strives for the confidentiality of all participants. Confidentiality refers to the researchers having identifiable information about participants but not disclosing it in the research results (Bhattacharjee, 2019). Oates et al. (2022) further press on the importance of pseudonymisation in cases where the participant might be easily identifiable in order to avoid negative consequences for the participants if identified. Since the participants were interviewed in their professional role, participants may feel pressure to present socially desirable responses to protect their professional reputation when discussing AI-policy experiences. To address this, all responses in the transcripts and reporting were anonymized, including names, organizations, and other identifiers. Participants were informed of this anonymity before the interviews and in the consent form to encourage candid disclosure.

The interviews were recorded and transcribed if the participant had consented to it, which presents a certain risk that needs to be taken into account. For the transcription, a digital tool will be leveraged, however, all transcriptions will be controlled by the researchers to ensure correctness. Furthermore, participants were able to read the interview transcripts afterwards. This ensures that they are being accurately represented. It also gives them the right to explain or correct any perceived misrepresentations from the interviews. More on the transcription process, is found in sections 3.5.1 *Transcription* and 3.5.3 *Transcript validation*. After the transcription process is done, the recorded interviews were deleted to ensure data privacy for the participants.

The projects data storing was done in accordance with the “Guidelines for the processing of personal data in student projects at Lund University School of Economics and Management”, this was done to ensure that each informant's personal and sensitive data was stored and used ethically and to minimize the risk of harm to the informant, as this study is dealing with voice recordings that according to the guidelines are classified as sensitive personal data (LUSEM, 2023). The guidelines are in alignment with GDPR, and state in short that the data should be accurate, stored securely, obtained lawfully and purposefully (LUSEM, 2023). This was achieved through the consent form, ensuring that the interview guide only asked what was necessary for the interview and also stated the purpose of the interview to each informant both on email and in-person. To ensure accuracy, transcript validation was performed and the recorded material was stored locally and deleted properly.

4 Findings

In this chapter the informant's answers will be compiled based on the three themes guiding the interview guide and the literature review: AI in organisations, Policy Formulation, Policy in practise and the six sub-themes identified during the interviews.

4.1 AI in organisations

In this section the responses coded as AI in Use will be presented.

4.1.1 AI in use

The AI tools used, according to the interviewees, were most commonly generative AI, with Microsoft Copilot being the most frequently mentioned tool.

I1 accounts for the use of Copilot because their organisation has Microsoft Office subscription. Beyond this, ChatGPT is also used, chosen based on perceived efficiency and with the goal to improve overall organisational efficiency.

I2 describes a broad AI usage within their organization, primarily Microsoft CoPilot for enhancing email communication. I2 also notes that their organisation has a GPT connected to OpenAI's ChatGPT, as well as other LLMs such as Gemini, CoPilot, and Claude. The use of AI is motivated by its ability to increase efficiency, and I2 experiences pressure to adopt AI in order to maintain their competitive edge.

I3 described generative AI and AI agents as the main tools in use, primarily from Photoshop, ChatGPT, Gemini and CoPilot, although CoPilot was considered very poorly. Since Copilot and Gemini were included in their Microsoft package they were used by default, while ChatGPT was chosen based on perceived best performance. I3 highlighted that the organisation's purpose of the AI use was to gain efficiency.

I4 explained that Copilot and Codex were chosen because they were perceived as the most efficient and cost-effective options. I4 stated that the purpose of AI was to improve efficiency.

I5 described the use of coding-specific tools such as GitHub Copilot, Claude Code, and Cursor. For the broader organisation, all employees have access to Microsoft CoPilot. Additionally, specific AI tools are sometimes integrated into other platforms in use, such as those used for HR. I5 noted that AI use cases within the organization are overall very broad. I5 also mentioned that tools such as DeepSeek, which is non-western based, are not used. Overall, I5s organisations strive to balance security and efficiency when selecting AI-tools and actively seeks employee input in the decision-making process.

“Then we also receive these requests from the business for new tools. And there you have to be humble and admit that we don't have the ability to keep track of all types of tools that come out with AI functionality. Instead, we're dependent on employees feeding that information to us. And then we'll often receive some requests, like: I want to

try this because it's popped up, or we have a customer who's talking about it.” (Appendix E #65)

I6 a&b described the use of generative AI such as Copilot, included through a deal with Microsoft as part of their package agreement. They also have an AI prompt interface with the option to select several different models. The goals of their AI use is centered on improved efficiency and alignment with organization's strategic objectives.

I7 describes the use of mainly LLMs for software-development, including tools such as GitHub Copilot and Cursor AI. Beyond software development I7 states that ChatGPT is used in marketing, finance, and other organisational functions. I7's organisation also develops its own AI solutions primarily used in their products for detection and sensors. The informant cited multiple reasons for the AI adoption, including performance, efficiency, and competence development. I7 also noted that AI use helps the organisation attract and retain skilled employees and remain a competitive employer.

“But also from an employer perspective it's relevant, if we want to have skilled software developers here, they expect to be able to use these tools as well. So if we don't offer that, I think quite a few would want to work somewhere else where they're allowed to.” (Appendix G #8)

Short analysis

Across all informants, generative AI dominates, with Microsoft Copilot being the single most commonly mentioned tool, largely due to its bundling with the Microsoft 365 enterprise agreements. Although all organisations used this tool, there was disagreement regarding its quality and usability. Increased efficiency was cited as a motivator for the AI use by nearly all informants; however, security, compliance and economic aspects were considered when choosing tools.

4.2 Policy formulation

In this section the responses related to Policy formulation process will be presented, these are coded as Purpose of Policy and Policy formulation.

4.2.1 Purpose of policy

I1 described their AI-policy as designed to ensure safe, ethical, and efficient AI use. I1 further stated that frameworks, security and responsibility are at the core of their AI-policy, while also expressing the difficulties of balancing these principles with a desire for innovation and experimentation. The policy was also built around core values rather than stating explicit tools and rules, as I1 felt that a more prescriptive policy would become obsolete too quickly. I1 also expressed concern about the loss of accountability in the absence of an AI-policy.

“But what is still very clear is that it is human in the loop. Because I was a little worried about that when we started. If we don't get the [AI-policy] in place. There is a risk that we develop a kind of 'blame the bot culture.'” (Appendix A #26)

The goal for I1 was to produce a concise, value-based policy rather than a rigid document of rules.

“There are definitely more comprehensive AI-policies out there. But the goal was really, what is actually the absolute minimum amount of information we need to have” (Appendix A #26)

When discussing which regulations most influence their AI-policy, I1 identified GDPR as the primary regulatory driver and noted that the organisation has begun examining other frameworks such as the EU AI Act, though regulatory compliance currently remains on a basic level.

I2 was not directly involved in formulating the AI-policy but described it as a bridge that enables AI use while ensuring security. A central concern for the organisation was balancing the benefits of AI adoption, particularly in terms of efficiency, with the protection of confidential business information.

I3 took a different approach to the purpose of the AI-policy, viewing primary as a tool for communication with customers and external stakeholders. For I3, the main concerns were honesty and transparency in AI use.

“Which is also more of a, perhaps not a policy but more of a promise, and we thought that was clear... communicative as well. Policy sounds boring but promise is more communicative” (Appendix C #13)

“...we didn't change how we worked after the policy, it was more that we clarified how we worked.” (Appendix C #19)

I3's AI-policy was developed in response to questions and concerns from stakeholders regarding the organisation's use of AI, and served as a statement, intended to address fears of job displacement, a sentiment I3 described as highly prevalent in their sector. The informant further noted a general mistrust of AI within the sector, stemming from previous privacy breaches by AI companies. I3 identified the core values of the AI-policy as ensuring consistency and clarity, summarizing them as: *“But it is clarity, reservation and feedback in dialogue”*. (Appendix C #15).

I4 stated that the policy exists primarily to comply with regulatory requirements concerning customer data, and that its overarching goal is to limit AI use. In an earlier iteration of the policy, AI was banned altogether within the organisation. However, the respondent acknowledged that this approach was not sustainable in the long-term. The informant also noted the difficulty of balancing the potential opportunities of AI with the associated risks and regulatory demands. Nevertheless, I4 suggested that the policy itself could serve as a competitive advantage.

“And unfortunately there are perhaps competitors who don't think about this perspective. Now we are trying to rather go with the approach of selling ourselves to larger organisations so they know that we don't do this.” (Appendix D #52)

I4's primary concern when developing the AI-policy was the loss of control of organisational data:

“Well it was a few years ago now and back then the fear on my part was that our information would go astray without us having control over it. That is it, that we lose control over our information, we are no longer the owner of it, because it is someone else who indirectly becomes the owner and can train their AI on our information and we have no control over where, when that happens” (Appendix D #18)

I5 articulated the primary goal of their AI-policy as follows:

“The primary goal is to ensure that we can have a high implementation of AI but that we do it in a responsible way that lives up to our expectations and our customers' expectations of us.” (Appendix E #16)

I5's focus is on ensuring that AI does not negatively impact the organisation's core objectives. This is, according to I5, achieved by establishing employee understanding of acceptable AI use through the policy. I5 further stated that they did not wish to prohibit AI entirely but rather viewed it as a powerful tool if used appropriately. Customer expectations and pressures were also cited as a key driver of the AI-policy, given that the organisation operates across a broad range of sectors and countries, each with distinct risk appetites and regulatory requirements.

“Yes but the expectations from customers affect us to a high degree. It is like for all businesses, it is the needs of the market that govern to a greater extent.” (Appendix E #30)

I6 a&b stated that their policy is designed to support the organisation's core objectives, while upholding transparency and credibility. I6b acknowledges the risks associated with AI use as a central driver of the AI-policy *“It can easily go wrong if we use AI tools in an incorrect, clumsy and immoral way”* (Appendix F #20). They further stated that the policy seeks to balance the adoption of emerging technologies with the need to maintain trust among stakeholders and customers.

I7 identified a dual purpose for the policy:

“Yes, we see it both as a way to protect the company when it comes to ownership around the data and the source code. At the same time it is actually somewhat of a communication tool to encourage AI use as well.” (Appendix G #34)

Regarding communication, I7 noted that this currently takes the form of internal communication aimed at encouraging AI adoption among employees. However, I7 also acknowledged that, in the future, having a well-defined AI-policy could build trust with buyers and investors, potentially functioning as a marketing tool.

Short analysis

With respect to the purpose of AI-policy, I1, I2, I4, and I5 were primarily concerned at managing the internal risks and ensuring that the AI adoption aligns with the organisational expectations regarding, quality, security and ethics. By contrast, both I3 and I7 viewed AI-policy as a tool of communication. While I3's AI-policy is viewed as a “promise” to external stakeholders, I7 saw it as a means of building trust with customers and a potential marketing asset. I5 additionally identified customer demands as a primary driver of AI-policy development.

What unites the motivations across all informants is an underlying tension between the control and adoption of AI, as organisations seek to harness its potential while simultaneously managing ethical, reputational and regulatory risks

4.2.2 Policy formulation process

I1 described their policy as emerging from internal demand for frameworks governing AI.

“In August 2025 I had developed a larger survey that went out to all employees where we asked exactly this, how much do you use AI today and what do you need in order to feel secure? And then it was precisely this: frameworks, training, what are we allowed to do, so that [the AI-policy] was a direct reaction to this” (Appendix A #20)

The policy was developed by I1 in response to this demand, together with the organisation CTO, management team, and legal department. During its creation they examined best practises and, in keeping with their goal of maintaining a short and simple policy, considered what the minimum necessary content would be for the AI-policy to remain valuable. When discussing the inclusion of legal and regulatory aspects, I1 stated that this did not feel relevant, as compliance is something the organization must adhere to regardless.

I2 had limited insight into the policy formulation process but describes the policy as having been developed at the global organisational level by a dedicated team, which sets boundaries and negotiates contracts. I2 was not aware of which specific departments drove decisions regarding AI-policy. From a user perspective, I2 noted that the policy and tools preceded the actual AI use with AI tools rolled out in a staged, gate-by-gate manner, first accessible to IT and digitalization teams and then gradually extended to the rest of the organisation.

I3 described that there was initially little interest in an AI-policy, however, being the communication chief, I3 began receiving questions from external stakeholders regarding the organisation's AI use. I3 stated that *“The [stakeholders] probably thought that we were doing things that were bad for them.”* (Appendix C #19) which formulated the basis of the AI-policy.

“So that was the starting point. To do a sort of, what is sometimes called a materiality analysis. A stakeholder analysis. And think about those who care about us. What questions do they have around AI? And how would we answer those questions?” (Appendix C #23)

When creating the policy, they also reviewed other organisations' AI-policies within their sector. I3 further stated that the policy is not designed to communicate specific rules, as this would cause it to lose relevance too quickly.

I4 stated that the policy was created as a response to concerns about information and data leakage. I4 argued that, prior to the rise of AI, there had been little incentive to share data externally, but this changed rapidly. The initial policy banned all AI use within the organisation but I4 realized that this was not sustainable in the long run, as it proved impossible to regulate their employees and coworkers in that manner. The informant concluded that it was thus more effective to provide guidance on how to use AI responsibly. Having reached this conclusion, I4 looked to other organizations' AI-policies for inspiration and adapted them to fit their organizational context. I4 also highlighted the issue of low employee engagement in the AI-

policy, and specifically created a shorter, more condensed version aimed for employees that do not read the official policy document.

I5 described their policy as a response to the release of ChatGPT and the broader impact of generative AI. This promoted a need to create employee guidelines on identifying prohibited and unprohibited AI use. The first draft of the AI-policy was produced in 2023 and was driven primarily from a compliance and information security perspective. The AI-policy was formulated by the national and scandinavian compliance-manager, the CISO (Chief Information Security Officer), alongside representatives from multiple departments.

I5 noted that the organisation initially had developed general guidelines regarding AI use but soon realized that employees needed explicit references to the specific tools they were using and requesting. However, the pace of technological development made it difficult to ensure such a policy was current and representative. As a result, the organisation transitioned to maintaining separate routines for each tool.

“So today we have replaced guidance around tools with separate routines. So you can read a routine for each respective tool that we say something about, and then the policy is more general. The pace of the policy is a bit lower and instead we can update these guidelines” (Appendix E #28)

I6 a&b both stated that their policy was developed in response to AIs global impact and a recognized need for an organisation-wide policy to maintain trust among consumers and stakeholders. In response, the organisation established an AI-committee responsible for overseeing decisions regarding AI-use, and an AI-council comprising employees from multiple departments produced the first draft of the AI-policy. The policy was created to ensure all members of the organisation use AI responsibly and in accordance with the organisation's core mission. I6a stated that *“So in that way I think it was so easy, in quotation marks, to make a policy. You just need to start from what our mission is as a [sector organisation].” (Appendix F #31)*

I6b further noted that during the policy formulation process, the organisation coordinated with other organisations within their sector to discuss approaches to AI use, which in turn informed their policy. *“It is not a set of rules, it is a good conversation between everyone. How do we maintain credibility and still use this new technology.” (Appendix F #33)*

I7 stated that the AI-policy was created after the informant introduced AI in the organisation and recognized the need for guidelines to ensure data privacy. There were concerns that employees were using AI tools outside of the organisational environment, and the associated risks promoted the development of an AI-policy. The policy was partially inspired by AI-Sweden and other organisations' AI-policies, which were adopted to suit I7's specific organisational context. I7 also conducted interviews both with management and employees to further develop the policy, particularly regarding risk identification and organisational needs. This was considered particularly important, as the informant emphasised that the policy should not only restrict use but actively encourage it. Regarding regulation and governance, I7 noted that these aspects were not explicitly included in the policy, in order to avoid unnecessary complexity.

“And there it also becomes regulatory in the end, you have to comply with it [the regulations] if you want to sell on the European market. We haven't made the AI-policy

more complex based on that. Because it's something that applies to us regardless and that we have to follow in parallel anyway.” (Appendix G #32)

Short analysis

One tension that emerged in relation to policy formulation is the trade-off between specificity and longevity. I1, I3, I7, and I5 all converge on the view that hard rules and tool-specific references cause policies to become outdated quickly given the pace of AI development. I7 and I1 both also excluded regulatory content for this same reason, not because regulation is irrelevant but because compliance is treated as a parallel obligation that exists regardless of whether it is codified in the policy. I5 encountered this issue directly when an initially tool-specific policy proved unworkable as tools evolved. Their solution was a two-tiered approach, supplementing a general AI-policy with separate routine documents for each tool, a potentially more scalable alternative to purely principles-based policies.

Furthermore, informants I3, I4, I6 a&b and I7 describe drawing inspiration from other organisations AI-policies when formulating their own, while I1 consulted best practise during policy development. Whereas I3, I4 and I7 adapted external policies to fit their specific organisational context, I6b described a more collaborative form of inspiration though sector-level dialogue with peer organisations.

4.3 Policy in practise

In this section the responses related to the Policy in practise will be presented, these are coded as Communication, Compliance, and Future Handling.

4.3.1 Communication

I1's AI-policy document is located in the organisation's intranet and has one AI-policy for all the departments. In I1's organisation, to ensure employees understand and make sure they are aware of the policy, the employees go through mandatory training. This is further underscored by management reminding the employees who have not yet completed the training. Other than this, if employees had questions, they could direct their questions to I1 directly. I1 accounts for having monthly digital meetings with all the employees where they have also previously announced the AI-policy when it was released. Following the introduction of new policies, additional announcements are made. I1 also describes plans to restructure the learning courses from one 20-minute course to many smaller separate courses, in order to make the learning process more digestible. I1 also states that there are plans in place to give management training in the AI-policy so that they are also confident in what it means to follow the guidelines. I1 highlights the importance of continuous learning and reminding of the AI-policy.

I2's AI-policy document is located on their internal platform. I2 accounts for having received information through emails, training courses and webinars when AI-policy documents have been released. I2 also notes that there is continuous communication of the AI-policy. I2's company has one and the same AI-policy spanning over the whole organisation.

I3's AI-policy document is located on their website. In I3's organisation, a meeting was conducted to inform about the AI-policy and the meeting opened up for feedback and contributions. I3 described this as a starting point for more discussions on the subject of AI use. After later versions of the policy were launched additional emails were sent out. In order to assist the employees to understand the AI-policy, a short training exercise on the AI-policy was sent out and done by the employees. I3's organisation has one and the same AI-policy for all departments as well as one AI-policy that spans over the entire corporate group.

I4's AI-policy document is located in their organisation's intranet. I4's organisation has lunch and learns where employees receive information. In addition to this, employees could also direct questions to I4 directly. When the AI-policy was launched it was introduced through the monthly newsletter, yearly calibration-meetings and through lunch and learns. I4 has one AI-policy for every department, but they have one section in the policy that is specific for their IT department. The primary support tools include lunch and learn, yearly calibration meetings, and direct access to I4 for questions. I4 also added a summary passage highlighting the AI-policy's key points, allowing the employees to grasp the essentials without reading the full document.

I5's AI-policy document is located in their organisational communication platform where all of the policy and governance documents are accessible to employees. When there are updates in policies the staff will be notified about the changes through this system. Regarding education, I5 states that they have educational initiatives in place regarding AI, focused on efficient use and value creation in line with the organisation's values. I5 states that dedicated compliance employees are available for policy-related questions and actively encourages staff to reach out to them. I5 also accounts for their organisation having an AI-council that decides on the AI guidelines and determines the tools available to support employees. When discussing compliance support I5 noted:

"Yes, you can always contact us who work with compliance issues. We try to be as visible and accessible as possible." (Appendix E #44)

I6 a&b's AI-policy document is located in their organisation's intranet which is accessible from outside of the organisation as well. When the AI-policy was launched it was introduced through the intranet and management emails and informational meetings for management. I6 a&b's AI-policy is reviewed, renewed, and communicated to the organisation and published on a yearly basis. Their AI-policy consists of one overarching policy, supplemented by department-specific policies with more detailed specifications. To help their employees understand how their AI-policy applies to practise, their organisation has adopted an AI model trained on their policy documents, in which employees can direct questions they would have otherwise brought to their manager. Beyond this, I6b identifies the employee's own manager as the most valuable resource for AI-policy guidance.

I7's AI-policy document is located on their internal HR site. When the AI-policy was launched it was announced through email and performance appraisals as well as in onboarding of new employees. When there are big changes made in the AI-policy, it is announced in the company's channels and emails to employees. Regarding the tools available to support policy understanding, during onboarding, employees are required to demonstrate their understanding of the policies. In addition to this, any further questions can be directed to their manager.

Short analysis

Across all organisations the AI-policy is stored in a centralized digital location such as their intranet as I1, I4, and I6 a&b or an organisation wide communication platform as I2 and I5 or HR platform as I7, one notable exception is I3, which has their policy published on their official website making it accessible outside the organisation as well. A majority of the organisations felt the need for support beyond policy awareness and adopted different approaches to ensuring policy understanding. These efforts span from training initiatives accounted for by I1, I3, and I5, to support from managers and responsible compliance personnel as mentioned by I5 where employees can direct questions regarding the policy. Notably I6 a&b stated that they have a dedicated AI model that can help in situations where employees are unsure how to act in accordance with the AI-policy.

4.3.2 Compliance

I1 describes keeping the AI-policy brief in order to encourage employees to read it.

“It's only one and a half A4 pages long. Because we were like, nobody reads these damn policy documents. So it has to be short.” (Appendix A #12)

I1 states that the digitalisation unit is partly responsible for ensuring compliance, along with the employees who should comply with the AI-policy and report any incidents, as well as the national steering group. According to I1 the digitalisation unit is responsible for managing website settings and controlling what information can be shared on external websites, ensuring that mishandling of information is prevented through technical constraints. Other than this, I1's equivalent of an AI office are the ones responsible for the follow-up in cases of any AI-policy breaches. In addition to this the nearest manager would also be involved in the follow-up process, as well as the relevant government authorities such as breaches involving personal data. I1 states that a lot of their follow-up process comes from and concerns laws and regulations in place such as GDPR among others. I1 states that the AI-policy concerns their daily work by regulating AI use. It also generates ongoing discussions and questions about how the policy should be interpreted and whether it should be revised in the future. I1 states that challenges concerning the AI-policy are, for example, that very few people read the policy documents unless properly motivated.

Concerning I1's perceived fulfilment of the policy's purpose, I1 perceives the AI-policy as fulfilling its purpose by addressing the organization's AI ethics and guidelines about safety and trust in a concise and accessible manner. I1 gives the example that accountability always rests with the human using the AI, rather than with the AI itself.

I2 states that everyone is responsible for upholding compliance with their policies. I2 also states that whenever a new policy is introduced, employees are required to complete a training exercise before accessing software functions related to that policy document. For example, in I2's organisation, AI can be used as a tool but not for validation purposes. The responsibility for validation always rests with the human overseeing the process. I2 sees the AI-policy as a good thing, as it allows the company to safely provide AI services to employees, which I2 agrees is the desired approach. I2 acknowledges that the company's AI services are limited but considers them sufficient for daily work. I2 also feels that the reduction of risk to the company justifies these limitations. I2 perceives that the AI-policy fulfils its purpose by protecting

them and the company from harm stemming from the AI usage. I2 states that they have not read the policy document themselves, but that they trust those responsible for formulating it. I2 also adds that since the policy regulations are built into the tools they use, and given the training received, they do not feel the need to read the full AI-policy document.

“Again it's a little bit sad to say that I have never looked for it. I trust the people who have written the policy ...” (Appendix B #40)

I3 states that their company doesn't have a dedicated AI-compliance role, as legal responsibility lies with the stakeholders rather than within the organisation itself. Also, I3 states that Sweden currently has no dedicated AI legislation, unlike some other countries, and therefore the organisation does not employ a dedicated AI lawyer. In the event of an AI-policy breach, I3 states that a meeting would be arranged with the nearest manager and the issue would be resolved internally. Regarding the day-to-day impact of the AI-policy, I3 says that the largest impact is that they are approached with questions concerning the policy document. I3 identifies the biggest challenge as keeping the policy relevant given the rapidly evolving AI landscape. I3 perceives the policy as fulfilling its purpose. I3 describes external pressures to abandon the use of AI but does not advocate for this approach and therefore feels that the policy fulfils its purpose by meeting those expectations halfway while still maintaining an increase in productivity by using AI. I3 also perceives that the broader corporate group-level AI-policy document is not widely read.

“Then we also have this corporate group policy. But it's in English and just, I don't know, I don't think anyone reads it.” (Appendix C #35)

I4 states that an initial AI-policy that forbade proved unenforceable, as it was impossible to govern colleagues and parts of the organisation in this way. I4 states that one main goal with the AI-policy is to preserve the privacy and security of the information their company handles and by constructing guidelines for safe AI use rather than forbidding it outright, it better supports employee compliance. Regarding compliance responsibility, I4 personally holds responsibility for maintaining it. In addition to this both internal and external auditors are responsible for ensuring compliance with the AI-policy. I4 states that in the event of a policy breach, the mistake is first corrected, followed by a disciplinary action and some form of educational measure. I4 states that the AI-policy affects their daily work in the form of productivity loss. I4 states that broader AI use would significantly increase productivity, but since it is not allowed in certain areas, it limits the productivity gains that could otherwise be achieved through AI use. I4 states that they have made a summarizing passage in order to incentivize people to actually read and take in the most important parts of the AI-policy.

“...how many will actually read all of this [the AI-policy]? No, it's not many and it's limited, and then I thought, no we need to have something for those who don't read anything.” (Appendix D #56)

I5 states that compliance with the policy is everyone's responsibility, but the main responsibility to ensure compliance is within a department solely dedicated to internal rules and compliance. In addition to the dedicated department, I5 states that they also conduct additional internal revisions. Other than this, each manager is also partly responsible for having their employees complying with the AI-policy guidelines. If the AI-policy was to be breached, I5 stated that there is a standard process where HR follows up on the proceedings. I5 states that

the AI-policy affects their daily work by restricting both the choice of tools and what information I5 are able to use and share within those tools.

“That means we need to be quite on our toes to keep track of what actually applies to each respective tool.” (Appendix E #56)

I5 states that keeping track of the AI models' licences as the models' change is considered a challenge with the AI-policy. I5 points out that their software must be closely monitored, as vendors may silently introduce AI functionalities, or on the contrary, label features as AI when they turn out not to be. I5 describes the core challenge they face as tracking and keeping pace with these changes. I5 identifies one challenge as ensuring the AI-policy is actually read by employees.

“Then you have to be a bit humble and admit that it's [policy documents] not what people are passionate about when they wake up in the morning and check if a policy has been updated. So it's a constant challenge to push out this type of information.” (Appendix E #64)

I5 says that they believe the AI-policy has fulfilled its purpose by having been widespread and absorbed by employees at I5's company. I5 perceives that the policy has fulfilled its purpose by having a wide reach and incentivising reflections surrounding use of AI and how the AI-policy should be revised. I5 also notes that the AI-policy has generated a large quantity of questions and suggestions, such as requests for specific AI tools to be approved for future use.

I6b stated that compliance is the shared responsibility of both managers and all employees. If the AI-policy was breached, I6b stated that it would be in the hands of the closest manager to handle it. I6a stated that the AI-policy affects their daily work by shaping the short-term and long-term goals set for the organisation. I6a described that, by allowing the organisation to utilize AI, they can set bigger goals concerning productivity. I6b stated that their daily work is affected by the AI-policy by having to continue keeping it relevant and useful for everyone complying with it and communicating within the organisation to gain feedback and reiterating the AI-policy. I6a stated that one challenge with the AI-policy is keeping it relevant to the needs of the various different departments within the organisation. I6b states that they find the AI-policy to fulfil its purpose by having it communicate the organisation's values, rules and ethical and responsible AI use. I6a also states that they find that the AI-policy fulfils its purpose by engaging and encouraging the organisation to reflect on AI use and its future direction.

I7 states that all managers are responsible for ensuring employee compliance with the AI-policy. In the event of an AI breach, I7 states that focus would be on correcting the error made. On a daily basis I7 perceives that the AI-policy affects them in a positive way, since it provides them with guidelines and tools to enhance productivity while actively encouraging AI use. I7 states that a key challenge they face is maintaining the AI-policy's relevance given how rapidly the AI landscape and the AI tools themselves evolve. I7 perceives the policy as fulfilling its purpose by meeting organisational needs at a justified cost, with AI tools delivering on their promise, especially regarding data privacy. I7 also perceives that the policy is fulfilling its purpose by being short and concrete, highlighting the guidelines most relevant to their specific industry and organisation.

“So rather, in this situation at least, the philosophy has been to keep the AI-policy simple for what we do in the company with those parts.” (Appendix E #32)

Short analysis

I1, I2, I3, I4, I5 and I7 all share the perception that employees lack motivation to read policy documents. I1, I4 and I7 also address this by deliberately keeping their policy documents concise, aiming to provide a clear overview that employees will actually read and absorb. The AI-policy is seen as both a hindrance of productivity by I4, as well as a motivator and guide for achieving productivity gains by I3, I6a and I7. A shared challenge among interviewees I3, I5, I6 a&b's, and I7 is keeping the AI-policy relevant, both in keeping pace with the AI models and their capabilities and tracking how those capabilities can be applied to benefit the company and its operations.

4.3.3 Future handling

I1 stated that they have a set schedule for updating the policy, currently in two-year increments, however, I1 recognize that significant shifts in AI technology or regulation may necessitate earlier revisions. I1 further noted that improving employee AI literacy will also be crucial going forward. *“The AI-policy is good, but it's the patience to continuously communicate it that will be what matters now.” (Appendix A #62)*

I2 stated that updating and maintaining a policy is highly challenging, comparing it to smartphone software updates that people tend to postpone until the last possible moment. I2 also noted the need to be able to make rapid changes to the policy, particularly in anticipation of further development within the regulatory landscape. I2 describes the situation as follows:

“They are just like sucking data and information out of us, and we don't have an overview of what's going on, but soon we will see the impact of that and soon we will get some slap from left and right from this technology, of course as much as they are nice and helpful, they could be harmful.” (Appendix B #6)

I3 identified one of the core challenges regarding the AI-policy as ensuring its continued relevance and keeping pace with technological advancements. I3 emphasised the importance of ensuring that the policy accurately reflects the realities of AI use:

“But the challenge remains to ensure that we have an AI-policy that reflects the technological development and how we want to work with the AI-policy” (Appendix C #47)

At present, the policy requires updating, as I3s organisation now has considerably more knowledge about AI and how it is used internally than when the AI-policy was first formulated.

I4 stated that, in terms of future handling, they strive to listen to their employees and monitor how other organizations are updating their policies. The informant noted that there is ongoing work to reformulate the policy, both to foster broader acceptance and to enable the introduction of new tools where appropriate. Looking ahead, I4 reflected on the prospect of more

trustworthy and valuable tools becoming available: *“That is my hope, that I have either found a service in Sweden that we dare to trust and that can also create value...”* (Appendix D #58). I4 also expressed the need for organisation-wide understanding of AI-technology, and more specifically how LLMs function, and considered how this understanding might be reflected in the AI-policy.

I5 discussed the rapid pace of AI development and its implication for AI use within the organisation. I5 noted that the tools currently in use differ from those used six months ago and will likely differ again in the six months ahead. In the same manner, I5 considered it unlikely that the AI-policy could remain relevant without revision over a one-to-two-year horizon. The pace of AI development also creates challenges in formulating an AI-policy that is both practical and concrete. Furthermore, I5 noted that a core aspect of their policy development is shaped by internal demands, where employees' adoption of, and need for new AI tools prompts the organisation to iteratively revise the AI-policy in order to establish regulated permissions for their uses.

I6b expressed that the organisation works continuously with feedback through ongoing interactions. Further I6b anticipated that the nature of the human-AI collaboration will change significantly in the future, with AI becoming more deeply integrated into everyday work and users coming to view it as a natural extension of themselves. I6b drew a comparison to the smartphone revolution, noting that this shift will have implications for how the AI-policy should be formulated. When discussing policy updates, I6b also stated that the organisation will need to remain reactive, as changes to the policy are contingent on development made by the major AI companies.

I7 identified that the fast-paced environment of AI is one of the core challenges when developing AI-policy, given the risk of it quickly losing relevance. I7 aims to keep the AI-policy aligned with the technical realities of the organisation and to ensure it remains open for new tools and use cases.

“Given that technology is moving so fast out there, we notice that a new way of doing things emerges that we don't actually cover with the policy. What we've said so far might not be enough.” (Appendix G #56)

I7 also reflected on the global political climate and how it may need to be considered in future AI-policy, particularly with regards to tools originating from certain countries and organisations.

Short analysis

A shared challenge emerges across all informants: as a result of the rapid pace of technological development, AI-policies risk becoming obsolete almost as soon as they are written. All informants acknowledged that their current policies already require updating or will do so in the next one to two years. I2, I3, I5, and I7 explicitly framed this core challenge as ensuring that policy actually represents both the technological and organisational realities, rather than lagging behind them.

5 Discussion

In this chapter the empirical results will be presented in relation to the previous research and to the Institutional theory. This chapter is divided into the three main themes presented in the literature and their respective sub-themes identified during the empirical data collection. This chapter aims to create a discussion that reflects both the empirical insights as well as the theoretical perspective.

5.1 AI in organisations

In this section the responses coded as AI in Use will be analysed and compared to the literature review.

5.1.1 AI in use

I2 accounts for external pressure stemming from a fear of becoming obsolete or losing competitive edge if the organisation falls behind in AI adoption. This reflects what DiMaggio & Powell (1983) describe as mimetic isomorphism, where adoption is driven by imitation of peers rather than deliberate evaluation. This is because the environment creates a pressure to conform. Van Giffen and Ludwig (2023) identify precisely this dynamic among managers fearing obsolescence. This suggests that mimetic pressure is a structural driver independent of organisational readiness and fit. When adoption is motivated mimetically, the tools risk being chosen out of imitation rather than design, thus creating conditions where strategic and value alignment (Mäntymäki et al., 2023) are difficult to achieve. This could be seen as all organisations use Co-Pilot but the perceived strategic alignment of the usage remains uneven across informants.

In the case of I7, they motivated their broad inclusion of AI models in their AI-policy as a way to stay competitive as an employer. This indicates that they view their AI-policy as a marketing tool to leverage a competitive edge and a way to gain legitimacy among the working force, both within their organisation and outside of it. This is supported by Papagiannidis et al. (2025) who claims that organisations can leverage their AI governance in order to gain legitimacy and reputation as a company as well as improving employee loyalty.

Since the AI-policy is written to appeal to the employees the connection could be made that the policy should be held in high regard by the employees and in continuation could indicate a coupling mechanism between policy and practise. This type of close coupling as a consequence of policy creation with the end users' desires and needs in mind is further supported by Flowerday & Tuyikeze (2016) who states that in order for a policy document to succeed, it should be supported by its users.

I5's account of tool selection being shaped by employee and customer demands illustrates a form of normative pressure within the organisation, where professional expectation and preferences of internal users drive the adoption in the organisation. This, however, differs from the normative isomorphism which is driven by professional networks and legitimation (DiMaggio & Powell, 1983). The finding of I5 shows a bottom-up stakeholder pressure as a

driver for change. This has direct implications for policy formulation, if a governing framework fails to account for these bottom-up pressures it risks being created around assumptions that no longer reflect the organisation's reality.

5.2 Policy formulation

In this section the responses related to the Policy formulation, coded as Purpose of Policy and Policy formulation, will be analysed and compared to the literature review.

5.2.1 Purpose of policy

The informants' accounts revealed two conceptualisations of the purpose of an AI-policy. The first and most prevalent, treats policy as a governing document and a formal mechanism for directing employee behaviour, managing risk and ensuring responsible AI use. This view is in alignment with Aphale et al. (2013) who defines policies as first and foremost a guideline for steering behavior. The second treats the policy as a communicative tool, aimed not at regulating internal practises but rather at conveying organisational values and building legitimacy with internal and external stakeholders. The two orientations are not mutually exclusive, and several informants reflect on both views of the AI-policy.

When viewing AI-policy through the lens of a governing standpoint, further distinction can be made regarding *strategic alignment* and *value alignment*. Strategic alignment refers to governing AI towards organisational goals (Mäntymäki et al., 2023), both I1 and I5, frame their AI-policies as an enabler for high adoption that is designed to give employees the framework and governance needed to use AI in way that aligns with the organisations overall strategic goals.

Value alignment refers to governing AI in accordance with ethical principles (Mäntymäki et al., 2023). This balance between ensuring organisational gain while simultaneously upholding ethical values was raised by informants I1, I2, I4, I5, and I6 a&b .

Both informants I1 and I4 highlighted the difficulties in ensuring value alignment and balancing ethical considerations and risks alongside the opportunities of AI. I1 expressed specifically a wish to innovate and experiment with the technology but expressed difficulties with aligning this with safe and ethical use. This corresponds with the findings from Mäntymäki et al. (2023), who argue that such decisions are often difficult because they frequently require organisations to trade off strategic alignment against value alignment. Moreover, de Laat (2021) and Fioravante (2024) similarly highlighted the difficulties in aligning ethics with profitability. At the same time there is a considerable pressure to adopt AI, I2 mentioned, in alignment with the findings from Van Giffen and Ludwig (2023), a fear of losing competitive edge by failing to capitalize on AI, which according to Fioravante (2024) often further exacerbate the value alignment problem.

It's worth noting, however, that value alignment and strategic alignment are not necessarily in opposition. For some informants, ethical and responsible AI use becomes a source of strategic value and a market differentiator that builds trust with clients and investors. This is particularly visible in I4's statement where the informant views their restrictive policy as a competitive advantage when targeting larger organisations. I7 also anticipates that a credible AI-

policy could function as a signal of trust towards future customers. Furthermore, I5 states that it's the expectations from customers in regard to responsible AI use that affects them and their AI-policy formulation.

This leads to a blurring of ethical and strategic rationale, and connects to the second conceptualization of policy purpose, namely policy as a communicative tool. As previously presented informants I4, I5, and I7 all mention their policy as a marketing tool, which aligns with the findings of Papagiannidis et al. (2025) who argued that developing responsible AI practises can lead to improvement of reputation and legitimacy. This is a part of the environmental pressures outlined by Mäntymäki et al. (2023) and illustrates how outside pressures affects internal AI-policy, such as I5's organisation who had a strong focus on meeting customers' own risk appetites with respect to AI use.

I3 took a different approach to policy as a communication tool, viewing it solely as a means of communicating with outside stakeholders. I3 stated that their AI-policy was more of a promise, another example of a response to Mäntymäki et al. (2023) environmental pressures, and to their stakeholder's general mistrust in AI. The policy was thus created as a way to respond to external questions regarding how the organisation used AI.

When discussing policy as a driver of legitimacy, it is important to consider the risk of policy-practise decoupling, as policies adopted for marketing or competitive edge face a larger risk of policy-practise decoupling (de Laat, 2021; Horneber, 2025). Furthermore, when ethics is used from a marketing standpoint it creates a risk for ethics washing where the organization uses ethics to market itself as responsible without anchoring these claims in practise (de Laat, 2021; Fioravante, 2024). In the cases of I4, I5, and I7 it is difficult to identify if such decoupling is present. In the case of I3, however, the risk of policy-practise decoupling is arguably reduced as I3 stated that policy did not change the organisation's processes but rather served to articulate them. This suggests that the policy rules were already embedded in the organization's practises and, as such, do not remain unimplemented or merely symbolic, avoiding the form of policy-practise decoupling outlined by Bromley and Powell (2012).

5.2.2 Policy formulation process

Butcher & Beridze (2019) argue that the field of AI governance is in disarray, a view that is further supported by Mäntymäki et al. (2022) who suggests that this disorder is present in organisations' own governance of AI. This disorder is reflected in varied ways in which organisations have approached the formulation of their AI-policies. During the creation of their respective AI-policies, two broad directions emerged: producing a detailed, comprehensive policy document with all specifics explicitly stated, or deliberately creating a short, high-level document. This lack of uniformity in policy formulation may indicate the very point raised by Mäntymäki et al. (2022) and Butcher & Beridze (2019) regarding organisational disorganisation in AI governance. This is further supported by de Laat (2021) who highlights that defining AI governance is further impeded by the lack of consensus on norms and best practises.

I1, I3, I5, and I7 deliberately kept their AI-policies short in order to ensure longevity. Anticipating future developments, they recognized that excessive specificity would require frequent revisions given the rapid pace of AI development. I1, I3, I5, and I7 further noted that employees are not naturally incentivized to read policy documents unprompted, which served as an additional motivation to keep the policy concise. I5, while opting for a short core policy chose

a middle ground through a two-tiered approach, separating the policy into two parts, a general AI-policy and a supplementary routine document. By doing so, I5 chose to specify the routine document in considerable detail with a shorter iteration cycle, while allowing the AI-policy itself to remain concise and relevant over a longer period. This allows I5's organisation to adopt a more scalable approach instead of relying on a solely principle-based AI-policy. I5 perceives that the level of detail in the routine document is necessary given the nature of their employee's professional roles. I4, by contrast, chose the opposite route, opting for a longer and more detailed policy document, though including a short summary at the beginning to provide an overview.

These approaches to optimizing the length and format of the AI-policy are related to the point raised by Flowerday & Tuyikeze (2016) who argue that a policy, in order to avoid redundancy, should be constructed in a manner that is supported by its uses. In this regard, I1, I3, I5, and I7 have identified, and acted on the assumption, that their users would not be supported by a long and demanding policy document, and therefore adapted it for maximal readability.

I4 expressed that their initial AI-policy banned all AI use entirely, but quickly realized that this approach was unworkable, a clear instance of decoupling. This could be understood as a consequence of what Berente et al. (2019) describes as weak institutional pressures to comply, combined with incongruent logics, which often result in resistance. Since the policy consisted of guidelines rather than enforceable rules, and its underlying logic was incongruent with the values and practises of the organisation, it was met with resistance from employees. I4 was therefore forced to reconsider the AI-policy formulation, in line with the principle suggested by Flowerday & Tuyikeze (2016), that a policy should be constructed in a manner supported by its users.

During the formulation process, both I1 and I7 chose to exclude existing regulations from their policy. They both chose this approach since they perceived that it would be unnecessary to include regulations that would apply regardless. This reflects a perception of external regulations as a reliable and effective institutional pressure. Rather than exemplifying the risk of a gap between law and practise identified by Horneber (2025), this approach suggests that I1 and I7 consider the gap to be insignificant enough that it does not need to be addressed within their own policies. Since I1 and I7 perceive the EU regulation to be sufficiently effective to achieve its intended purpose, they "outsource" that accountability to the relevant legislation and therefore perceive that the gap between law and practise is insignificant enough to have to complement in their own AI-policy. Treating regulatory compliance as a parallel obligation rather than something that needs to be codified in the policy. According to Berente et al. (2019), a strong pressure to comply, as is the case with laws and regulation, combined with congruent institutional logics, results in a faithful appropriation of the policy at hand. Since I1 and I7 both handle personal data, and are therefore subject to GDPR, their organisations' institutional logics are congruent with those underlying the GDPR. As a result, I1's and I7's organisations' AI-policy and related regulations are closely coupled, functioning as complementary rather than competing frameworks.

Since EU laws have been taken into account during the AI-policy creation, this suggests that AI-policy formulation in these organisations is not a fully-internally driven process but is rather shaped by the movements within the same field.

With regard to the content of the AI-policy, the approaches taken by the different organizations can be understood through the lens of isomorphism, specifically the three forms identified by DiMaggio and Powell (1983): mimetic, coercive and normative (DiMaggio & Powell, 1983).

I1, I3, I4, and I7 have, by creating their AI-policy partly based on other organisations' AI-policies, followed a distinct mimetic isomorphic pattern. Mimetic isomorphism occurs when an organisation has taken inspiration from other organisations, either authoritative or similar in profession (DiMaggio & Powell, 1983). This is a form of modelling, which often occurs when organisations are faced with uncertainty and choose to mimic other organisations that are perceived as more legitimate in order to gain legitimacy themselves (DiMaggio & Powell, 1983). This is precisely the case with I1, I3, I4 and I7 who have looked to best practises, AI-Sweden and other similar organisations for inspiration and thereafter adapting it for their own specific context. The fact that the majority of the participating companies in this study are undergoing mimetic isomorphism combined with DiMaggio and Powell (1983) observation that mimetic isomorphism is common when organisations are faced with uncertainty further supports de Laat (2021), Mäntymäki et al. (2022) and Butcher & Beridze (2019) arguments that the AI governance field is a relatively new and emerging field that is still in its early stages.

I5, can be seen as having undergone normative isomorphism, where knowledge from both professional experience and education influence the decisions being made (DiMaggio & Powell, 1983). Since I5 described having developed their AI-policy independently, this suggests that decisions were grounded in their own knowledge from their professional experiences or prior education, which is the basis of normative isomorphism (DiMaggio & Powell, 1983). A potential cause of I5 having undergone normative isomorphism instead of mimetic isomorphism, which is the most common response to uncertainty, could be attributed to the professionalization of the field of the company's employees, and their need for specific details on the AI tools in use (DiMaggio & Powell, 1983). This could further be explained as the environmental layer having the highest impact on the policy formulation process where the professional community is what defined the norm of AI governance (Mäntymäki et al., 2023). At the organisational level, individuals' expertise was informally translated into practises through the policy formulation process.

I6 a&b's policy formulation process, by contrast, reflects coercive isomorphism, as their organisation collaborated on shared guidelines in response to growing institutional pressures, both formal and informal, and sector wide expectations surrounding responsible AI governance, therefore maintaining legitimacy within their organisational field (DiMaggio & Powell, 1983). In the case of I2, where a significant body of laws and regulations governs their industry, navigating and integrating regulatory requirements has been a central focus of the policy formulation process. Policy formulation in response to these pressures, this could indicate on coercive isomorphism as the driving force. Therefore, the deduction can be drawn that the environmental layer of the organisation's logics, where the external pressures from competitors and external regulations lie, has played a significant role in I2's policymaking (Mäntymäki et al., 2023). By also responding to weaker institutional pressures, the organisation can further gain legitimacy, which may be especially important given their high-risk operating environment.

With regard to the broader regulatory environment, I1, I3, I4, I5, and I7 express that EU laws such as the GDPR and the AI Act fulfil their purpose but do not affect them in a close manner. Fioravante (2024) describes that when national and international policies fail to keep pace

with technological advancements, they risk putting the onus of AI governance and creating guidelines for ethical AI use onto organisations instead. They further state that by doing this, it risks creating a regulatory vacuum, in which organisations themselves have to balance profitability and ethics when governing AI. While I1, I3, I4, I5, and I7 describe a sense of vacuum of specifically AI-related regulations on both an EU and a national level, they perceive that other “neighbouring” regulations fill in the compliance gaps to a sufficient level. Interviewees I1, I3, I4, I5, and I7 therefore acknowledge the regulatory vacuum Fioravante (2024) describes but perceives that the surrounding regulations are enough to satisfy current needs. Since the EU has a relative AI regulatory vacuum, it makes it more difficult for organisations to gain reinforcement for their AI-policies from EU AI regulations, as pointed out by Horneber (2025) and Bromley and Powell (2012). In addition to this, Bromley and Powell (2012) also state that weak legal reinforcement may increase the likelihood of policy-practise decoupling.

5.3 Policy in practise

In this section the responses related to the Policy in practise, coded as Communication, Compliance, and Future Handling, will be analysed and compared to the literature review.

5.3.1 Communication

All organisations' AI-policy documents are stored on a centralized digital platform and are accessible to all members of the organisation. In the cases of I6 a&b and I3 the AI-policies are also accessible from outside the organisation. However, there was a general consensus among all informants that additional initiatives are needed to anchor the policy within the organisation, though the nature of these initiatives varies across organisations.

Mäntymäki et al. (2023) underscore the importance of developing internal capabilities to ensure *strategic alignment* when implementing AI, specifically, ensuring that employees understand the expectations and directions governing AI use. This also includes answering the *what, why* and *how* of AI use. This means fostering policy understanding, beyond policy awareness. Berente et al. (2019) argues that when incongruent institutional logics are present, where logic is embedded in a policy but not aligned with the logic that guides an employee's daily work, loose coupling or resistance is a likely outcome, as employees will not enact the policy as intended.

Communication of *what* the policy requires, *why* it exists and *how* it applies to an employee could potentially bridge this logic gap. Policy understanding may therefore be understood as a coupling mechanism. When employees are equipped with all three dimensions, the logic of the policy may become more legible and more congruent with their own working context. This tightens the connection between institutional intent and technical practise, thus according to Berente et al. (2019) reducing the conditions during which decoupling occurs. Furthermore Bromley & Powell (2012) note that policy-practise decoupling is likely to happen when organisations do not invest the time and resources needed for actual implementation, and lack of communication could be understood as a result of such a failure of insufficient investment.

Informants I1, I3, and I5 all stated that they implemented training initiatives to help users navigate the AI-policy. I1 noted that during the launch of the policy, all employees were required to complete mandatory training. This, currently, remains a one-time initiative though plans for more continuous reinforcement of the AI-policy exists, such as management training. I1 further stressed the importance of continuous training. I3 gave a similar account, describing how all employees underwent a short training session following the AI-policy launch. I5 also reported having training initiatives related to AI, however, these were more focused on training employees in the practical use of AI and value generation rather than the policy itself. Daugherty & Wilson (2018) stress the need for continuous staff training and retraining to ensure all employees in the organisation possess the necessary skills.

I1, I4, I5, I6 a&b, and I7 all accounted for the role of leadership in the communication of the AI-policy and how policy understanding is formed in practise. I7 noted that employees with questions about the policy could pose them to their manager, and I6 a&b similarly positioned the closest manager as the main support resource for policy related uncertainty. This positions leadership not only as an administrative function but as a communicator. Flowerday & Tuyikeze (2016) pointed out the importance of management support both to ensure policy awareness but also to actively ensure education. They further argue that committed leadership would motivate employees' acceptance and compliance with policy. Thus, actively affecting the risks of decoupling. Shneiderman (2020) also stresses the importance of leadership commitment and especially their role in handling responsibility and training.

Beyond relying on managers or static documents I6 a&b's organisation also developed a dedicated AI model trained on their own policy document, creating a tool designed specifically to help employees navigate policy-related uncertainty in real time. This reflects a level of resource commitments that signals strong management support. As previously mentioned, Bromley & Powell (2012), who argues that policy-practise decoupling occurs when organisations fail to make such resource allocations, investment in this model, alongside the management commitment, therefore functions as coupling mechanisms. Furthermore, this also indicates a dual alignment: not only is there value alignment, where policy represents the ethical and operational priorities, but also strategic alignment where resource allocation actively supports the realization of those in practise. This shows close connection between strategic priorities, and the stated values is particularly important given the organisation's ethical commitment to transparency and credibility are central to its core objectives.

5.3.2 Compliance

Horneber (2025) and Bromley and Powell (2012) argue that limited policy reinforcement weakens the structures needed to prevent policy-practise decoupling, identifying accountability fragmentation as a key compliance risk. Schneider et al. (2023) further state that policy without definition of responsibilities and accountabilities is structurally incomplete while Ayling and Chapman (2022) specifically argue that AI makes it inherently difficult to assign accountability in practise.

The finding reflects this fragmentation. I2, I3, and I5 all state that the responsibility of policy compliance rests entirely with the individual employees. When managers do not enforce compliance the conditions Horneber (2025) identifies as decoupling factors are more likely to emerge. I3 presents a notable case in which the responsibility does not reside within the organisation but is attributed to external stakeholders, potentially creating an accountability gap

beyond the organisational borders, where enforcement becomes difficult to sustain in practise. I6 a&b stated that the responsibility is shared among the managers and employees, while I7 identifies compliance as a managerial responsibility. This variation is consistent with Ayling & Chapman's (2022) observation that AI use obscures responsibility in systems outcome, making accountability assignment inherently difficult.

In contrast, I5 noted that, alongside the individual responsibility of employees, their organisation maintains a dedicated compliance department and internal revisions ensure compliance. Similarly, I1 has a team dedicated to compliance, who also ensured technical support and guardrails to ensure compliance. I4 described a more structured approach, including both internal and external audits. These mechanisms function as reinforcement structures that operationalize accountability concretely, directly reducing decoupling risk, which rises when responsibility is fragmented or unclear (Horneber, 2025).

A central finding regarding AI-policy was the consensus among I1, I3, I5, and I7 that employees are unlikely to read the policy documents, I2 also acknowledged not having read the AI-policy. This may be an indicator of what Bromley and Powell (2012) describe as policy-practise decoupling, where the policy exists only as a symbol without meaningfully affecting the organisation. I5 specifically highlighted the challenge of diffusing the policy content within the organisation given that employees rarely engage with policy documents voluntarily. It is important to note, however, that policy can be communicated through other means than document reading, such as training and support, as mentioned in the previous section 5.3.1 *Communication*. When employees are aware that policy exists but lack comprehension of its requirements, the policy cannot fulfil its intended purpose of steering behaviour as outlined by Aphale et al. (2013). This gap between awareness and understanding is a direct risk factor for decoupling.

Flowerday & Tuyikeze (2016) argue that policies risk becoming redundant or poorly constructed when they fail to be supported by its users, and further contended that both management and employee support is critical to successful policy implementation. This is illustrated by I1, whose policy formulation began with an organisation-wide survey designed to incorporate the employees needs in the policy. I7 in a similar manner included both management and employees during the crafting process. This may establish a foundation of co-ownership discussed by Flowerday & Tuyikeze (2016), making the policy more likely to be internalised. Flowerday & Tuyikeze (2016) argue that this ownership improves policy implementation success and policy understanding, through the creation of a more digestible policy. This could further be understood as a coupling mechanism, as co-ownership may impact the degree to which the policy aligns with the employees institutional logic, and Berente et al. (2019) argues that incongruent institutional logic leads to loose coupling or resistance of policy. Meaning that by establishing co-ownership and grounding the policy in employees own working logic the AI-policy has greater opportunity for real impact.

One illustration of this dynamic can be found in I4 experience, their initial AI-policy imposed a total ban on AI use, which I4 quickly recognized as unworkable, as it proved impossible to enforce. This outcome can be understood as the consequences of the policy being incongruent with the organisation and employees own working logic. As I4's employees were already preparing to integrate AI into their daily work, a total ban on AI use was fundamentally at odds with their existing practises. This is precisely the type of logical incongruence Berente et al. (2019) identifies as the cause of resistance. I4 subsequently revised the policy to support both organisation's objective and employees, reframing it as a set of guidelines on responsible use.

Several informants described a tension between AI-policy as an enabler or a constraint on productivity. This is an example of the strategic alignment discussed by Mäntymäki et al. (2023), where policy should be aligned to not only manage risk, but support the broader organisational objectives. When a policy is formulated in a way that fails to do this, it may indicate a misalignment between strategy and value alignment, undermining employees abilities and, as Flowerday & Tuyikeze (2016) warn, risk redundancy. I4 specifically noted that their current AI-policy restricts operational capacity and reduces productivity, suggesting precisely this form of misalignment. This echoes the findings of Ayling & Chapman (2022) who state that the current AI landscape has made it increasingly difficult for organisations to balance profitability with risk management. By contrast, I3 described a more balanced relationship between value and strategic alignment, stating that their AI-policy effectively addresses the environmental pressures regarding ethics while simultaneously enabling increased productivity.

A further compliance challenge identified across informants is maintaining a policy that remains relevant as the underlying technology advances at an accelerating pace. This mirrors the findings of de Laat (2021), Gianni et al. (2022) and Papagiannidis et al. (2025) all of whom argue that the pace of AI advancements and diffusion has intensified the tensions inherent in effective AI governance. This puts pressure on policy making as Van Giffen and Ludwig (2023) as well as Holmström (2024) similarly highlight the importance of ensuring that AI governance frameworks are capable of accommodating the technological change rather than functioning as static artefacts.

The findings reflect these structural pressures directly. I5 highlighted that licensing terms shift continuously, making it difficult for the organisation to maintain an AI-policy that accurately reflects the organisation's AI environment. When a policy is unable to keep pace with the technology it is designed to govern, it risks what Bromley & Powell (2012) and Horneber (2025) described as means-ends decoupling. This occurs when the policy remains formally in place while the environment it is set to govern has moved beyond its scope and the policy no longer fulfils its intended purpose.

I5 also struggled to define what constitutes AI in their organisational context. This is not surprising given, as Dignum (2019) notes, defining AI itself remains an unresolved conceptual problem. When organisations struggle to define the boundaries of AI, governing it becomes inherently challenging as this ambiguity makes it difficult to determine whether the AI-policy is achieving its intended purpose. The ambiguity of definitions ties in to the argument Horneber (2025) makes about the risk of means-ends decoupling in complex and uncertain fields, where it becomes difficult to compare intended and actual outcomes. As a result, organisations face challenges in evaluating whether an AI-policy has been successfully implemented and whether it has had its intended impacts.

5.3.3 *Future handling*

As discussed in 5.3.2 *Compliance*, ensuring policy relevance within a rapidly evolving technological landscape remains a complex challenge for organisations. All informants acknowledge that their current AI-policies either already need updating or will do so within the next one to two years. This indicates that maintaining applicability over time is a central challenge for the future of AI-policies, as a policy document that fails to accurately reflect an organisation's technical realities risks means-ends decoupling presented by Bromley &

Powell (2012) and Horneber (2025), where a document remains in place but fails to achieve its intended governance purpose. This is consistent with Holmström's (2024) argument regarding the importance of ensuring adaptability in order to handle the rapid change in the AI field. I1 stated that, going forward, a goal of the AI-policy is to maintain continuous communication and deliver learning initiatives throughout the organisation, in order to spread both AI literacy and awareness of the AI-policy. This is further supported by Van Giffen & Ludwig (2023), who state that AI governance cannot be seen as a one-time static event but rather a process of continuous learning and adaptation. For organisations this means the question not only concerns how to create an AI-policy but rather how to sustain relevance over time, particularly given the accelerating pace of AI development.

6 Conclusion

The purpose of this study was to explore how internal AI-policies are formulated and how these policies are translated into practise within organisations. Through a qualitative interview study, the aim is to answer the following research question: *How do organisations navigate AI-policy formulation and AI-policy in practise?*

Two distinct perspectives on the purpose of AI-policy emerged from the empirical findings: AI-policy acting as a governance document and AI-policy acting as a communication tool. This distinction shapes the starting point of the policy formulation process and signals objectives of the AI-policy. Regarding purpose, this study identified organisations experiencing difficulty in ensuring both strategic alignment and value alignment when implementing AI-policy. However, the study also found that ethical and value alignment can function as a strategic asset, where an AI-policy is positioned primarily as a marketing tool and a source of competitive advantage rather than as an internal governance document.

With respect to AI-policy formulation, this study identifies three isomorphic mechanisms that influence how organisations approach the process, namely: normative, coercive and mimetic. These three mechanisms emerged from the empirical data as the primary ways in which an organisation navigates AI-policy formulation. All of the organisations have been exposed to, and acted upon similar external pressures, yet they have responded to these pressures in distinctly different ways. External forces such as regulations, societal norms, or even guidelines agreed upon by organisations in the same industry are some of the forces affecting how as well as why an organisation takes on an AI-policy.

Regarding the translation of AI-policy into practise, the findings indicate that doing so effectively demands specific compliance efforts in order to avoid decoupling. The empirical findings show that organisations need to strive for AI-policy understanding rather than focusing solely on AI-policy awareness. Ensuring that employees are not only aware of the policy's existence but have the tools and comprehension necessary to ensure compliance. Among the respondents there was a consensus that the policy documents are rarely read, and several organisations thus strived for shorter and simpler AI-policies, supplementing them with other mechanisms to steer behaviour, such as through routine documents, training initiatives, tools, and direct management and compliance personnel support.

In answer to the research question, organizations navigate AI-policy formulation and practise through acting on isomorphic pressures and furthermore using the AI-policy to align with their strategic aims and organisational value. This study further suggests that organisations actively balance isomorphic pressures against their own strategic and value-driven objectives. This study highlights that AI-policy has become as much a social and communicative tool as a governance one, making the external pressures and norms more impactful for policy formulation. Organisations further navigate the AI-policy practise through efforts that ensure, not only, AI-policy awareness but AI-policy understanding to ensure the AI-policy achieves its intended objectives.

This study furthermore illuminates the gap between abstract ideals and their translation into practise where organisations face multiple institutional pressures, and further identifies AI-

policy formulation intent and AI-policy practise as two distinct processes that organisations often conflate, making it difficult to identify potential issues and their causes.

As more regulations regarding AI governance are imposed, such as the EU AI Act, the pressures on organisations and their operational AI-policies are further escalated. It is therefore increasingly important that these internal AI-governance processes are functioning effectively within organisations.

6.1 Future research

This study has been limited by time and resource constraints and further efforts should strive to broaden the scope of the data collection to enhance generalizability.

Further this study identifies the mechanisms that affect AI-policy creation and practise from the perspective of management personnel involved in the compliance or formulation of AI-policy. This limits the perspective gained especially in regard to AI-policy in practise where the insights of employees directly affected by the policy could gain a more nuanced view of how the AI-policy is understood by employees.

For future studies, another interesting continuation of this study would be to extend the scope of the research to also include the industries the organisations operate in as a factor. This would contribute to a deeper exploration of how the nature of different industries, for example industry specific regulations and norms, impacts both the formulation and implementation of the AI-policy. Furthermore, including the size of an organisation as a factor could further deepen the insights gained, since it could be an important factor both in AI- and policy-maturity, and societal pressures. For example, a larger organisation might be affected more by external societal pressures than a smaller company that doesn't face the same public scrutiny.

Furthermore, in the future another aspect might be taken into account, namely new regulations arising, for example for when the EU AI Act has taken force, it might bring more implications for AI governance than the existing regulations.

Another interesting aspect to further investigate is to broaden the scope of the geographical area that is covered in the study. Since this study only covers the Nordic countries, a Nordic perspective and way of operating companies may impact the way this study has found organisations navigate formulating their AI-policy and translating it into practise.

Appendix: AI contribution statement

The AI tools that have been utilized in this thesis are ChatGPT, Perplexity AI, Claude, NotebookLM and Whisper. These tools, with the exception of Whisper, have been used to obtain a quick overview of potential sources, helping us assess whether they are relevant before reading them in full. All sources used in this thesis were subsequently read and evaluated personally and in their entirety by us, AI was not used to read or interpret any sources on our behalf.

Furthermore, these AI tools were assisting us with proofreading, specifically regarding grammar and flow. The actual written content of the thesis was independently produced by us. AI did not generate any content included in the thesis. Claude and Perplexity AI were also used in the translation of quotes from informants, however, all quotes were checked by us to ensure accurate representation.

We have used prompts to sharpen our analysis and challenge us in our arguments, pushing us to think further in our analysis.

References

- Ågerfalk, P. J. (2020). Artificial Intelligence as Digital Agency, *European Journal of Information Systems*, vol. 29, no. 1, pp.1–8
<https://doi.org/10.1080/0960085X.2020.1721947>
- Alvehus, J. (2023). *Skriva Uppsats Med Kvalitativ Metod: En Handbok*, 3th edn, Stockholm: Liber
- Alvehus, J., Eldén, S., Wästerfors, D. & Öhlin, S. (2025). *Metodboken: Att Tänka, Undersöka, Analysera Och Skriva På Samhällsvetenskapliga Grunder*, Lund: Studentlitteratur
- Aphale, M. S., Norman, T. J. & Şensoy, M. (2013). Goal-Directed Policy Conflict Detection and Prioritisation, in H. Aldewereld & J. S. Sichman (eds), *Coordination, Organizations, Institutions, and Norms in Agent Systems VIII*, 2013, [e-book] Berlin, Heidelberg: Springer, <https://link.springer.com/book/10.1007/978-3-642-37756-3>
- Ayling, J. & Chapman, A. (2022). Putting AI Ethics to Work: Are the Tools Fit for Purpose?, *AI and Ethics*, vol. 2, no. 3, pp.405–429 <https://doi.org/10.1007/s43681-021-00084-x>
- Berente, N., Gu, B., Recker, J. & Santhanam, R. (2021). Managing Artificial Intelligence, *MIS Quarterly*, vol. 45, no. 3, pp.1433–1450
<https://doi.org/10.25300/MISQ/2021/16274>
- Berente, N., Lyytinen, K., Yoo, Y. & Maurer, C. (2019). Institutional Logics and Pluralistic Responses to Enterprise System Implementation: A Qualitative Meta-Analysis, *MIS Quarterly*, vol. 43, no. 3, pp.873–902 <https://doi.org/10.25300/MISQ/2019/14214>
- Bhattacharjee, A. (2019). *Social Science Research: Principles, Methods and practises* (Revised Edition), [e-book] University of Southern Queensland, Available Online: <https://usq.pressbooks.pub/socialscienceresearch/>
- Braun, V. & Clarke, V. (2006). Using Thematic Analysis in Psychology, *Qualitative Research in Psychology*, vol. 3, no. 2, pp.77–101
<https://doi.org/10.1191/1478088706qp063oa>
- Bromley, P. & Powell, W. W. (2012). From Smoke and Mirrors to Walking the Talk: Decoupling in the Contemporary World, *Academy of Management Annals*, vol. 6, no. 1, pp.483–530 <https://doi.org/10.5465/19416520.2012.684462>
- Bryman, A. (2018). *Samhällsvetenskapliga Metoder*, 3th edn., Stockholm: Liber
- Butcher, J. & Beridze, I. (2019). What Is the State of Artificial Intelligence Governance Globally?, *The RUSI Journal*, vol. 164, pp.88–96
<https://doi.org/10.1080/03071847.2019.1694260>
- Conboy, K., Crowston, K., Lundström, J. E., Jarvenpaa, S., Ram, S., Mikalef, P. & Ågerfalk, P. (2022). Artificial Intelligence in Information Systems: State of the Art and Research Roadmap, *Communications of the Association for Information Systems*, vol. 50, no. 1, <https://doi.org/10.17705/1CAIS.05017>
- Daugherty, P. R. & Wilson, H. J. (2018). *Human + Machine: Reimagining Work in the Age of AI*, La Vergne: Harvard Business Review Press
- de Laat, P. B. (2021). Companies Committed to Responsible AI: From Principles towards Implementation and Regulation?, *Philosophy & Technology*, vol. 34, no. 4, pp.1135–1193 <https://doi.org/10.1007/s13347-021-00474-3>

- Dignum, V. (2019). Responsible Artificial Intelligence: How to Develop and Use AI in a Responsible Way, [e-book] Cham: Springer International Publishing, <http://link.springer.com/10.1007/978-3-030-30371-6>
- DiMaggio, P. J. & Powell, W. W. (1983). The Iron Cage Revisited: Institutional Isomorphism and Collective Rationality in Organizational Fields, *American Sociological Review*, vol. 48, no. 2, pp.147–160 <https://doi.org/10.2307/2095101>
- Fioravante, R. (2024). Beyond the Business Case for Responsible Artificial Intelligence: Strategic CSR in Light of Digital Washing and the Moral Human Argument, *Sustainability*, vol. 16, no. 3, p.1232 <https://doi.org/10.3390/su16031232>
- Flowerday, S. V. & Tuyikeze, T. (2016). Information Security Policy Development and Implementation: The What, How and Who, *Computers & Security*, vol. 61, pp.169–183 <https://doi.org/10.1016/j.cose.2016.06.002>
- Gianni, R., Lehtinen, S. & Nieminen, M. (2022). Governance of Responsible AI: From Ethical Guidelines to Cooperative Policies, *Frontiers in Computer Science*, vol. 4, <https://doi.org/10.3389/fcomp.2022.873437>
- Goldkuhl, G. (2012). Pragmatism vs Interpretivism in Qualitative Information Systems Research, *European Journal of Information Systems*, vol. 21, no. 2, pp.135–146 <https://doi.org/10.1057/ejis.2011.54>
- Grøder, C. H., Parmiggiani, E. & Williams, R. (2025). Navigating the Future: Planning for AI in Public Services, *Scandinavian Journal of Information Systems*, vol. 37, no. 2, <https://doi.org/10.17705/3SJIS/037.15>
- Guest, G., MacQueen, K. & Namey, E. (2012). Applied Thematic Analysis, [e-book] California: SAGE Publications, Inc., Available Online: [https://meth-ods.sagepub.com/book/applied-thematic-analysis](https://methods.sagepub.com/book/applied-thematic-analysis)
- LUSEM. (2023) Guidelines for the Processing of Personal Data in Student Projects at Lund University School of Economics and Management. Available Online: [https://www.lusem.lu.se/sites/lusem.lu.se/files/2024-01/personal-data-in-student-pro-jects-2023.pdf](https://www.lusem.lu.se/sites/lusem.lu.se/files/2024-01/personal-data-in-student-projects-2023.pdf) [Accessed 16 April 2026]
- Holmström, J. (2024). A Layered Organizing Logic for Generative AI Evolution. Insights from Digital Infrastructure Theory, *Scandinavian Journal of Information Systems* vol. 36, Available Online: <https://aisel.aisnet.org/sjis/vol36/iss1/11>
- Horneber, D. (2025). Understanding the Implementation of Responsible Artificial Intelligence in Organizations: A Neo-Institutional Theory Perspective, *Communications of the Association for Information Systems*, vol. 57, pp.185–218 <https://doi.org/10.17705/1CAIS.05708>
- Jabbouri, R., Truong, Y., Dirk, S. & Palmer, M. (2019). Organizational Decoupling: A Systematic Literature Review and Directions for Future Research, EURAM Conference 2019, proceedings, <https://pure.qub.ac.uk/en/publications/organizational-decoupling-a-systematic-literature-review-and-dire/> [Accessed 17 April 2026]
- Jacobsen, D. I. (2002). Vad, Hur Och Varför: Om Metodval i Företagsekonomi Och Andra Samhällsvetenskapliga Ämnen, Lund: Studentlitteratur
- Jacobsen, D. I. (2017). Hur Genomför Man Undersökningar? Introduktion till Samhällsvetenskapliga Metoder, Lund: Studentlitteratur
- Janiesch, C., Zschech, P. & Heinrich, K. (2021). Machine Learning and Deep Learning, *Electronic Markets*, vol. 31, no. 3, pp.685–695 <https://doi.org/10.1007/s12525-021-00475-2>
- Kaplan, B. & Maxwell, J. (2005). Qualitative Research Methods for Evaluating Computer Information Systems, In: Anderson, J.G., Aydin, C.E. (eds) *Evaluating the Organizational Impact of Healthcare Information Systems*, pp.30–55

- King, J. L. & Kraemer, K. L. (2019). Policy: An Information Systems Frontier, *Journal of the Association for Information Systems*, pp.842–847
<https://aisel.aisnet.org/jais/vol20/iss6/2>
- Kirk, J. & Miller, M. L. (2011). Reliability and Validity in Qualitative Research, [e-book] SAGE Publications, Inc., Available Online: <https://meth-ods.sagepub.com/book/mono/reliability-and-validity-in-qualitative-research/toc>
- Lee, R. S. T. (2020). Artificial Intelligence in Daily Life, [e-book] Singapore: Springer Singapore, Available Online: <https://link.springer.com/10.1007/978-981-15-7695-9>
- Mäntymäki, M., Minkkinen, M., Birkstedt, T. & Viljanen, M. (2022). Defining Organizational AI Governance, *AI and Ethics*, vol. 2, no. 4, pp.603–609
<https://doi.org/10.1007/s43681-022-00143-x>
- Mäntymäki, M., Minkkinen, M., Birkstedt, T. & Viljanen, M. (2023). Putting AI Ethics into practice: The Hourglass Model of Organizational AI Governance, preprint,
<https://arxiv.org/abs/2206.00335>
- Maslej, N., Fattorini, L., Perrault, R., Gil, Y., Parli, V., Kariuki, N., Capstick, E., Reuel, A., Brynjolfsson, E., Etchemendy, J., Ligett, K., Lyons, T., Manyika, J., Carlos Niebles, J., Shoham, Y., Wald, R., Walsh, T., Hamrah, A., Santarlasci, L., Betts Loyufo, J., Rome, A., Shi, A. & Oak, S. (2025). The AI Index 2025 Annual Report, Stanford: Stanford University, <https://doi.org/10.48550/arXiv.2504.07139>
- McMullin, C. (2023). Transcription and Qualitative Methods: Implications for Third Sector Research, *International Journal of Voluntary and Nonprofit Organizations*, vol 34, pp.140–153 <https://doi.org/10.1007/s11266-021-00400-3>
- Minkkinen, M. & Mäntymäki, M. (2023). The Institutional Logics Underpinning Organizational AI Governance practises, *14th Scandinavian Conference on Information Systems*, <https://aisel.aisnet.org/scis2023/12>
- Oates, B. J., Griffiths, M. & McLean, R. (2022). Researching Information Systems and Computing, 2nd edn., Los Angeles: SAGE
- Olsson, T. (2008). Medievardagen: En Introduktion till Kvalitativa Studier, Malmö: Gleerup
- Orton, J. D. & Weick, K. E. (1990). Loosely Coupled Systems: A Reconceptualization, *The Academy of Management Review*, vol. 15, no. 2, pp.203–223
<https://doi.org/10.2307/258154>
- Papagiannidis, E., Mikalef, P. & Conboy, K. (2025). Responsible Artificial Intelligence Governance: A Review and Research Framework, *The Journal of Strategic Information Systems*, vol. 34, no. 2, <https://doi.org/10.1016/j.jsis.2024.101885>
- Patton, M. Q. (2015). Qualitative Research & Evaluation Methods: Integrating Theory and practise, 4th edn., California: SAGE
- Powell, W. W. & DiMaggio, P. J (1991). The New Institutionalism in Organizational Analysis, 2th edn, Chicago: University Of Chicago Press
- Powell, W. W., & Bromley, P. (2015). New Institutionalism in the Analysis of Complex Organizations. In International Encyclopedia of the Social & Behavioral Sciences [e-book], 2nd edn, pp.764-769. Elsevier <https://doi.org/10.1016/B978-0-08-097086-8.32181-X>
- Rakova, B., Yang, J., Cramer, H. & Chowdhury, R. (2021). Where Responsible AI Meets Reality: Practitioner Perspectives on Enablers for Shifting Organizational Practices, *Proceedings of the ACM on Human-Computer Interaction*, vol. 5, no. CSCW1, pp.1–23
<https://doi.org/10.1145/3449081>
- Ransbotham, S., Khodabandeh, S., Fehling, R., LaFountain, B. & Kiron, D. (2019). Winning With AI, *MIT Sloan Management Review*, Available Online: <https://sloanreview.mit.edu/projects/winning-with-ai/> [Accessed 8 May 2026]

- Recker, J. (2021). Scientific Research in Information Systems: A Beginner's Guide [e-book], 2nd edn., Heidelberg; Springer Berlin <https://doi.org/10.1007/978-3-642-30048-6>
- Rivera, A., Abhari, K. & Xiao, B. (2025). Responsible AI Design: The Authenticity, Control, Transparency Theory, *Journal of the Association for Information Systems*, vol. 26, no. 5, pp.1337–1389 <http://dx.doi.org/10.2139/ssrn.5509678>
- Saunders, M., Lewis, P. & Thornhill, A. (2009). Understanding Research Philosophies and Approaches, vol. 4, pp.106–135 In: Saunders, M.N.K. Lewis, P, and Thornhill, A (eds) *Research Methods for Business Students*, 8th edn. Harlow: Pearson Education, pp. 128–171.
- Schneider, J., Abraham, R., Meske, C., & Vom Brocke, J. (2023). Artificial Intelligence Governance For Businesses. *Information Systems Management*, vol. 40 no. 3, pp. 229–249. <https://doi.org/10.1080/10580530.2022.2085825>
- Schultze, U. & Avital, M. (2011). Designing Interviews to Generate Rich Data for Information Systems Research, *Information and Organization*, vol. 21, no. 1, pp.1–16 <https://doi.org/10.1016/j.infoandorg.2010.11.001>
- Shneiderman, B. (2020). Bridging the Gap Between Ethics and practise: Guidelines for Reliable, Safe, and Trustworthy Human-Centered AI Systems, *ACM Transactions on Interactive Intelligent Systems*, vol. 10, no. 4, p.1-31 <https://doi.org/10.1145/3419764>
- Van Giffen, B. & Ludwig, H. (2023). How Boards of Directors Govern Artificial Intelligence, *MIS Quarterly Executive*, vol. 22 no. 4, pp.251–272 <https://doi.org/10.17705/2msqe.00085>
- Vassilakopoulou, P., Parmiggiani, E., Shollo, A. & Grisot, M. (2022). Responsible AI: Concepts, Critical Perspectives and an Information Systems Research Agenda, *Scandinavian Journal of Information Systems*, vol. 34, pp.89–112 <https://aisel.aisnet.org/sjis/vol34/iss2/3>
- Williams, R., Anderson, S., Cresswell, K., Kannelønning, M. S., Mozaffar, H. & Yang, X. (2024). Domesticating AI in Medical Diagnosis, *Technology in Society*, vol. 76 <https://doi.org/10.1016/j.techsoc.2024.102469>