

TALLINN UNIVERSITY OF TECHNOLOGY
School of Information Technologies

Aaryan Nair 245683IVCM

**INTERACTION OF LAW AND TECHNOLOGY - AI
ACCOUNTABILITY IN CYBERSECURITY UNDER THE EU
DIGITAL SERVICES ACT**

Master's Thesis

Supervisor: Aleksi Oskar Johannes Kajander
Lawyer and Legal Researcher

Tallinn 2026

TALLINNA TEHNIKAÜLIKOOL
Infotehnoloogia teaduskond

Aaryan Nair 245683IVCM

**ÕIGUSE JA TEHNOLOOGIA KOOSTÖÖ -
TEHISINTELLEKTI VASTUTUS
KÜBERTURVALISUSES EUROOPA LIIDU
DIGITEENUSTE SEADUSE RAAMES**

Magistritöö

Juhendaja: Aleks Oskar Johannes Kajander
Lawyer and Legal Researcher

Tallinn 2026

Author's Declaration of Originality

I hereby certify that I am the sole author of this thesis. All the used materials, references to the literature and the work of others have been referred to. This thesis has not been presented for examination anywhere else.

On the Use of Artificial Intelligence

Throughout the preparation of this thesis, AI-based tools were used in a limited and supportive capacity. In particular, the tools Grammarly and Microsoft Copilot (ChatGPT-based) were utilized.

Grammarly was used for spellchecking, grammatical correction, and assistance with sentence formulation to improve the clarity and readability of the text.

Microsoft Copilot was used in several complementary ways. First, it was utilized as a supplementary support tool to gain an initial understanding of unfamiliar concepts, particularly in areas related to AI governance, cybersecurity, and systemic risk. This use was limited to background comprehension and did not replace the consultation of academic or legal sources. Second, it assisted in the formatting and generation of BibTeX entries for sources where such entries were not readily available. Fourth, it supported limited refinement of academic language and expression across the thesis.

All outputs generated with the assistance of these tools were critically reviewed, verified, and revised by the author. No AI tool was used to generate original research findings, legal analysis, or conclusions. All substantive research, interpretation, and argumentation presented in this thesis are the independent work of the author, who retains full responsibility for its content.

Author: Aaryan Nair

20.05.2026

Abstract

This thesis examines whether Article 34 of the European Union’s Digital Services Act (DSA) effectively addresses cybersecurity risks arising from AI-driven content moderation and recommender systems on large online platforms. While the DSA introduces systemic risk assessment as a core governance mechanism, the research evaluates whether this procedural framework is capable of operationalising meaningful accountability for cybersecurity risks linked to algorithmic system design and behavior. Using a doctrinal and analytical methodology, the thesis combines a structured literature review with regulatory gap analysis and scenario-based stress testing. The findings show that although Article 34 represents a significant conceptual shift towards systemic risk governance, its application to AI-driven cybersecurity risks remains largely abstract. The analysis demonstrates that the under-capture of such risks cannot be attributed solely to weak or inconsistent platform implementation, but instead reflects structural features of Article 34, including open-textured risk definitions, reliance on platform self-assessment, absence of AI-specific technical obligations, and limited audit ability. Scenario-based stress testing further reveals that Article 34 struggles to detect and mitigate foreseeable cybersecurity risks under realistic conditions of scale, coordination, and crisis. The thesis concludes that while Article 34 currently functions primarily as a coordination and reporting mechanism, targeted refinements could significantly strengthen its capacity to govern AI-driven systemic cybersecurity risks without undermining the DSA’s flexible, risk-based regulatory model.

Annotatsioon
ÕIGUSE JA TEHNOLOOGIA KOOSTÖÖ - TEHISINTELLEKTI
VASTUTUS KÜBERTURVALISUSES EUROOPA LIIDU
DIGITEENUSTE SEADUSE RAAMES

Käesolev lõputöö uurib, kas Euroopa Liidu Digitaalteenuste määruse (Digital Services Act, DSA) artikkel 34 käsitleb tõhusalt küberturvalisusega seotud riske, mis tulenevad tehisintellektil põhinevatest sisu modereerimise ja soovitusüsteemidest suurtes veebi-platvormides. Kuigi DSA käsitleb süsteemsete riskide hindamist keskse platvormide reguleerimise mehhanismina, analüüsitakse töös, kas see protseduuriline raamistik suudab praktikas tagada sisulise vastutuse algoritmiliste süsteemide ülesehituse ja toimimisega seotud küberturberiskide eest. Töö kasutab dogmaatilist ja analüütilist uurimismeetodit, ühendades struktureeritud kirjanduse ülevaate regulatiivse lünga analüüsi ja stsenaariumipõhise stressitestimisega. Uurimistulemused näitavad, et kuigi artikkel 34 kujutab endast olulist kontseptuaalset nihet süsteemse riskijuhtimise suunas, jääb selle kohaldamine tehisintellektil põhinevate küberturberiskide osas suuresti abstraktseks. Analüüs osutab, et selliste riskide alahõlmatuse ei tulene üksnes platvormide ebaühtlasest või nõrgast rakendamisest, vaid peegeldab artikli 34 struktuursete omadusi, sealhulgas avatud riskimõisteid, ulatuslikku enesehindamisele tuginemist, tehisintellekti-spetsiifiliste tehniliste nõuete puudumist ja piiratud auditeeritavust. Lisaks näitab stsenaariumipõhine stressitestimine, et artikkel 34 ei suuda realistlikes mastaabi-, koordineeritus- ja kriisiolukordades ette nähtavaid küberturberiske järjepidevalt tuvastada ega maandada. Lõputöö järeldab, et kuigi artikkel 34 toimib praegu peamiselt koordineerimise ja aruandluse mehhanismina, võivad sihipärased täiendused oluliselt parandada selle suutlikkust reguleerida tehisintellektil põhinevaid süsteemseid küberturberiske, kahjustamata samas DSA paindlikku ja riskipõhist regulatiivset ülesehitust.

List of Abbreviations and Terms

VLOP - Very Large Online Platform/s
DSA- Digital Services Act
AI - Artificial Intelligence
EU - European Union
GDPR - General Data Protection Regulation
CRA - Cyber Resilience Act
SLR - Systematic Literature Review
IT - Information Technology
IDS - Intrusion Detection System
IPS - Intrusion Protection System
RGA – Regulatory Gap Analysis
SBST – Scenario-Based Stress Testing
RQ – Research Question
SRQ – Sub-Research Question
DPIA – Data Protection Impact Assessment
RTO – Recovery Time Objective
RPO – Recovery Point Objective
CERT – Computer Emergency Response Team
CSIRT – Computer Security Incident Response Team
ICT – Information and Communication Technology
API – Application Programming Interface
ML – Machine Learning
CIB - Coordinated Inauthentic Behavior

Artificial Intelligence (AI) Systems or models that perform tasks normally requiring human intelligence, including learning, pattern recognition, decision-making, and content classification. In this thesis, AI primarily refers to algorithmic systems used for content moderation, recommender systems, and automated decision-making on online platforms.

Snowballing is a literature search technique in which relevant sources are identified by examining the references and citations of initially selected papers. Rather than relying solely on keyword-based database searches, snowballing allows the researcher to systematically expand the body of literature by following citation links between academic works, thereby uncovering influential or otherwise overlooked studies within a research field.

Backward Snowballing refers to the process of reviewing the reference lists of selected articles in order to identify earlier works that informed the study. By tracing cited sources, the researcher can locate foundational or seminal literature that may not have been captured through initial search queries.

Forward Snowballing involves identifying newer publications that cite a selected article. This method allows the researcher to track the development of a topic over time and to identify more recent contributions that build upon earlier work.

Algorithmic System A set of computational procedures that process input data to produce outputs such as rankings, classifications, recommendations, or moderation decisions. Algorithmic systems may incorporate machine learning techniques and operate autonomously at scale.

Content Moderation The automated or manual process by which online platforms detect, evaluate, remove, demote, or otherwise restrict content based on legal requirements or platform rules. In this thesis, content moderation is examined primarily in its AI-driven and automated forms.

Cybersecurity In the context of this thesis, cybersecurity refers to risks arising from the design, functioning, and exploitation of AI-driven platform systems that may compromise system integrity, reliability, or resilience and produce large-scale economic, societal, or public-security harm. The term is used narrowly to denote systemic, design-level risks rather than isolated technical incidents.

Digital Services Act (DSA) Regulation (EU) 2022/2065 establishing a regulatory framework for online intermediary services in the European Union, with particular emphasis on platform accountability, systemic risk governance, and transparency obligations.

Implementation Failure A situation in which legal obligations are inadequately fulfilled in practice, for example through superficial compliance, incomplete assessments, or weak execution of required procedures, despite the existence of formal regulatory requirements.

Platform Governance The set of legal, procedural, and technical mechanisms through which online platforms are regulated, including risk assessment, mitigation, transparency, and oversight obligations imposed by EU law.

Recommender System An automated system that determines the ranking, prioritization, or presentation of content to users based on engagement signals, behavioral data, or inferred preferences. Recommender systems are a central focus of this thesis due to their role in amplifying systemic risks.

Regulatory Gap A discrepancy between the risks addressed by a legal framework and the risks that arise in practice, where certain foreseeable harms are insufficiently regulated due to structural limitations, ambiguity, or absence of specific obligations.

Regulatory Gap Analysis (RGA) A methodological approach used to identify and analyze structural deficiencies in a legal framework by comparing normative regulatory expectations with observed practices and outcomes.

Scenario-Based Stress Testing (SBST) A methodological tool that evaluates the effectiveness of legal obligations under realistic, high-impact scenarios by testing whether regulatory mechanisms would detect, mitigate, or audit systemic risks under adverse conditions.

Structural Regulatory Gap A limitation arising from the design of a legal provision itself, where foreseeable risks are not required to be identified or mitigated even under

conditions of good-faith implementation.

Systemic Risk A risk that arises from the design, functioning, or scale of a service such that it can cause widespread or cascading harm to society, the economy, or public security. Under Article 34 DSA, systemic risks are associated with platform design choices rather than isolated incidents.

Very Large Online Platform (VLOP) An online platform designated under the DSA with at least 45 million monthly active users in the European Union, subject to enhanced systemic-risk assessment, mitigation, and audit obligations.

Table of Contents

1	Introduction	11
1.1	Background and Motivation	11
1.2	Overview of legal challenges in data protection and liability with the growth of AI and its accountability	12
1.3	Conceptual Definition of Cybersecurity in the Context of Article 34 DSA	12
2	Literature Review	14
2.1	Research Questions	14
2.2	Methodology and Reasoning	14
2.2.1	Method of selection also included	16
2.2.2	Reasoning for RQs	17
2.3	Expected Outcome	17
2.4	Selected Papers for the Systematic Literature Review	19
2.5	Observation	38
2.6	Summary of the Review	39
3	Discussion	47
3.1	Analysis and Discussion	47
4	Regulatory Gap Analysis	49
4.1	Regulatory Gap Analysis	49
4.1.1	Conceptual Gaps	49
4.1.2	Procedural Gaps	50
4.1.3	Technical Gaps	50
4.1.4	Enforcement Gaps	51
4.1.5	Interim Conclusion	51
5	Scenario Based Stress Testing	53
5.1	Scenario-Based Stress Testing of AI-Driven Systemic Risks	53
5.1.1	Purpose and Role of Scenario-Based Stress Testing	53
5.1.2	Selection of Scenarios for Stress Testing	53
5.1.3	Analytical Logic Applied Across Scenarios	54
5.1.4	Scenario 1: Meta Recommender Systems and the Cambridge Analytica Scandal	55
5.1.5	Scenario 2: YouTube Recommender Algorithms and Harmful Content Amplification	57

5.1.6	Scenario 3: Coordinated Inauthentic Behavior and AI Detection Failures at Meta	59
5.1.7	Cross-Scenario Findings	61
6	Strengthening the Systemic-Risk Framework of Article 34	62
6.1	Purpose and Scope of the Chapter	62
6.2	Clarifying AI-Driven Cybersecurity Risks as Systemic Risks	62
6.3	Introducing Minimum AI-Specific Risk Assessment Requirements	63
6.4	Strengthening the Link Between Risk Identification and Mitigation Obligations	63
6.5	Enhancing Audit-ability and Supervisory Oversight	64
6.6	Interim Reflections	64
7	Conclusion	65
	References	67

List of Figures

Table 1 : Keywords Used

Table 2 : Summary of Relevant Research Studies and Alignment with Research Questions

1. Introduction

1.1 Background and Motivation

The rapid growth of digital technology has significantly complicated the relationship between law and innovation. Legal systems now struggle to keep pace with fast-evolving tools such as artificial intelligence, blockchain, and autonomous systems, which challenge traditional legal assumptions surrounding human responsibility, jurisdiction, and transparency. These technologies often outpace the ability of existing legal frameworks to address critical concerns such as privacy, security, and accountability. In response, the law must not only regulate but also actively shape the development of emerging technologies, ensuring that they align with public interests, human rights, and ethical standards. This evolving intersection between law and technology demands flexible, interdisciplinary approaches to modern governance.

Moreover, the relationship between law and technology reflects a deeper inquiry into how societies govern risk, power, and human agency in increasingly digital environments. Technology is not merely a tool but a transformative force that reshapes how information is processed, how systems operate, and how vulnerabilities emerge and propagate. In this dynamic landscape, law functions not only as a regulatory mechanism but also as a framework for ensuring the security, resilience, and trustworthiness of technologically mediated systems. As artificial intelligence, platform infrastructures, and algorithmic decision-making become deeply embedded in critical digital ecosystems, cybersecurity concerns extend beyond isolated technical failures to encompass systemic risks inherent in the design and operation of these systems. Vulnerabilities in AI-driven platforms—whether arising from adversarial manipulation, automation biases, or large-scale coordination effects—can produce cascading consequences that affect economic stability, public security, and democratic processes. Traditional legal approaches to accountability, which focus on discrete incidents and identifiable actors, are increasingly insufficient to address these distributed and structural risks. Accordingly, the challenge for contemporary legal systems is not only to regulate technological misuse, but to develop frameworks capable of anticipating, mitigating, and governing cybersecurity risks that emerge at the systemic level, ensuring that digital infrastructures remain secure, accountable, and aligned with broader societal interests.

1.2 Overview of legal challenges in data protection and liability with the growth of AI and its accountability

The rapid advancement of artificial intelligence (AI) has introduced significant legal challenges in data protection and liability. AI systems often process vast amounts of personal data, raising complex questions about compliance with data protection laws such as the EU GDPR, DSA, NIS2. One major issue is ensuring transparency and explainability in AI-driven decision-making. Many AI models operate as "black boxes," making it difficult for individuals to understand how their data is used and for regulators to determine whether processing is lawful.

In terms of liability, AI challenges the traditional frameworks that rely on clear human agency. Autonomous systems can act unpredictably or in ways not foreseen by developers, complicating the assignment of legal responsibility. Key concerns include whether liability should fall on developers, users, or the AI system itself, particularly in high-risk areas like healthcare, finance, or cybersecurity. The EU proposed AI Liability Directive and AI Act aim to address these gaps by imposing stricter obligations on high-risk AI providers and by adapting liability rules for AI-related damage.

Ultimately, balancing innovation with accountability is crucial. Legal frameworks must evolve to ensure robust data protection and fair liability allocation without stifling technological progress or undermining public trust in AI.

1.3 Conceptual Definition of Cybersecurity in the Context of Article 34 DSA

The term *cybersecurity* is used in this thesis in a deliberately narrow and functional sense, tailored to the systemic-risk framework established by Article 34 of the Digital Services Act (DSA). Rather than adopting a general or infrastructure-oriented definition commonly associated with EU cybersecurity law (e.g., NIS2), this thesis conceptualizes cybersecurity as a category of *systemic risk arising from the design, functioning, and exploitation of AI-driven platform systems*.

In this context, cybersecurity refers to risks that affect the integrity, reliability, and resilience of AI-enabled information processing and decision-making systems—such as recommender systems, automated content moderation, and generative AI tools - where failures or manipulation of these systems can produce large-scale economic, societal, or public-security harms. These risks include, but are not limited to, adversarial exploitation

of algorithmic systems, coordinated manipulation of recommender mechanisms, degradation of automated moderation accuracy, and cascading failures caused by feedback loops within platform architectures.

Importantly, this definition distinguishes between three analytically distinct categories of risk. First, *technical system risks* concern vulnerabilities inherent in AI models and infrastructures, such as robustness failures or susceptibility to adversarial attacks. Second, *systemic operational risks* arise when AI systems function as designed but generate harmful outcomes at scale due to optimization choices, amplification dynamics, or interaction effects. Third, *societal and public-security risks* emerge when these technical and operational failures undermine economic stability, public order, or democratic processes.

The thesis excludes isolated cybersecurity incidents; such as insider threats or data breaches, unless they contribute to system-wide effects. The analytical focus is therefore on cybersecurity risks that are *systemic by nature*, meaning that they originate from or are materially exacerbated by platform design choices and AI-driven automation. By adopting this definition, the thesis aligns the concept of cybersecurity with the regulatory logic of Article 34 and clarifies that the analysis concerns structural risk governance rather than general IT security or compliance failures.

2. Literature Review

To better understand these dynamics, a systematic literature review (SLR) was conducted to identify, analyze, and synthesize academic and practical research on the intersection of law and technology, AI, and cybersecurity frameworks and regulations laid down.

2.1 Research Questions

The following research questions are designed to support the thesis and are designed to explore the growth of law and technology, and explore the growing controversial path of AI accountability in cybersecurity.

RQ1: How do the systemic risk assessment obligations under Article 34 of the Digital Services Act (DSA) account for cybersecurity risks originating from AI-driven content moderation and recommender systems?

- **SRQ1:** What legal and procedural mechanisms could operationalise AI accountability and liability for cybersecurity failures within the DSA framework?

2.2 Methodology and Reasoning

The literature review does not follow a strict systematic literature review protocol, but instead adopts a structured and thematic approach consistent with doctrinal legal research,[1, 2] with emphasis on main objectives, aimed at identifying how existing academic, regulatory, and policy-oriented sources conceptualize systemic risk and cybersecurity in the context of AI-driven digital platforms. This approach is appropriate given that the research focuses on the interpretation and evaluation of a legal framework — *Article 34 of the Digital Services Act* — rather than on empirical measurement or quantitative modeling.

The review is based on a structured selection of sources, including peer-reviewed journal articles, institutional reports, policy papers, and official EU documentation. Sources were selected according to three main criteria

- Relevance to the research questions,
- Availability and assurance of the authenticity of papers,
- Emphasis on AI accountability and its increasing misuse.
- Conceptual engagement with systemic risk or AI governance, and

- Recency in light of the rapidly evolving regulatory landscape.

Particular emphasis was placed on literature addressing recommender systems, algorithmic governance, and systemic risk assessment under EU law.

The analytical method applied in this review is thematic synthesis. Rather than summarizing sources individually, the literature is examined through recurring themes, including the procedural nature of systemic-risk regulation, the role of platform discretion, and the limitations of existing accountability mechanisms for AI-driven systems. These themes are then used to identify conceptual and regulatory gaps that inform the subsequent Regulatory Gap Analysis and Scenario-Based Stress Testing.

This methodological approach is consistent with established doctrinal legal research practices, which emphasize the systematic interpretation of legal norms in light of secondary sources and theoretical frameworks[3]. As noted in legal scholarship, doctrinal analysis remains a central method for evaluating the adequacy of regulatory structures and identifying areas where legal frameworks may fail to address emerging technological risks [4, 3].

By combining doctrinal interpretation with thematic analysis of interdisciplinary literature, the review provides a foundation for assessing whether Article 34 is capable of governing AI-driven cybersecurity risks in practice, and supports the evaluative framework developed in later chapters.

The literature search was conducted across several tools and academic databases, including:

- **IEEE Xplore**
- **ScienceDirect**
- **TalTech Digikogu**
- **Google/Google Scholar**

These tools and databases were selected based on prior knowledge of their usability and relevance to cybersecurity, information systems, and academic research. Tools and some databases need to be considered as search engines, rather than real databases, since they take their results from other databases.

2.2.1 Method of selection also included

- In addition to database searches, snowballing was used to ensure comprehensive coverage of the research domain.
- Both forward snowballing and backward snowballing [5] were employed to trace the conceptual evolution of cybersecurity law, AI accountability, and data protection frameworks.
- Snowballing was conducted using key articles identified through initial keyword-based searches on themes such as “systemic risk”, “Article 34 DSA”, “AI governance”, and “recommender systems”. Backward snowballing involved reviewing the reference lists of these core publications, while forward snowballing focused on identifying more recent studies citing them, thereby enabling the identification of influential and related literature beyond database search results.
- This approach enabled the identification of influential studies that might not be apparent through standard keyword-based queries.
- Furthermore, primary legal sources, including official EU regulatory texts such as the NIS2 Directive (2022) and the General Data Protection Regulation (GDPR), was incorporated to provide contextual and policy-oriented depth.
- These sources, though not peer-reviewed, were essential in analyzing the legal and institutional foundations underpinning academic discourse on cybersecurity and AI regulation.

The keywords selected are seen in **Table 1**

Article 34	DSA	AI	String / Keyword
Technology	AI	DSA	String / Keyword
Cybersecurity	Liability	DSA	String / Keyword
Systemic Risk	AI	DSA	String / Keyword
AI governance	Recommender systems	Article 34	String / Keyword

Table 1. Keywords Used

- The table above presents the different search strings and keywords that were used to identify relevant sources and data for this study. Each search string contributed to the identification of at least one source that was considered relevant to the research objectives. The use of structured string searches enabled the systematic retrieval of literature across multiple databases, ensuring that the selection process remained focused and aligned with the scope of the thesis.
- The search strings were designed to capture highly topic-specific literature related to

systemic risk, AI governance, and cybersecurity within the context of the Digital Services Act.

2.2.2 Reasoning for RQs

- This research question is not merely asking whether Article 34 DSA mentions cybersecurity. Rather, it probes a deeper structural issue:
 - Whether the DSA’s systemic risk framework is conceptually and practically capable of capturing cybersecurity threats that arise indirectly through AI-driven content moderation and recommender systems.
- Article 34 obliges Very Large Online Platforms (VLOPs) to assess systemic risks related to:
 - the dissemination of illegal content,
 - negative effects on fundamental rights,
 - civic discourse and public security.
- SRQ1 shifts the focus from risk identification to responsibility allocation. It asks how abstract DSA obligations can be translated into concrete accountability and liability mechanisms for AI-related cybersecurity failures. The underlying concern is that:
- AI systems diffuse responsibility across developers, deployers, and operators.
- The DSA relies heavily on procedural compliance (risk assessments, audits, mitigation measures), which may dilute causal attribution when cybersecurity failures occur.

2.3 Expected Outcome

The expected outcome of this research is to determine whether Article 34 of the Digital Services Act is capable of effectively identifying and governing cybersecurity risks arising from AI-driven content moderation and recommender systems on large online platforms. Based on the existing literature and the structure of the DSA, the analysis anticipates that while Article 34 provides an important conceptual framework for addressing systemic risk, its practical application to AI-driven cybersecurity risks will reveal significant limitations. More specifically, the research expects to show that the procedural nature of Article 34, combined with its reliance on platform-led risk assessment and the absence of clearly defined technical or AI-specific requirements, results in an under-identification of risks linked to algorithmic system design and behavior. The literature suggests that systemic risks arising from recommender systems, targeting architectures, and automated moderation are often diffuse, cumulative, and embedded within normal system operation, making them

difficult to detect through high-level, self-reported assessments.

The analysis further anticipates that the shortcomings of Article 34 cannot be explained solely by weak or inconsistent implementation practices, but instead reflect structural characteristics of the regulatory framework itself, including open-ended risk definitions and limited mechanisms for enforcing substantive risk evaluation. Through the application of Regulatory Gap Analysis and Scenario-Based Stress Testing, the thesis expects to demonstrate that AI-driven cybersecurity risks can emerge even where systems function as intended and in compliance with existing rules, thereby challenging the effectiveness of procedural risk governance.

Ultimately, the research aims to identify targeted and legally feasible refinements that could enhance Article 34's ability to operate as a genuinely preventive instrument. These refinements are expected to focus on improving the recognition of AI-driven systemic risk, strengthening the connection between risk assessment and mitigation, and enabling more meaningful oversight of platform behavior, while preserving the flexibility of the DSA's risk-based approach.

2.4 Selected Papers for the Systematic Literature Review

1) Balcioğlu, Çelik, Altındağ (2025) - A turning point in AI: Europe's human-centric approach to technology regulation

Introduction: Balcioğlu et al. (2025)[6] explore Europe's human-centric approach to technology regulation, focusing on aligning innovation with ethical principles and responsible AI oversight. They emphasize the European Union's leadership in ensuring that digital technologies reflect social and legal accountability.

Why chosen: This paper was chosen for its relevance to the thesis's emphasis on ethical-legal frameworks and its interest in ensuring safety when using AI in cybersecurity. It clarifies the EU's attempt to operationalise values—trust, fairness, and accountability—within regulatory instruments such as the AI Act.

Research gap: The paper lacks a sector-specific focus on AI-driven cybersecurity and does not address liability allocation when AI systems malfunction. The absence of DSA infrastructure integration limits its utility for accountability analysis.

2) Brinker (2024) - Identification and Demarcation—A General Definition and Method to Address Information Technology in European IT Security Law

Introduction: Brinker (2024)[7] defines and clarifies “information technology” within EU security law. The paper develops a legal-interpretive framework to delineate technological systems and ensure consistent application of IT legislation.

Why chosen: This work provides conceptual precision for identifying how AI-based systems fit within cybersecurity regulation. It helps frame the legal taxonomy underpinning AI liability debates.

Research gaps: The author does not connect definitional insights to operational or accountability contexts. There is minimal engagement with AI regulation or enforcement under evolving cybersecurity directives.

3) Beltrán, Marta (2025) — AI Algorithms under Scrutiny: GDPR, DSA, AI Act and CRA as Pillars for Algorithmic Security and Privacy in the EU

Introduction: This article [8] evaluates the combined effect of four primary EU frameworks (GDPR, Digital Services Act, AI Act, Cyber Resilience Act) on algorithmic security and privacy [8][9]. It assesses overlaps, potential redundancies, and complementarities. It proposes a conceptual model for how layered regulation can (or cannot) produce coherent safeguards for algorithmic systems in digital services and critical products.

Why chosen: Beltrán’s work explicitly treats that interaction and is therefore directly useful in showing where NIS2/AI Act/GDPR/CRA align or collide. It helps you articulate precise legal friction points (e.g., when a conformity assessment under CRA overlaps with risk-management duties under NIS2 or transparency duties under GDPR).[9]

Research gaps: The piece is synthetic and high-level, it diagnoses misalignment but stops short of proposing enforceable liability models or concrete regulatory drafting fixes. It also tends to treat “algorithmic systems” as an abstract class and provides limited operational guidance for AI used in cybersecurity operations (IDS/IPS automation) or for sectoral enforcement practice. This leaves room for your thesis to propose concrete amendments or contractual/legal mechanisms tailored to cyber-defensive AI.

4) Chiara (2024) — Towards a Right to Cybersecurity in EU Law? The Challenges Ahead

Introduction: Chiara (2024)[10] explores whether a distinct “right to cybersecurity” should be recognised under EU law. The author evaluates human-rights frameworks and the EU’s constitutional order, arguing that cybersecurity is essential to digital autonomy and societal resilience.

Why chosen: This paper extends the thesis toward the rights-based dimension of cybersecurity, framing protection not only as compliance but as entitlement. It provides a normative foundation for integrating human rights reasoning into AI governance discourse.

Research gaps: Chiara’s [10]proposal lacks a procedural pathway for enforcement and omits AI-specific liability contexts. The study is primarily conceptual, leaving the institutional and cross-border implications unexplored. The absence of operationalization guidance limits its applicability to cybersecurity regulation.

5) Cockfield (2003) - Towards a Law and Technology Theory CALT Conference: What Is Legal Knowledge

Introduction: Cockfield (2003)[11] presents an early theoretical inquiry into the interaction between law and technology. The article conceptualizes “technological determinism” and its legal responses, arguing that statutes both shape and adapt to technological change.

Why chosen: Included to provide historical and theoretical grounding, this paper contextualizes modern AI and cybersecurity debates within long-standing questions about regulatory adaptability. It underpins the thesis’s examination of how legal frameworks evolve alongside emerging technologies.

Research gaps: Being two decades old, the paper predates AI-specific challenges and cyber-governance instruments. It lacks empirical validation and a modern regulatory context. The framework’s generality offers limited guidance for addressing today’s AI liability and cybersecurity accountability concerns.

6) Cockfield and Pridmore (2007) - A Synthetic Theory of Law and Technology Symposium: Towards a General Theory of Law and Technology

Introduction: Cockfield and Pridmore (2007)[12] offer a unique perspective by synthesizing theories of law and technology, explaining how legal and technical systems co-evolve within dynamic social contexts. Their argument for an adaptive, interdisciplinary regulatory approach is a significant contribution to the field.

Why chosen: This work, with its theoretical continuity, explains why AI governance demands flexible, principle-driven legal responses. It enriches the conceptual base for analyzing cybersecurity accountability and engages the audience intellectually.

Research gap: The authors offer no specific applications to AI or cybersecurity. Their argument remains abstract, lacking empirical grounding or discussion of liability.

7) Charanjit Singh (2023) - European Cyber Security Law in 2023: A Review of the Advances in Network and Information Security Directive 2

Introduction: The 2023[13] review summarizes developments in European cybersecurity law after the adoption of the NIS2 Directive. It outlines expanded regulatory coverage and enhanced reporting obligations, providing practical insights into the evolving legal landscape.

Why chosen: This review establishes a baseline understanding of current EU cybersecurity legislation and clarifies the institutional mechanisms shaping AI-related regulation.

Research gap: The analysis remains descriptive and excludes AI implications, accountability, and liability allocation. It lacks comparative or evaluative analysis.

8) Fischer-Hübner et al. (2021) - Stakeholder Perspectives and Requirements on Cybersecurity in Europe

Introduction: Fischer-Hübner et al. (2021)[14] capture stakeholder views on cybersecurity policy through interviews and surveys. They identify calls for transparency, accountability, and cross-sector collaboration.

Why chosen: This paper adds empirical evidence on how institutional and societal actors perceive regulatory gaps, grounding the findings in stakeholder-based perspectives. This unique perspective is a valuable addition to the literature.

Research gap: It lacks doctrinal analysis or policy translation. AI-specific issues are unaddressed, and the connection to legal enforcement is underdeveloped.

9) Gifford (2007) - Law and Technology: Interactions and Relationships Symposium: Towards a General Theory of Law and Technology

Introduction: Gifford (2007)[15] examines mutual influences between law and technology, focusing on jurisprudential adaptability to innovation. The work highlights continuous reform as a condition for legal relevance.

Why chosen: It provides conceptual reinforcement of the idea that adaptive legal systems are not just important but essential for managing emerging cybersecurity risks and AI accountability, highlighting the urgency of the topic.

Research gap: The paper is theoretical and predates AI and digital rights frameworks. It omits any analysis of modern regulatory practice.

10) Jørgensen, Gunasekaran and Ma (2025) - Impact of EU Laws on AI Adoption in Smart Grids: A Review of Regulatory Barriers, Technological Challenges, and Stakeholder Benefits

Introduction: Jørgensen et al. (2025)[16] review legal and technical challenges facing AI adoption in smart-grid systems. They link data protection, cybersecurity, and operational efficiency concerns under EU regulation.

Why chosen: It demonstrates sectoral interaction between innovation and regulation, mirroring AI's role in cybersecurity operations.

Research gap: The study lacks comparative transposition or analysis of legal liability. It remains limited to the energy sector without broader AI-cybersecurity governance synthesis.

11) Mal'ko and Kostenko (2022) - Legal Techniques and Legal Technology: The Complex Nature of the Interaction

Introduction: Mal'ko and Kostenko (2022)[17] discuss the complex relationship between legal techniques and legal technology, analysing how technological processes reshape legislative and enforcement methods. They conceptualise "legal technology" as both a tool and an evolving social process that modifies legal interactions.

Why chosen: This article was included in the review for its examination of how technology alters legal reasoning and enforcement, and for offering a philosophical foundation for understanding technology's interaction in cybersecurity regulation.

Research gap: The paper remains largely theoretical, with minimal application to cybersecurity or AI systems. It also lacks empirical evidence on the integration of legal technologies into data protection or accountability mechanisms.

12) Manolopoulos (2003) - Raising 'Cyber-Borders': The Interaction Between Law and Technology

Introduction: Manolopoulos (2003)[18] explores the metaphor of "cyber-borders," addressing how law attempts to reassert control in cyberspace through regulatory frameworks. The article analyses how technological change challenges sovereignty and enforcement capacity.

Why chosen: This paper establishes a historical and jurisprudential link between legal jurisdiction and digital regulation, which is relevant to how current EU cybersecurity frameworks seek to manage cross-border AI systems.

Research gap: The work predates modern instruments like the GDPR and NIS2. It does not account for AI liability or global regulatory coherence, limiting its applicability to today's complex cyber-legal environment.

13) Mitrou (2021) - Data Protection, Artificial Intelligence and Cognitive Services

Introduction: Mitrou (2021)[19] evaluates the interaction between AI, data protection, and cognitive computing services under the GDPR. The paper investigates how automated systems challenge existing data-processing principles and accountability models.

Why chosen: It contributes directly to the AI–data protection interface, a crucial area of research, by analysing the sufficiency of regulatory frameworks in privacy and cybersecurity contexts.

Research gap: Although comprehensive, the study lacks a focused treatment of sectoral cybersecurity and omits liability mechanisms when AI systems cause security breaches.

14) Novelli et al. (2024) - Generative AI in EU Law: Liability, Privacy, Intellectual Property, and Cybersecurity

Introduction: Novelli et al. (2024)[20] examine AI’s intersection with EU law, focusing on liability, privacy, intellectual property, and cybersecurity. They propose legal harmonisation to address the multidimensional risks posed by generative AI.

Why chosen: This multidisciplinary paper is central to the paper because it directly addresses AI liability and cybersecurity governance. It synthesises multiple legal instruments to illustrate systemic fragmentation.

Research gap: Despite its breadth, it offers limited empirical data on enforcement or judicial interpretation. Its harmonisation proposal remains conceptual without implementation strategies.

15) Papakonstantinou (2022) - Cybersecurity as Praxis and as a State: The EU Law Path towards Acknowledgement of a New Right to Cybersecurity?

Introduction: Papakonstantinou (2022)[21] investigates the evolution of cybersecurity as a legal and conceptual construct within EU law, arguing that cybersecurity is emerging as a distinct legal right. The article links governance, enforcement, and citizen protection.

Why chosen: This study bridges cybersecurity policy and rights-based interpretation, offering a doctrinal basis for a proposal that accountability must align with legal entitlements.

Research gap: The paper lacks practical implementation recommendations and does not explicitly discuss AI systems. It omits aspects of enforcement and liability related to automated cybersecurity tools.

Relevance to research questions: It supports RQ1 by revealing conceptual progress toward recognizing cybersecurity as a legal norm. It strengthens SRQ4 by identifying gaps in translating this right into enforceable accountability mechanisms for AI-enabled systems.

16) Rampášek, Mesarčík and Andraško (2025) - Evolving Cybersecurity of AI-Featured Digital Products and Services: Rise of Standardization and Certification?

Introduction: Rampášek et al. (2025)[22] analyze the growing standardization and certification initiatives for AI-enabled digital products within EU cybersecurity law. They emphasize the importance of harmonized technical standards for ensuring trust and interoperability.

Why chosen: This study connects regulatory and technical governance, directly supporting examination of how law operationalization cybersecurity for AI systems through certification and compliance mechanisms.

Research gap: The paper remains focused on standardization without addressing legal liability or enforcement under NIS2 or the AI Act. It also neglects to assess whether certification guarantees effective accountability.

17) Tamò-Larrieux, Guitton, Mayer and Lutz (2024) - Regulating for Trust: Can Law Establish Trust in Artificial Intelligence?

Introduction: Tamò-Larrieux et al. (2024)[23] explore how law can cultivate public trust in AI systems. They propose "trust-by-design" approaches that integrate transparency, explainability, and human oversight into governance.

Why chosen: This paper enriches the paper's foundation by linking normative governance and legal accountability, highlighting how societal trust influences regulatory effectiveness in AI-driven cybersecurity.

Research gap: Although comprehensive, the study remains conceptual and lacks specific references to cybersecurity. It omits analysis of enforcement tools or liability distribution when AI systems cause harm.

18) Ufert (2020) - AI Regulation Through the Lens of Fundamental Rights: How Well Does the GDPR Address the Challenges Posed by AI?

Introduction: Ufert (2020)[24] assesses how well the GDPR addresses AI's challenges from a fundamental rights perspective. The article analyses whether the GDPR adequately addresses automation, bias, and algorithmic opacity.

Why chosen: It supports discussion of data protection and accountability, providing doctrinal depth on how privacy law adapts to AI and cybersecurity needs.

Research gap: The paper foreshadows the AI Act and NIS2 and lacks insights into establishing grounds in current frameworks. It does not cover sectoral cybersecurity applications and AI liability mechanisms.

19) Zeng, Yunis, Khalil and Mirza (2024) - Towards a Conceptual Framework for AI-Driven Anomaly Detection in Smart City IoT Networks for Enhanced Cybersecurity

Introduction: Zeng et al. (2024) [25] propose a conceptual framework for AI-driven anomaly detection in smart-city IoT networks. The paper highlights how AI enhances predictive cybersecurity through automation and adaptive learning.

Why chosen: It introduces the technical dimension, which is crucial for demonstrating how AI technologies operate in cybersecurity and how legal frameworks must adapt to their complexity. It misses key aspects covered in this paper, yet it provides a technical understanding of AI's functionality and the potential of AI in cybersecurity when adequately regulated.

Research gap: The article lacks legal analysis and does not discuss accountability or liability for automated cybersecurity actions.

20) Griffin (2025) Governing Platforms Through Corporate Risk Management: The Politics of Systemic Risk in the DSA [26]

Introduction: Griffin critically examines the governance logic of Article 34, arguing that systemic risk under the DSA is a socially and politically constructed concept that privileges technocratic expertise and delegates significant discretion to platforms.

Why Chosen: This work directly interrogates the institutional design of Article 34 and explains why platform-led risk assessments may downplay AI-driven cybersecurity risks.

Research Gap: The article does not address technical AI vulnerabilities such as adversarial attacks or model manipulation in content moderation systems.

21) Palumbo (2025) Systemic Risk Management and Constitutional Limits of Delegation under the DSA and AI Act [27]

Introduction: Palumbo analyzes Article 34 through EU constitutional law, arguing that the DSA delegates excessive normative discretion to platforms and regulators, raising concerns under principles such as the Meroni doctrine.

Why Chosen: Palumbo's work is very relevant for assessing whether Article 34 provides sufficiently constraining legal standards for platform risk assessments involving AI systems.

Research Gap: The analysis remains legal-constitutional and does not examine AI-specific cybersecurity threats or technical enforcement challenges.

22) Loi, Fabbri and Ferrario (2025) Regulating the Undefined: Addressing Systemic Risks in the DSA [28]

Introduction: The authors identify a ‘convergence problem’ in defining systemic risk and propose a dual-track framework distinguishing between research flexibility and compliance clarity.

Why chosen: It was chosen for its direct engagement with the conceptual shortcomings of Article 34 systemic-risk obligations.

Research Gap: The framework is not applied to AI-driven cybersecurity threats or automated moderation systems.

23) Kaesling and Wolf (2025) — Sustainability and Risk Management under the DSA [29]

Introduction: The authors clarify the legal meaning of ‘systemic’ risk in Article 34 by examining unlisted risk categories such as sustainability.

Why chosen: Its interpretive analysis supports extending Article 34 to technical AI-related harms.

Research Gap: The study does not engage directly with AI systems or cybersecurity vulnerabilities.

24) Torraco (2025) — Risk in the DSA and AI Act: Implications for Media Freedom and Disinformation [30]

Introduction: Torraco examines how recommender systems generate systemic risks to media pluralism and public discourse under the DSA.

Why chosen: Recommender-system analysis overlaps with AI-driven exploitation and manipulation relevant to cybersecurity risks.

Research gap: The paper focuses on informational harms rather than technical cyber vulnerabilities.

25) EU Commission (2025) — First Report on the Landscape of Systemic Risks [31]

Introduction: The European Commission’s First Report on the Landscape of Systemic Risks (2025) provides an overarching regulatory assessment of systemic risks posed by Very Large Online Platforms (VLOPs) under the Digital Services Act (DSA). It identifies recurring risk patterns, including disinformation, societal polarization, and emerging risks linked to generative AI and automated recommender systems. The report frames these risks as structural rather than isolated failures, emphasizing their cross-platform and cross-border nature.

Why chosen: This report is selected because it represents the Commission’s authoritative interpretation of what constitutes “systemic risk” under Article 34 DSA, effectively setting the baseline for compliance expectations. It signals regulatory priorities and enforcement focus areas for VLOPs, particularly regarding AI-driven content moderation and amplification systems. As such, it is a key reference point for understanding how regulators conceptualize and operationalizes systemic risk at the EU level.

Research gap: While the report offers a high-level taxonomy of systemic risks, it does not establish binding or standardized technical criteria for assessing AI systems’ contributions to those risks. It leaves unresolved how risks should be measured, compared, or audited in practice across different platform architectures. This gap creates uncertainty for both regulators and platforms when translating qualitative risk categories into concrete technical and governance controls.

26) Hacker, Kasirzadeh and Edwards (2025) — AI, Digital Platforms, and the New Systemic Risk [32]

Introduction: Hacker, Kasirzadeh, and Edwards (2025) propose a structured four-level framework for conceptualizing AI-related systemic risk, ranging from technical failures to broader societal and institutional harms. The paper situates AI-driven risks within the context of digital platforms, emphasizing how scale, automation, and feedback loops can transform localized failures into systemic threats. It also offers a comparative analysis of how the DSA and the AI Act address these risks from distinct regulatory angles.

Why chosen: This paper is selected because it provides one of the most comprehensive theoretical models for understanding AI-driven systemic risks, particularly in platform environments. Its multi-level framework helps bridge technical, organizational, and societal dimensions of risk, making it highly relevant for academic analysis of AI governance. The comparative treatment of the DSA and AI Act further enriches the discussion by highlighting regulatory complementarities and tensions.

Research gap: Despite its strong conceptual contribution, the paper does not engage in a detailed analysis of enforcement mechanisms under Article 34 DSA. In particular, it leaves unexplored how systemic risk assessments are operationalized, monitored, and validated by regulators in practice. This limits its ability to inform concrete compliance strategies for platforms deploying high-impact AI systems.

27) Panigutti et al. (2025) — Algorithm-Driven Risk Investigation in Online Platforms
[33]

Introduction: Panigutti et al. (2025) develop empirical audit methodologies for identifying and measuring risks arising from algorithmic recommender and content moderation systems. The paper focuses on observable harms such as bias amplification, misinformation spread, and unfair content suppression, using data-driven and experimental approaches. It situates algorithmic auditing as a necessary complement to platform self-assessment in large-scale digital environments.

Why chosen: This work is chosen because it offers concrete technical methodologies—such as behavioral testing and outcome-based audits—that are largely absent from the high-level risk assessment requirements under Article 34 DSA. It demonstrates how systemic risks generated by AI systems can be empirically detected rather than merely described. As such, it directly addresses the operational gap between regulatory obligations and technical implementation.

Research gap: Despite its strong methodological contribution, the paper does not explicitly map its proposed audit techniques onto specific DSA legal obligations or compliance workflows. It remains unclear how these methods could be institutionalized within mandatory risk assessments, regulatory reporting, or supervisory oversight. This limits its direct applicability for platforms seeking legally robust DSA compliance strategies.

28) Knight-Georgetown Institute (2025) — Systemic Risk Assessment under the DSA
[34]

Introduction: The Knight–Georgetown Institute’s 2025 report evaluates early systemic risk assessments submitted by VLOPs under Article 34 DSA. It finds that many assessments lack methodological rigor, rely on vague qualitative descriptions, and insufficiently analyze how platform design and algorithmic choices contribute to systemic harms. The report highlights a disconnect between statutory requirements and current industry practice.

Why chosen: This report is selected because it provides one of the first empirical evalu-

ations of how Article 34 is being implemented in practice. By systematically assessing real-world risk assessment documents, it demonstrates structural weaknesses in current compliance approaches. It therefore serves as an important reality check on the effectiveness of the DSA's systemic risk governance framework.

Research gap: While the report clearly documents deficiencies in risk assessment quality, it does not explore the downstream legal consequences of such failures. In particular, it leaves unexamined how weak assessments might trigger liability, sanctions, or enhanced regulatory intervention under the DSA. This omission limits its usefulness for analyzing enforcement dynamics and accountability mechanisms.

29) Eder, Niklas (2024) Making Systemic Risk Assessments Work: How the DSA Creates a Virtuous Loop to Address the Societal Harms of Content Moderation [35]

Introduction: Eder analyses systemic risk assessments as the central governance innovation introduced by the Digital Services Act. He argues that Article 34 does not aim to define systemic risks exhaustively, but instead establishes a procedural framework through which harms connected to platform design and content moderation can be progressively identified and mitigated. The paper frames systemic risk assessment as an iterative regulatory process rather than a static compliance obligation.

Why Chosen: This article is highly relevant as it provides one of the most detailed doctrinal explanations of how systemic risk assessments are supposed to function under the DSA. It helps clarify that the weakness of current assessments is not accidental but tied to the DSA's procedural regulatory philosophy.

Research gap: The article remains largely normative and procedural. It does not address AI-specific cybersecurity risks or examine whether purely procedural risk governance is adequate for technically complex systems such as AI-driven moderation and recommender algorithms.

30) Husovec, Martin (2024) General Risk Management under the Digital Services Act [36]

Introduction: Husovec offers a doctrinal analysis of Articles 34–35 DSA by comparing the DSA’s risk-management model with frameworks in the GDPR and financial regulation. He argues that the DSA adopts a deliberately limited, content-neutral risk-management approach constrained by legality and fundamental-rights protections.

Why chosen: This work is essential to understanding the legal limits of Article 34 enforcement, particularly why regulators cannot easily impose substantive risk-reduction measures on lawful content or algorithmic systems.

Research gap: The chapter does not examine how these doctrinal constraints interact with AI-driven cybersecurity risks, nor whether additional technical obligations could be legally justified within the DSA framework.

31) Tommasi, Sara (2023/2025) Risk-Based Approach in the Digital Services Act and the Artificial Intelligence Act [37]

Introduction: Tommasi compares the DSA’s systemic-risk framework with the AI Act’s risk-based model. She shows that the DSA treats systemic risks broadly and procedurally, while the AI Act adopts predefined categories of unacceptable and high-risk AI systems.

Why chosen: The paper is critical for analyzing regulatory inconsistency between the DSA and the AI Act, especially when AI systems deployed by platforms create cybersecurity or societal risks that fall between regulatory regimes.

Research gap: The analysis remains abstract and does not focus on platform-specific AI systems such as recommender engines or automated moderation tools, leaving implementation challenges unexplored.

32) Halil, Kollnig and Tamò-Larrieux (2025) Regulating Pressing Systemic Risks – but Not Too Soon? [38]

Introduction: This paper empirically studies the DSA's data-access provisions intended to support systemic-risk research. The authors analyze multiple data-access requests submitted across EU member states and find significant delays and procedural obstacles.

Why Chosen: It directly demonstrates that systemic-risk governance under the DSA fails without access to platform data, making it highly relevant for evaluating Article 34's practical effectiveness.

Research gap: The article focuses on data-access mechanisms rather than AI-specific risks and does not address how enforcement tools could be strengthened to support cybersecurity analysis.

33) Tamò-Larrieux, Aurelia (2024) Regulating for Trust: Can Law Establish Trust in Artificial Intelligence? [23]

Introduction: Tamò-Larrieux explores how law seeks to establish trust in AI systems through transparency, accountability, and procedural governance. The paper situates trust as a regulatory objective closely tied to risk management.

Why Chosen: This article provides a normative foundation for linking Article 34 risk assessments with public trust in AI-driven platforms.

Research gap: The work remains conceptual and does not analyze the DSA specifically or address cybersecurity vulnerabilities in AI systems.

34) Malgieri and Pasquale (2024) Licensing High-Risk Artificial Intelligence: Toward Ex Ante Justification [39]

Introduction: The authors argue for ex-ante regulatory controls over high-risk AI systems, proposing licensing and justification mechanisms prior to deployment.

Why Chosen: This paper is chosen because it introduces a robust ex-ante accountability model that sharply contrasts with the DSA's largely ex-post, procedural approach to systemic risk management. Its licensing and justification framework offers a normative alternative to risk self-assessment and mitigation duties. This contrast is particularly valuable for evaluating whether Article 34's design is sufficient for governing AI-driven

systemic risks on large platforms.

Research gap: The article does not examine platform regulation or Article 34 specifically, leaving open questions about integration with the DSA framework. It leaves unanswered how licensing-based approaches could be integrated with Article 34 risk assessments, mitigation obligations, or enforcement mechanisms.

35) Allen, Hilary J. (2024) A Harm-Focused Approach to Systemic Risk Regulation [40]

Introduction: Allen develops a harm focused theory of systemic risk regulation, originally grounded in financial and technological systems, arguing that systemic risk should be assessed not merely through probabilistic models or compliance checklists, but through an evaluation of how localized failures propagate into system-wide harms. The paper critiques technocratic and self-regulatory approaches to systemic risk, emphasizing that complex socio-technical systems such as AI enabled platforms exhibit emergent vulnerabilities that cannot be captured through isolated impact assessments.

Why Chosen: Her framework is transferable to platform governance and supports analyzing AI-driven cybersecurity risks as system-wide failures rather than isolated incidents.

Research gap: The paper does not engage with EU digital regulation or explain how harm-focused models could be operationalized under the DSA.

36) Kaushal et al. (2024) Automated Transparency: A Legal and Empirical Analysis of the DSA Transparency Database [41]

Introduction: This paper presents one of the first large-scale empirical evaluations of the Digital Services Act's transparency regime, analyzing over 130 million "statements of reasons" submitted by platforms under the DSA. The authors combine legal analysis with computational methods to examine how platforms operationalise transparency obligations in practice. The authors argue that transparency under the DSA risks becoming performative rather than functional, particularly when data is provided in formats that resist aggregation, auditing, or comparative analysis. This finding has direct relevance for Article 34 systemic-risk assessments, which depend on transparency data to identify patterns of harm, content moderation failures, and algorithmic bias at scale.

Why Chosen: It provides empirical evidence that current DSA transparency mechanisms, which underpin systemic-risk assessment, are insufficient for meaningful oversight.

Research Gap: The paper focuses on transparency rather than systemic-risk mitigation and does not assess Article 34 directly

37) Vandezande, Niels (2024) Cybersecurity in the EU: How the NIS2 Directive Stacks up Against Its Predecessor [42]

Introduction: Vandezande offers a comprehensive legal analysis of the NIS2 Directive, focusing on its strengthened enforcement mechanisms, expanded sectoral scope, and enhanced risk-management obligations. The paper highlights how NIS2 operationalizes cybersecurity governance through explicit technical and organizational requirements, mandatory incident reporting, and clear supervisory powers for regulators. Unlike the Digital Services Act, NIS2 adopts a prescriptive model that directly addresses technical system resilience. This analysis is highly relevant for assessing Article 34 DSA by contrast. While both instruments pursue systemic resilience, the DSA adopts a procedural and platform-led risk-assessment model, whereas NIS2 imposes concrete cybersecurity obligations with defined benchmarks.

Why Chosen: The paper provides a cybersecurity benchmark against which the DSA's weaker approach to AI-driven cyber risks can be evaluated.

Research gap: The analysis does not examine platform-specific systemic risks or interactions with Article 34 DSA.

38) Bodó, Balázs (2025) Governance by Trust Mediators in the Digital Society [43]

Introduction: Bodó examines the transformation of digital intermediaries into “trust mediators,” entities that increasingly shape social, informational, and economic environments through algorithmic decision-making and private governance arrangements. The paper argues that trust in digital systems is no longer primarily established through public law, but through platform-designed procedures, metrics, and automated moderation systems that operate largely outside democratic oversight.

This perspective is directly relevant to Article 34 systemic-risk assessments, which rely heavily on platforms' internal representations of risk and mitigation effectiveness. Bodó's analysis highlights how trust mediation creates structural incentives for platforms to prioritize reputational risk over societal harm reduction.

Why Chosen: The paper supports arguments about corporate influence and risk prioritization, complementing Griffin's political-economy critique.

Research Gap: The article does not focus on EU regulation or AI-specific cybersecurity risks.

39) Halil and Kollnig (2024) Platform Data Access and Systemic Risk Research in the EU[44]

Introduction: Halil and Kollnig investigate the practical functioning of the DSA's data-access provisions for researchers, focusing on whether these mechanisms enable meaningful study of systemic risks. Through empirical analysis of multiple data-access requests and regulator responses, the authors show that procedural complexity, administrative delays, and fragmented national implementation significantly obstruct timely access to platform data.

Why Chosen: It reinforces the conclusion that Article 34 cannot function effectively without enforceable data-access rights.

Research gap: Does not develop legal reform proposals specific to Article 34.

40) Eder, Griffin and Husovec (2024–2025) Procedural Risk Governance in EU Digital Law [45]

Introduction: This body of scholarship collectively conceptualises risk governance in EU digital regulation as procedural, iterative, and intentionally open-ended. Across multiple peer-reviewed contributions, the authors argue that the DSA represents a paradigmatic shift from liability-based regulation toward governance through processes, audits, and continuous assessment. Systemic risk, within this model, is not a fixed legal category but an evolving construct shaped by institutional practice, stakeholder input, and regulatory interpretation.

Why Chosen: It provides the theoretical backbone for analyzing why Article 34 struggles to address AI-driven cybersecurity risks.

Research Gap: This literature broadly acknowledges does not solve the problem of technical opacity in AI systems.

2.5 Observation

The literature reviewed reveals a consistent and significant gap between the regulatory ambition of Article 34 of the Digital Services Act and its practical capacity to account for cybersecurity risks arising from AI-driven content moderation and recommender systems. Foundational works in law and technology theory (e.g., Cockfield, Gifford) establish that legal frameworks often lag behind technological innovation, a pattern that persists in contemporary digital regulation. While these works provide important conceptual grounding, they offer limited operational guidance for addressing AI-specific risks.

More recent scholarship directly addressing the DSA highlights that systemic risk under Article 34 is framed primarily as a procedural governance obligation rather than a technically enforceable standard. Authors such as Griffin, Palumbo, Loi et al., and Husovec demonstrate that platforms retain significant discretion in defining, assessing, and reporting systemic risks, resulting in assessments that are high-level, narrative-driven, and insufficiently focused on algorithmic system design. This structural reliance on platform self-assessment creates blind spots for cybersecurity risks that originate from AI architectures, feedback loops, and automated decision-making.

Empirical and implementation-focused studies further reinforce this conclusion. Analyses by the Knight-Georgetown Institute, Kaushal et al., and Halil et al. show that existing systemic risk assessments and transparency mechanisms lack methodological rigor, standardized metrics, and enforceable data-access pathways. As a result, risks associated with adversarial manipulation, recommender-system gaming, and cascading AI failures remain largely under-detected and under-mitigated.

At the same time, comparative and theoretical contributions (e.g., Hacker et al., Allen, Malgieri and Pasquale) suggest that AI-driven cybersecurity failures should be understood as systemic phenomena rather than isolated incidents. These works underscore the inadequacy of purely procedural risk governance for complex socio-technical systems and point toward the need for clearer legal and procedural mechanisms—such as algorithmic audits, standardized risk metrics, and ex-ante accountability tools—to operationalize AI accountability under the DSA.

Overall, the literature indicates that while Article 34 formally acknowledges the role of platform design and algorithmic systems in generating systemic risks, it does not currently provide sufficient legal or procedural mechanisms to ensure meaningful accountability for AI-driven cybersecurity failures.

2.6 Summary of the Review

The literature collectively demonstrates that contemporary EU digital regulation struggles to adequately govern cybersecurity risks arising from AI-driven content moderation and recommender systems. Foundational law-and-technology scholarship establishes that legal frameworks historically lag behind technological innovation, creating persistent governance gaps when regulation confronts complex socio-technical systems [11, 15]. This dynamic remains evident in current EU instruments, where AI systems introduce opacity, automation, and scale that challenge traditional accountability models [19, 24]. Technical studies further confirm that AI-based cybersecurity mechanisms, while improving detection capabilities, generate new risks by automating decision-making without clearly attributable responsibility [25].

DSA-focused scholarship reveals that Article 34 conceptualizes systemic risk primarily as a procedural obligation rather than a substantively enforceable standard. Multiple doctrinal analyses show that platforms retain wide discretion in defining, assessing, and reporting systemic risks, resulting in high-level and narrative-driven assessments that insufficiently interrogate algorithmic architectures [26, 27, 28]. Empirical evaluations of VLOP systemic-risk reports corroborate this concern, finding that platform submissions lack methodological rigor, standardized metrics, and meaningful engagement with AI system design [34]. As a result, cybersecurity risks stemming from recommender-system manipulation, adversarial exploitation, and automated moderation failures remain under-assessed. A further strand of the literature emphasizes that AI-driven cybersecurity harms should be understood as system-wide risks, rather than isolated technical incidents. Conceptual frameworks developed by Hacker et al. (2025)[32] and Allen (2024)[40] demonstrate how AI failures propagate through interconnected platform ecosystems, producing cascading societal and security harms. However, Article 34 currently lacks explicit procedural mechanisms—such as algorithmic auditing, robustness testing, and enforceable data-access rights—necessary to translate systemic-risk identification into accountability [33, 38]. Comparative analyses further show that, unlike NIS2’s prescriptive cybersecurity obligations, the DSA relies heavily on platform self-assessment, reinforcing concerns about corporate capture and weak enforcement [42, 43].

Overall, the literature indicates that while Article 34 formally recognizes the role of platform design and algorithmic systems in generating systemic risks, it does not yet provide sufficient legal or procedural tools to ensure effective AI accountability for cybersecurity failures. This regulatory gap directly motivates the evaluative and reform-oriented analysis undertaken in the subsequent chapters.

Table 2. Summary of Relevant Research Studies and Alignment with Research Questions

Nr	Title	Year	Research Domain	Problem Discussed	Method/Strategy	Relevant RQs	Description
1	Cockfield – Towards a Law and Technology Theory	2003	Law and Technology Theory	Legal frameworks lag behind technological change	Theoretical legal analysis	RQ1	Establishes foundational theory on legal adaptability relevant to systemic AI risks.
2	Gifford – Law and Technology Interactions	2007	Legal Theory	Rigid legal structures fail to adapt	Jurisprudential analysis	RQ1	Reinforces need for adaptive regulation in technologically complex systems.
3	Mitrou – Data Protection, AI and Cognitive Services	2021	AI & Data Protection	GDPR limitations for AI accountability	Regulatory analysis	RQ1	Identifies opacity and accountability gaps relevant to AI cybersecurity.
4	Manolopoulos – Raising Cyber-Borders	2003	ICT Law	Jurisdictional limits in cyberspace	Doctrinal analysis	RQ1	Highlights regulatory challenges in cross-border digital systems.
5	Brinker – Identification and Demarcation	2024	Cybersecurity Law	Definitional ambiguity of IT systems	Legal demarcation framework	RQ1	Clarifies scope of IT systems relevant for AI risk classification.

Continued on next page

Nr	Title	Year	Research Domain	Problem Discussed	Method/Strategy	Relevant RQs	Description
6	Cockfield – What is Legal Knowledge	2003	Legal Theory	Law inadequately reflects technical complexity	Conceptual synthesis	RQ1	Supports reconceptualization of legal reasoning for AI governance.
7	Fischer-Hübner et al. – Stakeholder Perspectives	2021	Cybersecurity Governance	Fragmented EU cybersecurity coordination	Empirical stakeholder study	RQ1	Demonstrates governance gaps affecting systemic risk management.
8	Singh – European Cybersecurity Law in 2023	2023	EU Cyber Law	NIS2 integration challenges	Legal review	RQ1	Provides baseline for EU cybersecurity regulatory evolution.
9	Jørgensen et al. – AI in Smart Grids	2025	AI Regulation	Regulatory uncertainty for AI systems	Policy review	RQ1	Illustrates sectoral AI governance challenges.
10	Mal’ko & Kostenko – Legal Techniques and Technology	2022	Legal Theory	Legal-technical integration gaps	Doctrinal analysis	RQ1	Provides conceptual insight into legal-technical co-evolution.
11	Papakonstantinou – Cybersecurity as Praxis	2022	EU Cyber Law	Lack of right to cybersecurity	Doctrinal analysis	RQ1	Frames cybersecurity as normative foundation for risk governance.

Continued on next page

Nr	Title	Year	Research Domain	Problem Discussed	Method/Strategy	Relevant RQs	Description
12	Rampášek et al. – AI Certification	2025	AI Cybersecurity	Certification gaps for AI systems	Regulatory analysis	RQ1	Highlights limits of technical certification for accountability.
13	Tamò-Larrieux et al. – Trust in AI	2024	AI Governance	Trust deficit in AI regulation	Normative analysis	RQ1	Links trust to accountability in AI governance.
14	Zeng et al. – AI Anomaly Detection	2024	Technical Cybersecurity	Automated detection lacks legal framing	Technical framework	—	Provides technical context without legal accountability focus.
15	Balcioğlu et al. – Human-Centric AI Regulation	2025	AI Ethics	Ethics insufficiently operationalized	Policy analysis	RQ1	Supports human-centric framing of AI accountability.
16	Novelli et al. – Generative AI in EU Law	2024	AI Law	Fragmented liability frameworks	Doctrinal synthesis	RQ1	Integrates AI liability and cybersecurity governance.
17	Beltrán – AI Algorithms under Scrutiny	2025	EU Governance	Regulatory overlap across AI laws	Comparative analysis	RQ1	Exposes fragmentation affecting AI risk regulation.

Continued on next page

Nr	Title	Year	Research Domain	Problem Discussed	Method/Strategy	Relevant RQs	Description
18	Chiara – Right to Cybersecurity	2024	EU Human Rights	Cybersecurity not codified as a right	Normative inquiry	RQ1	Supports rights-based interpretation of systemic risk.
19	Vandezande – NIS2 Directive Assessment	2024	EU Cyber Law	AI accountability gaps under NIS2	Comparative analysis	RQ1	Provides cybersecurity benchmark for DSA comparison.
20	Griffin – Corporate Risk Management under DSA	2025	DSA Governance	Platform control of systemic risk framing	Institutional analysis	RQ1, SRQ1	Demonstrates corporate discretion limiting AI accountability.
21	Palumbo – Constitutional Limits of Delegation	2025	EU Constitutional Law	Over-delegation under Article 34	Doctrinal analysis	RQ1, SRQ1	Highlights enforcement limits of platform self-assessment.
22	Loi et al. – Regulating the Undefined	2025	Platform Governance	Undefined systemic risk concept	Conceptual framework	RQ1	Explains inconsistency in AI risk assessment.
23	Kaesling & Wolf – Sustainability and DSA Risk	2025	DSA Interpretation	Scope of systemic risk unclear	Legal interpretation	RQ1	Supports inclusion of AI cybersecurity harms.

Continued on next page

Nr	Title	Year	Research Domain	Problem Discussed	Method/Strategy	Relevant RQs	Description
24	Torraco – Media and Disinformation Risks	2025	Platform Governance	Recommender-driven systemic risks	Policy analysis	RQ1	Links algorithmic amplification to systemic harm.
25	EU Commission – Systemic Risk Landscape	2025	EU Regulation	Inconsistent platform risk reporting	Institutional analysis	RQ1, SRQ1	Provides regulator benchmark while exposing gaps.
26	Hacker et al. – AI and Systemic Risk	2025	AI Governance	AI risks propagate system-wide	Conceptual model	RQ1, SRQ1	Frames AI cybersecurity failures as systemic.
27	Panigutti et al. – Algorithmic Audits	2025	Algorithmic Accountability	Lack of audit tools under DSA	Empirical methods	RQ1, SRQ1	Supplies operational tools for AI accountability.
28	Knight-Georgetown – DSA Risk Assessments	2025	DSA Enforcement	Superficial platform assessments	Empirical review	RQ1, SRQ1	Shows failure to assess AI architecture risks.
29	Eder – Making Risk Assessments Work	2024	DSA Governance	Procedural governance limits	Doctrinal analysis	RQ1, SRQ1	Explains weaknesses in procedural risk regulation.
30	Husovec – Risk Management under DSA	2024	EU Platform Law	Legal limits on risk mitigation	Doctrinal comparison	RQ1	Explains why AI obligations are weak under Article 34.

Continued on next page

Nr	Title	Year	Research Domain	Problem Discussed	Method/Strategy	Relevant RQs	Description
31	Tommasi – DSA and AI Act Risk Models	2023	EU AI Regulation	Regulatory misalignment	Comparative analysis	RQ1, SRQ1	Shows AI cybersecurity risks fall between regimes.
32	Halil et al. – Data Access and Risk Research	2025	DSA Transparency	Lack of enforceable data access	Empirical study	RQ1, SRQ1	Demonstrates dependency of accountability on transparency.
33	Tamò-Larrieux – Trust and Governance	2024	AI Governance	Procedural trust deficits	Normative analysis	RQ1	Connects trust erosion to weak AI oversight.
34	Malgieri & Pasquale – Licensing High-Risk AI	2024	AI Liability	Ex-post governance failures	Normative proposal	SRQ1	Proposes ex-ante accountability mechanisms.
35	Allen – Harm-Focused Systemic Risk	2024	Systemic Risk Theory	Cascading harms unaddressed	Theoretical framework	RQ1, SRQ1	Supports system-wide accountability for AI failures.
36	Kaushal et al. – Automated Transparency	2024	DSA Transparency	Performative transparency	Empirical analysis	RQ1	Shows transparency limits undermine AI risk detection.
37	Bodó – Trust Mediators	2025	Platform Governance	Platforms control risk narratives	Institutional analysis	RQ1	Reinforces corporate capture critique.

Continued on next page

Nr	Title	Year	Research Domain	Problem Discussed	Method/Strategy	Relevant RQs	Description
38	Halil & Kollnig – Platform Data Access	2024	DSA Enforcement	Research barriers persist	Empirical analysis	RQ1, SRQ1	Shows Article 34 depends on enforceable access rights.
39	Eder, Griffin & Husovec – Procedural Risk Governance	2025	EU Digital Law	Procedural governance insufficient	Theoretical synthesis	RQ1, SRQ1	Explains structural limits of DSA risk governance.
40	Eder, Griffin & Husovec – Procedural Risk Governance in EU Digital Law	2025	EU Digital Regulation Theory	Procedural risk governance insufficient for technically complex AI systems	Cross-doctrinal legal synthesis	RQ1, SRQ1	Explains why Article 34's procedural and platform-led risk governance model struggles to ensure accountability for AI-driven cybersecurity risks.

3. Discussion

3.1 Analysis and Discussion

The findings of the literature review indicate that the Digital Services Act, and Article 34 in particular, represents a significant conceptual shift in EU digital regulation by explicitly recognizing systemic risk as a regulatory object. However, the reviewed scholarship consistently demonstrates that this recognition remains largely procedural and abstract when confronted with cybersecurity risks originating from AI-driven content moderation and recommender systems. While Article 34 formally requires platforms to assess risks arising from the design and functioning of their services, the absence of substantive benchmarks results in risk assessments that prioritize compliance narratives over technical scrutiny, raising questions as to whether such limitations stem from regulatory design or from implementation practices.

Across doctrinal and governance-oriented studies, a recurring theme is the extensive discretion afforded to platforms in defining what constitutes a systemic risk. Scholars analyzing Article 34 argue that this delegation enables platforms to frame risks in ways that align with corporate risk-management incentives rather than public cybersecurity concerns. As a result, AI-driven vulnerabilities—such as adversarial exploitation of recommender systems, cascading failures in automated moderation, and manipulation through coordinated inauthentic behavior—are rarely examined as design-level risks. Instead, they are treated as isolated incidents or behavioral problems attributable to users, rather than as structural cybersecurity issues embedded in platform architectures. At this stage, however, the literature does not yet allow a definitive attribution of this outcome to either deficient practice or structural under-specification of legal obligations.

Empirical evaluations of early systemic risk assessments reinforce this concern. Studies examining VLOP submissions show that risk reports tend to lack methodological rigor, standardized metrics, and meaningful engagement with algorithmic system design. Transparency mechanisms intended to support systemic-risk identification further suffer from fragmentation and limited accessibility, constraining independent verification and external scrutiny. While these observations may initially appear to reflect weak compliance, the literature also suggests that procedural variability is facilitated by the open-ended structure of Article 34 itself.

The literature also highlights a broader regulatory tension between procedural governance and technical accountability. Comparative analyses demonstrate that, unlike sector-specific cybersecurity regimes such as NIS2, the DSA does not impose concrete technical obligations related to system robustness, resilience testing, or failure mitigation. Consequently, AI-driven cybersecurity harms are governed through a model that emphasizes process over outcomes.

Taken together, the reviewed scholarship suggests that while Article 34 establishes an important governance framework for systemic risk, it does not yet provide the legal or procedural capacity required to operationalise AI accountability for cybersecurity failures.

However, this section does not seek to determine conclusively whether the identified shortcomings are attributable primarily to regulatory design or to insufficient implementation. That distinction is examined systematically in the Regulatory Gap Analysis presented in the following section.

4. Regulatory Gap Analysis

4.1 Regulatory Gap Analysis

This section applies a Regulatory Gap Analysis (RGA) to assess whether Article 34 of the Digital Services Act (DSA) effectively addresses cybersecurity risks originating From AI-driven content moderation and recommender systems. The analysis compares the normative expectations embedded in Article 34 with platform practices and the shortcomings identified in academic and institutional literature, thereby addressing RQ1 and SRQ1.

Before examining specific deficiencies, it is necessary to distinguish analytically between two potential sources of inadequacy: *implementation failures*, where platforms comply only superficially with existing obligations, and *structural regulatory gaps*, where the legal framework itself does not require the identification, assessment, or mitigation of foreseeable risks. While instances of weak implementation are referenced where relevant, this section focuses primarily on structural regulatory gaps, evaluating whether Article 34 is capable of operationalizing AI accountability even under conditions of good-faith compliance.

4.1.1 Conceptual Gaps

At the conceptual level, Article 34 establishes systemic risk as a core regulatory concern but deliberately leaves it underspecified. While the provision requires platforms to assess risks arising from the design and functioning of their services, it does not define systemic risk in a way that clearly captures AI-driven cybersecurity threats. Scholarly analyses show that this ambiguity enables platforms to frame systemic risks narrowly, often prioritizing content-related harms over structural vulnerabilities embedded in algorithmic architectures (Griffin, 2025; Loi et al., 2025)[26][28].

Legal interpretation further demonstrates that the catalog of systemic risks under Article 34 is non-exhaustive, requiring interpretive expansion beyond enumerated harms (Kaesling & Wolf, 2025)[29]. However, in the absence of explicit AI- or cyber-specific risk categories, risks such as adversarial attacks on moderation systems, recommender manipulation, and cascading failures caused by automated feedback loops are not consistently treated as systemic by default. As a result, cybersecurity risks are acknowledged abstractly but remain weakly operationalized within the conceptual framework of Article 34.

As a result, cybersecurity risks are acknowledged abstractly but remain weakly operationalized within the conceptual framework of Article 34. These limitations persist independently of platform compliance practices, as Article 34 does not require AI-driven cybersecurity risks to be classified as systemic by default.

4.1.2 Procedural Gaps

Procedurally, Article 34 primarily relies on platform-led self-assessment to identify and evaluate systemic risks. Although platforms are required to conduct annual assessments, the DSA does not mandate standardized methodologies, performance metrics, or analytical frameworks. Empirical evaluations of early VLOP systemic risk reports reveal that many assessments lack methodological transparency, clear indicators, and analysis of platform design choices (Knight-Georgetown Institute, 2025)[34].

This procedural flexibility reflects the DSA’s governance-oriented regulatory philosophy, which deliberately constructs systemic risk assessment as a procedural obligation rather than a substantive standard (Eder, 2024)[35]. However, doctrinal analyses indicate that such reliance on self-assessment, combined with legally constrained audit requirements, allows platforms to satisfy formal compliance obligations without engaging in meaningful cybersecurity evaluation of AI-driven systems (Husovec, 2024)[36]. While such deficiencies may appear to reflect implementation weaknesses, the literature indicates that they are enabled by the procedural design of Article 34 itself, which does not mandate standardized methodologies or AI-specific cybersecurity assessment criteria.

4.1.3 Technical Gaps

The regulatory gap is most pronounced at the technical level. Article 34 does not impose explicit obligations related to AI system robustness, adversarial resilience, or exploitable testing, despite these being well-established concerns in cybersecurity research. Conceptual work on AI-platform systems demonstrates that AI failures often manifest as cascading and collective harms rather than isolated incidents, underscoring their systemic nature (Hacker et al., 2025)[32].

Nevertheless, current DSA obligations do not require platforms to conduct adversarial testing, evaluate susceptibility to model poisoning, or disclose algorithmic performance metrics relevant to cybersecurity risk. Empirical research on algorithmic auditing shows that, without such technical evaluations, significant risks remain invisible to governance mechanisms (Panigutti et al., 2025)[33]. Technical studies further illustrate how AI-driven

security systems introduce new vulnerabilities when training data or model behavior is compromised, complicating accountability in the absence of regulatory oversight (Zeng et al., 2024).[25]. These omissions constitute a clear case of structural regulatory deficiency rather than implementation failure, as Article 34 does not impose any legal obligation to evaluate AI system robustness, resilience, or exploitability.

4.1.4 Enforcement Gaps

Finally, enforcement-related gaps further weaken the effectiveness of Article 34. The DSA does not establish clear sanctions for vague, incomplete, or superficial systemic risk assessments, reducing incentives for substantive engagement. Constitutional analyses of the DSA highlight that significant normative discretion is delegated to private actors, creating enforcement and accountability concerns (Palumbo, 2025)[27].

Thus discretion asymmetry is therefore embedded in the regulatory architecture of Article 34 and cannot be resolved solely through enhanced enforcement or supervisory practice.

Currently in practice, this delegation produces discretion asymmetries in which platforms define risks and mitigation priorities, while regulators operate reactively and with limited technical access. Governance scholarship further shows that platforms act as trust mediators, shaping how risks are framed and prioritized through internal governance mechanisms (Bodó, 2025)[43]. The absence of binding links between risk identification (Article 34), mitigation (Article 35), and auditing (Article 37) therefore undermines the operationalization of AI accountability for cybersecurity failures.

4.1.5 Interim Conclusion

Taken together, the Regulatory Gap Analysis demonstrates that the primary source of inadequacy in the operation of Article 34 lies in its structural design rather than in insufficient implementation alone. Conceptual ambiguity, procedural reliance on platform-led self-assessment, the absence of AI-specific technical obligations, and embedded enforcement asymmetries collectively limit the provision's capacity to operationalize accountability for AI-driven cybersecurity risks.

While instances of weak platform practice are observable, the analysis shows that even robust and good-faith implementation would not address several foreseeable risks arising from the design and functioning of AI-driven systems. Article 34 therefore functions

predominantly as a coordination and reporting mechanism rather than as a substantive instrument of systemic risk governance. These findings provide the analytical foundation for the scenario-based stress testing and comparative legal benchmarking undertaken in the subsequent sections.

5. Scenario Based Stress Testing

5.1 Scenario-Based Stress Testing of AI-Driven Systemic Risks

5.1.1 Purpose and Role of Scenario-Based Stress Testing

Scenario-Based Stress Testing (SBST) is employed in this thesis as a qualitative and evaluative method to examine whether Article 34 of the Digital Services Act (DSA) is capable of identifying, assessing, and mitigating AI-driven cybersecurity risks before systemic harm materializes. Unlike doctrinal analysis, which interprets legal provisions in abstraction, SBST operationalizes Article 34 by testing its regulatory logic against empirically documented real-world scenarios in which AI-driven platform systems produced large-scale harm and subsequently required system-level intervention.

In my assessment, SBST is particularly appropriate for analyzing Article 34 because systemic risks associated with AI-driven platforms rarely manifest as discrete unlawful events. Instead, they emerge from lawful system behavior operating at scale, often becoming visible only after societal harm has already occurred. Stress-testing the regulation against documented cases, therefore, allows this thesis to evaluate whether Article 34 functions as a genuinely preventive governance mechanism or whether it remains predominantly reactive.

5.1.2 Selection of Scenarios for Stress Testing

The selection of scenarios for the Scenario-Based Stress Testing (SBST) in this thesis follows a structured and purposive approach, aimed at identifying empirically documented cases in which AI-driven platform systems have produced systemic risks resulting in observable harm and subsequent system-level modification. Rather than relying on hypothetical examples, the analysis focuses exclusively on real-world scenarios to ensure that the stress test reflects conditions under which systemic risks have demonstrably materialized.

The selection process is guided by four main criteria. First, each scenario must involve a clear and central role of AI-driven systems, such as recommender systems, algorithmic targeting, or automated moderation tools. This ensures that the analysis directly engages with the types of technological systems governed by Article 34 of the Digital Services Act.

Second, the selected scenarios must exhibit systemic characteristics, meaning that the risk extends beyond individual instances of harm and generates broader effects on users, information ecosystems, or societal processes. This criterion aligns the analysis with the concept of systemic risk as defined within the DSA.

Third, each scenario must be supported by authoritative sources, including regulatory investigations, academic research, or platform transparency reports, and must involve documented platform responses or system modifications. The inclusion of such cases allows the analysis to trace a clear link between risk emergence, harm realization, and subsequent corrective action, thereby providing a robust basis for stress testing.

Fourth, the scenarios are selected to represent distinct categories of AI-driven risk, including algorithmic amplification, data-driven targeting, and adversarial exploitation. This diversity enables the analysis to test Article 34 across different dimensions of platform functionality and risk generation.

In the authors assessment, this selection methodology ensures that the stress testing remains both analytically rigorous and empirically grounded, while avoiding speculative conclusions. It also allows the evaluation to focus on whether Article 34 would have required the identification and mitigation of these risks under conditions in which their occurrence was not only foreseeable but also repeatedly observed in practice.

5.1.3 Analytical Logic Applied Across Scenarios

Across all scenarios, the stress test follows a consistent analytical logic. For each case, the analysis examines:

- The role of AI system implementation in producing or amplifying risk;
- The systemic nature of the harm generated;
- Whether Article 34 would plausibly require ex-ante identification of the risk;
- Whether mitigation obligations under Article 35 or oversight under Article 37 would realistically be triggered; and
- The extent of harm to users, data integrity, and information ecosystems. This approach ensures analytical consistency while allowing scenario-specific detail and interpretation.

5.1.4 Scenario 1: Meta Recommender Systems and the Cambridge Analytica Scandal

Empirical Background and AI System Context

The Cambridge Analytica scandal[46] represents one of the most prominent cases of data-driven political manipulation on social media platforms. While the immediate legal infringement concerned unauthorized third-party access to user data, subsequent investigations demonstrated that systemic harm arose from the interaction between algorithmic profiling, behavioral targeting, and content recommendation systems optimized for engagement [46, 47]. Facebook’s recommender and targeting architectures enabled political messaging to be delivered at scale with unprecedented precision, amplifying the reach and impact of coordinated influence campaigns.

In response to regulatory scrutiny, Meta implemented extensive system-level changes, including restricting API access, reducing targeting granularity, redesigning internal data governance structures, and modifying recommender-system functionalities [48, 49][50]. These modifications reflect an implicit acknowledgment that the risk originated from platform architecture rather than from isolated misuse.

Stress-Test Application to Article 34

Applying Article 34 retrospectively to this scenario involves assessing whether a platform would have been required to identify algorithmic targeting and recommendation architectures as sources of systemic risk *ex ante*. Although Article 34 refers broadly to risks arising from the design and functioning of platform services, it does not explicitly mandate analysis of how recommender systems and micro-targeting tools may facilitate coordinated exploitation.

In my assessment, this ambiguity would likely allow the risk to be characterized narrowly as a data protection or third-party misuse issue, rather than as a systemic cybersecurity risk embedded in AI-driven platform design. Article 34 does not require platforms to assess how lawful AI systems might be repurposed strategically to generate large-scale societal harm.

User Harm, Data Distortion, and Systemic Effects

The harm produced in this scenario extended beyond individual privacy violations.[51] Users were exposed to opaque behavioral targeting, undermining autonomy, informed political participation, and trust in digital platforms. At the data level, engagement signals

generated by manipulated content polluted optimization processes, reinforcing algorithmic feedback loops that increased the effectiveness of manipulation. From a societal perspective, democratic integrity and public trust suffered significant harm.

Evaluative Conclusion and Article 34 Gap Alignment

The author concludes that Article 34 would not have reliably captured this risk prior to harm materializing. The scenario reveals multiple regulatory gaps: a conceptual gap in classifying AI-driven political amplification as systemic risk; a procedural gap in requiring analysis of the exploitability of recommender systems; a technical gap in assessing amplification feedback loops; and an enforcement gap in which mitigation occurred only after external intervention. This illustrates Article 34's limited capacity to address AI-driven systemic manipulation *ex ante*.

Findings from the Stress Test

The analysis of the Cambridge Analytica case shows that the systemic risk did not arise solely from unlawful data access, but from the way Meta's AI-driven targeting and recommender systems operated at scale. Algorithmic profiling and content distribution mechanisms enabled political messages to be delivered with high precision and reach, amplifying the impact of data misuse far beyond individual users.

The harm materialized through predictable interactions between platform design, engagement optimization, and coordinated exploitation. These dynamics indicate that the risk was embedded in the platform's architecture rather than arising from accidental misuse. Importantly, corrective action occurred only after regulatory intervention and public scrutiny, leading Meta to redesign its data access rules and targeting systems.

When stress-tested against Article 34, this scenario reveals that the regulation does not clearly require platforms to treat algorithmic targeting and recommender-system amplification as systemic risks *ex ante*. The risk could easily be framed as a compliance or data protection issue, thereby allowing its systemic dimension to remain inadequately addressed within the procedural risk assessment framework.

Author's Reflection

In the author's assessment, this scenario demonstrates that systemic risk can emerge even when AI systems operate exactly as designed. It is particularly concerning that platform responses occurred only after significant societal and psychological harm had already materialized, despite the fact that such risks were foreseeable given the scale and precision of algorithmic targeting in recommender systems.

This situation can be characterized as indicative of a structural limitation of Article 34. Although the provision refers to risks arising from system design, it does not require platforms to critically assess how recommender and targeting systems may be exploited for coordinated manipulation. As a consequence, accountability within the current framework tends to remain reactive rather than preventive.

According to the author’s analysis, preventing similar risks in the future would require clearer constraints under Article 34, particularly regarding how AI-driven profiling and amplification mechanisms interact with broader political and societal contexts. In the absence of such requirements, the regulatory approach risks continuing to rely on post-hoc enforcement rather than anticipatory risk governance.

5.1.5 Scenario 2: YouTube Recommender Algorithms and Harmful Content Amplification

Empirical Background and AI System Context

A substantial body of empirical research has demonstrated that YouTube’s engagement-optimized recommender system systematically amplified harmful but lawful content, including health misinformation and content affecting minors [52, 53]. These outcomes were driven not by content illegality but by optimization strategies that prioritized watch time and user engagement, producing algorithmic “rabbit holes”.

Following public and regulatory scrutiny, YouTube modified its recommender system by downgrading borderline content, limiting recommendation chains, and introducing friction mechanisms designed to curb amplification dynamics [54]. These changes indicate that systemic risk was recognized only after harm had become visible.

Stress-Test Application to Article 34

The stress test examines whether Article 34 would require platforms to identify recommender-system amplification dynamics as systemic risks when harm arises from lawful optimization choices. While recommender systems fall within the general scope of Article 34, the provision does not mandate analysis of engagement-driven feedback loops or require platforms to reassess optimization objectives themselves.

In my view, this omission allows systemic harm to remain under-analyzed so long as content remains lawful, thereby limiting Article 34’s preventive potential.

User Harm and Information Ecosystem Degradation

Users experienced repeated exposure to misleading or harmful content without transparency or meaningful choice. Algorithmically distorted engagement data fed back into optimization processes, entrenching amplification bias. The resulting harm extended beyond individual users to broader information ecosystems, particularly in sensitive domains such as public health.

Evaluative Conclusion and Article 34 Gap Alignment

I argue that this scenario exposes a structural incapacity of Article 34 to govern lawful AI behavior that produces harmful outcomes. Identified gaps include the failure to treat lawful amplification as systemic risk, the absence of obligations to analyze feedback loops, discretionary mitigation under Article 35, and audits under Article 37 that focus on documentation rather than system behavior.

Findings from the Stress Test

The YouTube scenario illustrates how lawful content can generate systemic harm through AI-driven recommendation systems. Empirical evidence demonstrates that engagement-optimized algorithms repeatedly amplified harmful and misleading material, particularly in areas such as health information, without violating content legality standards.

The harm emerged from algorithmic feedback loops that continuously promoted content attracting user interaction, thereby distorting information exposure at scale. Users were repeatedly directed toward sensational or misleading narratives, while engagement signals reinforced these recommendation patterns. This, in turn, was a misalignment in the recommender system's recognition of content and of harm to users.

Stress-testing this scenario against Article 34 reveals that the regulation does not require platforms to assess the risks associated with lawful recommendation dynamics. Mitigation obligations were not structurally triggered, and system-level corrections occurred only after sustained external pressure.

Author's Reflection

In the author's assessment, this scenario highlights a fundamental weakness in how Article 34 approaches systemic risk. In this instance, the harm did not arise from illegal content, but from the manner in which the recommender system prioritized engagement without sufficient consideration of downstream consequences. As demonstrated in the findings, the issue extended beyond harm to users, reflecting a broader failure in systemic

risk management and irregularities in the implementation and evaluation of the system.

It is further observed that Article 34 does not adequately address situations in which AI systems produce harmful outcomes through lawful optimization choices. The regulatory framework allows platforms to focus primarily on content legality, thereby overlooking the wider societal effects generated by algorithmic amplification processes.

This raises significant concerns regarding user protection and the integrity of information ecosystems, particularly in sensitive contexts such as public health. From the author's perspective, a more effective systemic-risk framework would require platforms to assess how recommender systems influence user exposure and behavior, even in the absence of clear legal violations. Without such an approach, systemic harms are likely to remain insufficiently addressed within the current regulatory structure.

5.1.6 Scenario 3: Coordinated Inauthentic Behavior and AI Detection Failures at Meta

Empirical Background and AI System Context

Meta's Adversarial Threat Reports document repeated campaigns of coordinated inauthentic behavior (CIB) that exploited weaknesses in AI-driven detection and moderation systems [55]. These campaigns targeted elections, public health debates, and geopolitical discourse. Adversarial actors repeatedly adapted to enforcement signals, identifying detection thresholds and exploiting system blind spots.

Meta has repeatedly redesigned its AI detection models and moderation pipelines in response to these failures, indicating that the risk stemmed from predictable exploitability in system design rather than from isolated misconduct.

Stress-Test Application to Article 34

This scenario stress-tests whether Article 34 would require platforms to treat adversarial exploitation of AI moderation systems as a systemic cybersecurity risk. Article 34 does not mandate adversarial robustness testing, continuous exploitability monitoring, or performance degradation analysis for AI-driven moderation systems.

In my assessment, this omission allows coordinated manipulation to be framed as episodic abuse rather than as a structural vulnerability intrinsic to automated systems.

Harm to Users, Data Integrity, and Public Trust

Users were exposed to manipulated narratives and synthetic engagement patterns that distorted public discourse and undermined trust in platform governance. Polluted data environments degraded the long-term effectiveness of AI moderation models, reinforcing cycles of ineffective enforcement. Harm therefore extended to both content outcomes and AI system integrity.

Evaluative Conclusion and Article 34 Gap Alignment

I conclude that Article 34 structurally fails to capture adversarial exploitability as a systemic risk category. Gaps include the absence of technical robustness requirements, procedural misclassifications of coordinated manipulation as user behavior, and reliance on reactive system redesign rather than preventive governance.

Findings from the Stress Test

The analysis of coordinated inauthentic behavior demonstrates that AI-driven content moderation and detection systems are repeatedly exploited by organized actors. Platform transparency reports show that adversarial groups systematically adapt to automated enforcement thresholds, allowing harmful content and networks to persist at scale.

The systemic risk arises not from individual violations, but from predictable weaknesses in AI systems used to detect manipulation. Platforms have repeatedly redesigned detection models after discovering these vulnerabilities, indicating that the exploitability of AI moderation is a recurring and foreseeable issue.

When assessed under Article 34, this scenario reveals that adversarial exploitability is not treated as a default systemic risk. The regulation does not require adversarial robustness testing or continuous evaluation of AI system performance, allowing mitigation to occur only after harm becomes visible.

Author's Reflection

In the author's assessment, this scenario demonstrates how failures in AI systems are frequently interpreted as instances of user misconduct rather than as vulnerabilities arising from system design. It is particularly notable that platform responses tend to occur only after coordinated manipulation has been exposed, despite the fact that such forms of exploitation are both widely documented and foreseeable.

It is further observed that Article 34 does not sufficiently address this category of cyberse-

curity risk. By emphasizing procedural risk reporting, the framework allows adversarial weaknesses inherent in AI-driven systems to remain insufficiently examined until external pressure necessitates intervention.

From the author's perspective, a stronger emphasis on system robustness and exploitability assessment would help reduce harm to users and enhance trust in platform governance. In the absence of such measures, Article 34 risks functioning primarily as a reporting mechanism rather than as an effective preventive safeguard.

5.1.7 Cross-Scenario Findings

Across all scenarios, consistent patterns emerge. AI systems produced systemic risk while functioning as designed; harm escalated before regulatory mechanisms were activated; platform interventions occurred only post hoc; and Article 34 lacked sufficient ex-ante regulatory force. These findings support the conclusion that observed failures stem primarily from structural limitations of Article 34 rather than from deficient implementation.

6. Strengthening the Systemic-Risk Framework of Article 34

6.1 Purpose and Scope of the Chapter

This chapter proposes targeted options to strengthen Article 34 of the Digital Services Act (DSA) in light of the structural limitations identified in the preceding analysis. The objective is not to introduce alternative regulatory regimes or to expand the scope of the thesis beyond platform systemic-risk governance. Rather, the chapter explores how targeted and legally feasible refinements could reduce ambiguity, improve accountability, and enhance Article 34's capacity to address cybersecurity risks arising from AI-driven content moderation and recommender systems.

All proposals developed in this chapter are derived directly from the Regulatory Gap Analysis and the Scenario-Based Stress Testing. References to adjacent EU legal frameworks are used selectively and instrumentally, solely to illustrate regulatory techniques capable of addressing the identified gaps, without displacing Article 34 as the primary object of analysis.

6.2 Clarifying AI-Driven Cybersecurity Risks as Systemic Risks

A central finding of the Regulatory Gap Analysis is that Article 34 does not specify with sufficient clarity how cybersecurity risks originating from AI-driven platform systems should be conceptualized as systemic risks. Although the provision broadly refers to risks arising from the design and functioning of platform services, it does not indicate whether vulnerabilities such as adversarial manipulation, algorithmic amplification cascades, or degradation of automated moderation accuracy fall within the core scope of systemic risk assessment.

This conceptual openness enables platforms to frame AI-driven cybersecurity risks as isolated incidents or user-behavior problems rather than as design-level vulnerabilities embedded in system architectures. A first strengthening option would therefore consist of explicit interpretative clarification, through guidance or implementing acts, confirming that cybersecurity risks arising from AI-driven platform systems qualify as systemic risks by default when capable of producing large-scale economic, societal, or public security

effects.

Such clarification would not expand the material scope of Article 34 but would reduce uncertainty by anchoring AI-driven cybersecurity risks firmly within its systemic-risk framework, thereby limiting excessive reliance on platform discretion.

6.3 Introducing Minimum AI-Specific Risk Assessment Requirements

The analysis further demonstrates that Article 34 imposes a duty to conduct systemic-risk assessments without defining minimum content requirements for those assessments, particularly with respect to AI-driven systems. As a result, platforms may formally comply with Article 34 while omitting foreseeable cybersecurity risks linked to algorithmic design choices.

A targeted regulatory response would be to introduce minimum AI-specific assessment elements within the Article 34 risk-assessment process. Such elements could include the evaluation of susceptibility to coordinated manipulation and adversarial exploitation, analysis of amplification dynamics and feedback loops in recommender systems, and consideration of system performance degradation under stress conditions.

The purpose of these requirements would not be to impose uniform technical solutions, but to ensure that foreseeable AI-driven cybersecurity risks cannot be excluded from systemic-risk assessments altogether. By strengthening the procedural baseline, Article 34 would move closer to operationalizing meaningful accountability while preserving its flexible, risk-based governance model.

6.4 Strengthening the Link Between Risk Identification and Mitigation Obligations

Scenario-Based Stress Testing indicates that even where systemic risks are identified, mitigation obligations under Article 35 may not be triggered in a timely or consistent manner. This reflects the absence of predefined indicators linking risk identification under Article 34 to mandatory mitigation measures.

To address this disconnect, Article 34 could be complemented by predefined systemic-risk indicators for AI-driven cybersecurity risks, particularly in economically or socially sensitive contexts. Where such indicators are met, platforms would be required to implement mitigation measures under Article 35 rather than relying solely on discretionary risk

management responses.

By introducing clearer activation conditions for mitigation, Article 34 would shift from a predominantly descriptive reporting mechanism toward a more preventive form of systemic-risk governance, reducing the regulatory gap identified in earlier chapters.

6.5 Enhancing Audit-ability and Supervisory Oversight

The effectiveness of systemic-risk governance depends on the credibility of oversight mechanisms. As established in the Regulatory Gap Analysis, audits conducted under Article 37 currently focus primarily on procedural compliance and documentation, offering limited insight into the actual behavior of AI-driven systems under conditions of stress.

A further strengthening option would be to clarify audit expectations so that supervisory verification extends to whether AI-specific cybersecurity risks were meaningfully assessed and mitigated. This could include the examination of internal monitoring mechanisms for system degradation, verification of mitigation measures following risk identification, and evaluation of whether platform practices align with identified systemic risks.

Such an approach would remain consistent with the procedural nature of the DSA while reducing the gap between formal compliance and substantive accountability, thereby reinforcing the effectiveness of Articles 34, 35, and 37 as a coherent governance framework.

6.6 Interim Reflections

Taken together, the options explored in this chapter demonstrate that the limitations of Article 34 identified in earlier chapters do not necessitate its replacement or fundamental redesign. Rather, they point toward the need for targeted refinements that reduce ambiguity, constrain excessive discretion, and introduce AI-specific expectations where systemic cybersecurity risks are foreseeable.

By clarifying the scope of systemic risk, strengthening minimum assessment requirements, linking risk identification to mitigation, and enhancing audit ability, Article 34 could evolve into a more effective instrument for governing AI-driven cybersecurity risks. These refinements preserve the DSA's risk-based and flexible character while addressing the structural weaknesses revealed through the Regulatory Gap Analysis and Scenario-Based Stress Testing.

7. Conclusion

This thesis set out to examine how the systemic risk assessment obligations under Article 34 of the Digital Services Act (DSA) account for cybersecurity risks arising from AI-driven content moderation and recommender systems. In addressing this question, the thesis combined doctrinal legal analysis with empirical insights from early platform practice, regulatory gap analysis, and scenario-based stress testing to evaluate whether Article 34 is capable of operationalising meaningful accountability for such risks.

The analysis demonstrates that Article 34 represents an important conceptual development in EU platform governance by explicitly recognizing systemic risk as a regulatory concern. However, the findings show that this recognition remains largely procedural and abstract when applied to cybersecurity risks originating from the design and functioning of AI-driven systems. While platforms are formally required to conduct systemic risk assessments, the absence of defined benchmarks, AI-specific assessment requirements, and enforceable technical expectations limits the practical effectiveness of this obligation.

A key contribution of the thesis lies in its distinction between implementation failures and structural regulatory gaps. Although instances of weak or inconsistent platform practice can be observed, the Regulatory Gap Analysis establishes that many of the identified shortcomings persist independently of compliance quality. Rather, they are enabled by the open-textured and discretionary design of Article 34, which leaves the identification and characterization of AI-driven cybersecurity risks largely to platform judgment. As a result, foreseeable risks such as adversarial exploitation, amplification cascades, and cascading system failures often remain under-captured even under conditions of good-faith implementation.

Scenario-Based Stress Testing further reinforces this conclusion by demonstrating that Article 34 struggles to detect, mitigate, and audit AI-driven cybersecurity risks under realistic conditions of scale, coordination, and crisis. The stress-testing scenarios illustrate that the regulatory framework is particularly weak where harmful outcomes emerge from system behavior that is technically lawful and formally compliant, but systemically destabilizing.

In response to these findings, the thesis does not advocate a fundamental reconfiguration of the DSA's systemic-risk governance model. Instead, it identifies targeted and legally feasible refinements capable of strengthening Article 34 while preserving its risk-based and

flexible structure. These include clarifying the treatment of AI-driven cybersecurity risks as systemic risks by default, introducing minimum AI-specific assessment requirements, strengthening the link between risk identification and mitigation obligations, and enhancing audit ability through more substantive supervisory verification.

This thesis contributes to the emerging intersection of cybersecurity, AI governance, and EU platform regulation by demonstrating that systemic risks on large online platforms frequently arise not from technical failure or unlawful conduct, but from the normal operation of AI-driven systems designed to optimize engagement and scale. Through the combined application of doctrinal analysis, Regulatory Gap Analysis, and Scenario-Based Stress Testing, the research shows that Article 34 of the Digital Services Act is structurally limited in its ability to identify and mitigate such risks *ex ante*. In particular, the findings reveal that the provision's reliance on procedural self-assessment, coupled with the absence of AI-specific analytical obligations, allows foreseeable harms linked to recommender systems, targeting architectures, and automated moderation to remain under-captured until after they materialize. By exposing these limitations and proposing targeted refinements that remain consistent with the DSA's risk-based governance model, this thesis advances a more anticipatory understanding of systemic risk regulation, contributing to improved protection of users, enhanced data integrity, and greater resilience of digital information ecosystems.

Overall, the thesis concludes that Article 34, in its current form, functions more effectively as a coordination and reporting mechanism than as a robust tool of cybersecurity accountability for AI-driven platform systems. While it provides an essential governance framework, its capacity to address systemic cybersecurity risks depends on reducing ambiguity, constraining excessive discretion, and aligning procedural obligations more closely with the technical realities of AI-driven risk.

References

- [1] Hannah Snyder. “Literature review as a research methodology: An overview and guidelines”. In: *Journal of Business Research* 104 (2019), pp. 333–339.
- [2] Virginia Braun and Victoria Clarke. “Using thematic analysis in psychology”. In: *Qualitative Research in Psychology* 3.2 (2006), pp. 77–101.
- [3] Christopher McCrudden. “Legal Research and the Social Sciences”. In: *Law Quarterly Review* 122 (2006), pp. 632–650.
- [4] Jan M. Smits. *The Mind and Method of the Legal Academic*. Edward Elgar, 2012.
- [5] Claes Wohlin. “Guidelines for snowballing in systematic literature studies and a replication in software engineering”. In: *Proceedings of the 18th International Conference on Evaluation and Assessment in Software Engineering* (2014).
- [6] Yavuz Selim Balcioğlu, Ahmet Alkan Çelik, and Erkut Altındağ. “A Turning Point in AI: Europe’s Human-Centric Approach to Technology Regulation”. In: *Journal of Responsible Technology* 23 (2025), p. 100128. DOI: 10.1016/j.jrt.2025.100128.
- [7] Nils Brinker. “Identification and Demarcation—A General Definition and Method to Address Information Technology in European IT Security Law”. In: *Computer Law & Security Review* 52 (2024), p. 105927. DOI: 10.1016/j.clsr.2023.105927.
- [8] Marta Beltrán. “AI Algorithms under Scrutiny: GDPR, DSA, AI Act and CRA as Pillars for Algorithmic Security and Privacy in the European Union”. In: *Computers & Security* 158 (2025), p. 104628. DOI: 10.1016/j.cose.2025.104628.
- [9] *Regulation (EU) 2016/679 of the European Parliament and of the Council of 27 April 2016 on the protection of natural persons with regard to the processing of personal data and on the free movement of such data, and repealing Directive 95/46/EC (General Data Protection Regulation) (Text with EEA relevance)*. May 2016.
- [10] Pier Giorgio Chiara. “Towards a Right to Cybersecurity in EU Law? The Challenges Ahead”. In: *Computer Law & Security Review* 53 (2024), p. 105961. DOI: 10.1016/j.clsr.2024.105961.
- [11] Arthur J. Cockfield. “Towards a Law and Technology Theory CALT Conference: What Is Legal Knowledge”. In: *Manitoba Law Journal* 30.3 (2003), pp. 383–416.

- [12] Arthur Cockfield and Jason Pridmore. “A Synthetic Theory of Law and Technology Symposium: Towards a General Theory of Law and Technology”. In: *Minnesota Journal of Law, Science & Technology* 8.2 (2007), pp. 475–514.
- [13] Dr. Charanjit Singh, European Commission. “European Cyber Security Law in 2023: A Review of the Advances in Network and Information Security Directive 2”. In: *European Cyber Security Review* (2023).
- [14] Simone Fischer-Hübner, Cristina Alcaraz, Afonso Ferreira, et al. “Stakeholder Perspectives and Requirements on Cybersecurity in Europe”. In: *Journal of Information Security and Applications* 61 (2021), p. 102916. DOI: 10.1016/j.jisa.2021.102916.
- [15] Daniel J. Gifford. “Law and Technology: Interactions and Relationships Symposium: Towards a General Theory of Law and Technology”. In: *Minnesota Journal of Law, Science & Technology* 8.2 (2007), pp. 571–588.
- [16] Bo Nørregaard Jørgensen, Saraswathy Shamini Gunasekaran, and Zheng Grace Ma. “Impact of EU Laws on AI Adoption in Smart Grids: A Review of Regulatory Barriers, Technological Challenges, and Stakeholder Benefits”. In: *Energies* 18.12 (2025), p. 3002. DOI: 10.3390/en18123002.
- [17] A. V. Mal’ko and M. A. Kostenko. “Legal Techniques and Legal Technology: The Complex Nature of the Interaction”. In: *Gosudarstvo i Pravo* 7 (2022), pp. 22–29. DOI: 10.31857/S102694520020992-2.
- [18] A. Manolopoulos. “Raising ‘Cyber-Borders’: The Interaction Between Law and Technology”. In: *International Journal of Law and Information Technology* 11.1 (2003), pp. 40–58. DOI: 10.1093/ijlit/11.1.40.
- [19] Lilian Mitrou. “Data Protection, Artificial Intelligence and Cognitive Services”. In: *European Journal of Law and Technology* (2021).
- [20] Claudio Novelli et al. “Generative AI in EU Law: Liability, Privacy, Intellectual Property, and Cybersecurity”. In: *Computer Law & Security Review* 55 (2024), p. 106066. DOI: 10.1016/j.clsr.2024.106066.
- [21] Vagelis Papakonstantinou. “Cybersecurity as Praxis and as a State: The EU Law Path towards Acknowledgement of a New Right to Cybersecurity?” In: *Computer Law & Security Review* 44 (2022), p. 105653. DOI: 10.1016/j.clsr.2022.105653.
- [22] Michal Rampášek, Matúš Mesarčík, and Jozef Andraško. “Evolving Cybersecurity of AI-Featured Digital Products and Services: Rise of Standardisation and Certification?” In: *Computer Law & Security Review* 56 (2025), p. 106093. DOI: 10.1016/j.clsr.2024.106093.

- [23] Aurelia Tamò-Larrieux et al. “Regulating for Trust: Can Law Establish Trust in Artificial Intelligence?” In: *Regulation & Governance* 18.3 (2024), pp. 780–801. DOI: 10.1111/rego.12568.
- [24] Fabienne Ufert. “AI Regulation Through the Lens of Fundamental Rights: How Well Does the GDPR Address the Challenges Posed by AI?” In: *European Papers: A Journal on Law and Integration* 5.2 (2020), pp. 1087–1097. DOI: 10.15166/2499-8249/394.
- [25] Heng Zeng et al. “Towards a Conceptual Framework for AI-Driven Anomaly Detection in Smart City IoT Networks for Enhanced Cybersecurity”. In: *Journal of Innovation & Knowledge* 9.4 (2024), p. 100601. DOI: 10.1016/j.jik.2024.100601.
- [26] Rachel Griffin. “Governing Platforms Through Corporate Risk Management: The Politics of Systemic Risk in the Digital Services Act”. In: *European Law Open* (2025). DOI: 10.1017/elo.2025.5. URL: <https://www.cambridge.org/core/journals/european-law-open/article/governing-platforms-through-corporate-risk-management/>.
- [27] Andrea Palumbo. “Systemic Risk Management and the Constitutional Limits of Delegation under the Digital Services Act and the AI Act”. In: *European Journal of Risk Regulation* (2025). DOI: 10.1017/err.2025.6. URL: <https://www.cambridge.org/core/journals/european-journal-of-risk-regulation/article/systemic-risk-management/>.
- [28] Michele Loi, Matteo Fabbri, and Andrea Ferrario. “Regulating the Undefined: Addressing Systemic Risks in the Digital Services Act”. In: *Philosophy & Technology* (2025). DOI: 10.1007/s13347-025-00903-7. URL: <https://link.springer.com/article/10.1007/s13347-025-00903-7>.
- [29] Katharina Kaesling and Annegret Wolf. “Sustainability and Risk Management under the Digital Services Act”. In: *GRUR International* (2025). DOI: 10.1093/grurint/ikaf009. URL: <https://academic.oup.com/grurint/article/74/2/119/7943600>.
- [30] Rosa Maria Torraco. “Risk in the Digital Services Act and the AI Act: Implications for Media Freedom and Disinformation”. In: *Computer Law & Security Review* (2025). DOI: 10.1016/j.clsr.2024.106120. URL: <https://www.sciencedirect.com/science/article/pii/S0267364924001205>.

- [31] European Commission. *Digital Services Act: First Report on the Landscape of Systemic Risks*. Tech. rep. Official EU institutional report; no DOI assigned. European Commission, 2025. URL: <https://digital-strategy.ec.europa.eu/en/news/digital-services-act-report-lays-out-landscape-systemic-risks-online>.
- [32] Philipp Hacker, Atoosa Kasirzadeh, and Lilian Edwards. “AI, Digital Platforms, and the New Systemic Risk”. In: *arXiv preprint* (2025). arXiv: 2509.17878 [cs.CY]. URL: <https://arxiv.org/abs/2509.17878>.
- [33] Cecilia Panigutti et al. “Algorithm-Driven Risk Investigation in Online Platforms”. In: *Proceedings of the ACM Conference on Fairness, Accountability, and Transparency*. 2025. DOI: 10.1145/3715275.3732052. URL: <https://dl.acm.org/doi/10.1145/3715275.3732052>.
- [34] Knight-Georgetown Institute. *Systemic Risk Assessment under the Digital Services Act*. Tech. rep. Academic policy report; no DOI assigned. Georgetown University, 2025. URL: <https://kgi.georgetown.edu/research/systemic-risk-assessment-under-the-digital-services-act/>.
- [35] Niklas Eder. “Making Systemic Risk Assessments Work: How the Digital Services Act Creates a Virtuous Loop”. In: *German Law Journal* (2024). DOI: 10.1017/glj.2024.18. URL: <https://www.cambridge.org/core/journals/german-law-journal/article/making-systemic-risk-assessments-work/>.
- [36] Martin Husovec. “General Risk Management under the Digital Services Act”. In: *Oxford Handbook of Platform Law*. Peer-reviewed book chapter; no DOI assigned. Oxford University Press, 2024. URL: <https://academic.oup.com/edited-volume/56034/chapter/441610289>.
- [37] Sara Tommasi. “Risk-Based Approaches in the Digital Services Act and the Artificial Intelligence Act”. In: *SpringerBriefs in Law*. Springer, 2023. DOI: 10.1007/978-3-031-31646-5_3. URL: https://link.springer.com/chapter/10.1007/978-3-031-31646-5_3.
- [38] Defne Halil, Konrad Kollnig, and Aurelia Tamò-Larrieux. “Regulating Pressing Systemic Risks under the Digital Services Act”. In: *Internet Policy Review* (2025). DOI: 10.14763/2025.1.1761. URL: <https://policyreview.info/articles/analysis/regulating-pressing-systemic-risks>.
- [39] Gianclaudio Malgieri and Frank Pasquale. “Licensing High-Risk Artificial Intelligence: Toward Ex Ante Justification”. In: *Computer Law & Security Review* (2024). DOI: 10.1016/j.clsr.2024.106093. URL: <https://www.sciencedirect.com/science/article/pii/S0267364924000935>.

- [40] Hilary J. Allen. “A Harm-Focused Approach to Systemic Risk Regulation”. In: *Computer Law & Security Review* (2024). DOI: 10.1016/j.clsr.2024.106081. URL: <https://www.sciencedirect.com/science/article/pii/S0267364924000819>.
- [41] R. Kaushal et al. “Automated Transparency: A Legal and Empirical Analysis of the Digital Services Act Transparency Database”. In: *Proceedings of the ACM Conference on Fairness, Accountability, and Transparency*. 2024. DOI: 10.1145/3630106.3658934. URL: <https://dl.acm.org/doi/10.1145/3630106.3658934>.
- [42] Niels Vandezande. “Cybersecurity in the EU: How the NIS2-Directive Stacks up against Its Predecessor”. In: *Computer Law & Security Review* 52 (2024), p. 105890. DOI: 10.1016/j.clsr.2023.105890.
- [43] Balázs Bodó. “Governance by Trust Mediators in the Digital Society”. In: *Journal of Trust Research* (2025). DOI: 10.1080/21515581.2025.2311024. URL: <https://www.tandfonline.com/doi/full/10.1080/21515581.2025.2311024>.
- [44] Defne Halil and Konrad Kollnig. “Platform Data Access and Systemic Risk Research in the European Union”. In: *Internet Policy Review* (2024). DOI: 10.14763/2024.4.1749. URL: <https://policyreview.info/articles/analysis/platform-data-access>.
- [45] Niklas Eder, Rachel Griffin, and Martin Husovec. “Procedural Risk Governance in European Digital Law”. In: *European Law Open* (2025). DOI: 10.1017/elo.2025.14. URL: <https://www.cambridge.org/core/journals/european-law-open/article/procedural-risk-governance/>.
- [46] Information Commissioner’s Office. *Investigation into the use of data analytics in political campaigns*. Accessed 2026-02-23. UK Information Commissioner’s Office, July 2018. URL: <https://ico.org.uk/media/action-weve-taken/2260125/investigation-into-the-use-of-data-analytics-in-political-campaigns-final-20180711.pdf>.
- [47] Carole Cadwalladr and Emma Graham-Harrison. “Revealed: 50 million Facebook profiles harvested for Cambridge Analytica in major data breach”. In: *The Guardian* (Mar. 2018). URL: <https://www.theguardian.com/news/2018/mar/17/cambridge-analytica-facebook-influence-us-election>.
- [48] Federal Trade Commission. *In the Matter of Facebook, Inc.* Consent order imposing data governance and system redesign obligations. Federal Trade Commission, 2019. URL: https://www.ftc.gov/system/files/documents/cases/182_3109_facebook_decision_order_7-24-19.pdf.

- [49] Meta Platforms, Inc. *Facebook’s Plan to Prevent Abuse of Data*. Official platform response following Cambridge Analytica scrutiny. 2019. URL: <https://about.fb.com/news/2019/03/data-protection/>.
- [50] Nicholas Confessore. “Cambridge Analytica and Facebook: The scandal and the fall-out so far”. In: *The New York Times* (Apr. 2018). URL: <https://www.nytimes.com/2018/04/04/us/politics/cambridge-analytica-scandal-fallout.html>.
- [51] Ann Cavoukian. *Privacy by Design: The 7 Foundational Principles*. Tech. rep. Information and Privacy Commissioner of Ontario, 2009.
- [52] Mozilla Foundation. *YouTube Regrets: A crowdsourced investigation into YouTube’s recommendation algorithm*. Empirical evidence of systemic recommendation failures. Mozilla Foundation, 2021. URL: <https://foundation.mozilla.org/en/campaigns/youtube-regrets/>.
- [53] Manoel Horta Ribeiro, Raphael Ottoni, and Robert West. “Auditing radicalization pathways on YouTube”. In: *Proceedings of the ACM Conference on Fairness, Accountability, and Transparency* (2020), pp. 131–141. DOI: 10.1145/3351095.3372879.
- [54] YouTube. *How YouTube Addresses Harmful Content Recommendations*. Official description of recommender-system changes. 2022. URL: <https://blog.youtube/inside-youtube/how-youtube-addresses-harmful-content/>.
- [55] Meta Platforms, Inc. *Adversarial Threat Report*. Platform documentation of coordinated inauthentic behaviour and AI exploitation. Meta, 2023. URL: <https://transparency.fb.com/policies/community-standards/adversarial-threats/>.