

A Framework for Resolving AI Compliance Conflicts

Ahmed Taha

IU International University of Applied Sciences, Germany

ahmed.taha@iu.org

Abstract—The European Union Artificial Intelligence Act (EU AI Act) introduces binding requirements for robustness, security, and accountability in high-risk AI systems. In practice, these requirements often conflict with data protection obligations under the General Data Protection Regulation (GDPR), where security mechanisms depend on access to sensitive data restricted by privacy regulations. This creates a “compliance trilemma” where making systems secured, protecting privacy, and being transparent can conflict with each other. Instead of treating this as purely a legal issue, as it is often approached in regulations and standards, this paper looks at it as a security engineering challenge. In real world use cases, these requirements behave like interconnected security services: some support each other, while others create conflicts and constraints. As a result, organizations can face situations where meeting one requirement makes it harder to satisfy another, leading to design-time compliance challenges. To address this challenge, this work proposes a dependency based security governance framework that explicitly accounts for dependencies between requirements. It combines dependency modeling with the Analytic Hierarchy Process (AHP) to support structured and transparent decision making. This allows conflicting requirements to be identified and systematically prioritized based on their relative importance. The approach is demonstrated through a financial fraud detection use case, illustrating how such conflicts can be managed in a practical and defensible manner.

Keywords—EU AI Act, High-Risk AI, AI Governance, AHP, Security Dependencies

I. INTRODUCTION

Artificial Intelligence systems have evolved from experimental tools to core components of safety-critical and security-sensitive infrastructures [1]. Modern deployments in domains such as healthcare, finance, and public administration expose AI systems to adversarial threats including data poisoning, model inversion, and evasion attacks. These threats extend the traditional attack surface beyond infrastructure into data pipelines and model behavior, necessitating a shift toward “Security-by-Design”.

The European Union Artificial Intelligence Act (EU AI Act) [8] establishes the first binding regulatory framework to mitigate AI-related risks through mandates for robustness, transparency, and human oversight. Simultaneously, the General Data Protection Regulation (GDPR) [3] imposes rigorous constraints on data processing, such as data minimization and storage limits. This intersection creates an inherent regulatory friction: while the EU AI Act necessitates extensive data access for adversarial testing and traceability to ensure security, the GDPR strictly limits such practices. Consequently,

organizations must navigate a “compliance grey zone” where the security and safety requirements for AI collide with established privacy regulations and measures.

Recent studies indicate that a significant proportion of organizations lack a unified governance approach to manage the overlap between AI safety and data privacy, leaving them vulnerable to both regulatory penalties and technical exploitation [4]. Small and Medium-sized Enterprises (SMEs), in particular, face a “Compliance Trilemma” where they cannot simultaneously maximize Security (Robustness), Privacy (Minimization), and Transparency without a structured decision-support mechanism. These observations suggest that AI Act compliance cannot be addressed through legal interpretation alone, but requires a security engineering perspective.

A. Problem Statement: Regulatory Dependency Deadlock

Current compliance approaches rely on qualitative interpretation and checklist-based governance, which fail to capture dependencies between security controls and regulatory constraints. In high-risk AI systems, this leads to a Regulatory Dependency Deadlock, where required security controls are structurally blocked by privacy constraints.

For example:

- 1) Adversarial robustness (AI Act Art. 15) [8] requires access to training and attack data, which is restricted by GDPR data minimization.
- 2) Transparency (AI Act Art. 13) [8] depends on explaining model logic, which increases vulnerability to Model Inversion attacks (violating confidentiality).
- 3) Data minimization (GDPR Art. 5) [3] limits retention of data required for robustness and auditing.

These conflicts cannot be resolved through qualitative reasoning alone and require auditable decision mechanisms.

B. Contributions

To address this challenge, this paper introduces a Dependency Based Security Governance Framework that treats compliance as an engineering problem. The proposed framework includes:

- 1) Extending the service dependency framework, proposed in [5], to model the intersection of the EU AI Act and GDPR as a directed graph. This allows for the automated detection of compliance deadlocks at the design time.
- 2) Presenting a quantitative mechanism for resolving conflicting objectives using the Analytic Hierarchy Process

(AHP) [6], allowing trade-offs to be expressed through an auditable decision ratio aligned with the regulatory requirements.

- 3) Proposing a Regulatory Collision Matrix that links AI Act obligations to ISO/IEC 27001 [2] technical controls and GDPR constraints (Table I). This artifact bridges the gap between legal intent and security operations, explicitly identifying where standard security implementations trigger privacy violations.
- 4) Validating the framework against a high-risk financial fraud detection scenario, demonstrating how the model successfully prioritizes adversarial robustness over the right to erasure in a manner compatible with EU regulatory requirements for critical infrastructure.

C. Paper Organization

The remainder of this paper is structured as follows. Section II reviews related work on regulatory conflicts in AI governance, security service dependencies, and quantitative decision-making approaches. Section III introduces the proposed Dependency-Aware Security Governance Framework, including service formalization, dependency modeling, and risk-informed prioritization using AHP. Section IV validates the framework through a high-risk financial fraud detection case study. Finally, Section V concludes the paper and outlines directions for future research.

II. RELATED WORK

The interaction between the EU AI Act and the GDPR presents a complex compliance challenge. Recent studies describe this tension as a form of regulatory misalignment, where compliance with one framework may undermine compliance with another [15], [16]. This challenge is especially evident in high-risk AI systems that require continuous retraining, bias detection, and adversarial testing [11]. Although prior work provides valuable legal and policy-level interpretations, it does not address how these conflicts can be resolved at the system design or engineering level. Consequently, the legal requirements directly influence the feasibility of implementing core mechanisms such as logging, monitoring, and data retention.

Furthermore, the growing body of work on adversarial machine learning highlights the need for robust and continuously evaluated AI systems. Research has demonstrated that machine learning models are vulnerable to a range of attacks, including adversarial examples, data poisoning, model inversion, and membership inference [9], [18]. Mitigating these threats requires ongoing validation, monitoring, and, in many cases, the retention of training and evaluation data. At the same time, these defensive practices frequently involve processing sensitive or personal data, creating a direct conflict with privacy-preserving requirements. While this trade-off between security and privacy has been acknowledged in prior surveys [19], existing approaches largely focus on technical mitigation strategies and do not address the regulatory implications of these conflicts.

In parallel, security engineering research has shown that modern systems are characterized by complex dependencies between security mechanisms. In distributed and cloud environments, such interdependencies have been shown to introduce systemic risk, particularly when they are not explicitly modeled or managed [5], [17]. Despite advancements across legal analysis, security engineering, and quantitative governance, these domains remain largely disconnected. Legal scholarship provides normative interpretations but lacks operational implementation models. Security research captures technical dependencies but does not incorporate regulatory constraints. Quantitative approaches offer structured decision support, yet have not been systematically applied to resolving regulatory conflicts in AI systems.

Consequently, there is no unified framework that models regulatory obligations as interdependent services, identifies compliance deadlocks at design time, and resolves such conflicts using quantitative mechanisms. This paper addresses this gap by integrating a risk based dependency model, providing a structured and engineering-driven approach to AI governance in high-risk environments.

III. METHODOLOGY: THE DEPENDENCY-AWARE GOVERNANCE FRAMEWORK

To address the Compliance Trilemma identified in Section I, we propose a Dependency-Aware Security Assurance Framework that treats regulatory compliance as a security engineering problem rather than a static legal checklist. The framework models legally mandated trust properties as security assurance services whose misconfiguration, absence, or conflict constitutes an exploitable system weakness.

The proposed methodology is explained and demonstrated through a representative case study to illustrate its practical applicability in Section IV. The methodology consists of three phases as shown in Fig. 1:

- 1) Security Assurance Service Formalization, where legal obligations are mapped to technical security assurance services;
- 2) Dependency Conflict Modeling, where functional and blocking dependencies are identified; and
- 3) Risk Prioritization, where conflicting security objectives are weighted using the Analytic Hierarchy Process (AHP) to support defensible architectural decisions within a legally permissible design space.

A. Phase 1: Security Assurance Service Formalization

We propose transforming the security- and safety-relevant technical obligations of the EU AI Act and GDPR into a Security Assurance Service Catalog (S). Let $S = \{s_1, s_2, \dots, s_n\}$ be the set of security assurance services required for the secured operation of a High-Risk AI system. Each service s_i is defined by a tuple:

$$s_i = \langle ID, Capability, ConstraintSet \rangle \quad (1)$$

where each service is defined by three elements: an *ID* that links it to a specific regulatory requirement, a *Capability*

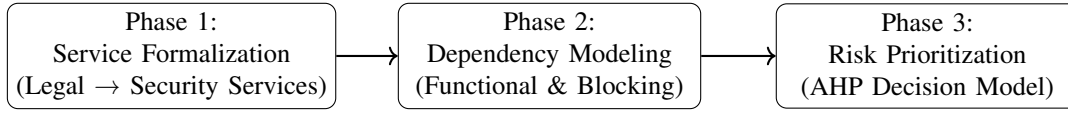


Fig. 1. Dependency-Aware Security Governance Framework. The model transforms regulatory obligations into security services, identifies dependency conflicts, and resolves them through quantitative risk-informed prioritization.

that describes what the system does in practice to meet that requirement, and a *ConstraintSet* (C_i) that captures the constraints limiting how this capability can be implemented, such as data protection rules.

B. Phase 2: Dependency Conflict Modeling

The core contribution of the framework lies in modeling security dependency conflicts between assurance services. Building on the service dependency model proposed by Taha et al. [5], we define a dependency relation D over the service set S , where $s_i \rightarrow s_j$ implies that the secure execution of service s_i requires the correct functioning of s_j . We identify two types of dependencies:

- 1) Functional Dependency ($s_i \xrightarrow{F} s_j$): The output of s_j is a necessary input for s_i (e.g., $s_{robustness} \xrightarrow{F} s_{data_access}$, robustness testing cannot be securely performed without access to representative training or evaluation data).
- 2) Blocking Dependency ($s_k \xrightarrow{B} s_j$): A constraint in service s_k prevents the execution of s_j (e.g., $s_{GDPR_privacy} \xrightarrow{B} s_{data_access}$, privacy constraints block unrestricted data access when sensitive personal data is involved).

A “Security Compliance Deadlock” occurs if there exists a path where a Functional Dependency is severed by a Blocking Dependency along a direct or transitive path in the dependency graph:

$$\exists(s_i, s_j, s_k) : (s_i \xrightarrow{F} s_j) \wedge (s_k \xrightarrow{B} s_j) \quad (2)$$

For example, if $s_{robustness} \xrightarrow{F} s_{data_access}$ but $s_{GDPR_privacy} \xrightarrow{B} s_{data_access}$, the system is flagged as “Non-Compliant by Design.”

To support dependency analysis, we develop a Regulatory Collision Matrix (Table I) that maps EU AI Act requirements to corresponding security assurance controls, ISO/IEC 27001 domains, and GDPR constraints.

The matrix operationalizes the framework by linking legal obligations to their technical implementations and associated constraints. This enables the identification of blocking dependencies, where enforcing a regulatory constraint (e.g., GDPR data minimization) restricts the implementation of a required security control (e.g., adversarial robustness).

These interactions reveal design-time compliance deadlocks, which serve as inputs to the risk prioritization process described in Phase 3.

C. Phase 3: Risk Prioritization

When a deadlock is detected, we apply the AHP methodology to mathematically derive a justifiable decision. We

structure this process into three explicit steps as shown in Fig. 2:

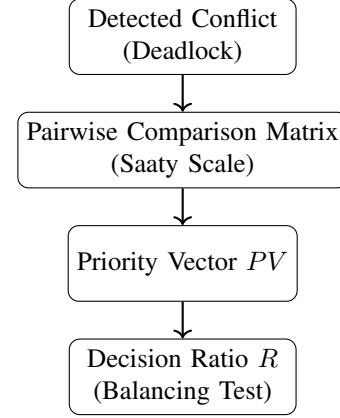


Fig. 2. AHP-Based Decision Flow for resolving regulatory conflicts.

1) *Step 1. The Pairwise Comparison Matrix (A)*: This step quantifies the relative importance of conflicting compliance objectives based on the organization’s risk appetite. To achieve this, a pairwise comparison matrix $A \in \mathbb{R}^{n \times n}$ is constructed, where each element a_{ij} represents the relative importance of objective i over objective j , using the Saaty fundamental scale (1–9), where 1 denotes equal importance, 3 moderate, 5 strong, 7 very strong, and 9 extreme importance [20].

$$A = \begin{bmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{n1} & a_{n2} & \dots & a_{nn} \end{bmatrix} \quad (3)$$

The matrix is governed by the Rule of Reciprocity, defined as:

$$a_{ji} = \frac{1}{a_{ij}} \quad (4)$$

This ensures logical consistency within the matrix; for example, if objective i is assigned a value of 3 (moderate importance) over objective j , then objective j is automatically assigned a value of $1/3$ relative to objective i .

2) *Step 2. Deriving the Priority Vector (PV)*: In this step, the pairwise comparisons in matrix A are converted into a single vector PV that represents the relative importance of each objective. The priority vector is obtained by solving the eigenvalue equation:

$$A \cdot PV = \lambda_{max} \cdot PV \quad (5)$$

Where:

TABLE I
THE REGULATORY COLLISION MATRIX: AI ACT GOALS, SECURITY CONTROLS, ISO 27001 DOMAINS, AND GDPR CONSTRAINTS

AI Act Requirement (Compliance Objective) [8]	Security Assurance Control (Implementation)	ISO/IEC 27001 Domain [2]	GDPR Constraint [3]	Blocking Dependency / Deadlock Description
Art. 15 (Accuracy, Robustness & Cybersecurity) [8]	Adversarial Training (Retention of adversarial inputs and attack patterns for regression and robustness validation) [9]	A.14 - System Acquisition, Development and Maintenance [2]	Arts. 5(1)(c), 17 (Data Minimization, Erasure) [3]	Retention Deadlock: Validating robustness requires retaining adversarial inputs that may violate data minimization and erasure mandates.
Art. 10 (Data Governance) [8]	Bias and Fairness Detection (Processing demographic attributes to assess discriminatory outcomes) [10], [11]	A.12 - Operations Security (Data Management) [2]	Art. 9 (Special Categories of Data) [3]	Classification Deadlock: Detecting bias requires processing sensitive demographic attributes generally prohibited under GDPR special category rules.
Art. 14 (Human Oversight) [8]	Explainable AI (XAI) (Disclosure of decision rationale to operators and users) [12]	A.12.4 - Logging and Monitoring [2]	Arts. 5(1)(f), 32 (Integrity, Confidentiality) [3]	Reconstruction Deadlock: High-fidelity explanations risk "model inversion," where internal logic exposure compromises system confidentiality.
Art. 12 (Record-Keeping) [8]	Immutable Logging (WORM-based audit trails for traceability and accountability) [14]	A.12.4 - Logging and Monitoring [2]	Art. 5(1)(e) (Storage Limitation) [3]	Lifecycle Deadlock: Mandatory traceability logs for AI accountability often conflict with GDPR storage limitation periods.
Art. 13 (Transparency) [8]	Automated Model Fact Sheets (User-facing disclosures and system transparency mechanisms) [13]	A.13 - Communications Security [2] A.18 - Compliance [2]	Arts. 5(1)(f), 32 (Integrity, Confidentiality) [3]	Exposure Deadlock: Public transparency regarding model parameters or logic expands the attack surface for evasion or probing attacks.

- *PV* (Priority Vector): The normalized vector (summing to 1) representing the weight of each objective, used for decision-making.
- λ_{max} (Principal Eigenvalue): A scalar value used to assess the consistency of the pairwise comparisons. Ideally, $\lambda_{max} \approx n$. If λ_{max} deviates significantly from n , the judgments are inconsistent and should be revised.

The output of this step is a normalized priority vector representing the relative weights for all n objectives:

$$PV = [w_1, w_2, \dots, w_n] \quad (6)$$

Where each w_i corresponds to the weight of objective i , such that $\sum_{i=1}^n w_i = 1$. Thus, each element of *PV* directly corresponds to a compliance objective defined in Step 1, providing a quantitative representation of their relative importance.

3) *Step 3. The Algorithmic Balancing Test*: A balancing test is required to compare competing objectives, where the benefit of a given objective is weighed against the impact on another conflicting objective.

In this work, the balancing test is modeled by mapping these competing concerns to the compliance objectives defined in Step 1. Using the priority vector $PV = [w_1, w_2, \dots, w_n]$ derived in Step 2, each objective is assigned a quantitative weight reflecting its relative importance.

$$R = \frac{w_i}{w_j} \quad (7)$$

where w_i and w_j represent the weights of two objectives.

If $R > \tau$ (where $\tau \geq 1.0$ is the organizational risk threshold), the framework supports prioritizing objective i over objective j . Otherwise, objective j is prioritized.

The resulting value of R serves as a transparent and auditable decision artifact [7].

The methodology is applied and validated through the case study in Section IV, demonstrating how the proposed approach enables structured and auditable decision-making in practice.

IV. CASE STUDY: REGULATORY CONFLICT RESOLUTION IN AI-BASED FRAUD DETECTION

To demonstrate the practical application of the proposed dependency-aware security assurance framework, we apply the methodology to a High-Risk AI system deployed for financial fraud detection and anti-money laundering (AML) monitoring. In this domain, security robustness requirements frequently collide with privacy and transparency obligations imposed by regulatory frameworks, such as the EU AI Act and GDPR.

The case study follows the three phases defined in Section III: (1) Security Assurance Service Formalization, (2) Dependency Conflict Modeling, and (3) Risk Prioritization.

A. Phase 1: Security Assurance Service Formalization

The financial institution operates a machine learning system that analyzes transaction records to detect fraudulent financial activity. The system continuously retrains its detection model using historical fraud patterns in order to remain robust against evolving evasion techniques.

Based on the Regulatory Collision Matrix (Table I), three relevant security assurance services are identified:

- $s_1 = \langle AIAct_Art15, \text{Adversarial Training}, C_1 \rangle$

This service implements robustness by training the model using past attack patterns and suspicious transactions (i.e., Adversarial Training). While this improves security, it requires storing data that may be restricted by regulations, captured in the constraint set (C_1).

- $s_2 = \langle \text{GDPR_Art17}, \text{Personal Data Deletion}, C_2 \rangle$
This service represents privacy-related requirements, specifically the enforcement of the Right to Erasure, which ensures that personal data is deleted upon request. It ensures compliance with data protection requirements by removing stored user data from the system. However, this capability is constrained by operational and security requirements (C_2), such as the need to retain data for auditing, fraud detection, or model improvement.
- $s_3 = \langle \text{AIAct_Art13}, \text{Decision Explainability}, C_3 \rangle$
This service implements transparency requirements by enabling explanations of automated decisions for auditors or affected users. It supports understanding and accountability of model behavior. However, this capability is constrained by limitations (C_3) such as the risk of exposing sensitive information, protecting proprietary models, or the inherent complexity of certain machine learning techniques.

The resulting Security Assurance Service Catalog is therefore:

$$S = \{s_1, s_2, s_3\}$$

Each service corresponds to a regulatory obligation and introduces technical capabilities as well as operational constraints.

B. Phase 2: Dependency Conflict Modeling

Next, we model the functional and blocking dependencies between the identified services.

Functional Dependency: Robust fraud detection requires access to historical transaction data in order to perform adversarial training. Therefore:

$$s_1 \xrightarrow{F} s_{data}$$

where s_{data} represents the data access capability required for model training.

Blocking Dependency: The GDPR Right to Erasure may remove the same transaction records required for adversarial training. This creates the blocking relation:

$$s_2 \xrightarrow{B} s_{data}$$

This produces the following deadlock condition:

$$(s_1 \xrightarrow{F} s_{data}) \wedge (s_2 \xrightarrow{B} s_{data})$$

which indicates that the system becomes *Non-Compliant by Design* if the two services are enforced simultaneously without a prioritization mechanism.

A second conflict emerges between transparency (s_3) and robustness (s_1). Revealing detailed model explanations may expose the detection logic to adversaries and enable evasion attacks. Thus transparency can indirectly weaken robustness

guarantees. These conflicts illustrate how regulatory requirements originating from different frameworks can produce design-time compliance deadlocks in High-Risk AI systems.

C. Phase 3: Risk Prioritization

To resolve these conflicts, the framework applies the Analytic Hierarchy Process (AHP) to determine the relative importance of the competing assurance services. The institution evaluates the services based on its regulatory and operational risk posture in the financial sector.

- Robustness (s_1) vs Privacy (s_2):
Preventing money laundering and financial fraud is considered critically important for financial stability and regulatory compliance. A Saaty score of 7 (very strong importance) is assigned.
- Robustness (s_1) vs Transparency (s_3):
Excessive transparency may expose detection logic to adversaries. Robustness is therefore strongly prioritized (score 5).
- Transparency (s_3) vs Privacy (s_2):
Transparency is moderately prioritized over erasure to support auditing and accountability (score 2).

Each element in the matrix A (cf., equation 3) represents the relative importance of service s_i over service s_j . For example:

$$\begin{aligned} a_{11} &= \frac{s_1}{s_1} = 1 && \text{(Robustness vs Robustness)} \\ a_{12} &= \frac{s_1}{s_2} = 7 && \text{(Robustness vs Privacy)} \\ a_{13} &= \frac{s_1}{s_3} = 5 && \text{(Robustness vs Transparency)} \end{aligned}$$

The remaining elements are derived using the reciprocity property $a_{ji} = 1/a_{ij}$, and diagonal elements are equal to 1.

Thus, the pairwise comparison matrix is defined as:

$$A = \begin{matrix} & \begin{matrix} s_1 & s_2 & s_3 \end{matrix} \\ \begin{matrix} s_1 \\ s_2 \\ s_3 \end{matrix} & \begin{pmatrix} 1 & 7 & 5 \\ 1/7 & 1 & 1/2 \\ 1/5 & 2 & 1 \end{pmatrix} \end{matrix}$$

Solving the eigenvector equation yields the normalized priority vector:

$$PV = \begin{pmatrix} w_1 & w_2 & w_3 \\ 0.73 & 0.08 & 0.19 \end{pmatrix}$$

The result indicates that security robustness represents 73% of the system's governance priority, while privacy obligations represent 8%.

Algorithmic Balancing Test

When a data subject invokes the Right to Erasure, the framework evaluates the competing objectives using a Legitimate Interest balancing ratio:

$$R = \frac{w_1}{w_2}$$

Substituting the calculated weights yields:

$$R = \frac{0.73}{0.08} = 9.12$$

Since the ratio is significantly greater than the threshold value of 1, the model concludes that the legitimate interest in fraud prevention outweighs the individual's request for data deletion. Based on this evaluation, the institution can lawfully retain the transaction records under the Legitimate Interest legal basis defined in GDPR Article 6(1)(f). The computed ratio $R = 9.12$ serves as an auditable artifact demonstrating that the organization performed a structured proportionality assessment when resolving the regulatory conflict.

This example illustrates how the proposed framework enables organizations to detect compliance deadlocks at design time and resolve them through a transparent and mathematically grounded prioritization process.

V. CONCLUSION AND FUTURE WORK

The EU Artificial Intelligence Act introduces a paradigm shift in the regulation of AI systems by embedding security, robustness, and accountability as enforceable legal obligations. However, as demonstrated throughout this paper, these requirements do not operate in isolation. Instead, they intersect with existing regulatory approach such as the GDPR in ways that create structural conflicts at the system design level.

The analysis demonstrates that these conflicts cannot be effectively addressed through purely qualitative compliance approaches. Rather, they must be treated as security engineering problems characterized by interdependent services. Thus, we proposed a Dependency based Security Governance Framework that formalizes regulatory requirements as security assurance services, models their functional and blocking dependencies, and resolves resulting deadlocks using a quantitative prioritization mechanism based on the Analytic Hierarchy Process.

Furthermore, we demonstrated the validity of the proposed model using a financial fraud detection case study, in which the proposed approach enables organizations to detect "Non-Compliant by Design" conditions early in the system life-cycle and to justify trade-offs between competing regulatory objectives through transparent and auditable metrics. From a practical perspective, the model supports security architects, compliance officers, and system designers in aligning AI safety requirements with data protection obligations without relying on ad hoc or purely legal reasoning. More broadly, it contributes to bridging the gap between regulatory intent and technical implementation, which remains one of the central challenges in AI governance.

Future work will focus on extending the framework to incorporate additional regulatory instruments such as the NIS2 Directive and sector-specific compliance requirements, as well as automating dependency detection within AI system pipelines through policy-aware tooling.

REFERENCES

- [1] B. Biggio and F. Roli, "Wild patterns: Ten years after the rise of adversarial machine learning," *Pattern Recognition*, vol. 84, pp. 317–331, 2018.
- [2] International Organization for Standardization, *ISO/IEC 27001:2022 Information security, cybersecurity and privacy protection, Information security management systems Requirements*, ISO, Geneva, Switzerland, 2022.
- [3] European Parliament and Council, "Regulation (EU) 2016/679 on the protection of natural persons with regard to the processing of personal data and on the free movement of such data (General Data Protection Regulation)," *Official Journal of the European Union*, L 119/1, 2016.
- [4] N. Hussain, "Navigating the EU AI Act: Strategic Compliance and Governance Recommendations for Organisations Adopting High Risk AI systems," Master Thesis, IU International University of Applied Sciences, 2025.
- [5] A. Taha, P. Metzler, J. Luna, R. Trapero, and N. Suri, "Identifying and Utilizing Dependencies Across Cloud Security Services," in *Proceedings of the 11th ACM on Asia Conference on Computer and Communications Security (ASIA CCS '16)*, 2016, pp. 297–308.
- [6] A. Taha, R. Trapero, J. Luna, and N. Suri, "AHP-Based Quantitative Approach for Assessing and Comparing Cloud Security," *IEEE International Conference on Trust, Security and Privacy in Computing and Communications (TrustCom)*, 2014.
- [7] Article 29 Data Protection Working Party, "Opinion 06/2014 on the notion of legitimate interests of the data controller under Article 7 of Directive 95/46/EC," 844/14/EN, 2014.
- [8] N. Smuha, "Regulation 2024/1689 of the European Parliament & Council of June 13, 2024 (EU Artificial Intelligence Act)," *International Legal Materials*, vol. 64, no. 5, pp. 1234–1381, 2025.
- [9] I. Goodfellow, J. Shlens, and C. Szegedy, "Explaining and Harnessing Adversarial Examples," in *Proc. ICLR*, 2015.
- [10] S. Barocas and A. Selbst, "Big Data's Disparate Impact," *California Law Review*, vol. 104, no. 3, 2016.
- [11] N. Mehrabi et al., "A Survey on Bias and Fairness in Machine Learning," *ACM Computing Surveys*, vol. 54, no. 6, 2021.
- [12] R. Guidotti et al., "A Survey of Methods for Explaining Black Box Models," *ACM Computing Surveys*, vol. 51, no. 5, 2018.
- [13] M. Mitchell et al., "Model Cards for Model Reporting," in *Proc. FAT**, 2019.
- [14] A. Behl and K. Behl, "Cybersecurity and Cyberwar: What Everyone Needs to Know," Oxford University Press, 2021.
- [15] M. Veale and F. Borgesius, "Demystifying the Draft EU Artificial Intelligence Act," *Computer Law Review International*, 2021.
- [16] L. Edwards, "Regulating AI in Europe: Four Problems and Four Solutions," *European Journal of Law and Technology*, 2022.
- [17] J. Somorovsky, "Systematic Analysis of Security Protocols," 2017.
- [18] R. Shokri et al., "Membership Inference Attacks Against Machine Learning Models," *IEEE S&P*, 2017.
- [19] N. Papernot et al., "SoK: Security and Privacy in Machine Learning," *IEEE EuroS&P*, 2018.
- [20] T. Saaty, "Decision making with the analytic hierarchy process," *International Journal of Services Sciences*, vol. 1, no. 1, pp. 83–98, 2008.