

Translating EU AI Act Transparency Requirements into Software Engineering Specifications

Eduardo Almentero^{1,4}[0000–0002–5664–1080], Henrique P. S. Sousa², Roxana L. Q. Portugal³[0000–0001–7693–5353], Tadeu M. Classe²[0000–0001–9849–5133], and Marcos Kalinowski⁴[0000–0003–1445–3425]

¹ UFRRJ, Seropédica, Brazil. almentero@ufrjr.br

² UNIRIO, Rio de Janeiro, Brazil. {hsousa,tadeu.classe}@uniriotec.br

³ UNSAAC, Cuzco, Peru. roxana.portugal@ifkw.lmu.de

⁴ PUC-Rio, Rio de Janeiro, Brazil. kalinowski@inf.puc-rio.br

Abstract. The European Union Artificial Intelligence Act introduces strict obligations for high-risk AI systems, yet a significant translation gap remains between its legal mandates and the precise specifications that software engineering practice demands. This paper proposes a five-phase operationalization process grounded in Design Science Research, integrating the Non-Functional Requirements Framework with the ISO/IEC 25059:2023 quality model to bridge this gap for Article 13 transparency requirements. Applied analytically to a credit scoring AI system as an illustrative high-risk domain, the process extracts atomic legal obligations, maps them to standardized quality characteristics, instantiates context-specific softgoals, models interdependencies through Softgoal Interdependency Graphs, and produces measurable acceptance criteria in a bidirectional traceability matrix, yielding nine traceable requirements across nine quality dimensions. While Article 13 can be largely operationalized through this approach, some provisions remain technically indeterminate and require accompanying organizational governance mechanisms. The paper contributes a structured vocabulary and roadmap for translating AI regulation into verifiable software requirements, intended for multidisciplinary teams of requirements engineers, legal experts, and domain specialists.

Keywords: EU AI Act · NFR Framework · ISO/IEC 25059 · Requirements Engineering · AI Transparency · Software Quality.

1 Introduction

The proliferation of artificial intelligence systems across high-stakes domains, including clinical decision support, credit risk assessment, predictive policing, and automated recruitment, has accelerated at a pace that outstrips existing governance frameworks. As AI systems increasingly mediate consequential decisions affecting individuals' health, liberty, and economic opportunities, concerns regarding their opacity, accountability, and potential for systematic harm

have intensified among regulators, civil society, and the scientific community [1], leaving deployers without sufficient guidance to evaluate AI-generated outputs and affected individuals without meaningful recourse when algorithmic decisions cause harm [2].

The European Union Artificial Intelligence Act introduces a risk-based approach to ensure the safety and trustworthiness of AI systems [3], with Article 13 mandating that high-risk AI systems provide clear instructions for use detailing their characteristics, capabilities, and limitations [4]. These obligations aim to enable deployers to appropriately interpret AI outputs, thereby mitigating risks to health, safety, and fundamental rights [5, 6]. The problem is that Article 13 is written at a level of abstraction that accommodates judicial interpretation but does not translate naturally into engineering specifications. Software engineering demands precision, measurability, and verifiability, and the qualitative language of legal texts does not provide these [7].

This translation gap has been recognized in related domains. The requirements engineering literature has produced structured approaches for operationalizing legal obligations [10–12, 33], and Explainable AI (XAI) research has developed techniques for making model outputs interpretable [13]. However, GDPR approaches were designed for well-bounded data protection obligations and do not account for the dynamic, context-sensitive nature of AI system behavior. XAI frameworks similarly fall short, as their focus on post-hoc explanation diverges from the system-level documentation and operational constraints that Article 13 concerns [13, 8]. A review of relevant quality standards, including ISO/IEC 25059:2023 and IEEE 7001:2021, further reveals that while these frameworks establish pertinent quality characteristics, they do not provide explicit mappings from specific legal articles to verifiable engineering requirements.

This paper addresses this gap by proposing a five-phase operationalization process, grounded in Design Science Research, that integrates the Non-Functional Requirements Framework and Softgoal Interdependency Graphs with ISO/IEC 25059:2023 to translate Article 13 into verifiable software engineering specifications [14–16]. The paper is organized around three research questions: RQ1: How can Article 13 obligations be decomposed into discrete, verifiable software engineering requirements? RQ2: To what extent do the NFR Framework and ISO/IEC 25059:2023 provide sufficient coverage for this operationalization? RQ3: What are the practical implications for AI system providers and deployers seeking regulatory compliance?

The contributions are: (i) a systematic mapping of Article 13 obligations onto software engineering requirements organized by quality dimension; (ii) a five-phase operationalization process combining the NFR Framework with ISO/IEC 25059:2023; (iii) verifiable acceptance criteria enabling conformity assessment by developers and market surveillance authorities; and (iv) a critical analysis of Article 13 provisions that resist direct technical operationalization. The application yields nine software engineering requirements across nine quality dimensions. The analysis reveals that while the majority of Article 13 obligations can be systematically operationalized, a subset remains technically indeterminate due to

dependence on deployment context and judicial interpretation, suggesting that full compliance requires accompanying organizational governance mechanisms beyond technical specification alone.

2 Theoretical Background

AI Regulation and the EU AI Act. The EU AI Act establishes a risk-based regulatory framework that classifies AI systems into four risk tiers, with high-risk systems subject to the most stringent requirements [17]. Systems listed in Annex III of the EU AI Act (the annex of the regulation itself, not of this paper) including those used in credit scoring, recruitment, and clinical decision support, must comply with obligations covering data governance, human oversight, and transparency. Article 13 specifically mandates that high-risk AI systems be accompanied by instructions for use that disclose their intended purpose, performance characteristics, known limitations, and conditions under which human oversight is required.

AI Transparency and Explainability. Transparency in AI systems is a system-level property concerned with making relevant aspects of a system, including its behavior, decision processes, and operational conditions, accessible and understandable to stakeholders [18]. Explainability, in contrast, focuses on making individual model outputs interpretable by exposing how specific decisions are produced, typically through post-hoc techniques such as LIME, SHAP, or attention visualization [19–21]. Although closely related, explainability is best understood as a contributing component of transparency rather than an equivalent concept, as it operates at the level of individual decisions rather than the system as a whole [18]. This distinction is consequential for Article 13 compliance: the Act’s obligations are primarily system-level transparency requirements addressed to providers and deployers, not user-facing explanation mechanisms.

Requirements Engineering for Legal Compliance. The challenge of translating legal texts into software specifications has been extensively studied across various regulatory domains. Structured methodologies have been developed to derive verifiable requirements from regulatory prose, employing techniques such as natural language processing, goal modeling, and ontology-based approaches, among others [11, 12, 22–24]. These approaches demonstrate that normative obligations can be systematically decomposed into technical specifications amenable to verification and validation. Villamizar et al. [31] propose a complementary vision framework integrating artefact-based and perspective-based methods to support the specification of trustworthy AI requirements, illustrating how RE practices can be structured to address trustworthiness concerns including transparency and accountability. Habiba et al. [32] confirm, through a systematic mapping of 126 RE4AI studies, that requirements specification and explainability remain among the most prevalent challenges in requirements engineering for AI-based systems, underscoring the practical need for structured operationalization approaches such as the one proposed here.

Non-Functional Requirements and the NFR Framework. Non-functional requirements (NFRs) express quality properties of software systems

that cannot be reduced to discrete functional behaviors, including performance, security, and transparency [25]. The NFR Framework, introduced by [26], provides a principled mechanism for modeling NFRs as softgoals, which are goals that are satisfied to an acceptable degree rather than strictly satisfied in a binary sense. Softgoal Interdependency Graphs (SIGs) extend this formalism by capturing hierarchical decompositions and contribution relationships among softgoals, including positive contributions such as Help and Make, and negative contributions such as Hurt and Break. This capacity to model partial satisfaction and inter-goal trade-offs makes the NFR Framework particularly appropriate for transparency obligations, which are inherently matters of degree and frequently in tension with competing properties such as privacy and security.

AI Quality Standards. ISO/IEC 25059:2023 extends the SQuaRE quality model to AI systems, defining quality characteristics relevant to AI transparency, including Explainability, Understandability, Controllability, and Auditability [27]. IEEE 7001:2021 similarly establishes measurability criteria for transparency in autonomous systems across five stakeholder levels [28]. While both standards provide a valuable vocabulary for characterizing AI transparency properties, neither was designed to map specific legislative obligations to verifiable engineering requirements. They therefore serve as necessary but insufficient foundations for Article 13 compliance, a limitation this paper addresses by integrating ISO/IEC 25059:2023 into a structured operationalization process.

3 Methodology

This study adopts a Design Science Research (DSR) paradigm [29], which is appropriate when the research objective is the design and evaluation of a novel artifact intended to solve a practical engineering problem. In this context, the artifact is a five-phase operationalization process that translates the transparency obligations of Article 13 of the EU AI Act into verifiable software engineering requirements. DSR was selected over purely analytical or empirical approaches because the translation gap identified in Section 2 is not merely a theoretical problem but an engineering one that requires a prescriptive solution applicable in practice. Concretely, practitioners tasked with demonstrating compliance with Article 13 currently lack a structured method for translating legal provisions into engineering specifications, a shared vocabulary linking legal concepts to technical quality characteristics, and auditable artifacts connecting regulatory obligations to measurable criteria, precisely the type of engineering problem that Design Science Research (DSR) is intended to address.

The process was constructed through an iterative design grounded in three DSR cycles. The relevance cycle motivated the design by establishing the practical problem of Article 13 operationalization and the requirements any solution must satisfy. The design cycle produced the process artifact through successive refinements, drawing on the NFR Framework, Softgoal Interdependency Graphs, and ISO/IEC 25059:2023 as theoretical foundations. The rigor cycle evaluated the process against established requirements engineering principles and assessed its outputs through an illustrative application to a credit scoring AI system,

examining requirement coverage, mapping consistency, and practical applicability from a developer perspective.

It is important to note that the credit scoring application presented in Section 4 is an illustrative application rather than a case study in the methodological sense. It demonstrates how the pipeline operates when applied to a specific domain and serves to make the process more concrete and understandable, but it does not constitute empirical validation. The process relies on structured analytical reading of the legal text and holistic judgment when mapping legal obligations to ISO quality characteristics. Key steps, including legal requirement extraction, mapping to ISO characteristics, and softgoal instantiation, involve interpretive decisions that depend on the analyst’s understanding of both the legal text and the standard. This subjectivity is an acknowledged limitation and is discussed further in Section 7.

The theoretical foundations were selected based on their complementary contributions to the operationalization problem. The NFR Framework and Softgoal Interdependency Graphs provide a principled mechanism for modeling non-functional requirements as satisficeable goals rather than binary conditions, which is essential given that transparency is inherently a matter of degree. ISO/IEC 25059:2023 serves as a semantic bridge between legal language and engineering terminology, grounding the process in internationally recognized AI quality characteristics and preventing the proliferation of ad hoc terminology that arises when legal concepts are translated informally [30].

4 The Proposed Process

The process operationalizes Article 13 obligations through five phases that transform qualitative legal provisions into measurable software engineering specifications (Fig. 1). Each phase produces a concrete artifact that serves as input to the subsequent phase, and the entire pipeline is unified by a bidirectional traceability matrix. The first two phases operate directly on the legal text and ISO/IEC 25059:2023 and are therefore domain independent. From Phase 3 onwards, the process requires domain-specific instantiation. To illustrate these context-dependent phases, a credit scoring AI system is used as a representative high-risk domain, as defined in Annex III of the EU AI Act. The outputs presented for these phases are illustrative instances of what the process produces, not prescriptive requirements for credit scoring systems specifically.

4.1 Phase 1: Legal Requirement Extraction.

The first phase extracts atomic, engineering-relevant legal requirements from the text of Article 13 through structured analytical reading. The extraction process identifies discrete obligation units within the legal text, namely clauses that impose specific and independently addressable duties on providers or deployers, and classifies each according to its transparency type. The process systematically distinguishes between two categories of obligation: information transparency requirements, which concern what the system must disclose to deployers, and

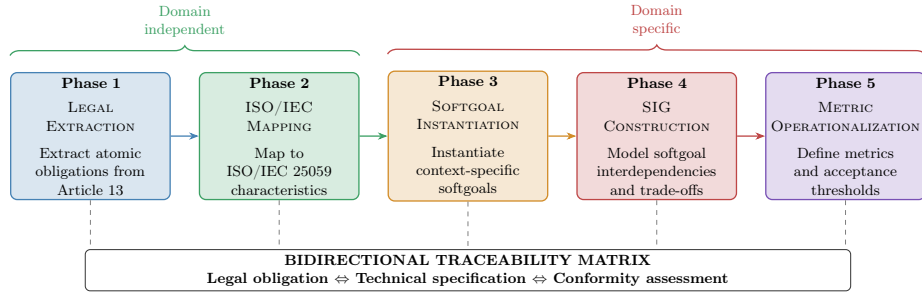


Fig. 1. Five-Phase Operationalization Process.

Table 1. Illustrative extraction of atomic legal requirements from Article 13 of the EU AI Act for a credit scoring AI system.

ID	Legal Requirement	Type
LR1	The system must identify itself as an AI system to the deployer	Information
LR2	The system must disclose its intended purpose and use conditions	Information
LR3	The system must describe its level of accuracy and performance metrics	Information
LR4	The system must disclose known limitations and foreseeable misuse scenarios	Information
LR5	The system must describe the characteristics of the input data it requires	Information
LR6	The system must provide instructions enabling deployers to interpret its outputs	Process
LR7	The system must indicate when human oversight is required	Process
LR8	The system must support the ability of deployers to monitor its operation	Process
LR9	The system must describe any foreseeable changes in performance over time	Information

process transparency requirements, which concern how the system must behave during operation. This classification is performed through careful reading and interpretation of the legal text, without automated parsing, and reflects the analyst’s judgment regarding the engineering relevance of each clause. The involvement of legal or compliance experts to validate the extracted requirements is strongly recommended to mitigate interpretive subjectivity at this stage. The output consists of a set of discrete legal requirements, each identified by a unique identifier and assigned a transparency type, as illustrated in Table 1. This phase is independent of the deployment domain.

4.2 Phase 2: Mapping to ISO/IEC 25059.

The second phase classifies each extracted legal requirement according to the standardized AI quality characteristics defined in ISO/IEC 25059:2023, as illustrated in Table 2. This classification is performed through a holistic analytical

process: the analyst reads each legal requirement, identifies the core engineering concern it expresses, and associates it with the ISO quality characteristic or characteristics that most closely capture that concern. The process is manual and relies on the analyst’s combined understanding of the legal text and the standard’s quality model. A single legal requirement may map to more than one quality characteristic where its obligations span multiple engineering concerns. This mapping reduces interpretive ambiguity and prevents the proliferation of ad hoc vocabulary that arises when legal concepts are translated informally into technical ones [30]. As this phase operates directly on the legal text and the standard, the resulting mappings are domain independent. Alternative mappings are possible, and involvement of a requirements engineer familiar with ISO/IEC 25059:2023 is recommended to ensure consistency.

Table 2. Illustrative mapping of extracted legal requirements to ISO/IEC 25059:2023 quality characteristics for a credit scoring AI system.

ID	Legal Requirement	ISO/IEC 25059 Characteristic
LR1	AI system self-identification	Transparency
LR2	Intended purpose disclosure	Transparency, Controllability
LR3	Accuracy and performance metrics	Performance Efficiency, Reliability
LR4	Limitations and misuse scenarios	Risk Management, Transparency
LR5	Input data characteristics	Data Quality, Transparency
LR6	Output interpretation instructions	Explainability, Understandability
LR7	Human oversight indication	Controllability, Human Agency
LR8	Operational monitoring support	Maintainability, Auditability
LR9	Performance change disclosure	Reliability, Transparency

4.3 Phase 3: Softgoal Instantiation.

The third phase contextually instantiates each ISO quality characteristic as a softgoal using the NFR syntax Type [Topic], as illustrated in Table 3. The Type component identifies the quality characteristic, while the Topic component anchors the softgoal to the specific system aspect under consideration. The Topic component is inherently context dependent and must be defined by the system provider and domain expert in accordance with the specific operational context, system architecture, and deployment domain.

4.4 Phase 4: SIG Construction.

Credit scoring is explicitly listed in Annex III of the EU AI Act as a high-risk domain subject to Article 13 obligations. It is used here because the domain creates particularly salient trade-offs: the obligation to explain credit decision logic may conflict with proprietary model protection, making the SIG’s capacity to represent negative contribution links especially valuable in this context.

Table 3. Illustrative contextual instantiation of ISO/IEC 25059:2023 quality characteristics as softgoals for a credit scoring AI system.

ISO Characteristic	Softgoal Instantiation
Transparency	Transparency [System Identity]
Explainability	Explainability [Credit Decision Logic]
Understandability	Understandability [Output Interpretation]
Controllability	Controllability [Deployment Conditions]
Reliability	Reliability [Performance Stability]
Data Quality	Data Quality [Input Feature Disclosure]
Risk Management	Risk Management [Limitation Disclosure]
Auditability	Auditability [Operational Monitoring]
Human Agency	Human Agency [Oversight Triggers]

The fourth phase models the hierarchical relationships among the instantiated softgoals using a Softgoal Interdependency Graph. The SIG captures AND/OR decompositions and Contribution Links representing the directional influence one softgoal exerts on another. Four link types are employed: Make (full positive enablement), Help (partial positive contribution), Hurt (partial negative impact), and Break (full prevention). The specific interdependencies are inherently domain dependent; the SIG must be constructed by the system provider with knowledge of the operational environment, as illustrated in Fig. 2.

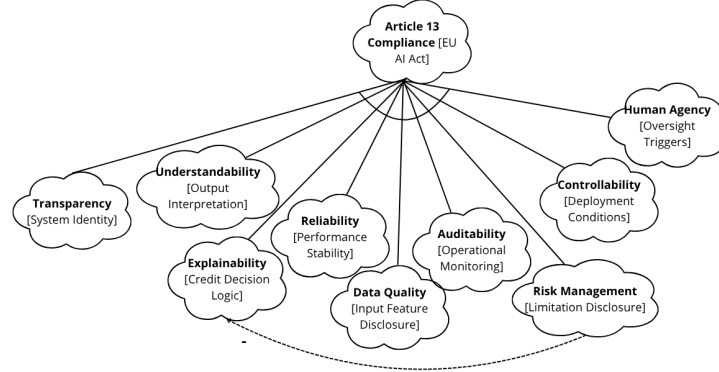


Fig. 2. Softgoal Interdependency Graph for Article 13 compliance.

4.5 Phase 5: Metric Operationalization.

The fifth phase closes the pipeline by defining measurable proxies and evaluation scales for the leaf nodes of the SIG, enabling the degree of softgoal satisfaction to be assessed empirically. For each leaf node, the process specifies a metric, measurement scale, and minimum acceptance threshold, as illustrated in Table 4.

Important: The metrics and thresholds presented in Table 4 are illustrative proposals only. They are not empirically validated standards, nor are they intended to be prescriptive for any specific domain. They demonstrate the type of operationalization this phase produces. System providers must define and calibrate their own metrics and thresholds in accordance with their specific operational context, applicable standards, and the requirements of the relevant conformity assessment body. The process treats transparency as a satisficeable property, meaning metrics capture degrees of compliance rather than binary pass/fail conditions.

Table 4. Illustrative measurable proxies, evaluation scales, and acceptance thresholds for leaf-level softgoals in a credit scoring AI system.

Softgoal	Metric	Scale	Threshold
Transparency [System Identity]	Presence of AI disclosure statement in user documentation	Binary	Must be present
Explainability [Credit Decision Logic]	Number of decision factors disclosed with relative weight	Ordinal (0–5+)	Min. 3 factors
Understandability [Output Interpretation]	Flesch reading ease score of output interpretation guide	Continuous (0–100)	Min. score 60
Controllability [Deployment Conditions]	Completeness of use condition checklist	Percentage (0–100%)	Min. 90%
Reliability [Performance Stability]	Frequency of performance reporting updates	Nominal (never, annually, quarterly, monthly)	Min. quarterly
Data Quality [Input Feature Disclosure]	Ratio of documented to total input features	Percentage (0–100%)	Min. 100%
Risk Management [Limitation Disclosure]	Number of documented limitation scenarios	Ordinal (0–5+)	Min. 3 scenarios
Auditability [Operational Monitoring]	Availability of structured operational logs	Binary	Must be present
Human Agency [Oversight Triggers]	Number of defined conditions triggering human review	Ordinal (0–5+)	Min. 2 conditions

4.6 Bidirectional Traceability Matrix.

The process is unified by a bidirectional traceability matrix consolidating the outputs of all five phases, as presented in Table 5. The matrix can be traversed from a legal requirement to its technical implementation (supporting compliance demonstration) and from a technical specification back to its legal basis (supporting conformity assessment). The ISO characteristic column is domain independent; the softgoal and metric columns are context dependent.

Table 5. Bidirectional traceability matrix linking Article 13 legal requirements to ISO/IEC 25059:2023 characteristics, softgoal instantiations, and operationalization metrics for a credit scoring AI system.

LR	Legal Requirement	ISO Characteristic	Softgoal	Metric
LR1	AI system self-identification	Transparency	Transparency [System Identity]	AI disclosure statement
LR2	Intended purpose disclosure	Transparency, Controllability	Controllability [Deployment Conditions]	Use condition checklist
LR3	Accuracy and performance metrics	Performance Efficiency, Reliability	Reliability [Performance Stability]	Performance reporting frequency
LR4	Limitations and misuse scenarios	Risk Management, Transparency	Risk Management [Limitation Disclosure]	Documented limitation scenarios
LR5	Input data characteristics	Data Quality, Transparency	Data Quality [Input Feature Disclosure]	Input feature documentation ratio
LR6	Output interpretation instructions	Explainability, Understandability	Explainability [Credit Decision Logic]	Decision factors disclosed
LR7	Human oversight indication	Controllability, Human Agency	Human Agency [Oversight Triggers]	Defined oversight conditions
LR8	Operational monitoring support	Maintainability, Auditability	Auditability [Operational Monitoring]	Operational log availability
LR9	Performance change disclosure	Reliability, Transparency	Reliability [Performance Stability]	Performance reporting frequency

5 Discussion

The illustrative application of the process to Article 13 in the credit scoring domain produced nine atomic legal requirements mapped across nine ISO/IEC 25059:2023 quality characteristics, instantiated as nine softgoals, and operationalized through nine measurable metrics. The bidirectional traceability matrix consolidates these outputs into a single auditable artifact.

The information/process distinction established in Phase 1 proved to be structurally consequential: information obligations map more readily to binary or ratio-scale metrics, whereas process obligations require more demanding operationalization and context-sensitive thresholds. A residual compliance tension is observable in LR4 and LR6, where explainability obligations may conflict with proprietary model protection, a relationship explicitly represented in the SIG through a Hurt contribution link.

ISO/IEC 25059:2023 and the NFR Framework proved complementary and jointly sufficient for Article 13 coverage. This complements findings from the broader RE4AI literature: Habiba et al. [32] identify requirements specification and explainability as among the most persistent challenges in AI requirements engineering, and the process proposed here directly addresses both by providing a structured vocabulary and traceable operationalization path. Villamizar et al. [31] similarly argue for the integration of artefact-based and perspective-based RE methods to operationalize trustworthiness requirements, a vision that aligns with the approach taken here.

However, a subset of Article 13 provisions remains technically indeterminate due to dependence on deployment context and judicial interpretation, confirming

that full compliance requires organizational governance beyond technical specification alone. This finding converges with Ayala-Rivera and Pasquale [11], whose GDPR operationalization work also identifies residual normative indeterminacy as a structural limitation of any purely technical compliance approach.

The principal methodological difference between this work and GDPR-focused approaches lies in the nature of the obligations. GDPR obligations are largely data-centric and well-bounded, amenable to relatively straightforward mapping. Article 13 obligations are behavioral, system-level, and context-sensitive, requiring the NFR Framework’s satisficing semantics and SIG trade-off modeling to be adequately represented. This constitutes the primary methodological advance of this work over prior legal RE approaches.

6 Related Work

Ayala-Rivera and Pasquale [11] propose a structured process to operationalize GDPR obligations into software requirements, making it the closest structural analog to this work. Their approach was designed for well-bounded data protection obligations and does not account for the dynamic, context-sensitive nature of AI system behavior. The present process extends this paradigm to the AI domain by grounding transparency obligations in ISO/IEC 25059:2023 and the NFR Framework, producing softgoal-based specifications with explicit metric-level acceptance thresholds. A key structural difference is that GDPR obligations are primarily data-centric and amenable to binary verification, whereas Article 13 obligations are behavioral and degree-sensitive, requiring satisficing semantics and SIG trade-off modeling that GDPR operationalization approaches do not employ. This distinction motivates the NFR Framework integration and represents the primary methodological advance over prior legal RE work.

Sovrano et al. [6] directly address the EU AI Act by examining how explainability metrics can support compliance. Their focus remains on post-hoc explainability techniques rather than system-level documentation and operational constraints. The present approach complements this by operationalizing transparency obligations at the requirements specification level, producing traceable engineering artifacts.

Kelly et al. [9] propose a methodological approach to EU AI Act compliance for safety-critical products, addressing compliance broadly across multiple provisions without producing verifiable metric-level specifications. The present process focuses specifically on Article 13 and delivers quantitative acceptance thresholds, SIGs, and a bidirectional traceability matrix.

Villamizar et al. [31] present a complementary vision for operationalizing trustworthy AI requirements by integrating AMDiRE with PerSpecML. While their framework is broader in scope, addressing multiple trustworthiness dimensions, the present work is more narrowly focused on Article 13 transparency obligations and goes further in producing quantitative metric-level specifications and a bidirectional traceability matrix. The two approaches are complementary: Villamizar et al. provide the concern identification and artefact structuring layer, while the present process provides the legal-to-metric operationalization layer.

Conclusion

This paper addressed the translation gap between the qualitative transparency obligations of Article 13 of the EU AI Act and the precise, verifiable specifications required for software engineering practice. The proposed five-phase process integrates the NFR Framework, Softgoal Interdependency Graphs, and ISO/IEC 25059:2023 into an operationalization pipeline that transforms legal mandates into measurable engineering requirements with explicit traceability. The resulting bidirectional traceability matrix constitutes a compliance artifact usable by both system providers seeking to demonstrate conformity and market surveillance authorities conducting assessments.

Intended users and practical adoption. The process is designed for multidisciplinary teams. A requirements engineer should lead Phases 1 through 5, providing the methodological expertise to structure the extraction, mapping, and modeling steps. A legal or compliance expert should validate the outputs of Phase 1, since the extraction of atomic legal requirements involves interpretive decisions that benefit from legal knowledge. A domain expert should calibrate the outputs of Phases 3 through 5, since softgoal instantiation, SIG construction, and metric thresholds are context-sensitive and require knowledge of the target deployment environment. The process does not assume legal expertise on the part of the requirements engineer, nor engineering expertise on the part of the legal expert; its value lies in providing a structured interface between these knowledge domains. It is intended to be adopted during the requirements engineering phase of AI system development, before system design, so that transparency specifications can inform architectural and implementation decisions.

Limitations. Several limitations should be acknowledged. First, the process has not been empirically validated with practitioners. The credit scoring application is an illustrative exercise demonstrating process feasibility, not evidence of practical effectiveness or generalizability. Second, key steps involve interpretive subjectivity — the extraction, mapping, and metric definition all depend on analyst judgment, and no inter-rater agreement procedures were applied. Third, the metrics and thresholds in Table 4 are illustrative proposals not validated against real compliance assessment practice. Fourth, the process focuses exclusively on Article 13 and a single domain, leaving open questions about scalability.

Future Work. Empirical validation involving real multidisciplinary teams across multiple Annex III domains, including healthcare and automated recruitment, is the most immediate priority. Inter-rater agreement studies on the extraction and mapping activities would provide evidence regarding reproducibility. Structured derivation guidelines for Phases 1 and 2 would help reduce reliance on implicit judgment. The integration of stakeholder-stratified measurability criteria from IEEE 7001:2021 could further enrich Phase 5. Tool support to automate portions of the pipeline, particularly traceability matrix generation and mapping activities, would improve scalability and facilitate integration into existing development workflows. Finally, extending the process

to other articles of the EU AI Act would broaden its practical applicability and impact.

Acknowledgments. We express our gratitude to the Brazilian Research Council - CNPq (Grant 312275/2023-4), Rio de Janeiro State's Research Agency - FAPERJ (Grant E-26/204.256/2024), the Coordination for the Improvement of Higher Education Personnel (CAPES), and Kunumi Institute for their generous support.

References

1. Cheong, B. C. (2024). Transparency and accountability in AI systems: safeguarding wellbeing in the age of algorithmic decision-making. *Frontiers in Human Dynamics*, 6. <https://doi.org/10.3389/fhumd.2024.1421273>
2. Stahl, B. C., Antoniou, J., Bhalla, N., Brooks, L., Jansen, P., Lindqvist, B., Kirichenko, A. S., Marchal, S., Rodrigues, R., Santiago, N., Warso, Z., & Wright, D. (2023). A systematic review of artificial intelligence impact assessments. *Artificial Intelligence Review*, 56(11), 12799. Springer Science+Business Media. <https://doi.org/10.1007/s10462-023-10420-8>
3. Hermanns, H., Lauber-Rönsberg, A., Meinel, P., Sterz, S., & Zhang, H. (2024, October). AI Act for the working programmer. In *International Conference on Bridging the Gap between AI and Reality* (pp. 74-98). Cham: Springer Nature Switzerland.
4. Teodorescu, M., Sun, Y., Bhatia, H. N., & Makridis, C. (2025). An Analysis of the New EU AI Act and A Proposed Standardization Framework for Machine Learning Fairness. *arXiv preprint arXiv:2510.01281*.
5. Butt, J. (2024). Analytical study of the world's first EU Artificial Intelligence (AI) Act. *International Journal of Research Publication and Reviews*, 5(3), 7343-7364.
6. Sovrano, F., Lognoul, M., & Vilone, G. (2024). Aligning XAI with EU Regulations for Smart Biomedical Devices: A Methodology for Compliance Analysis. *FRONTIERS IN ARTIFICIAL INTELLIGENCE AND APPLICATIONS*, 392, 826-833.
7. Busuioc, M., Curtin, D., & Almada, M. (2023). Reclaiming transparency: contesting the logics of secrecy within the AI Act. *European Law Open*, 2(1), 79-105.
8. Sovrano, F., Sapienza, S., Palmirani, M., & Vitali, F. (2022). Metrics, explainability and the European AI act proposal. *J*, 5(1), 126-138.
9. Kelly, J., Zafar, S. A., Heidemann, L., Zacchi, J. V., Espinoza, D., & Mata, N. (2024, June). Navigating the EU AI Act: A methodological approach to compliance for safety-critical products. In *2024 IEEE Conference on Artificial Intelligence (CAI)* (pp. 979-984). IEEE.
10. Abualhaija, S., Ceci, M., & Briand, L. (2025). Legal requirements analysis: A regulatory compliance perspective. In *Handbook on natural language processing for requirements engineering* (pp. 209-242). Cham: Springer Nature Switzerland.
11. Ayala-Rivera, V., & Pasquale, L. (2018, August). The grace period has ended: An approach to operationalize GDPR requirements. In *2018 IEEE 26th International Requirements Engineering Conference (RE)* (pp. 136-146). IEEE.
12. Negri-Ribalta, C., Lombard-Platet, M., & Salinesi, C. (2024). Understanding the GDPR from a requirements engineering perspective—a systematic mapping study on regulatory data protection requirements. *Requirements Engineering*, 29(4), 523-549.
13. Cambria, E., Malandri, L., Mercurio, F., Mezzanzanica, M., & Nobani, N. (2023). A survey on XAI and natural language explanations. *Information Processing & Management*, 60(1), 103111.
14. Kesari, A., Breaux, T., Santos, S., Norton, T., & Singhal, A. (2025, June). From legal text to tech specs: Generative ai's interpretation of consent in privacy law. In *Proceedings of the Twentieth International Conference on Artificial Intelligence and Law* (pp. 159-168).

15. Shah, S. T. U., Hussein, M., Barcomb, A., & Moshirpour, M. (2025, September). Explainability as a compliance requirement: What regulated industries need from AI tools for design artifact generation. In 2025 IEEE 33rd International Requirements Engineering Conference Workshops (REW) (pp. 13-22). IEEE.
16. Hassani, S., Sabetzadeh, M., & Amyot, D. (2026). From law to Gherkin: A human-centred quasi-experiment on the quality of LLM-generated behavioural specifications from food-safety regulations. *Information and Software Technology*, 108122.
17. European Parliament & Council of the European Union. (2024). Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act). <http://data.europa.eu/eli/reg/2024/1689/oj>
18. Larsson, S., & Heintz, F. (2020). Transparency in artificial intelligence. *Internet policy review*, 9(2), 1-16.
19. Doshi-Velez, F., & Kim, B. (2017). Towards a rigorous science of interpretable machine learning. arXiv preprint arXiv:1702.08608.
20. Molnar, C. (2020). Interpretable machine learning. Lulu. com.
21. Ribeiro, M. T., Singh, S., & Guestrin, C. (2016, August). " Why should i trust you?" Explaining the predictions of any classifier. In *Proceedings of the 22nd ACM SIGKDD international conference on knowledge discovery and data mining* (pp. 1135-1144).
22. Abualhajja, S., Ceci, M., Sannier, N., Bianculli, D., Lannier, S., Siclari, M., ... & Tosza, S. (2025, September). LLM-assisted extraction of regulatory requirements: A case study on the GDPR. In 2025 IEEE 33rd International Requirements Engineering Conference (RE) (pp. 142-154). IEEE.
23. Kabzhan, Z., Shakhnovich, A., Gorshkov, S., Yemenov, Y., Gorshkov, F., & Shogelova, N. (2025). Semantic and ontology-based analysis of regulatory documents for construction industry digitalization. *Frontiers in Built Environment*, 11, 1575913.
24. Kosenkov, O., Unterkalmsteiner, M., Mendez, D., Fucci, D., Gorschek, T., & Fischbach, J. (2024, June). On developing an artifact-based approach to regulatory requirements engineering. In 2024 IEEE 32nd International Requirements Engineering Conference Workshops (REW) (pp. 262-271). IEEE.
25. IEEE Computer Society. (2024). Guide to the software engineering body of knowledge (SWEBOK guide): Version 4.0. (Bourque, P., & Fairley, R. E., Eds.). IEEE.
26. Chung L., Nixon, J. M. B. and Yu, A.: *Non-functional Requirements in Software Engineering*. Springer, Reading, Massachusetts, 2000.
27. International Organization for Standardization. (2023). Software engineering — Quality requirements and evaluation of AI systems (SQuaRE) — AI system quality model (ISO/IEC 25059:2023). <https://www.iso.org/standard/80655.html>
28. IEEE. (2022). IEEE standard for transparency of autonomous systems (IEEE Std 7001-2021). IEEE Computer Society. <https://ieeexplore.ieee.org/document/9726145>
29. Wieringa, R. (2014). *Design science methodology for information systems and software engineering*. Springer-Verlag Berlin Heidelberg.
30. Oviedo, J., Rodriguez, M., Trenta, A., Cannas, D., Natale, D., & Piattini, M. (2024). ISO/IEC quality standards for AI engineering. *Computer Science Review*, 54, 100681.
31. Villamizar, H., Mendez, D., & Kalinowski, M. (2025, September). Towards a Framework for Operationalizing the Specification of Trustworthy AI Requirements. In 2025 IEEE 33rd International Requirements Engineering Conference Workshops (REW) (pp. 533-537). IEEE.
32. Habiba, U. E., Haug, M., Bogner, J., & Wagner, S. (2024). How mature is requirements engineering for AI-based systems? A systematic mapping study on practices, challenges, and future research directions. *Requirements Engineering*, 29(4), 567-600.
33. ACM - Sousa, H.P.D.S., Almentero, E.K., Classe, T.M.D., Santos, R.J.D., & Leite, J.C.(2023). Uma abordagem baseada no Catálogo de Requisitos Não Funcionais para conformidade à LGPD. In 26th Workshop on Requirements Engineering (WER2023), Brazil. 10.29327/1298356.26-8