

OntoCRA-NS: A Neuro-Symbolic Framework for Explainable and Auditable Cybersecurity Compliance Aligned with EU Regulatory Frameworks

Jasmin Cosic, Admir Jukan, and Miroslav Baca

Abstract—At the EU level, the NIS2 Directive, Cyber Resilience Act (CRA), and AI Act (AIA) form a strategic regulatory triangle designed to strengthen Europe’s digital sovereignty and cyber resilience. Together, they should ensure that we have common level of cyber-security (NIS2), that products with digital elements are built securely (CRA), and AI systems are trustworthy, robust, safe and aligned with fundamental rights and EU values (AIA). This integrated approach creates a unified cybersecurity framework that protects citizens, businesses, and critical infrastructure from emerging digital and AI-driven threats across the EU. This paper addresses the growing need for automated, explainable, and auditable cybersecurity compliance mechanisms under the CRA. It introduces OntoCRA-NS, a neuro-symbolic framework that integrates neural language understanding with formal ontological reasoning to support automated compliance assessment. The system combines a neural extractor that interprets unstructured documentation and vulnerability data with an OWL-based ontology (OntoCRA) containing formalized requirements encoded as SWRL rules. Extracted evidence is transformed into ontology instances and processed by the Pellet reasoner to infer Declaration of Conformity (DoC) or Banned from Market decisions, accompanied by human-readable justifications. A proof-of-concept implementation demonstrates end-to-end reasoning using representative products from the CURIUM EU project, achieving deterministic, traceable, and legally interpretable compliance outcomes. The results confirm that hybrid neuro-symbolic reasoning enables explainable, trustworthy, and reproducible CRA compliance evaluation, contributing to the emerging paradigm of regulatory AI under the EU cybersecurity and AI governance frameworks.

Keywords: Cyber Resilience Act, Ontology, SWRL, OWL, Neuro-Symbolic AI, Explainable AI, Compliance Automation, EUCC, Trustworthy AI, Regulatory Reasoning.

I. INTRODUCTION

The European Union’s Cyber Resilience Act (CRA) [1] establishes a comprehensive legal framework to ensure that all products with digital elements placed on the EU market meet minimum cybersecurity and vulnerability management requirements throughout their lifecycle. Manufacturers must demonstrate compliance with essential requirements covering secure design, vulnerability disclosure, software maintenance, post-market surveillance and other requirements. On another hand, the mission of the NIS2 Directive [2], is to strengthen the overall level of cybersecurity and resilience across the European Union by setting common standards for risk management, incident reporting, supply-chain security,

governance, etc. Finally, AI Act [3], should ensure that artificial intelligence in the EU is secure, trustworthy, and aligned with fundamental rights and EU values. It establishes a risk-based framework that promotes innovation while preventing harmful uses of AI, ensuring systems are transparent, robust, and secure. Together, these regulations balance technological progress with safety, ethics, and (cyber)security. While the CRA aims to increase trust and accountability in digital products, its practical enforcement poses major challenges. Current conformity assessments are largely manual, expert-driven, and highly heterogeneous across stakeholders [4]. Few EU funded projects are engaged with these tasks [5]. The methods are mostly time-consuming, error-prone, and difficult to scale across complex software supply chains. Automated solutions are emerging, yet existing machine learning and natural-language processing (NLP) models. Although being capable of understanding unstructured documentation, they lack formal reasoning and explanation, both of which are essential for legal defensibility and regulatory transparency. Conversely, traditional ontology-based reasoning systems provide deterministic and auditable inference but cannot autonomously interpret free-form textual or semi-structured data sources.

To address this gap, this paper proposes OntoCRA-NS, a neuro-symbolic framework that fuses neural and symbolic AI to automate and explain CRA compliance assessment.

II. BACKGROUND AND RELATED WORK

The CRA is the first horizontal EU regulation introducing mandatory cybersecurity requirements for all products with digital elements. It complements other European initiatives such as the NIS2 Directive, the EU Cybersecurity Certification Framework (EUCC), and the forthcoming AI Act, forming a cohesive regulatory ecosystem focused on security-by-design, vulnerability management, and post-market accountability [2]. Under the CRA, manufacturers must ensure that their products are developed according to state-of-the-art cybersecurity practices, maintain patching support for at least five years, and report actively exploited vulnerabilities to ENISA and national authorities. Compliance must be demonstrated through technical documentation and a Declaration of Conformity (DoC) prior to market access. These obligations highlight the urgent need for automated, explainable compliance-assessment mechanisms capable of translating natural-language legal text into verifiable, machine-interpretable rules.

A. Ontology-Based Compliance Reasoning

Ontologies, explicit formal specification of the domain [6], provide a formal, semantic representation of regulatory domains, enabling reasoning engines to infer compliance or non-compliance from structured facts. Prior research in security ontologies has explored risk management [7], privacy and data-protection reasoning [8], ISO/IEC 27001 control verification [9], cyber-security-digital forensic [10], and cyber-security certification [11], [12]. Several works have modelled cybersecurity certification processes in OWL, capturing dependencies between controls, assets, and evidence [11],[12]. However, most existing systems operate on predefined datasets and lack mechanisms to ingest unstructured documentation or to explain decisions in natural-language form. Ontologies mostly stay on the level of static definition and without enriched rules. Rule-based extensions such as the Semantic Web Rule Language (SWRL) [13], allow executable compliance checks, yet they rely on complete, structured input - a condition rarely met in industrial settings. The OntoCRA_NS ontology presented in this work builds on these foundations, encoding CRA Essential and Vulnerability Requirements as SWRL rules while remaining open for integration with data extracted from textual or semi-structured sources.

B. Neuro-Symbolic Artificial Intelligence

Neuro-Symbolic AI (NSAI) [14] combines the data-driven learning capability of neural networks with the interpretability and logical structure of symbolic reasoning. Foundational contributions such as Neural-Symbolic Learning Systems [15], DeepProbLog [16], and Logic Tensor Networks [17] demonstrate how hybrid models can integrate probabilistic inference with formal logic to achieve both adaptability and explainability. Recent studies have applied NSAI to domains including visual reasoning, robotic planning, and legal text understanding [18], [19]. In regulatory contexts, neuro-symbolic architectures offer unique potential: neural models can parse complex policy language, while ontological reasoners ensure deterministic and legally sound conclusions. The research [20] also discusses the different role that explicit knowledge can play in Explainable AI. OntoCRA-NS advances this frontier by applying the neuro-symbolic paradigm to the CRA, integrating natural-language interpretation with an ontology-based reasoning core to produce transparent, auditable compliance decisions.

C. Research Gaps and Contribution

From the reviewed literature, three main limitations in compliance check process, persist:

- Lack of automation: Existing ontology frameworks depend on manual evidence input.
- Limited explainability: Machine-learning classifiers cannot justify decisions according to legal norms.

- Poor integration: There is little coupling between neural text analysis and symbolic compliance models.

OntoCRA-NS addresses these gaps by (1) introducing a reusable, CRA-aligned ontology with executable rules; (2) integrating a neural extraction component that populates the ontology automatically; and (3) providing human-readable explanations consistent with EU regulatory principles of transparency and accountability.

D. SoTa

Artificial Intelligence (AI) has undergone three significant paradigm shifts: from symbolic logic systems, through data-driven neural models, toward the emerging convergence known as neuro-symbolic AI [21]. Early AI research relied on symbolic reasoning - formal logic, ontologies, and rule-based inference - to represent knowledge in interpretable, structured ways. Ontology languages such as OWL and rule systems like SWRL have been used extensively in domains requiring explicit reasoning, such as semantic web technologies and knowledge representation [22]. However, symbolic systems are brittle when faced with unstructured, noisy, or incomplete real-world data. The second wave of AI, driven by deep neural networks [23], achieved remarkable performance in perception tasks, natural language processing, and pattern recognition. Models such as BERT and GPT can extract implicit relationships from vast unstructured corpora but lack transparency, traceability, and explicit reasoning capability. Their “black-box” nature makes them unsuitable for high-assurance or regulated environments such as cybersecurity, safety, and digital product compliance. Recent research therefore focuses on neuro-symbolic integration, combining the perceptual power of neural models with the interpretability of symbolic reasoning [24], [21]. Hybrid frameworks such as Logic Tensor Networks, DeepProbLog, and NeuralLP integrate logical constraints into neural training or translate neural predictions into formal reasoning. [25] These approaches enable AI systems to reason about rules, hierarchies, and relationships that pure deep learning cannot express. In parallel, the trustworthy and explainable AI (XAI) community has emphasized the importance of accountability, fairness, and transparency in AI systems. The European Union’s AI Act and the ENISA AI Cybersecurity Certification Framework explicitly demand interpretability and regulatory traceability in high-risk AI applications. This has led to emerging studies combining ontology-based reasoning with neural models for explainable legal or compliance-related inference. Works such as LegalBERT [26] and LexGLUE [27], demonstrate that neural architectures can understand legal semantics, but they do not provide verifiable logical reasoning. Despite these advances, current neuro-symbolic systems remain mostly generic and lack explicit alignment with regulatory frameworks such as the EU Cyber Resilience Act (CRA). Existing solutions do not yet model CRA’s essential and vulnerability requirements in a computable form, nor do they integrate compliance reasoning with adaptive learning or human oversight. This

gap is particularly relevant to cybersecurity governance, where explainability, accountability, and formal traceability are essential for certification and market authorization. To address this gap, the proposed OntoCRA-NS framework introduces domain-specific (cyber-security) neuro-symbolic architecture for automated and explainable conformity assessment under the CRA. OntoCRA-NS combines:

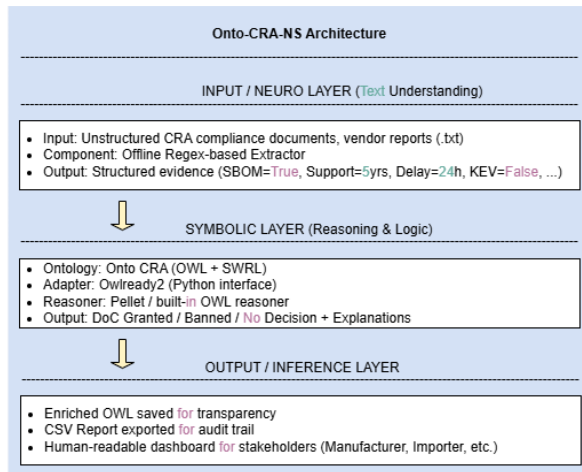
- a neural evidence extraction layer that interprets technical documentation, SBOMs, and vulnerability reports.
- a symbolic reasoning layer (OntoCRA_NS ontology in OWL/SWRL) that formalizes CRA requirements and infers decisions such as Declaration of Conformity (DoC) or Ban from Market.
- a feedback loop enabling human auditors to validate and improve the neural extraction process.

This approach positions OntoCRA-NS at the intersection of neuro-symbolic AI, trustworthy regulatory automation, and cybersecurity compliance. It aligns with ongoing ENISA and EU Commission efforts to establish verifiable, explainable AI infrastructures for critical digital systems. Despite recent advances in regulatory AI, no prior system explicitly models CRA requirements as executable ontological rules integrated with neural extraction. OntoCRA-NS fills this gap.

III. ONTOCRA-NS ARCHITECTURE

The OntoCRA-NS framework is designed as a modular, layered architecture that integrates neural evidence extraction and symbolic reasoning into a single, explainable compliance-assessment workflow (Figure 1). Its structure allows interoperability between data-driven interpretation and formal regulatory logic, ensuring that outputs are both machine-processable and legally auditable.

Figure 1. High-level OntoCRA-NS architecture



A. Overview

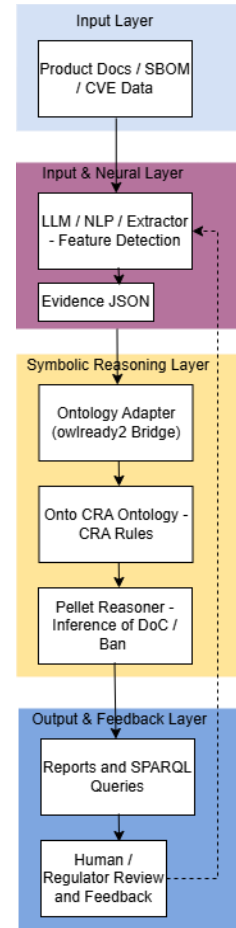
OntoCRA-NS follows a five-layer architecture as described below (Figure 2).

This separation of concerns provides flexibility: the neural components can evolve independently (e.g., LLM → fine-tuned model), while the ontology and rule base remain stable and transparent, ensuring consistent assessment procedures.

B. Input Layer

The input layer collects product documentation, software bills of materials (SBOMs), vulnerability reports (CVE/CISA KEV), and support-policy statements. Each item may exist in free-text or semi-structured form (PDF, JSON, XML). These documents contain implicit evidence about CRA obligations such as the existence of an SBOM, vulnerability handling processes, or support duration.

Figure 2. Proposed Architecture



C. Neural Extraction Layer

The neural extraction layer acts as the perception component. It employs an NLP/LLM model (e.g., GPT-based extractor, DistilBERT, or domain-specific transformer) to detect and classify CRA-relevant features. For lightweight

deployments, a rule-based regex extractor can simulate the same process (read from uploaded documentation).

The output is a structured Evidence JSON object representing each product’s compliance attributes, for example:

```
{
  "product_id": "CAC", "hasSBOM": true,
  "hasTechnicalDocumentation": true,
  "hasVulnerabilityHandlingProcess": true,
  "knownExploitedVulnerability": false, "supportYears": 5
}
```

Each attribute may also include a confidence score if produced by a neural model, allowing future fuzzy reasoning extensions.

D. Ontology Adapter Layer

This intermediary layer bridges the neural and symbolic components. It parses the Evidence JSON and programmatically populates the ontology using Owlready2. [28] New individuals of the class, i.e. *ProductWithDigitalElements* are created or updated with corresponding data and object properties (e.g., *hasSBOM*, *hasTechnicalDocumentation*, *supportPeriodYears*). The adapter enforces data-type consistency, validates ranges, and triggers reasoning tasks via the Pellet API. [29]

This mechanism transforms unstructured textual knowledge into a semantically valid, machine-interpretable knowledge graph aligned with CRA terminology.

E. Symbolic Reasoning Layer

At the core of OntoCRA-NS lies the OntoCRA v3 ontology, which formalizes CRA Essential Requirements (ER) and Vulnerability Requirements (VU) using OWL classes and SWRL rules. The Pellet reasoner executes these rules to infer compliance decisions. Example logic includes:

1. Missing SBOM → BannedFromMarket
2. Missing Technical Documentation → BannedFromMarket
3. Known exploited vulnerability → BannedFromMarket
4. All requirements met and support ≥ 5 years → DoC_Granted, etc.

This reasoning is deterministic and explainable: each inference is linked to explicit rule premises, enabling traceability.

F. Output and Feedback Layer

After reasoning, OntoCRA-NS generates a compliance decision (DoC_Granted or BannedFromMarket) together with explanations in natural language or SPARQL-based form. For instance, an inferred rule might output:

“Product “CAC” banned due to Known Exploited Vulnerability (CVE-2024-XXXX).”

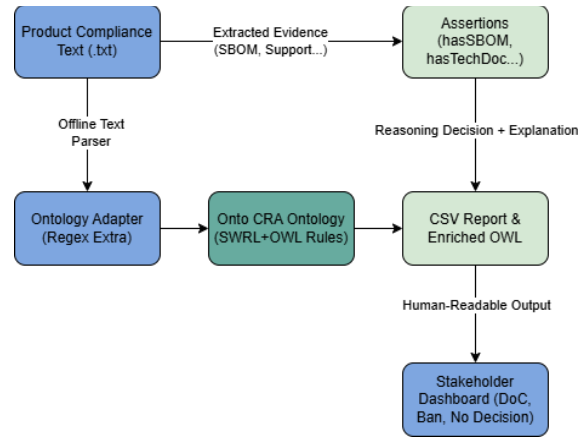
Results can be exported as structured reports, visual dashboards, or regulatory forms. A human auditor or regulator

reviews the explanation, validating or correcting the decision what brings us auditability capability. This feedback is then fed back into the neural extractor for continual improvement, forming the neuro-symbolic learning loop.

G. Advantages of Architecture

This approach ensures *Explainability* - every decision is backed by formal rule reasoning, *Reusability* - OntoCRA v3 can be extended to other EU regulations (EUCC, NIS2D, AI Act), *Scalability* - the neural extractor allows automatic ingestion of large document sets, *Interoperability* - standardized Evidence JSON - enables integration with external tools or digital-compliance platforms. Overall, this architecture operationalizes the CRA’s legal obligations within an *auditable*, hybrid reasoning environment—linking machine learning, symbolic logic, and regulatory semantics in a single, unified workflow.

Figure 3. Data flows in the framework



IV. IMPLEMENTATION OF THE PROTOTYPE

A proof-of-concept (PoC) prototype of OntoCRA-NS was implemented in Protégé [30], Python using the Owlready2 library for ontology management and the Pellet reasoner for SWRL rule execution. The OntoCRA ontology serves as the symbolic knowledge base, comprising 65 classes, 42 object properties, 15 data properties, and 28 SWRL rules encoding the Essential and Vulnerability Requirements defined by the EU Cyber Resilience Act (CRA). A concise summary of the quantitative evaluation metrics and protocol is provided in Section V-A; full PoC results are reported in the companion paper. [31]

The neural layer was simulated through a rule-based text extractor acting as a lightweight natural-language processing (NLP) component. This extractor reads product-related documentation and security information, automatically detecting the presence of key compliance attributes (e.g., SBOM availability, secure design process, vulnerability handling, technical documentation, patch support duration, etc.). Extracted evidence is structured into JSON format and injected into the ontology as individual instances of the *ProductWithDigitalElements* class. The reasoning engine performs automated inference over the populated ontology,

producing compliance decisions such as Declaration of Conformity (DoC) or Banned from Market, together with natural-language explanations. In the PoC setup, five representative products (CyReA, CAC, DPMA, DEMF, and PSTV) were modeled to simulate different CRA compliance levels.

TABLE I. PoC SETUP AND FRAMEWORK DECISION

Product	CRA Evidence	Decision	Reasoning Outcome
CyReA	Complete documentation, SBOM, 5-year support	DoC Granted	Fully compliant
CAC	Missing patch management, exploited CVE	Banned	Violation of VU-3
DPMA	Missing secure update mechanism	Banned	Non-conformance to ER-5
DEMF	Partial documentation	No Decision	Insufficient evidence
PSTV	All essential and vulnerability requirements met	DoC Granted	Compliant

The symbolic reasoning process was executed on a standard workstation (Intel i7, 32 GB RAM), with inference time averaging below 2 seconds per product, demonstrating the computational feasibility of the approach for regulatory - scale assessments. Each decision is traceable to the underlying rule premises, providing deterministic justification paths. For instance, the system infers that *“Product CAC is banned due to a known exploited vulnerability”* through a direct SWRL rule mapping between vulnerability status and CRA Article 10. This transparency enables human auditors to verify every inference step, fulfilling both explainability and auditability criteria mandated under the CRA and AI Act. While the current PoC uses rule-based extraction, the same pipeline can seamlessly integrate a neural LLM-based extractor to process larger document collections, providing confidence-weighted evidence for reasoning. This enhancement is part of ongoing work described in the companion paper focused on the practical deployment of OntoCRA-NS in an industrial environment.

Overall, the prototype demonstrates the technical soundness of the OntoCRA-NS architecture and validates the feasibility of integrating neural evidence interpretation with formal regulatory reasoning. The resulting hybrid pipeline enables automated, explainable, and reproducible CRA compliance evaluation—an essential capability for future ENISA - led conformity-assessment frameworks. The implementation described in this paper focuses on validating the OntoCRA-NS architecture and its core reasoning workflow. A companion study, currently under preparation, extends this work with a detailed proof-of-concept (PoC) and industrial use-case analysis conducted within the scope of the CURIUM and CUSTODES projects. That paper will present empirical evaluation results, dataset descriptions, and stakeholder feedback demonstrating the practical applicability

of the proposed framework in real-world CRA compliance assessment scenarios.

V. RESULTS, EVALUATION AND DISCUSSION

The achieved results demonstrate that OntoCRA-NS have the potential to operationalize not only CRA, but also all other EU legal requirements. The experiment (PoC) shows that formalizing CRA rules in an ontology and combining it with automated evidence extraction can produce compliance decisions that are legally interpretable and consistent. This directly addresses the CRA’s call for verifiable compliance processes, potentially reducing the burden on conformity assessment bodies. Unlike prior “ontology-only” systems that assumed structured inputs, OntoCRA-NS brings a hybrid neural-symbolic solution that can handle unstructured documentation (plain text or .pdf files) via AI, a capability not seen in previous cybersecurity certification ontologies.

We highlight three main innovations of OntoCRA-NS over related work: automatic evidence ingestion, executable legal semantics (direct mapping of law text to rules), and explainable outputs. The system provides transparency (clear rule-based decision paths), traceability (each conclusion is linked to specific requirements, aiding audits), and modularity (the neural and symbolic components can be updated independently). Moreover, the ontology and rules are reusable; the core model could be extended to related regulations (e.g. it could incorporate security requirements from EUCC or NIS2) with minimal changes. The low computational cost and deterministic nature make it feasible for integration into industry compliance toolchains.

Some limitations of the framework are: (i) the neural extraction layer in the PoC is simplified (regex-based), which may miss nuanced information; deploying a more robust LLM-based extractor is the next step; (ii) the ontology covers only a subset of CRA requirements (focused on core security and vulnerability management duties) and would need expansion to fully cover all CRA clauses; (iii) the reasoning is binary (compliant vs not); in practice, a degree of uncertainty or partial compliance might need representation, possibly via confidence scores or fuzzy logic, remediation path or conditional approval scenarios will be addressed in following, updated version of the ontology; (iv) rule maintenance requires ongoing expert oversight. (e.g., if new interpretations or standards emerge, the knowledge base must be updated), as well as ontology, repository updating and maintaining is challenging.

OntoCRA-NS itself would likely be considered a high-risk AI system under the EU AI Act since it aids in regulatory decisions. OntoCRA-NS is meant to assist, not replace, human compliance officers and evaluators in the Laboratory. The feedback loop allows human experts to review AI-extracted evidence and intervene if necessary, ensuring accountability remains with the appropriate stakeholders.

A. Quantitative Evaluation (Metrics & Protocol Summary)

To assess the extraction-to-reasoning pipeline while keeping this manuscript focused on architecture, we report here the metrics and evaluation protocol only; datasets, numbers, and result tables are provided in a companion PoC paper [31]. The PoC [31] instantiates OntoCRA-NS end-to-end (evidence extraction → ontology population → OWL/SWRL reasoning with Pellet) on a small set of CRA dossiers representative of typical documentation artefacts (technical documentation, SBOM, vulnerability/CVE & CISA Known Exploited Vulnerabilities (KEV) [32] indicators, support-policy statements).

1. Evidence-extraction metrics.

Precision, Recall, F1 (field-level) for CRA-relevant binary attributes: `hasSBOM`, `hasTechnicalDocumentation`, `hasVulnerabilityHandlingProcess`, `knownExploitedVulnerability`.

Numeric accuracy for continuous attributes (e.g., `supportPeriodYears`), reported as Mean Absolute Error (MAE) and tolerance-based accuracy (within $\pm k$ years).

2. End-to-end reasoning metrics.

Decision Agreement between the reasoner’s outcome (DoC Granted, Banned from Market, Assessment Pending) and a rule-derived ground truth computed from the annotated evidence using the same OntoCRA v3 rules; this isolates extraction/representation errors from legal interpretation.

False-Compliant Rate (safety metric): share of cases inferred Compliant/DoC by the system while ground truth is Non-Compliant—a primary risk metric for regulatory triage.

3. Runtime (descriptive feasibility).

Per-dossier breakdown of extraction, ontology population, and reasoning time, reported as descriptive statistics; in the PoC, symbolic reasoning remains lightweight relative to document ingestion/extraction.

4. Uncertainty & error analysis (concise).

We include a brief qualitative error analysis of typical extraction failures (e.g., ambiguous support statements, ungrounded evidence) and conservatively return Assessment Pending when decisive predicates are unknown—avoiding default compliance under incomplete evidence (open-world semantics).

5. Baselines & benchmarking.

The PoC uses a deterministic (regex/pattern) extractor to validate the pipeline and traceability; LLM-assisted extraction and comparative metrics will be reported separately. The reasoning layer (OntoCRA v3 + Pellet) remains unchanged to preserve auditability across extractor variants.

(Full protocol, dossier design, and numerical results are provided in the PoC paper.) [31]

6. Governance & Maintenance.

OntoCRA (v3) and the SWRL rule set are maintained under version control with tagged releases. Proposed changes follow a lightweight change-control process (proposal → expert review → tagged release) aligned to CRA delegated acts and ENISA guidance. This ensures traceable evolution of requirements while keeping the reasoning layer (OntoCRA v3 + Pellet) stable and auditable across extractor variants.

VI. CONCLUSION AND FUTURE WORK

This paper presented OntoCRA-NS, a novel framework integrating symbolic ontological reasoning with neural text processing to automate CRA compliance assessment. The approach was motivated by the pressing need for scalable, explainable compliance solutions in the face of new EU cybersecurity regulations. We demonstrated that OntoCRA-NS can successfully determine compliance or non-compliance for sample digital products, producing outcomes that are both machine-verifiable and legally transparent.

OntoCRA-NS contributes to the emerging paradigm of “Regulatory AI” – AI systems designed to enforce or verify compliance with laws. By providing a working example in the cybersecurity domain, it opens avenues for similar neuro-symbolic approaches in other regulatory areas (data protection, safety certification, etc.) where explainability is paramount. First, expanding the ontology to cover the full spectrum of CRA requirements (and keeping it updated as official guidance evolves) will be crucial. Second, integrating a more advanced LLM-based extractor and continuously training it on CRA-related documents could improve the system’s robustness and reduce the need for manual input formatting. Third, conducting more extensive evaluations with real industry data and in collaboration with certification bodies (e.g., using actual product documentation from manufacturers) is planned to validate the tool’s effectiveness in practical settings. We also plan to interface OntoCRA-NS with live vulnerability databases (e.g., CVE feeds) so that known exploits are flagged in real-time – aligning with CRA’s ongoing vulnerability disclosure obligations. OntoCRA-NS demonstrates a promising step toward trustworthy AI-driven compliance. As enforcement of the CRA “clock is ticking” (partial application by 2026, full application by 2027), tools that can automate compliance checking while preserving explanation and accountability will be invaluable to both industry and regulators. The companion case study paper further illustrates how this framework can be applied in an industrial context, highlighting the practical benefits and stakeholder perspectives of deploying OntoCRA-NS for real-world compliance assessment. OntoCRA-NS also demonstrates a promising path toward trustworthy, explainable AI-driven compliance aligned with emerging EU cybersecurity frameworks. Future work, presented in a companion paper, will describe the full, detailed, proof-of-concept

implementation and industrial use-case evaluation of OntoCRA-NS within the CUSTODES and CURIUM project environment, using real products with full set of technical documentation.

ACKNOWLEDGMENT

The research leading to these results received funding from the European Union's HORIZON Innovation Action program GA No. 101120684 - project CUSTODES, as well as DIGITAL-ECCC-2024-DEPLOY-CYBER-06 under proposal number 101190372, project CURIUM. Jasmin Cosic from DEKRA SE (Germany) participates as a CUSTODES project partner and Admir Jukan and Miroslav Baca from Cyber-security, Ltd. (Croatia) as CURIUM project partners.

We want to thank to all project partners, as well as specially to our colleagues, experts from DEKRA AI Team - Angels Aragones Martín and Xavier Valero González, for constructive suggestions on Framework and valuable comments.

Early linguistic polishing was performed with ChatGPT (OpenAI), and final grammar and style adjustments with Microsoft Copilot (M365 Copilot). AI tools were used exclusively for language editing; no technical content, modeling, analyses, or conclusions were produced by AI. All scientific content originates from the authors and the project real pilots/use-cases and was fully reviewed and approved by the authors.

REFERENCES

- [1] Regulation (EU) 2024/2847 of the European Parliament and of the Council of 5 December 2024 on horizontal cybersecurity requirements for products with digital elements (Cyber Resilience Act), OJ L, 2024/2847, Dec. 10, 2024. [Online]. Available: EUR-Lex [Accessed: 20.11.2025].
- [2] European Parliament and Council, "Directive (EU) 2022/2555 of the European Parliament and of the Council of 14 December 2022 on measures for a high common level of cybersecurity across the Union," Official Journal of the European Union, vol. L 333, pp. 80-152, Dec. 27, 2022. [Online]. Available: <https://eur-lex.europa.eu/eli/dir/2022/2555/oj>. [Accessed: 20.10.2025].
- [3] European Union, "Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonized rules on artificial intelligence," Official Journal of the European Union, July 12, 2024. [Online]. Available: <https://eur-lex.europa.eu/eli/reg/2024/1689/oj>. [Accessed: 20.12.2025].
- [4] European Cyber Security Organisation (ECISO), "Challenges of the industry to implement the CRA," Working Group 1 (WG1), Tech. Rep., Mar. 2024. [Online]. Available: https://ecs-org.eu/ecs-uploads/2024/03/ECISO_Survey_CRA_2024.pdf.
- [5] EU funded project, CRA-AI (<https://cybercertlabs.com/cra-ai/>), CURIUM (<https://curium-project.eu/>), CRACY (<https://cra-cy.eu/>), CYBERFORT (<https://www.i-energylink.eu/cyberfort-project/>), accessed 09.02.2026
- [6] Noy, N. F., McGuinness, D. L. & others (2001). Ontology development 101: A guide to creating your first ontology Stanford knowledge systems laboratory technical report KSL-01-05 and Stanford medical informatics technical report SMI-2001-0880,
- [7] M. Rossi et al., "Ontology-Driven Risk Management for IoT Systems," Computers & Security, 2022.
- [8] A. Smith and K. Patel, "Privacy Reasoning with Semantic Web Technologies," IEEE Trans. Knowledge and Data Engineering, 2021.
- [9] L. Chen et al., "Formal Verification of ISO/IEC 27001 Controls Using OWL," Proc. ICE, 2022.
- [10] J. Cosic, Z. Cosic, M. Baca, An Ontological Approach to Study and Manage Digital Chain of Custody of Digital Evidence, JIOS, 2011
- [11] J. Cosic, C. Argyro, M. M. Sánchez, A. D. Vizcaino Gomez and D. Morog, "Towards Automated Certification Framework of Composite Systems: A SWRL-Based Approach," 2025 IEEE International Conference on Engineering, Technology, and Innovation (ICE/ITMC), Valencia, Spain, 2025, pp. 1-8,
- [12] J. Cosic and A. Jukan, "Deciphering Cyber-Security Certifications: An Ontological Journey through Composite Systems and their certification," 2024 IEEE International Conference on Engineering, Technology, and Innovation (ICE/ITMC), Funchal, Portugal, 2024, pp. 1-6
- [13] I. Horrocks et al., "SWRL: A Semantic Web Rule Language Combining OWL and RuleML," W3C, 2004.
- [14] Sheth, A., Roy, K., & Gaur, M. (2023). Neurosymbolic AI--Why, What, and How. arXiv preprint arXiv:2305.00813.
- [15] A. d'Avila Garcez et al., "Neural-Symbolic Learning and Reasoning," IEEE Proc., 2019.
- [16] R. Manhaeve et al., "DeepProbLog: Neural Probabilistic Logic Programming," NeurIPS, 2018.
- [17] L. Serafini and A. Garcez, "Logic Tensor Networks," Artificial Intelligence, 2016.
- [18] S. Pirovano et al., "Hybrid Neural-Symbolic Models for Legal Text Reasoning," Applied AI Letters, 2023.
- [19] C. Li et al., "Explainable AI for Policy Compliance Using Ontologies," IEEE Access, 2022.
- [20] R. Confalonieri, G. Guizzardi, On the Multiple Roles of Ontologies in Explainable AI, 2023
- [21] L. Serafini and A. G. d'Avila Garcez, "Logic Tensor Networks: Deep Learning and Logical Reasoning from Data and Knowledge," J. Artif. Intell. Res., 2016., G. De Raedt, S. Dumančić, and T. Manhaeve, "Neuro-Symbolic AI: The 3rd Wave of AI,"
- [22] S. Staab and R. Studer, Handbook on Ontologies, Springer, 2010.
- [23] Y. LeCun, Y. Bengio, and G. Hinton, "Deep Learning," Nature, vol. 521, pp. 436-444, 2015. S. Besold, A. Garcez, and L. Serafini, "Neural-Symbolic Learning and Reasoning: A Survey and Interpretation," Frontiers in Artificial Intelligence, 2021
- [24] S. Besold, A. Garcez, and L. Serafini, "Neural-Symbolic Learning and Reasoning: A Survey and Interpretation," Frontiers in Artificial Intelligence, 2021
- [25] Yang, F., Yang, Z., and Cohen, W. W., "Differentiable Learning of Logical Rules for Knowledge Base Reasoning," Advances in Neural Information Processing Systems (NeurIPS), pp. 2319-2328, 2017. doi:10.48550/arXiv.1702.08367.
- [26] I. Chalkidis, M. Fergadiotis, P. Malakasiotis, N. Aletras, and I. Androustopoulos, "LEGAL-BERT: The Muppets Straight out of Law School," in Findings of EMNLP 2020, pp. 2898-2904, 2020. doi:10.18653/v1/2020.findings-emnlp.261.
- [27] N. Chalkidis et al., "LexGLUE: A Benchmark Dataset for Legal Language Understanding in the EU and US," ACL, 2022.
- [28] J.-B. Lamy, "Owlready: Ontology-oriented Programming in Python with Automatic Classification and High-Level Constructs for Biomedical Ontologies," Artificial Intelligence in Medicine, vol. 80, pp. 11-28, 2017. doi:10.1016/j.artmed.2017.07.002.
- [29] E. Sirin, B. Parsia, B. C. Grau, A. Kalyanpur, and Y. Katz, "Pellet: A Practical OWL-DL Reasoner," Web Semantics: Science, Services and Agents on the World Wide Web, vol. 5, no. 2, pp. 51-53, 2007. doi:10.1016/j.websem.2007.03.004.
- [30] Musen MA; Protégé Team. The Protégé Project: A Look Back and a Look Forward. AI Matters. 2015 Jun;1(4):4-12. doi: 10.1145/2757001.2757003. PMID: 27239556; PMCID: PMC4883684, <https://protege.stanford.edu/>
- [31] J. Cosic and M. Baca, Ontology-driven Compliance Reasoning for the EU Cyber Resilience Act: A Proof-of-Concept, 2026 (In press)
- [32] Cybersecurity and Infrastructure Security Agency (CISA). (n.d.). *Known Exploited Vulnerabilities Catalog*. Retrieved February 9, 2026, from <https://www.cisa.gov/known-exploited-vulnerabilities-catalog>