

A New Framework for Formal Reasoning over Symbolic and Sub-Symbolic Knowledge

(Extended Abstract)

Gianluca Cima¹, Marco Console¹ and Laura Papi¹

¹Sapienza University of Rome, Italy

Abstract

This discussion paper summarizes our work [1], recently accepted at the 40th International AAAI Conference on Artificial Intelligence (AAAI 2026). We propose a novel framework for combining symbolic and sub-symbolic knowledge. The framework is based on the notion of Hybrid Knowledge Base (HKB), a formalism that integrates ontologies with Machine Learning (ML) classifiers. In an HKB, an ontology provides a conceptual and logical description of the application domain, while ML classifiers operating over a feature space implicitly populate the instances of concepts and roles used in the ontology. We introduce a model-theoretic semantics for HKBs and study the computational complexity of query answering under two alternative semantics.

Keywords

Neuro-Symbolic AI, Hybrid Knowledge Bases, Query Answering, Computational Complexity

Machine Learning (ML) classifiers are widely used in many application domains, often supporting decision-making processes by providing data-driven predictions. They have demonstrated strong empirical performance in various tasks, turning out to be particularly effective in areas such as finance, healthcare, and law. However, due to the high complexity of their learned representations, ML classifiers are often treated as black-box models. As a result, while they capture useful inductive knowledge acquired from data, their applicability is limited in settings that require enforcement of global constraints, explicit multi-step reasoning, and logically verifiable interpretations of their predictions [2].

To better support such processes, it would be highly beneficial to contextualize the inductively acquired knowledge encoded in these ML models and enable formal reasoning over it, ideally by combining such knowledge with a logical specification of the domain of the data on which these classifiers operate. In this way, the learned knowledge could be interpreted and reasoned about within its application context, allowing domain constraints and logical inference to complement purely data-driven predictions. Despite the substantial advances made in Neuro-Symbolic AI in recent years, this particular challenge has received comparatively little attention.

To contribute to this line of research, we propose a novel logical framework that enables the integration of sub-symbolic knowledge induced by ML classifiers with symbolic knowledge specified by logic-based formalisms. Specifically, regarding the sub-symbolic component, we assume the availability of a set of already trained ML classifiers (possibly treated as black-box models) operating over a shared feature space. The symbolic component is specified by means of an ontology, which provides a formal description of the application domain corresponding to the data represented in that feature space. The framework is then based on the novel notion of *Hybrid Knowledge Base (HKB)*, consisting of two components: an ontology and a set of ML binary classifiers. Formally, an HKB is defined as follows.

Definition 1. A Hybrid Knowledge Base (HKB) $\mathcal{K}$ is a pair $\mathcal{K} = (\mathcal{O}, \Psi)$, where $\mathcal{O}$ is an ontology and Ψ is a set of classifiers based on a tuple $\mathcal{A}$ of attributes. More precisely, Ψ contains:

- a binary classifier κ_C , for each $C \in \text{sig}(\mathcal{O})$;
- a binary classifier over pairs λ_R , for each $R \in \text{sig}(\mathcal{O})$.

AIxLA 2026 Discussion Papers, 24th International Conference of the Italian Association for Artificial Intelligence, Perugia, Italy, 2026

✉ cima@diag.uniroma1.it (G. Cima); console@diag.uniroma1.it (M. Console); papi@diag.uniroma1.it (L. Papi)

id 0000-0003-1783-5605 (G. Cima); 0009-0004-5526-019X (M. Console); 0009-0003-2281-9500 (L. Papi)



© 2026 Copyright for this paper by its authors. Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

In the above definition, an ontology $\mathcal{O}$ is simply a first-order theory over a signature $\text{sig}(\mathcal{O})$ consisting of a finite set of concepts (unary predicates) and roles (binary predicates). The classifiers in Ψ are assumed to operate over a shared feature space induced by the tuple of attributes $\mathcal{A} = (A_1, \dots, A_n)$, where each attribute A_i is associated with a domain D_{A_i} of values, whose Cartesian product defines the feature space $\mathbb{F}(\mathcal{A}) = D_{A_1} \times \dots \times D_{A_n}$. In other words, an HKB can be seen as a Knowledge Base whose extensional component is represented by the set Ψ of classifiers, and in particular a classifier $\kappa_C: \mathbb{F}(\mathcal{A}) \rightarrow \{0, 1\}$, for each atomic concept $C \in \text{sig}(\mathcal{O})$, and a classifier $\lambda_R: \mathbb{F}(\mathcal{A}) \times \mathbb{F}(\mathcal{A}) \rightarrow \{0, 1\}$ (called classifier over pairs), for each atomic role $R \in \text{sig}(\mathcal{O})$. In what follows, we let $\text{sig}(\Psi) = \mathcal{A}$.

Example 1. Consider a medical scenario in which each patient is represented by an element of the feature space $\mathbb{F}(\mathcal{A})$, with $\mathcal{A} = (\text{age}, \text{bmi}, \text{ant}, \text{a1c}, \text{hb})$ denoting a tuple of patient attributes, where *age* is the patient age, with $D_{\text{age}} = \{0, 1, 2\}$ (0 corresponds to young, 1 to middle-aged, and 2 to senior); *bmi* is the body mass index, with $D_{\text{bmi}} = \{0, 1, 2\}$ (0 corresponds to low BMI, 1 medium, and 2 high); *ant* is the number of ABO blood group antigens, with $D_{\text{ant}} = \{0, 1, 2\}$ (0 corresponds to blood type O, 1 to either A or B, and 2 to AB); *a1c* indicates the value of blood glucose, with $D_{\text{a1c}} = \{0, 1, 2\}$ (0 corresponds to low glucose level, 1 medium, and 2 high); and *hb* contains the hemoglobin levels according to blood tests, with $D_{\text{hb}} = \{0, 1, 2\}$ (0 corresponds to low hemoglobin level, 1 medium, and 2 high).

Suppose also that the doctors exploit classifiers to determine whether a patient is diabetic (classifier κ_D), male (κ_M), or pregnant (κ_P), as well as whether a pair of patients are compatible blood donors (classifier over pairs λ_{cD}) or share a compatible blood type (λ_{cB}). Let them be defined as follows:

- For each $\bar{a} = (a_1, a_2, a_3, a_4, a_5) \in \mathbb{F}(\mathcal{A})$:
 - $\kappa_D(\bar{a}) = 1$ if and only if $a_2 + a_3 - 3 \geq 0$;
 - $\kappa_M(\bar{a}) = 1$ if and only if $3a_2 + a_5 - 5 \geq 0$;
 - $\kappa_P(\bar{a}) = 1$ if and only if $-a_2 + \frac{1}{2}a_3 + a_4 - \frac{5}{2} \geq 0$;
- For each pair $(\bar{a}, \bar{b}) \in \mathbb{F}(\mathcal{A}) \times \mathbb{F}(\mathcal{A})$, where $\bar{a} = (a_1, a_2, a_3, a_4, a_5)$ and $\bar{b} = (b_1, b_2, b_3, b_4, b_5)$:
 - $\lambda_{cD}(\bar{a}, \bar{b}) = 1$ if and only if $-a_3 - a_4 + b_3 + b_4 - 1 \geq 0$;
 - $\lambda_{cB}(\bar{a}, \bar{b}) = 1$ if and only if $a_3 - b_3 \geq 0$.

The inductive knowledge contained in these classifiers can be contextualized within a semantic context. Specifically, we model the above scenario through an ontology $\mathcal{O}$, comprising the atomic concepts D , M , and P for diabetics, males, and pregnant patients, respectively, and the atomic roles cD and cB for pairs of patients that are compatible donors and pairs of patients with compatible blood types, respectively. Furthermore, $\mathcal{O}$ states that a patient cannot be both pregnant and male, i.e. $\mathcal{O} = \{P \sqsubseteq \neg M\}$.

We can model the above scenario through the HKB $\mathcal{K} = (\mathcal{O}, \Psi)$, where $\mathcal{O}$ is the ontology with $\text{sig}(\mathcal{O}) = \{D, M, P, cD, cB\}$ and which contains the axiom $P \sqsubseteq \neg M$ stating that a patient cannot be both pregnant and male, while $\Psi = \{\kappa_D, \kappa_M, \kappa_P, \lambda_{cD}, \lambda_{cB}\}$.

The semantics of HKBs is formalized in logical terms through first-order interpretations. In particular, given a HKB $\mathcal{K} = (\mathcal{O}, \Psi)$ with $\text{sig}(\Psi) = \mathcal{A}$, an *interpretation* for $\mathcal{K}$ is a pair $\mathcal{I} = (\mathcal{I}, f)$, where:

- $\mathcal{I} = (\Delta^{\mathcal{I}}, \cdot^{\mathcal{I}})$ is a first-order interpretation for the ontology $\mathcal{O}$;
- $f: \mathbb{F}(\mathcal{A}) \rightarrow \mathcal{P}(\Delta^{\mathcal{I}})$ associates to each element $\bar{a} \in \mathbb{F}(\mathcal{A})$ a set $f(\bar{a})$ of objects from $\Delta^{\mathcal{I}}$;
- $f(\bar{a}) \cap f(\bar{b}) = \emptyset$ holds for every pair $(\bar{a}, \bar{b}) \in \mathbb{F}(\mathcal{A}) \times \mathbb{F}(\mathcal{A})$ such that $\bar{a} \neq \bar{b}$.

The function f acts as a bridge from the raw input vectors to the symbolic representation embodied by the objects in the interpretation domain $\Delta^{\mathcal{I}}$ of $\mathcal{I}$. Notably, the above notion of interpretation allows modeling situations in which a combination of attribute values $\bar{a} \in \mathbb{F}(\mathcal{A})$ does not correspond to any real-world entity (in this case, $f(\bar{a}) = \emptyset$) or represents one or more real entities sharing exactly those characteristics. The third bullet point ensures that distinct elements of the feature space are mapped to distinct objects of the interpretation domain (similarly to the Unique Name Assumption). In what follows, we denote by $\text{disc}(\mathcal{I})$ the set of $\mathcal{I}$ -discarded elements, i.e. $\text{disc}(\mathcal{I}) = \{\bar{a} \in \mathbb{F}(\mathcal{A}) \mid f(\bar{a}) = \emptyset\}$. We also say that $\mathcal{I}$ is *trivial* if $\text{disc}(\mathcal{I})$ coincides with the feature space $\mathbb{F}(\mathcal{A})$, i.e. if $\text{disc}(\mathcal{I}) = \mathbb{F}(\mathcal{A})$.

An interpretation $\mathcal{I} = (\mathcal{I}, f)$ for a HKB $\mathcal{K} = (\mathcal{O}, \Psi)$ is a model of $\mathcal{K}$ if both $\mathcal{I} \models \mathcal{O}$ and $\mathcal{I} \models \Psi$. As for the condition $\mathcal{I} \models \mathcal{O}$, we simply require that $\mathcal{I}$ satisfies all the axioms in $\mathcal{O}$. As for $\mathcal{I} \models \Psi$, once the universe of discourse $\Delta^{\mathcal{I}}$ and the function f are fixed, we treat the classifiers in Ψ as *exact mappings*, requiring that the extension of atomic concepts and roles faithfully follows the classifiers in Ψ . Thus, like the disjointness axioms in the ontology, the inclusions axioms also act as constraints in a given HKB. The next definition formalizes the above notion of model for a HKB.

Definition 2. Let $\mathcal{K} = (\mathcal{O}, \Psi)$ be a HKB with $\text{sig}(\Psi) = \mathcal{A}$, and let $\mathcal{I} = (\mathcal{I}, f)$ be an interpretation for $\mathcal{K}$. We say that $\mathcal{I}$ is a model of $\mathcal{K}$ if $\mathcal{I} \models \mathcal{O}$ and $\mathcal{I} \models \Psi$, i.e. the next conditions hold:

1. for each $C \in \text{sig}(\mathcal{O})$: $C^{\mathcal{I}} = \{o \mid \exists \bar{a}. \kappa_C(\bar{a}) = 1 \text{ and } o \in f(\bar{a})\}$;
2. for each $R \in \text{sig}(\mathcal{O})$: $R^{\mathcal{I}} = \{(o, o') \mid \exists \bar{a}, \bar{b}. \lambda_R(\bar{a}, \bar{b}) = 1 \text{ and } o \in f(\bar{a}) \text{ and } o' \in f(\bar{b})\}$.

Example 2. Recall Example 1. Let $\mathcal{I}$ be such that $\Delta^{\mathcal{I}} = \{o_{\bar{a}} \mid \kappa_P(\bar{a}) = 0 \vee \kappa_M(\bar{a}) = 0\}$, $D^{\mathcal{I}} = \{o_{\bar{a}} \mid \kappa_D(\bar{a}) = 1\}$, $P^{\mathcal{I}} = \{o_{\bar{a}} \mid \kappa_P(\bar{a}) = 1\}$, $M^{\mathcal{I}} = \{o_{\bar{a}} \mid \kappa_M(\bar{a}) = 1\}$, $cD^{\mathcal{I}} = \{(o_{\bar{a}}, o_{\bar{b}}) \mid \lambda_{cD}(\bar{a}, \bar{b}) = 1\}$, and $cB^{\mathcal{I}} = \{(o_{\bar{a}}, o_{\bar{b}}) \mid \lambda_{cB}(\bar{a}, \bar{b}) = 1\}$. Let $f: \mathbb{F}(\mathcal{A}) \rightarrow \mathcal{P}(\Delta^{\mathcal{I}})$ be the function such that $f(\bar{a}) = \emptyset$ if both $\kappa_P(\bar{a}) = 1$ and $\kappa_M(\bar{a}) = 1$; otherwise, $f(\bar{a}) = \{o_{\bar{a}}\}$. Consider now the interpretation $\mathcal{I} = (\mathcal{I}, f)$ for $\mathcal{K}$. By construction, we have $\mathcal{I} \models \Psi$. Furthermore, since $\mathcal{I}$ discards all feature space elements classified positively by both κ_P and κ_M , we can conclude that $\mathcal{I} \models \mathcal{O}$. It follows that $\mathcal{I}$ is a model of $\mathcal{K}$.

Rather than considering all models of a HKB, it is reasonable to restrict the attention to those that retain maximal information from the application of the classifiers over the sub-symbolic data. In this paper, we focus on models that discard feature space elements in a minimal fashion, according to set inclusion. In particular, we say that $\mathcal{I}$ is a *minimally-discarding model* of $\mathcal{K}$ if (i) $\mathcal{I}$ is a model of $\mathcal{K}$ and (ii) there is no model $\mathcal{I}' = (\mathcal{I}', f')$ of $\mathcal{K}$ such that $\text{disc}(\mathcal{I}') \subsetneq \text{disc}(\mathcal{I})$. Given a HKB $\mathcal{K}$, we denote by $\text{MinDisc}(\mathcal{K})$ the set of minimally-discarding models of $\mathcal{K}$. We point out that the notion of minimally-discarding model shares similar characteristics with the notion of repair, widely studied in the literature on consistent query answering.

Interestingly, a HKB may admit minimally-discarding models that discard different subsets of the feature space.

Example 3. Recall Example 1. Let $\mathcal{K}' = (\mathcal{O}', \Psi)$ be the HKB obtained from $\mathcal{K}$ by adding to $\mathcal{O}$ the axiom: $cD \sqsubseteq cB$. Consider $\bar{a} = (1, 1, 1, 2, 0)$ and $\bar{b} = (1, 0, 2, 2, 1)$. Note that (i) $\lambda_{cD}(\bar{a}, \bar{b}) = 1$ and $\lambda_{cB}(\bar{a}, \bar{b}) = 0$, and (ii) $\kappa_M(\bar{a}) = \kappa_M(\bar{b}) = 0$ and $\lambda_{cB}(\bar{a}, \bar{a}) = \lambda_{cB}(\bar{b}, \bar{b}) = 1$. Due to (i), we derive that there can be no model $\mathcal{I}$ of $\mathcal{K}'$ such that both $\bar{a} \notin \text{disc}(\mathcal{I})$ and $\bar{b} \notin \text{disc}(\mathcal{I})$. Due to (ii), we derive that there can be models $\mathcal{I}$ of $\mathcal{K}'$ such that $\bar{a} \notin \text{disc}(\mathcal{I})$ and models $\mathcal{I}'$ of $\mathcal{K}'$ such that $\bar{b} \notin \text{disc}(\mathcal{I}')$. As a result, there exist at least two distinct minimally-discarding models of $\mathcal{K}'$ that discard different subsets of the feature space.

Since the semantics of an HKB is based on the notion of model, we can define reasoning tasks over HKBs in terms of model-theoretic entailment. We focus on the following two key reasoning tasks:

- *non-trivial consistency checking*, which consists in determining whether a given HKB $\mathcal{K}$ admits at least one non-trivial model, i.e. a model $\mathcal{I}$ of $\mathcal{K}$ such that $\text{disc}(\mathcal{I}) \subsetneq \mathbb{F}(\mathcal{A})$;
- *query answering*, which consists in determining whether a given query is entailed by a given HKB with respect to a predefined semantics.

For query answering, we first adopt a cautious approach by considering only minimally-discarding models. This form of reasoning is reminiscent of classical *skeptical reasoning* over all repairs of a KB, a common strategy for handling query answering with respect to inconsistent KBs [3, 4]. Formally, let $\mathcal{K} = (\mathcal{O}, \Psi)$ be a HKB, let $q = \{\bar{x} \mid \varphi(\bar{x})\}$ be a first-order query over $\mathcal{O}$ of arity m , and let $\bar{t} = (\bar{a}_1, \dots, \bar{a}_m)$ be an m -tuple of elements from $\mathbb{F}(\mathcal{A})$, i.e. $\bar{t} \in \mathbb{F}(\mathcal{A})^m$. We say that $\bar{t}$ is a *skeptical-answer* to q over $\mathcal{K}$ if, for every $\mathcal{I} = (\mathcal{I}, f) \in \text{MinDisc}(\mathcal{K})$, the following condition holds: there exists an m -tuple $\bar{o} = (o_1, \dots, o_m)$ of objects from $\Delta^{\mathcal{I}}$ such that $\mathcal{I} \models \varphi(\bar{x}/\bar{o})$ and $o_i \in f(\bar{a}_i)$, for each $i \in [m]$.

Example 4. Recall Example 2, and consider the union of conjunctive queries (UCQ[≠]) $q = \{(x, y) \mid cD(x, y) \wedge x \neq y \wedge (P(x) \wedge P(y)) \vee (D(x) \wedge D(y))\}$ over $\mathcal{O}$, asking for pairs of distinct patients that are compatible blood donors and are either both pregnant or both diabetic. Consider now $\bar{a} = (1, 2, 1, 0, 2)$ and $\bar{b} = (0, 2, 1, 1, 2)$. It is not hard to verify that $(\bar{a}, \bar{b})$ is a skeptical answer to q over $\mathcal{K}$.

We start by studying the computational complexity of the above tasks under the setting in which the ontology $\mathcal{O}$ is specified in $DL\text{-}Lite_{RDFS}^\neg$ [5], i.e. the Description Logic counterpart of RDFS (called $DL\text{-}Lite_{RDFS}$ [6]) equipped with disjointness axioms, each classifier is a *Multi-Layer Perceptron* (MLP) [7], and the query (for the query answering problem) is either a *conjunctive query* (CQ) or a UCQ[≠]. For the classifiers, we let the domain of each attribute be $\{0, 1\}$. This is only for the sake of presentation, since all our results can be easily extended to the case in which the domains are finite sets of arbitrary size. We further assume that each binary classifier over pairs $\lambda_R: \{0, 1\}^n \times \{0, 1\}^n \rightarrow \{0, 1\}$ is treated as a standard classifier $\kappa_R: \{0, 1\}^{2n} \rightarrow \{0, 1\}$ operating on concatenated input pairs, i.e. $\lambda_R(x, x') = 1$ if and only if $\kappa_R(x \parallel x') = 1$. Finally, we mention that our investigation focuses on the *classifiers complexity*, which is the complexity where only the set Ψ of classifiers is regarded as the input, while the ontology and the query (if applicable) are fixed. This complexity measure is analogous to the *data complexity* [8] of reasoning tasks over KBs, where the input is the set of facts.

Theorem 1. For $DL\text{-}Lite_{RDFS}^\neg$ ontologies and MLP classifiers, we have that:

- Non-trivial consistency checking is NP-complete in classifiers complexity;
- Skeptical query answering is CONEXPTIME-complete in classifiers complexity for UCQ[≠]s.

For query answering, the hardness holds already for CQs. In light of such a negative result, we further investigate the query answering problem over HKBs by (i) considering a fragment of the ontology language $DL\text{-}Lite_{RDFS}^\neg$ and (ii) adopting an alternative semantics for query answers that avoid skeptically reasoning over all the minimally-discarding models. As for (i), we observe that the main source of complexity in the above result is the presence of axioms involving atomic roles. This naturally leads us to consider the fragment $\mathcal{H}^\neg$ of $DL\text{-}Lite_{RDFS}^\neg$ which forbids axioms that mention atomic roles. Note that $\mathcal{H}^\neg$ remains a useful ontology language, as it allows to express taxonomies (i.e. hierarchy of atomic concepts) as well as disjointness between atomic concepts. For instance, the ontology $\mathcal{O}$ of Example 1 is a $\mathcal{H}^\neg$ ontology, whereas the ontology $\mathcal{O}'$ of Example 3 is not. While non-trivial consistency checking remains NP-complete in classifiers complexity also in this restricted setting, for skeptical query answering we get the following result.

Theorem 2. For $\mathcal{H}^\neg$ ontologies, MLP classifiers, and UCQ[≠]s, the skeptical query answering problem is NP-complete in classifiers complexity. The hardness holds already for CQs.

We conclude our investigation by studying a semantics that soundly-approximate the skeptical answers semantics. Such an alternative semantics is inspired by the well-behaved *IAR semantics* adopted for query answering over inconsistent KBs [3], hence following the *WIDTIO* (When In Doubt Throw It Out) approach of the belief revision and update research area. Specifically, we say that a model $\mathcal{J}$ for $\mathcal{K}$ is a *minimally-discarding WIDTIO-model* of $\mathcal{K}$ if, for every $\bar{a} \in \mathbb{F}(\mathcal{A})$, the following holds: $\bar{a} \in \text{disc}(\mathcal{J})$ if and only if there exists a $\mathcal{J}' \in \text{MinDisc}(\mathcal{K})$ such that $\bar{a} \in \text{disc}(\mathcal{J}')$. Let now $q = \{\bar{x} \mid \varphi(\bar{x})\}$ be a UCQ[≠] over $\mathcal{O}$ of arity m and $\bar{t} = (\bar{a}_1, \dots, \bar{a}_m)$ be an m -tuple of elements from $\mathbb{F}(\mathcal{A})$. We say that $\bar{t}$ is a *WIDTIO-answer* to q over $\mathcal{K}$ if, for every minimally-discarding WIDTIO-model $\mathcal{J} = (\mathcal{I}, f)$ of $\mathcal{K}$, the following holds: there exists an m -tuple $\bar{o} = (o_1, \dots, o_m)$ of objects from $\Delta^\mathcal{I}$ such that $\mathcal{I} \models \varphi(\bar{x}/\bar{o})$ and $o_i \in f(\bar{a}_i)$, for each $i \in [m]$.

Theorem 3. For $DL\text{-}Lite_{RDFS}^\neg$ ontologies, MLP classifiers, and UCQ[≠]s, the WIDTIO query answering problem is Σ_2^P -complete in classifiers complexity. The hardness holds already for CQs.

Declaration on Generative AI

During the preparation of this work, the authors used OpenAI ChatGPT-4 in order to: Grammar and spelling check; Paraphrase and reword. After using this tool/service, the authors reviewed and edited the content as needed and take full responsibility for the publication's content.

References

- [1] G. Cima, M. Console, L. Papi, Foundations of formal reasoning over knowledge bases combining symbolic and sub-symbolic knowledge, in: Proceedings of the Fortieth AAAI Conference on Artificial Intelligence (AAAI 2026), 2026, pp. 18994–19002.
- [2] C. Thames, Y. Sun, A survey of artificial intelligence approaches to safety and mission-critical systems, in: 2024 Integrated Communications, Navigation and Surveillance Conference (ICNS), 2024, pp. 1–12.
- [3] D. Lembo, M. Lenzerini, R. Rosati, M. Ruzzi, D. F. Savo, Inconsistency-tolerant query answering in ontology-based data access, *Journal of Web Semantics* 33 (2015) 3–29.
- [4] M. Bienvenu, C. Bourgaux, Inconsistency-tolerant querying of description logic knowledge bases, in: J. Z. Pan, D. Calvanese, T. Eiter, I. Horrocks, M. Kifer, F. Lin, Y. Zhao (Eds.), Proceedings of the Twelfth International Summer School on Reasoning Web: Logical Foundation of Knowledge Graph Construction and Query Answering (RW 2016), volume 9885 of *Lecture Notes in Computer Science*, Springer, 2016, pp. 156–202.
- [5] G. Cima, M. Console, R. M. Delfino, M. Lenzerini, A. Poggi, Answering conjunctive queries with safe negation and inequalities over RDFS knowledge bases, in: Proceedings of the Thirty-Ninth AAAI Conference on Artificial Intelligence (AAAI 25), 2025, pp. 14824–14831.
- [6] B. Cuenca Grau, A possible simplification of the semantic web architecture, in: Proceedings of the Thirteenth International World Wide Web Conference (WWW 2004), 2004, pp. 704–713.
- [7] P. Barceló, M. Monet, J. Pérez, B. Subercaseaux, Model interpretability through the lens of computational complexity, in: Proceedings of the Thirty-Third Annual Conference on Neural Information Processing Systems (NeurIPS 2020), 2020.
- [8] M. Y. Vardi, The complexity of relational query languages (extended abstract), in: Proceedings of the Fourteenth Annual ACM Symposium on Theory of Computing (STOC 1982), 1982, pp. 137–146.