

# Neuro-Symbolic methods for Trustworthy AI: a systematic review with a focus on interpretability

Cyprien Michel--Del  tie  <sup>a,b,\*</sup> and Md Kamruzzaman Sarker  <sup>a,c,\*\*</sup>

<sup>a</sup> *Computing Sciences Department, University of Hartford, CT, USA*

<sup>b</sup> *Computer Science Department, ENS de Lyon, France*

<sup>c</sup> *Computer Science Department, Bowie State University, MD, USA*

## Abstract.

Recent advances in Artificial Intelligence (AI) and especially in deep learning have manifested an increasing concern in trustworthiness. Neuro-symbolic methods, which mix some elements of neural networks with some elements of symbolic reasoning, have shown great potential for some aspects of trustworthiness, particularly for interpretability. In this paper, we provide an overview of the various ways neuro-symbolic methods have been used to increase the trustworthiness, in the latest literature of the leading conferences. In particular, we focus on the contributions of the recent articles that achieve better interpretability thanks to NeSy systems, while also considering contributions in a broader sense, such as safety, fairness, and privacy. We also did a categorization of the existing contributions along several key dimensions related to the symbolic structures they are exploiting, and the type of interpretability they provide.

Keywords: Neuro-Symbolic, Trustworthiness, Interpretability

## 1. Introduction

The field of Artificial Intelligence (AI) is in a continuous state of exploration, with its potential applications appearing to be endless. AI decision-making systems have demonstrated superior performance, frequently outperforming humans. However, this comes with a notable drawback: the decision processes of these systems lack transparency and are often incomprehensible. This issue becomes increasingly critical as AI systems begin to handle sensitive data and make crucial decisions in various sectors, ranging from autonomous driving to criminal justice. As a result, the demand for trustworthiness in AI systems is escalating. Particularly, the subject of interpretability has seen a significant rise in interest in recent years, and is now a major research focus. This increase is a direct consequence of recognizing that many top-tier AI systems are non-transparent and difficult to interpret, leading them to be labeled as “black boxes”. A common trend observed is that the larger the AI model, the more challenging it is to interpret its internal workings. These complex models pose a problem, as it becomes increasingly difficult to identify errors or biases within the system. Shifting towards more interpretable systems would cultivate greater trust in their decisions, enhance social acceptance, and encourage stakeholder discussions about their implementation [48].

Neuro-Symbolic AI (NeSy), which combines Machine Learning (ML) with mechanisms related to Knowledge Discovery and Data Mining (KDD), seeks to integrate neural networks with symbolic processing techniques. This field attracts interest from two distinct perspectives [35]. From a cognitive science perspective, while human brains exhibit connectionist characteristics like neural networks, they are also able to process complex symbolic structures. This capability is believed to play a crucial role in the superiority of human intelligence over other animals. Additionally, from a conceptual perspective, it appears that symbolic and neural approaches complement each other, each with their own strengths and weaknesses. For example, deep learning systems, trained on raw data, show robustness against outliers, a feature less prominent in symbolic systems. In contrast, symbolic systems can directly utilize expert knowledge and are generally more self-explanatory compared to their neural counterparts.

---

\*Corresponding author. E-mail: cyprien.michel--deletie@ens-lyon.fr.

\*\*Currently working at Bowie State University. E-mail: msarker@bowiestate.edu.

The self-explanatory nature of neuro-symbolic methods is especially relevant when considering trustworthiness and especially interpretability. Indeed, one of the main criticism toward the current neural models is their lack of transparency, but symbolic systems do not have this issue. Therefore, in tasks where the state-of-the-art is largely consisting of neural methods, developing a neuro-symbolic approach offers the opportunity to exploit the interpretability that the symbolic aspect provides.

In this paper, we present a systematic review of recent literature (from 2021 to 2022) on neuro-symbolic approaches with a focus on achieving high trustworthiness. The survey methodology was designed to encompass methods that could be considered as neuro-symbolic even-though they were not labeled as such; it is thus presenting a more representative overview and highlighting the amount of existing work that can be classified as neuro-symbolic. We considered different trustworthiness dimensions: privacy, fairness, safety, or interpretability. However, all but two of these papers focused on interpretability, thus interpretability became a primary focus of this work. These papers were further categorized using a traditional taxonomy in three dimensions: global versus local methods, self-explainable versus post-hoc explainability methods, and model-agnostic versus model-specific methods. We also reviewed papers dealing with interpretability based on the symbolic structures used. This review provides an overview of the current trends in this domain, highlighting areas that have been thoroughly explored, revealing common challenges and designs across multiple domains and pinpointing promising directions for future research.

The paper is organized as follows. Section 2 provides an extensive background of the core concepts involved in this survey, as well as grounding for our taxonomy and the presentation of related works. Section 3 presents the survey’s methodology, its framing and some observations on the papers found. Section 4 develops on the different types of symbolic knowledge used by the selected papers and presents each paper’s contribution. Section 5 analyzes the different types of interpretability provided by these papers. Section 6 presents additional learnings from our review and further discussions.

## 2. Background

### 2.1. History of Neuro-Symbolic AI

The genesis of neuro-symbolic (NeSy) research is deeply intertwined with the history of Artificial Intelligence (AI), its roots arguably dating back to a seminal 1943 paper by McCulloch and Pitts [62]. This pioneering work used propositional logic to model neural connections, setting the foundation for what would evolve into NeSy. Historically, the field of AI has been bifurcated into two primary paradigms: symbolism and connectionism. Symbolism approached intelligence through the lens of logic and rules, while connectionism favored learning driven by probabilistic methods. From the mid-1950s to the late 1980s, symbolic models dominated the early AI landscape, as researchers predominantly pursued this approach to create problem solving systems [91]. However, the field encountered unexpected hurdles, leading to the infamous “AI winter” of the 1980s, marked by a significant decline in AI interest and funding [85]. Despite this setback, research on symbolic AI persisted, albeit overshadowed by the resurgence of connectionist AI in the early 2010s. This revival, fueled by the impressive capabilities of deep learning in areas such as image classification, brought new attention to the field. Nevertheless, alongside these advancements came increasing concerns over the limitations of connectionist systems, such as vulnerability to adversarial attacks, low interpretability, challenges in integrating expert knowledge, limited reasoning capabilities, and inherent biases.

NeSy emerged as a promising avenue to address these challenges. Although its conceptual roots span several decades, it was not until the 1990s that NeSy began to crystallize as a distinct field of study, gaining more structured research attention in the early 2000s [77]. NeSy aims to synthesize the strengths of both symbolic and neural elements, striving to create systems that exhibit robust learning capabilities (able to improve from raw data) and strong reasoning prowess (capable of abstraction and combinatorial reasoning). Although neural networks have shown impressive performances, logic remains a cornerstone in modeling thought and behavior [13]. The integration of these paradigms holds the promise of retaining their respective strengths while mitigating their weaknesses. However, this integration is challenging due to their fundamentally different methodologies: statistical inductive

learning and distributed representations in connectionism, contrasted with logical deductive reasoning and localist representations in symbolism [91].

NeSy has shown its utility in various ways, such as leveraging symbolic knowledge bases and metadata to enhance deep learning systems, providing greater explainability through background knowledge, and solving complex problems that benefit from symbolic reasoning structures [35]. Furthermore, NeSy has found successful applications in diverse industrial contexts, including business process modeling, trust management in e-commerce, coordination in large-scale multi-agent systems, and multimodal processing and applications [13].

## 2.2. Background on Trustworthiness

The concept of trustworthiness is paramount in any decision-making system. At its core, a system is deemed trustworthy if it can be relied on for high-stakes decisions with minimal or no supervision. Although this certainly includes performance, as a high-performing system is a prerequisite for trustworthiness, in the AI field, trustworthiness encompasses several additional dimensions: *interpretability*, *fairness*, *robustness*, *privacy*, and *safety* [4, 55, 64].

*Fairness* focuses on ensuring that AI models do not harbor biases that could lead to discrimination against certain groups [87]. This is especially pertinent in AI applications involving the classification of people, such as risk assessments in criminal recidivism or automatic resume screening, both of which are rapidly gaining traction [73]. Studies have uncovered biases in some deployed systems against racial minorities, even in the absence of explicit racial data input [43]. Addressing these biases to ensure fairness towards all groups is a critical concern.

*Privacy* relates to safeguarding the private data used to train AI models [64]. There is a risk that interaction with the deployed models or analysis of those could inadvertently expose sensitive training data, a situation that raises significant privacy concerns.

*Robustness* is about the system’s ability to function correctly in scenarios that deviate from its training data distribution [87]. This is vital across AI applications, as it is often impossible to anticipate all potential scenarios a system may encounter. The susceptibility of deep neural networks to adversarial attacks is particularly concerning, when subtle manipulations of input data can lead to incorrect interpretations by the AI, despite being obvious to humans. *Since robustness is closely intertwined with performance, it is often unclear when research specifically focuses on robustness; hence, papers primarily addressing robustness were not included in our review.*

*Safety* is a critical aspect of trustworthiness that focuses on preventing accidents and unintended harmful behaviors in machine learning systems [4]. These issues can arise due to errors in the specification of objectives, oversights in the learning process, or other implementation mistakes. As AI systems are increasingly deployed in complex environments with real stakes, ensuring their safety becomes paramount. This involves creating scalable solutions to mitigate risks and avoid potential adverse impacts on society, making AI systems not only effective but also reliable and secure.

## 2.3. Background on Interpretability

*Interpretability* is the most extensively addressed aspect of trustworthiness, experiencing exponential growth as a research domain [6, 82]. There is little consensus on the precise definition of interpretability, but it can be broadly defined as the extent to which a system’s operations can be understood by users [22, 48]. This includes access to mechanisms or reasoning that underpin the system’s predictions. Simpler systems are naturally more interpretable, which is why this was not a major topic in earlier AI systems that used simpler methods like decision trees. However, with the complexity of deep neural networks, interpretability has become a critical concern, for societal acceptance as well as regulatory compliance; indeed, both United States and the European Union have imposed a right to explanation for consumers [6]. We also argue that interpretability is crucial for a better understanding of the systems, which will help to develop them further and to overcome their flaws.

The characterization and approach to interpretability in AI is a subject of ongoing debate. Although many papers use explainability and interpretability interchangeably, some argue that explainability is a stronger concept than interpretability [6, 10, 28]. Since the frontiers between these two concepts is quite vague, our review is based on the terminology used by the authors of the papers, which may not always align with this distinction; thus explainable and interpretable should be understood as interchangeable in our paper. Generally, interpretability is self-assessed

by researchers, leading to calls for more rigorous taxonomies and evaluations [9, 22, 40, 80]. It is also important to note that explainability is not the “silver bullet” for AI trustworthiness. Studies have shown that while explainability can enhance AI collaboration with novices, it does not necessarily do so with experts [63]; a combination of AI and human decision-making can be quicker but less accurate when AI provides explanations [40]; and there is a risk that explanations, even if not particularly useful, can unduly increase public acceptance, leading to over-reliance on AI [23, 93].

Interpretability in AI systems has been tackled through a variety of methods, which can usually be divided into two categories depending on what kind of interpretability they provide: either *ante-hoc* or *post-hoc* [9, 80, 82]. *Ante-hoc explainability* is a characteristic of systems which are inherently designed to be easily interpretable from the inside (also known as *self-explainable* systems). These systems are structured in order to make their internal processes straightforward and clear. *Post-hoc explainability* is a characteristic of systems that are not inherently transparent but present an additional layer that attempts to give more insights about the underlying decision process. This method is particularly versatile, as it can be applied to virtually any system, allowing for the continued use of high-performance models. However, a drawback of post-hoc explainability is that the explanations it provides might not always accurately reflect the true workings of the system. This concern is highlighted by Rudin [71], who argues against the use of such explainability, suggesting that it can be misleading. Conversely, Gilpin et al. [28] argued that when using post-hoc explanations, it is crucial to clearly inform users about their potential limitations. There are also approaches that fall somewhere between these two extremes [26, 47]. These methods aim to train systems in a way that makes their decision-making processes easier to interpret, without fundamentally altering their core structure.

Interpretability systems can also be categorized based on the scope of explanations, they can range from *local* to *global* [9, 80, 82]: *Local explanations* are explanations tailored to individual instances, providing insights on specific decisions or similar cases. A well-known example of a local explanation method is LIME [70], which is designed to offer explanations for specific data points. *Global explanations* aim to shed light on the system’s behavior as a whole, irrespective of individual inputs. Some methods [81, 103] provide explanations for a specific category of inputs, allowing a more targeted understanding of the system’s decisions in particular scenarios. These diverse approaches to interpretability demonstrate the complexity and varied nature of making AI systems transparent and understandable.

#### 2.4. Link between Interpretability and NeSy

Neuro-Symbolic AI (NeSy) and interpretability are intrinsically connected, primarily because symbols serve as an effective medium for explanations. Common practices in generating explanations include the use of decision trees or logic rules, which are inherently symbolic. Kambhampati et al. [46] have even suggested that symbols are essential for effective communication between humans and AI systems. Although visual representations like saliency maps are also popular for explanations, these may not be adequate for complex human-AI interactions that require a blend of implicit and explicit task knowledge. Since NeSy inherently involves dealing with symbols within decision systems, it naturally possesses a strong potential for high interpretability.

Another perspective on the connection between NeSy and interpretability is their shared role as intermediaries linking deep learning with neuroscience. As Angelov et al. [6] have pointed out, a key objective of explainability is to mimic human-like reasoning in a manner that elucidates the predictions made by AI systems. This goal aligns closely with the principles of NeSy, which integrate aspects of human cognitive processes and neural network-based learning. Therefore, the synergy between NeSy and interpretability is not only practical in terms of implementing symbolic representations for explanations but also fundamental in achieving a deeper, more human-like understanding of AI decision-making processes.

#### 2.5. Related Works

Trustworthiness, being a broad and multifaceted concept in AI, encompasses a diverse range of studies and reviews, often focused on specific domains within the field. A notable comprehensive survey by Liu, Wang et al. [59] addresses recent techniques for enhancing AI trustworthiness. This work examines trustworthiness over six dimensions: explainability, robustness, accountability /auditability, privacy, fairness, and environmental well-being.

Another notable review in the field of trustworthy Machine Learning was conducted by Serban et al. [75], providing valuable insights into methods to foster trust in AI systems.

Interpretability methods in AI have received considerable attention, with numerous reviews dedicated to this topic. In the latest reviews, the focus is usually expressed with the word *explainability* (XAI), but as discussed in Section 2.3, the core idea is the same as *interpretability*. For instance, Speith et al. [82] conducted an analysis of various taxonomies used to categorize interpretability methods. Their study revealed that these taxonomies are based on different criteria, such as the methods used, the type of explainability produced, or the conceptual approach, sometimes combining several of these aspects. They argued that the choice of taxonomy should align with the user’s needs and proposed a unified taxonomy to guide users. Similarly, the very comprehensive review by Barredo Arrieta et al. [9] analyzed and synthesized the existing taxonomies of XAI, and proposed to assess explainability based on the targeted audience. They also proposed both a review of existing transparent (ie. self-interpretable) methods and a review of post-hoc explainability methods. Lastly, they extended their review to what they call “*Responsible AI*”, a concept roughly equivalent to Trustworthy AI. Another relevant review by Weller [93] explored the concept of transparency, putting it in perspective with the different stakeholders of AI. This consideration of the different stakeholders is also of key importance in the frameworks of Kasizadeh [48] and Langer et al. [53]. Vilone et al. [89] have performed extensive classifications of explainable artificial methods, focusing on the formats of their outputs. This approach is highly beneficial for users seeking the most suitable system for their specific requirements.

Reviews specifically focusing on NeSy learning have also been published [12, 13, 31, 72, 91]. Sarker et al. [72] provided a systematic review of neuro-symbolic methods presented in leading conference proceedings, applying two different taxonomies to categorize these methods and noting a recent increase in their popularity. Besold et al. [13] presented a more conceptual review of the neural-symbolic field, discussing its foundations, current applications, and future challenges. Berlot-Attwell [12] explored the use of NeSy AI in Visual Question Answering (VQA), while Hamilton et al. [31] offered a detailed analysis of NeSy methods in Natural Language Processing (NLP), highlighting the challenges in classifying papers as NeSy due to the term’s ambiguity. Wang et al. [91] conducted a systematic overview of recent advances in neuro-symbolic computing and described a taxonomy in four dimensions, inspired by Bader and Hitzler [7]. While we considered insights from existing reviews, we found the intersection of NeSy and interpretability to be a compelling and underexplored area for a survey. To our knowledge, this is the first paper to systematically review neuro-symbolic methods through the lens of trustworthy AI.

### 3. Survey

#### 3.1. Methodology

The aim of this survey was to capture the current state of research in the application of neuro-symbolic tools for enhancing trustworthiness in AI. We focused on papers published in top academic venues from 2021 to 2022 (latest publications at the time of writing). While we are aware of emerging venues such as the NeSy AI Journal and those hosted by IOS Press and Sage, many of them were either very recent or not yet fully established at the time of our review. Additionally, although several well-regarded AI journals exist, their primary focus did not align specifically with the neuro-symbolic AI domain. As a result, we focused our literature review on top-tier AI conference proceedings, where substantial and timely research in this area was more readily accessible. We selected papers from the following conferences: *NeurIPS*, *AAAI*, *IJCAI*, *IJCL*, *ICML*, *NeSy*, *AACM FAccT*, and *KDD*. The volume of papers presented at these conferences, exceeding 10,000 over the last two years, required a more strategic approach to identify relevant papers, rather than reviewing each one individually.

In our commitment to transparency, we employed a detailed and systematic methodology. Using the *dblp* database, papers were initially filtered based on titles that contained keywords indicative of neuro-symbolic methods. These keywords included *symbol*, *logic* (excluding derivatives like *biologic* or *topologic*), *reason*, *inducti(on)*, *abducti(on)*, *concept*, *hybrid*, *ontolog(y)*, *relational*, *compositional*, and *rule*. Search-in-page tools were then used to determine if these papers frequently mentioned key terms related to trustworthiness, such as *interpretab(le)*, *explaina(ble)*, *explanat(ion)*, *trust*, *fair*, *faithful*, *priva(cy)*, *tractab(le)*, *safe* and *understandab(le)*. Papers meeting these criteria were examined in more detail to assess their relevance to our focus.

Additionally, to ensure that we did not overlook papers that explicitly mentioned the use of NeSy methods (but not in the title), we screened all papers with titles suggesting a focus on trustworthiness. We then reviewed papers that contained multiple mentions of the keyword *symbol* (and thus every derived terms such as *symbolic*) for further evaluation. This comprehensive approach was designed to capture a wide range of relevant research, ensuring a thorough overview of the intersection of neuro-symbolic research and trustworthiness in AI.

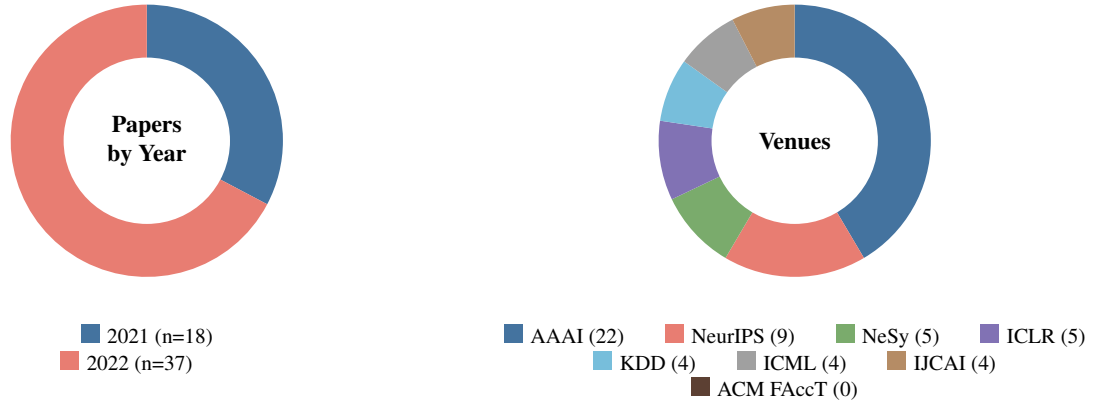


Fig. 1. Distribution of selected papers.

### 3.2. Framing the survey

Determining whether a paper’s approach qualifies as neuro-symbolic presented a significant challenge due to the broadness and ambiguity surrounding the definition of NeSy. To address this, we established specific criteria: a paper was included in our review only if it involved some form of symbolic knowledge manipulation (such as logic propositions, rules, action models, or graphs) directly contributing to trustworthiness. We specifically looked for papers where this symbolic knowledge played an active role in the process, rather than being a mere output. For example, if the explanations were presented in the form of a graph that was neither used nor executed in the system, we did not consider the symbolic role to be sufficiently integrated in the framework for it to be considered as neuro-symbolic.

While we recognize that this approach might have excluded some relevant papers, our objective was to minimize any systematic bias in our selection process that could lead to a skewed representation of the field. We noticed that many papers treated interpretability as a beneficial by-product rather than a primary focus, without substantial discussion or emphasis. To maintain the relevance and specificity of our survey, we chose to include only the papers where trustworthiness was a central motivation of the research. This decision inevitably introduced a degree of subjectivity into the selection of papers, but it was a necessary step to ensure the focus and coherence of our survey.

### 3.3. Some Statistics

Our comprehensive review yielded a total of 54 papers that employed neuro-symbolic methods with a clear emphasis on trustworthiness. An interesting pattern emerged from our analysis: the vast majority of these papers, except for two (one focusing on fairness [101] and another on safety [97]), focused on interpretability. This trend was notable despite our efforts to encompass a broader range of trustworthiness aspects such as fairness and privacy. This observation suggests that, currently, NeSy may not be widely utilized to address trustworthiness concerns beyond interpretability. Another observation is that of a significant increase in relevant publications in 2022, with 37 out of the 54 papers coming from this year alone (Figure 1), indicating a growing interest and expansion in this domain. The distribution of these papers across various conferences reveals that AAAI is the predominant venue for this type of research. We also noted that no papers from ACM FAccT ended up in our survey, despite the fact that

this conference specifically focuses on trustworthiness issues. This is because our keyword search found very few matches in FAccT papers, and the only papers that matched were not proposing a neuro-symbolic approach in our view. This was quite surprising, but it is worth mentioning that this is also a conference with quite few papers compared to the other conferences considered here.

### 3.4. Applications of NeSy Systems

While exploring interpretability contributions of the NeSy systems, we found that they were being used for different applications. Many of the proposed systems were working with visual data: either image classification [5, 39, 47, 81], action recognition in videos [36], agent communication about images [19], handwritten mathematical expression recognition [94], visual relation detection [15], or visual reasoning [21]. Equally many of the systems dealt with natural language settings: fake news detection [14, 42, 99], question answering [18, 104], unspecified NLP [103], text classification [101], commonsense reasoning [45], medical diagnosis through dialogue [60], text fiction tasks [66], or news recommendation [58]. A few applications were entirely based on graphs: knowledge graph completion [17, 102], query answering on knowledge graphs [105], graph classification [34], or imitating algorithms on graphs [27]. Some researchers worked in settings where an agent has to make different decisions (often trained with reinforcement learning) [25, 41, 83, 88]. In some cases, the methods were explicitly suited for multiple settings [8, 94]. Lastly, many other unique settings were explored: adaptive management [24], time series analysis [69], congestion control [76], safe execution of programs [97], or computer algebra [65]. The wide range of applications shows how versatile NeSy methods can be.

## 4. Analysis based on the type of symbolic knowledge

### 4.1. The different symbolic structures used

In our categorization of the papers, we found that 16 of them presented **rule-learning** approaches [20, 29, 30, 44, 52, 56, 61, 67, 74, 78, 86, 92, 95, 96, 98, 100]. These papers typically utilize deep learning to generate logic rules or decision trees for classification purposes. In these instances, machine learning techniques allow the creation of a symbolic model, which offers a high degree of transparency and interpretability. Beyond rule-learning approaches, we also analyzed the types of symbolic data structures employed in other NeSy systems. We identified that these systems could be broadly categorized into three types based on the symbolic data structures they manipulate: *logic*, *graphs*, and *other structures*. This classification, depicted in Figure 2, provides insight into the varied approaches within the NeSy field, highlighting the diversity of methods being explored to improve trustworthiness in AI systems.

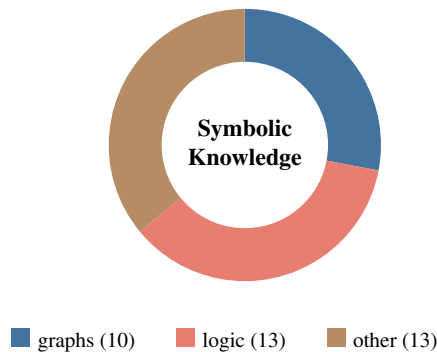


Fig. 2. Form of the symbolic knowledge (in papers that do not fall into the rule-learning category)

*Logic*, as used in various papers [5, 8, 14, 17, 27, 34, 45, 47, 54, 69, 76, 81, 101], typically involve logic propositions, often in the form of logic rules (e.g., *precondition*  $\rightarrow$  *class*). This approach usually uses symbolic reasoning

as a means to interpret and classify data. *Graphs* are another prevalent structure in NeSy research, encompassing a variety of types. Knowledge graphs are commonly used [58, 66, 102, 105], but the category also includes other types of graphs [18, 24, 42, 94, 99, 104], such as scene graphs, proof graphs, or Abstract Meaning Representation (AMR) graphs. These graphical structures are instrumental in representing relationships and dependencies in a visual and often intuitive manner. The third category, labeled as “*other*”, encompasses a variety of other symbolic data forms [15, 19, 21, 25, 26, 36, 39, 41, 60, 65, 83, 88, 103]. This includes, for example, symbolic descriptions of objects or symbolic programming languages. This category is diverse and encompasses a wide range of approaches in which symbolic representations take on various forms.

#### 4.2. Description of the reviewed papers

Interestingly, each of these three categories—logic structures, graphs, and other structures—encompasses a similar number of papers. These varied methodologies highlight the versatility of symbolic representations in AI and their potential to address different aspects of trustworthiness in sophisticated and nuanced ways. *Note that we chose not to elaborate about rule-learning papers [20, 29, 30, 44, 52, 56, 61, 67, 74, 78, 86, 92, 95, 96, 98, 100] in the following to maintain appropriate scope and manuscript length.*

##### 4.2.1. Approaches using logic

A common approach for interpretability is to extract symbolic rules that explain the system’s prediction. This can often involve modifying the system’s structure so that the rule extraction from it can be more feasible and faithful. For instance, Barbiero et al. [8] proposed a new approach in which the classifier is designed in a way that allows the extraction of logic rules to explain its predictions. This approach can be related to Lee et al.’s framework [54], which upgrades a deep model into a self-explainable version by naturally integrating human priors and rule generation into its predictions. Acting on the training step, Sharan et al. [76] proposed a method to train a deep model, then extracts from it symbolic rules. Similarly, Kasioumi et al. [47] proposed a new learning method (Elite BackProb) which promotes activation sparsity of the filters of a convolutional neural network, so that a rule extraction algorithm can be used to approximate its predictions. In the graph processing domain, Himmelburger et al. [34] made a framework which extract rules for post-hoc explanation of graph neural networks (GNN). Georgiev et al. [27] proposed concept-bottleneck GNNs, which are variants of GNNs with a new readout which allows the production of explanations in propositional logic based on inferred concepts. Cucala et al. [17] proposed a new class of knowledge graphs transformations that are always equivalent to the application of symbolic rules. Rajapaksha and Bergmeir [69] proposed a model to produce rule-based explanations of a black-box Global Forecasting Model on several time series.

Other approaches used logic in original ways, usually involving it in the system’s decision to make it more interpretable. For example, Chen et al. [14] proposed an approach that decomposes texts into phrases and uses aggregation logic to classify them as fake or not in an interpretable way. Yao et al.’s framework [101] parses advices on what a language model is wrongfully using for its prediction into First Order Logic (FOL) and use it to refine the weights of the language model. In this case the main goal is about increasing fairness and not interpretability. Ribeiro and Leite’s framework [81] maps a neural network’s internal states to concepts from an ontology, making it possible to build explanations from these concepts. Kalyanpur et al. [45] proposed a novel reasoner that combines symbolic reasoning with statistical functions for fuzzy unification and dynamic rule generation. In the image classification domain, An et al. [5] proposed a novel rule-guided method called dynamic ablation to provide explanations as well as visual highlights.

##### 4.2.2. Approaches using graphs

Among the works using graphs, a few approaches share the common characteristic of using knowledge graphs. For instance, Zha et al. [102] proposed a method for knowledge graph completion that outputs a pattern in the graph to explain the predictions, using BERT. Zhu et al. [105] developed an approach that converts logical queries into circuits including graph neural networks to answer them based on a knowledge graph. Liu et al. [58] proposed a new method for news recommendation: small anchor graphs are generated via reinforcement learning so that the similarity of two articles can be estimated by computing the number of paths connecting the two anchor graphs.

Peng et al. [66] designed a reinforcement learning agent that uses a knowledge graph to represent its belief about the world alongside an attention mechanism to be able to explain its reasoning.

Various approaches also used different types of graphs in original ways. For instance, Ferrer-Mestres et al. [24] proposed an approach to extract policies of a fixed length from policies or arbitrary lengths so that they are small enough to be interpretable, in the Mixed Observability Markov Decision Process Setting (policies are represented as graphs). Wu et al.'s framework [94] decomposes images of mathematical formulas into graphs to make the process of inferring the formula more interpretable. Zhong et al. [104] proposed a method that produces hybrid chains (mixing text and table data) and reason on those with a transformer to provide an answer, in a question answering (QA) setting. For the same task, Deng et al. [18] proposed a method that parses questions into AMR (abstract meaning representation) graphs and reason on those graphs to answer the questions. For the task of fake-news identification, Jin et al. [42] took inspiration from human's information-processing model to make a model that builds claim-evidence graphs to identify fake news. For a similar task but focusing on the propagation network of fake-news, Yang et al. [99] designed a framework that reveals which subgraphs of the news propagation network are the most important in a model's decision process.

#### *4.2.3. Approaches using other forms of symbolic knowledge*

Many papers used various forms of symbolic data, usually specific to their application. Some of them worked with autonomous agents, often trained with reinforcement learning agents. For instance, Sreedharan et al. [83] proposed a method to provide contrastive explanations with user-specified concepts in sequential decision-making settings, by building partial symbolic models of a local approximation of the task. In a similar setting, Jin et al. [41] developed a framework to learn action models and symbolic options with a symbolic planner, using reinforcement learning. Also considering agents trained with reinforcement learning, Finkelstein et al. [25] designed a protocol to apply transformations to the environment model of an autonomous agent in order to produce textual explanations of it. Targeting a broader range of models, Verma et al. [88] proposed a new approach based on query answering to estimate a black-box autonomous agent as an interpretable relational model.

In the medical domain, Jang et al. [39] proposed an approach that extracts symbolic representations from images and rules on these representations to diagnose some diabetes. Another paper in the medical domain, by Liu et al. [60], presented a method for medical diagnosis, that uses a Bayesian Network as well as conditional probability and mutual information matrices to direct an inquiry of the symptoms in order to identify a disease.

In the visual domain, Chen et al. [15] proposed a method that combines deep learning with analogical learning on visual relation detection, using object information and spatial information between objects, so that the relation identification relies on an interpretable algorithm. For dynamic visual reasoning in videos, Ding et al. [21] proposed a method that identifies objects and related physics concepts such as speed, then gives them as input to a physical simulator to predict what will happen next. Hua et al. [36] developed a method that consists of decomposing videos into object-relation chains, which allows both the classification and the production of explanations based on this representation.

Lastly, we observed a few original papers which are either generic or have a domain of focus which is not shared with other papers in this category. For instance, Geiger et al. [26] proposed a training method that allows the alignment of the neural network to a high-level causal model. Dessi et al. [19] presented a protocol to train two deep neural networks with a way of communicating using symbols, which is partially interpretable. Zhang et al.'s framework [103] extracts from a deep Natural Language Processing (NLP) model the interaction between the words that influenced the embedding, and outputs a tree structure. Peng et al. [65] proposed a new framework for symbolic computation that decomposes computations in fundamental transformations, performed by deep models.

### *4.3. Synthesis and Observations*

Overall, across the three categories of symbolic knowledge, several common trends and challenges emerge. First, the diversity of symbolic structures used across the three categories reflects the versatility of symbolic representations in AI, and suggests that the choice of symbolic structure is often driven by the specific application domain. For instance, graphs are more prevalent in natural language processing tasks, while logic rules are more commonly used in image classification and graph processing. Second, a notable observation is that many of the papers in these categories are domain-specific, with a significant number focusing on natural language processing, visual reasoning, and

autonomous agents. This suggests that the development of NeSy methods for trustworthiness is still largely driven by specific application needs, rather than by a unified theoretical framework. Finally, while the majority of papers across all three categories focus on interpretability, the approaches differ in the type of interpretability they provide. Logic-based approaches tend to provide more formal and precise explanations, while graph-based approaches often provide more intuitive and visual explanations. The “other” category, being more diverse, encompasses a wider range of interpretability types, from contrastive explanations to causal models. These observations suggest that, despite the variety of symbolic structures and application domains, there is a common underlying goal of bridging the gap between the high performance of neural models and the transparency of symbolic systems.

## 5. Classification based on the types of interpretability

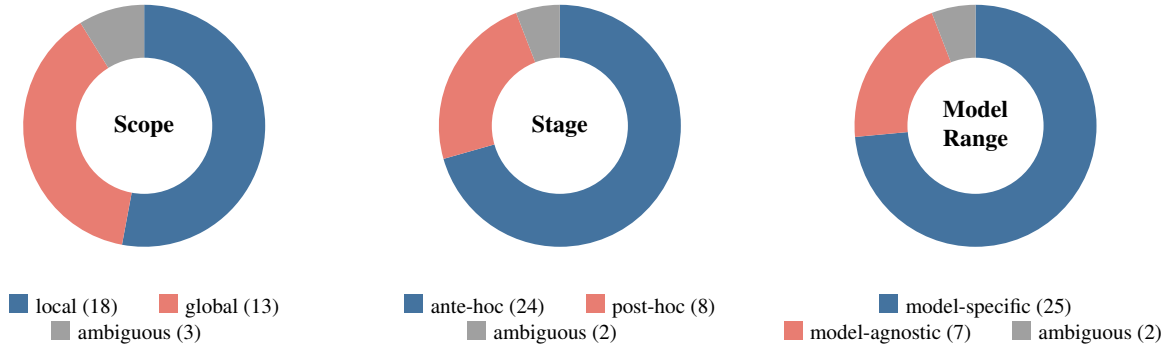


Fig. 3. Distribution of the different types of interpretability contributions.

Since interpretability represents the central theme of the majority of the collected studies, we systematically categorize these works to better analyze prevailing methodological trends. Specifically, the papers are organized along three widely recognized dimensions in interpretability research: (1) the **scope** of explainability, (2) the **stage** at which the interpretability method is applied, and (3) the degree of dependence on the underlying model architecture (**model range**). This taxonomy provides a structured view of how interpretability techniques are designed and deployed across the literature. The resulting distribution of studies across these dimensions is illustrated in Figure 3, while the detailed classification of individual papers is summarized in Table 1. These dimensions are frequently used to analyze papers in this field [9, 80, 82] and the description of those concepts is mentioned in Section 2.3. *In our analysis, we excluded rule-learning methods as they inherently fall into the ante-hoc, model-specific, and usually global scope categories.*

The first dimension of our taxonomy concerns the scope of explainability, which distinguishes between explanations that focus on individual predictions and those that describe the behavior of the model as a whole. Accordingly, interpretability methods in the literature are commonly categorized into local explanations, which explain a specific prediction or instance, and global explanations, which aim to capture the overall decision-making logic of the model. **Local** explanations are specific to a given input, providing insights into why a particular decision was made. In this case, the explanations are different for each individual processing; in other words, running the model several times on different inputs will give different explanations. On the other hand, **global** explanations offer a broader understanding, characterizing the behavior of the entire model. In between, an intermediate “*cohort*” scope is sometimes considered, for methods applicable to a subset of inputs rather than just one<sup>1</sup>. Our review found a relatively balanced distribution: 18 papers (53%) focus on local explanations, 13 papers (38%) on global explanations, and 3 papers (9%) with an ambiguous scope. This shows that both strategies are well studied in the literature.

<sup>1</sup>labeled in the figure as “ambiguous” for consistency with other dimensions

Regarding the **stage** of explanation, methods can be categorized as either *ante-hoc* (also known as *self-explainable*) or *post-hoc*. *Ante-hoc* interpretability contributions are frameworks designed to be inherently explainable, usually by making part of the decision process structured and comprehensible by humans. On the other hand, *post-hoc* methods generate explanations after or in parallel to the decision, often for decisions made by an opaque, black-box model. *Post-hoc* explanations can take various forms, such as a textual justification of a decision or a simplified model that mirrors the original model’s decisions. The fundamental difference of this in contrast with *ante-hoc* interpretability is that the explanation does not necessarily align with the real reasons why a particular output was produced by the framework. This means that it wouldn’t help to uncover potential bias or imperfections in the reasoning process, and might even hide problems in the decision process with justifications *a posteriori*. Therefore, it is questionable how such papers actually contribute to Trustworthiness which was our initial theme. In our study, we found that most of the papers covered provided ante-hoc interpretability contributions (24 ; 71%), and only 8 of them (24%) provided post-hoc interpretability contributions. Our survey found no clear correlation between the scope of explanations and the stage at which the method is applied, indicating a wide range of approaches addressing interpretability in AI systems.

The third dimension (**model range**) concerns whether the interpretability method is *model-agnostic* or *model-specific*. *Model-agnostic* methods can be applied universally across different models, while *model-specific* methods are tailored to a particular model. In this dimension, we observed almost the same distribution than for the previous categorization: 25 papers (74%) proposed a model-specific interpretability contribution, while 7 of them (20%) proposed a model agnostic contribution. In fact, this similarity is not due to chance, since we can see in table 1 that all ante-hoc explainability contributions fall into the model-specific category. This makes sense, because providing ante-hoc interpretability requires designing an inherently explainable method, so they are naturally model-specific. On the contrary, post-hoc explainability methods have the flexibility to be model-agnostic, and we observed that all but one of the methods in the post-hoc category were model-agnostic. An interesting exception we noticed is the work by Seungeon Lee [54], which involved modifying the final layer and training process of a deep model to enhance explainability.

From this classification, we derive several key insights. First, local explanations slightly outnumber global ones, suggesting a focus on instance-level interpretability in neuro-symbolic research. Second, ante-hoc methods dominate, reflecting the field’s emphasis on intrinsic interpretability rather than post-hoc explanations. Finally, the prevalence of model-specific methods indicates that neuro-symbolic interpretability often requires custom design rather than generic tools.

Table 1  
Classification of papers by interpretability dimensions (see Section 5)

|                                   |                                                                                                                     |                                                                      |                                 |
|-----------------------------------|---------------------------------------------------------------------------------------------------------------------|----------------------------------------------------------------------|---------------------------------|
| <b>Scope</b><br>(Local vs Global) | <b>Local:</b> [5, 8, 14, 24, 25, 36, 39, 41, 42, 45, 58, 65, 66, 69, 83, 94, 104, 105]                              | <b>Global:</b> [14, 17, 18, 21, 26, 27, 34, 44, 47, 60, 76, 88, 102] | <b>Ambiguous:</b> [19, 81, 103] |
| <b>Stage</b><br>(Ante vs Post)    | <b>Ante-hoc:</b> [8, 14, 15, 17–19, 21, 27, 36, 39, 41, 42, 44, 45, 54, 58, 60, 65, 66, 76, 94, 100, 102, 104, 105] | <b>Post-hoc:</b> [5, 25, 34, 69, 81, 83, 88, 103]                    | <b>Ambiguous:</b> [26, 47]      |
| <b>Type</b><br>(Spec. vs Agn.)    | <b>Model-specific:</b> [8, 14, 15, 17–19, 21, 27, 36, 39, 41, 42, 44, 45, 58, 60, 65, 66, 76, 81, 94, 100, 102–105] | <b>Model-agnostic:</b> [5, 25, 34, 54, 69, 83, 88]                   | <b>Ambiguous:</b> [26, 103]     |

## 6. Discussion

### 6.1. Lack of applications to Fairness, Privacy and Safety

The primary objective of this review was to examine how neuro-symbolic (NeSy) systems are applied to address different dimensions of AI trustworthiness. As expected, interpretability emerged as the dominant research focus

in the literature. However, the limited number of studies addressing other trustworthiness aspects—particularly fairness and privacy—was notable. Although some neuro-symbolic approaches targeting safety, privacy, or fairness may exist, their absence from the prominent conferences/journals (from where we collected the papers) suggests that these topics currently have limited visibility or adoption within the NeSy research landscape.

This suggests that there may be unexploited potential. In the context of privacy, the potential benefits of incorporating NeSy systems are not immediately apparent. It could be suggested that NeSy may not offer significant advantages for enhancing privacy in AI systems. However, caution is advised before making definitive statements about NeSy’s limitations in this area. Regarding safety, we did find one paper [97], so there seems to be at least some potential. The scarcity of NeSy approaches to safety could be attributed to the fact that safety is often seen as a broad concept with a lot of overlap with other dimensions, such as robustness for instance, and it is this dimension which appears as a clear goal in the publications. Safety, which is the study of worse-case scenarios and respect of critical constraints, encompasses considerations that are very specific to the target applications. Looking for the keyword “safe” in AAAI accepted papers, we observed that it was much less common than what we could find for “interpretab(le)” and “explainab(le)” combined, and that it appeared mostly in publications about reinforcement learning settings. We could hypothesize that neuro-symbolic approaches are less popular in these settings, and that it contributes to the rarity of NeSy approaches to safety.

Considering fairness, there seems to be untapped potential for NeSy integration. In addition to the paper we mentioned previously [101], we found another work by Wang et al. [90], which approached fairness by imposing rule-like constraints during the training process. Although this approach was deemed too narrow to qualify as a comprehensive NeSy integration in our review, it indicates the integration of fairness constraints could be facilitated by NeSy models. This suggests that further investigation into NeSy’s potential to address fairness in AI should be pursued. Moreover, there is a close link between interpretability and fairness, since having more understanding of a system’s decision could help to detect potential biases, as pointed out by Barredo Arrieta et al. [9]. This idea is supported by some works [1, 79] that address both explainability and fairness simultaneously. However, post-hoc explainability introduces a new risk of “fair-washing”, which is to give in appearance fair explanations to decisions that were taken because of biases [2]. Overall, since interpretability and fairness are related and NeSy designs have a high potential for interpretability, we believe that they have a high potential for fairness as well, which definitely requires further research.

## 6.2. *Lack of grounding based on common taxonomies*

Another insight from our review is the wide range of methods encompassed under the NeSy umbrella and the absence of clear categorization for these methods. The term “neuro-symbolic” itself is often not explicitly used in many papers. Although review papers like those by Sarker et al. [72] and Wang et al. [91] proposed conceptual taxonomies for NeSy systems, these classifications are not universally adopted in the literature. This lack of standardized taxonomy makes it challenging to categorize papers without a deep dive into their methodologies. Consequently, there is a need for more consensus in the research community regarding the taxonomy and terminology of NeSy systems. Our survey regrouped papers of different NeSy categories according to Wang et al. [91], and these categories regrouped papers of different application domains, suggesting that several researchers could face the same challenges without awareness of the insights from other fields. However, we could not always be sure of each framework’s category. More clear grounding on taxonomies would facilitate the identification and comparison of works with similar methodologies, regardless of their specific applications.

Regarding the interpretability, there are numerous existing works on the taxonomies (see Section 2), but the papers are too rarely providing clear grounding of their work on those. For instance, how consistent are the explanations with the actual decision-making is not often actually assessed; one usually needs to understand in depth the proposed method to evaluate the explanations consistently, despite this issue being of key importance. Another takeaway from our review is that it is very hard to quantify the degree of explainability and thus compare the different methods. Although some papers provided user studies [54, 60, 83, 99, 102], it is far from a universal practice, and user studies are not standardized. More systematic assessment, and standardised metrics, would make it possible to deduce actionable insights from the comparison of the different methods.

### 6.3. Common Challenges

One of the main challenges faced by interpretability methods is the trade-off between interpretability and performance. The reason for the domination of “black-box” deep models over the state-of-the-art in many domains is that they have shown empirically to be the best performing. While post-hoc explainability approaches keep the neural model untouched and thus do not alter their performance, self-interpretable frameworks are designed specifically for interpretability, thus their performance may not be optimal. However, as it was pointed out in earlier studies [10, 71], a trade-off between interpretability and performance is not systematic. Indeed, we found that a large part of our surveyed papers claimed new state-of-the-art results for their task, and a few others claimed competitive results with SOTA (most of the time neural approaches). Even if some of the papers showed performances below the SOTA or did not provide comparison with non-interpretable approaches, we can definitely say that neuro-symbolic approaches have the potential to be the best performing while providing interpretability, in a lot of domains. Therefore, we strongly recommend the exploration of NeSy strategies.

A common challenge for NeSy approaches in general is to design the interface between neural components and symbolic parts. When the symbolic knowledge is altering the training of the neural components (often through the loss), finding the right integration and the right weight is very challenging. When neural components need to output intermediate symbolic parts, the model cannot be trained end-to-end, and a main challenge is to choose the right symbolic space, as well as to get the model to provide outputs in this space. When the symbolic knowledge is restraining the output space of neural networks or acting directly on the inference process, a main challenge is to ensure good design so that the learning process is not hampered. On top of these difficulties, ensuring interpretability should be considered when designing the framework, since NeSy systems are not always interpretable. This often implies giving a strong enough role to the symbolic structures, especially at the steps leading to the final outputs. The actual explainability provided by the symbolic representations is also to be considered. All of these design challenges are common to methods from different domains using similar methodologies, therefore leveraging insights from designs in other domains is essential when approaching a new task.

Another challenge faced by NeSy systems is that of scalability. Regarding the symbolic structures, this means scaling the symbolic space to increase the potential expressivity and cover more cases, involving richer knowledge. However, more expressive symbolic spaces make the symbolic integration more difficult. For instance, many methods build rules in Propositional Logic, which lacks the expressivity of First-Order Logic. It is often difficult to scale the methods to First-Order Logic, one of the reasons being that it introduces a combinatorial quantity of possible formulas and possible reasoning patterns. For the neural components, the problem of scalability is mainly about data. Indeed, many NeSy approaches require datasets with additional structural information, or datasets of intermediate symbolic representations. It is hard to find such data in large quantities, while neural models often require a lot of data with a diverse enough distribution to be reliable and robust. Future research should explore ways to enhance the scalability of NeSy approaches.

## 7. Conclusion

This research work provides a comprehensive examination of the most recent advancements in neuro-symbolic (NeSy) methods, specifically focusing on their role in enhancing the trustworthiness of AI systems. Our findings reveal that the primary application of current NeSy methods for trustworthiness is centered around improving interpretability. By converging the fields of AI trustworthiness and NeSy integration, this study proposed a new unified analysis of these two intertwined domains.

The papers included in our survey were reviewed based on the symbolic structures they were exploiting. They were also systematically categorized on the basis of the scope, stage, and adaptability of the interpretability methods they developed. A key insight from our study is the recognition of the immense potential NeSy integration holds for interpretability. This potential is not restrained to any specific domain or application, indicating a broad and versatile utility of NeSy approaches.

However, this study also highlights a significant imbalance in the focus of current NeSy research. While a substantial part of this research is dedicated to enhancing interpretability, there is a noticeably smaller portion of works

aimed at improving other aspects of AI trustworthiness, such as security. This observation underscores an opportunity for future research to broaden the scope of NeSy applications, extending its benefits to other critical dimensions of AI trustworthiness, including but not limited to fairness, privacy, and safety. This study also uncovered a lack of grounding on existing taxonomies, and the lack of standardized assessment of interpretability.

In conclusion, this review provides insights on how the emerging research trend of improving trustworthiness via NeSy methods can be analyzed and structured. It should facilitate a comprehensive understanding of the field, and open avenues for future exploration in expanding the application of NeSy methods to address a wider array of trustworthiness concerns in AI systems.

## References

- [1] Ahn, Y., Lin, Y.: Fairsight: Visual analytics for fairness in decision making. *IEEE Trans. Vis. Comput. Graph.* **26**(1), 1086–1095 (2020). , <https://doi.org/10.1109/TVCG.2019.2934262>
- [2] Aivodji, U., Arai, H., Gambs, S., Hara, S.: Characterizing the risk of fairwashing. In: Ranzato, M., Beygelzimer, A., Dauphin, Y., Liang, P., Vaughan, J.W. (eds.) *Advances in Neural Information Processing Systems*. vol. 34, pp. 14822–14834. Curran Associates, Inc. (2021), [https://proceedings.neurips.cc/paper\\_files/paper/2021/file/7caf5e22ea3eb8175ab518429c8589a4-Paper.pdf](https://proceedings.neurips.cc/paper_files/paper/2021/file/7caf5e22ea3eb8175ab518429c8589a4-Paper.pdf)
- [3] Almazrouei, E., Alobeidli, H., Alshamsi, A., Cappelli, A., Cojocaru, R., Debbah, M., Goffinet, E., Heslow, D., Launay, J., Malartic, Q., Noun, B., Pannier, B., Penedo, G.: Falcon-40B: an open large language model with state-of-the-art performance (2023)
- [4] Amodei, D., Olah, C., Steinhardt, J., Christiano, P.F., Schulman, J., Mané, D.: Concrete problems in AI safety. *CoRR* **abs/1606.06565** (2016), <http://arxiv.org/abs/1606.06565>
- [5] An, J., Lai, Y., Han, Y.: Logic rule guided attribution with dynamic ablation. In: Thirty-Sixth AAAI Conference on Artificial Intelligence, AAAI 2022, Thirty-Fourth Conference on Innovative Applications of Artificial Intelligence, IAAI 2022, The Twelveth Symposium on Educational Advances in Artificial Intelligence, EAAI 2022 Virtual Event, February 22 - March 1, 2022. pp. 77–85. AAAI Press (2022), <https://ojs.aaai.org/index.php/AAAI/article/view/19881>
- [6] Angelov, P.P., Soares, E.A., Jiang, R., Arnold, N.I., Atkinson, P.M.: Explainable artificial intelligence: an analytical review. *WIREs Data Mining Knowl. Discov.* **11**(5) (2021). , <https://doi.org/10.1002/widm.1424>
- [7] Bader, S., Hitzler, P.: Dimensions of neural-symbolic integration - A structured survey. In: Artémov, S.N., Barringer, H., d’Avila Garcez, A.S., Lamb, L.C., Woods, J. (eds.) *We Will Show Them! Essays in Honour of Dov Gabbay*, Volume One. pp. 167–194. College Publications (2005)
- [8] Barbiero, P., Ciravegna, G., Giannini, F., Lió, P., Gori, M., Melacci, S.: Entropy-based logic explanations of neural networks. In: Thirty-Sixth AAAI Conference on Artificial Intelligence, AAAI 2022, Thirty-Fourth Conference on Innovative Applications of Artificial Intelligence, IAAI 2022, The Twelveth Symposium on Educational Advances in Artificial Intelligence, EAAI 2022 Virtual Event, February 22 - March 1, 2022. pp. 6046–6054. AAAI Press (2022), <https://ojs.aaai.org/index.php/AAAI/article/view/20551>
- [9] Barredo Arrieta, A., Díaz-Rodríguez, N., Del Ser, J., Bénéttot, A., Tabik, S., Barbado, A., García, S., Gil-Lopez, S., Molina, D., Benjamins, R., Chatila, R., Herrera, F.: Explainable artificial intelligence (xai): Concepts, taxonomies, opportunities and challenges toward responsible ai. *Information Fusion* **58**, 82–115 (Jun 2020). , <http://dx.doi.org/10.1016/j.inffus.2019.12.012>
- [10] Bell, A., Solano-Kamaiko, I., Nov, O., Stoyanovich, J.: It’s just not that simple: An empirical study of the accuracy-explainability trade-off in machine learning for public policy. In: *FAccT ’22: 2022 ACM Conference on Fairness, Accountability, and Transparency*, Seoul, Republic of Korea, June 21 - 24, 2022. pp. 248–266. ACM (2022). , <https://doi.org/10.1145/3531146.3533090>
- [11] Belle, V.: Symbolic logic meets machine learning: A brief survey in infinite domains. *CoRR* **abs/2006.08480** (2020), <https://arxiv.org/abs/2006.08480>
- [12] Berlot-Attwell, I.: Neuro-symbolic VQA: A review from the perspective of AGI desiderata. *CoRR* **abs/2104.06365** (2021), <https://arxiv.org/abs/2104.06365>
- [13] Besold, T.R., d’Avila Garcez, A.S., Bader, S., Bowman, H., Domingos, P.M., Hitzler, P., Kühnberger, K., Lamb, L.C., Lowd, D., Lima, P.M.V., de Penning, L., Pinkas, G., Poon, H., Zaverucha, G.: Neural-symbolic learning and reasoning: A survey and interpretation. *CoRR* **abs/1711.03902** (2017), <http://arxiv.org/abs/1711.03902>
- [14] Chen, J., Bao, Q., Sun, C., Zhang, X., Chen, J., Zhou, H., Xiao, Y., Li, L.: LOREN: logic-regularized reasoning for interpretable fact verification. In: Thirty-Sixth AAAI Conference on Artificial Intelligence, AAAI 2022, Thirty-Fourth Conference on Innovative Applications of Artificial Intelligence, IAAI 2022, The Twelveth Symposium on Educational Advances in Artificial Intelligence, EAAI 2022 Virtual Event, February 22 - March 1, 2022. pp. 10482–10491. AAAI Press (2022), <https://ojs.aaai.org/index.php/AAAI/article/view/21291>
- [15] Chen, K., Forbus, K.D.: Visual relation detection using hybrid analogical learning. In: Thirty-Fifth AAAI Conference on Artificial Intelligence, AAAI 2021, Thirty-Third Conference on Innovative Applications of Artificial Intelligence, IAAI 2021, The Eleventh Symposium on Educational Advances in Artificial Intelligence, EAAI 2021, Virtual Event, February 2-9, 2021. pp. 801–808. AAAI Press (2021), <https://ojs.aaai.org/index.php/AAAI/article/view/16162>
- [16] Chen, K., Forbus, K.D.: Visual relation detection using hybrid analogical learning. In: Thirty-Fifth AAAI Conference on Artificial Intelligence, AAAI 2021, Thirty-Third Conference on Innovative Applications of Artificial Intelligence, IAAI 2021, The Eleventh Symposium on Educational Advances in Artificial Intelligence, EAAI 2021, Virtual Event, February 2-9, 2021. pp. 801–808. AAAI Press (2021), <https://ojs.aaai.org/index.php/AAAI/article/view/16162>

- [17] Cucala, D.J.T., Grau, B.C., Kostylev, E.V., Motik, B.: Explainable gnn-based models over knowledge graphs. In: The Tenth International Conference on Learning Representations, ICLR 2022, Virtual Event, April 25–29, 2022. OpenReview.net (2022), <https://openreview.net/forum?id=CrCvGNHAIrz>
- [18] Deng, Z., Zhu, Y., Chen, Y., Witbrock, M., Riddle, P.: Interpretable amr-based question decomposition for multi-hop question answering. In: Raedt, L.D. (ed.) Proceedings of the Thirty-First International Joint Conference on Artificial Intelligence, IJCAI 2022, Vienna, Austria, 23–29 July 2022. pp. 4093–4099. [ijcai.org \(2022\)](https://doi.org/10.24963/ijcai.2022/568), <https://doi.org/10.24963/ijcai.2022/568>
- [19] Dessì, R., Kharitonov, E., Baroni, M.: Interpretable agent communication from scratch (with a generic visual processor emerging on the side). In: Ranzato, M., Beygelzimer, A., Dauphin, Y.N., Liang, P., Vaughan, J.W. (eds.) Advances in Neural Information Processing Systems 34: Annual Conference on Neural Information Processing Systems 2021, NeurIPS 2021, December 6–14, 2021, virtual. pp. 26937–26949 (2021), <https://proceedings.neurips.cc/paper/2021/hash/e250c59336b505ed411d455abaa30b4d-Abstract.html>
- [20] Dhaoui, A., Bertoncello, A., Gourvénec, S., Garnier, J., Pennec, E.L.: Causal and interpretable rules for time series analysis. In: Zhu, F., Ooi, B.C., Miao, C. (eds.) KDD '21: The 27th ACM SIGKDD Conference on Knowledge Discovery and Data Mining, Virtual Event, Singapore, August 14–18, 2021. pp. 2764–2772. ACM (2021), <https://doi.org/10.1145/3447548.3467161>
- [21] Ding, M., Chen, Z., Du, T., Luo, P., Tenenbaum, J., Gan, C.: Dynamic visual reasoning by learning differentiable physics models from video and language. In: Ranzato, M., Beygelzimer, A., Dauphin, Y.N., Liang, P., Vaughan, J.W. (eds.) Advances in Neural Information Processing Systems 34: Annual Conference on Neural Information Processing Systems 2021, NeurIPS 2021, December 6–14, 2021, virtual. pp. 887–899 (2021), <https://proceedings.neurips.cc/paper/2021/hash/07845cd9aefaf6cde3f8926d25138a3a2-Abstract.html>
- [22] Doshi-Velez, F., Kim, B.: Towards a rigorous science of interpretable machine learning (2017)
- [23] Ferrario, A., Loi, M.: How explainability contributes to trust in AI. In: FAccT '22: 2022 ACM Conference on Fairness, Accountability, and Transparency, Seoul, Republic of Korea, June 21 – 24, 2022. pp. 1457–1466. ACM (2022), <https://doi.org/10.1145/3531146.3533202>
- [24] Ferrer-Mestres, J., Dietterich, T.G., Buffet, O., Chades, I.: K-n-momdps: Towards interpretable solutions for adaptive management. In: Thirty-Fifth AAAI Conference on Artificial Intelligence, AAAI 2021, Thirty-Third Conference on Innovative Applications of Artificial Intelligence, IAAI 2021, The Eleventh Symposium on Educational Advances in Artificial Intelligence, EAAI 2021, Virtual Event, February 2–9, 2021. pp. 14775–14784. AAAI Press (2021), <https://ojs.aaai.org/index.php/AAAI/article/view/17735>
- [25] Finkelstein, M., Schlot, N.L., Liu, L., Kolumbus, Y., Parkes, D.C., Rosenschein, J.S., Keren, S.: Explainable reinforcement learning via model transforms. In: NeurIPS (2022), [http://papers.nips.cc/paper\\_files/paper/2022/hash/dbef234be68d8b170240511639610fd1-Abstract-Conference.html](http://papers.nips.cc/paper_files/paper/2022/hash/dbef234be68d8b170240511639610fd1-Abstract-Conference.html)
- [26] Geiger, A., Wu, Z., Lu, H., Rozner, J., Kreiss, E., Icard, T., Goodman, N.D., Potts, C.: Inducing causal structure for interpretable neural networks. In: Chaudhuri, K., Jegelka, S., Song, L., Szepesvári, C., Niu, G., Sabato, S. (eds.) International Conference on Machine Learning, ICML 2022, 17–23 July 2022, Baltimore, Maryland, USA. Proceedings of Machine Learning Research, vol. 162, pp. 7324–7338. PMLR (2022), <https://proceedings.mlr.press/v162/geiger22a.html>
- [27] Georgiev, D., Barbiero, P., Kazhdan, D., Velickovic, P., Lió, P.: Algorithmic concept-based explainable reasoning. In: Thirty-Sixth AAAI Conference on Artificial Intelligence, AAAI 2022, Thirty-Fourth Conference on Innovative Applications of Artificial Intelligence, IAAI 2022, The Twelfth Symposium on Educational Advances in Artificial Intelligence, EAAI 2022 Virtual Event, February 22 – March 1, 2022. pp. 6685–6693. AAAI Press (2022), <https://ojs.aaai.org/index.php/AAAI/article/view/20623>
- [28] Gilpin, L.H., Bau, D., Yuan, B.Z., Bajwa, A., Specter, M.A., Kagal, L.: Explaining explanations: An overview of interpretability of machine learning. In: Bonchi, F., Provost, F.J., Eliassi-Rad, T., Wang, W., Cattuto, C., Ghani, R. (eds.) 5th IEEE International Conference on Data Science and Advanced Analytics, DSAA 2018, Turin, Italy, October 1–3, 2018. pp. 80–89. IEEE (2018), <https://doi.org/10.1109/DSAA.2018.00018>
- [29] Glanois, C., Jiang, Z., Feng, X., Weng, P., Zimmer, M., Li, D., Liu, W., Hao, J.: Neuro-symbolic hierarchical rule induction. In: Chaudhuri, K., Jegelka, S., Song, L., Szepesvári, C., Niu, G., Sabato, S. (eds.) International Conference on Machine Learning, ICML 2022, 17–23 July 2022, Baltimore, Maryland, USA. Proceedings of Machine Learning Research, vol. 162, pp. 7583–7615. PMLR (2022), <https://proceedings.mlr.press/v162/glanois22a.html>
- [30] Glauer, M., West, R., Michie, S., Hastings, J.: Esc-rules: Explainable, semantically constrained rule sets. In: d’Avila Garcez, A.S., Jiménez-Ruiz, E. (eds.) Proceedings of the 16th International Workshop on Neural-Symbolic Learning and Reasoning as part of the 2nd International Joint Conference on Learning & Reasoning (IJCLR 2022), Cumberland Lodge, Windsor Great Park, UK, September 28–30, 2022. CEUR Workshop Proceedings, vol. 3212, pp. 94–103. CEUR-WS.org (2022), <https://ceur-ws.org/Vol-3212/paper7.pdf>
- [31] Hamilton, K., Nayak, A., Bozic, B., Longo, L.: Is neuro-symbolic AI meeting its promise in natural language processing? A structured review. CoRR **abs/2202.12205** (2022), <https://arxiv.org/abs/2202.12205>
- [32] Hao, S., Gu, Y., Ma, H., Hong, J.J., Wang, Z., Wang, D.Z., Hu, Z.: Reasoning with language model is planning with world model. CoRR **abs/2305.14992** (2023), <https://doi.org/10.48550/arXiv.2305.14992>
- [33] Heist, N., Paulheim, H.: The caligraph ontology as a challenge for OWL reasoners. In: Singh, G., Mutharaju, R., Kapanipathi, P. (eds.) Proceedings of the Semantic Reasoning Evaluation Challenge (SemREC 2021) co-located with the 20th International Semantic Web Conference (ISWC 2021), Virtual Event, October 27th, 2021. CEUR Workshop Proceedings, vol. 3123, pp. 21–31. CEUR-WS.org (2021), <https://ceur-ws.org/Vol-3123/paper3.pdf>
- [34] Himmelhuber, A., Zillner, S., Grimm, S., Ringsquandl, M., Joblin, M., Runkler, T.A.: A new concept for explaining graph neural networks. In: d’Avila Garcez, A.S., Jiménez-Ruiz, E. (eds.) Proceedings of the 15th International Workshop on Neural-Symbolic Learning and Reasoning as part of the 1st International Joint Conference on Learning & Reasoning (IJCLR 2021), Virtual conference, October 25–27, 2021. CEUR Workshop Proceedings, vol. 2986, pp. 1–5. CEUR-WS.org (2021), <https://ceur-ws.org/Vol-2986/paper1.pdf>
- [35] Hitzler, P., Eberhart, A., Ebrahimi, M., Sarker, M.K., Zhou, L.: Neuro-symbolic approaches in artificial intelligence. National Science Review **9**(6) (03 2022), <https://doi.org/10.1093/nsr/nwac035>, nwac035

- [36] Hua, H., Li, D., Li, R., Zhang, P., Renz, J., Cohn, A.G.: Towards explainable action recognition by salient qualitative spatial object relation chains. In: Thirty-Sixth AAAI Conference on Artificial Intelligence, AAAI 2022, Thirty-Fourth Conference on Innovative Applications of Artificial Intelligence, IAAI 2022, The Twelveth Symposium on Educational Advances in Artificial Intelligence, EAAI 2022 Virtual Event, February 22 - March 1, 2022. pp. 5710–5718. AAAI Press (2022), <https://ojs.aaai.org/index.php/AAAI/article/view/20513>
- [37] Huang, J., Chang, K.C.: Towards reasoning in large language models: A survey. CoRR **abs/2212.10403** (2022). , <https://doi.org/10.48550/arXiv.2212.10403>
- [38] Ignatiev, A., Marques-Silva, J., Narodytska, N., Stuckey, P.J.: Reasoning-based learning of interpretable ML models. In: Zhou, Z. (ed.) Proceedings of the Thirtieth International Joint Conference on Artificial Intelligence, IJCAI 2021, Virtual Event / Montreal, Canada, 19-27 August 2021. pp. 4458–4465. ijcai.org (2021). , <https://doi.org/10.24963/ijcai.2021/608>
- [39] Jang, S., Girard, M.J.A., Thiéry, A.H.: Explainable diabetic retinopathy classification based on neural-symbolic learning. In: d’Avila Garcez, A.S., Jiménez-Ruiz, E. (eds.) Proceedings of the 15th International Workshop on Neural-Symbolic Learning and Reasoning as part of the 1st International Joint Conference on Learning & Reasoning (IJCLR 2021), Virtual conference, October 25-27, 2021. CEUR Workshop Proceedings, vol. 2986, pp. 104–114. CEUR-WS.org (2021), <https://ceur-ws.org/Vol-2986/paper8.pdf>
- [40] Jesus, S., Belém, C., Balayan, V., Bento, J., Saleiro, P., Bizarro, P., Gama, J.: How can i choose an explainer? an application-grounded evaluation of post-hoc explanations (2021). , <https://arxiv.org/abs/2101.08758>
- [41] Jin, M., Ma, Z., Jin, K., Zhuo, H.H., Chen, C., Yu, C.: Creativity of AI: automatic symbolic option discovery for facilitating deep reinforcement learning. In: Thirty-Sixth AAAI Conference on Artificial Intelligence, AAAI 2022, Thirty-Fourth Conference on Innovative Applications of Artificial Intelligence, IAAI 2022, The Twelveth Symposium on Educational Advances in Artificial Intelligence, EAAI 2022 Virtual Event, February 22 - March 1, 2022. pp. 7042–7050. AAAI Press (2022), <https://ojs.aaai.org/index.php/AAAI/article/view/20663>
- [42] Jin, Y., Wang, X., Yang, R., Sun, Y., Wang, W., Liao, H., Xie, X.: Towards fine-grained reasoning for fake news detection. In: Thirty-Sixth AAAI Conference on Artificial Intelligence, AAAI 2022, Thirty-Fourth Conference on Innovative Applications of Artificial Intelligence, IAAI 2022, The Twelveth Symposium on Educational Advances in Artificial Intelligence, EAAI 2022 Virtual Event, February 22 - March 1, 2022. pp. 5746–5754. AAAI Press (2022), <https://ojs.aaai.org/index.php/AAAI/article/view/20517>
- [43] Johnson, T.L., Johnson, N.N., McCurdy, D., Olajide, M.S.: Facial recognition systems in policing and racial disparities in arrests. Government Information Quarterly **39**(4), 101753 (2022). , <https://www.sciencedirect.com/science/article/pii/S0740624X22000892>
- [44] Kakadiya, A., Natarajan, S., Ravindran, B.: Relational boosted bandits. In: Thirty-Fifth AAAI Conference on Artificial Intelligence, AAAI 2021, Thirty-Third Conference on Innovative Applications of Artificial Intelligence, IAAI 2021, The Eleventh Symposium on Educational Advances in Artificial Intelligence, EAAI 2021, Virtual Event, February 2-9, 2021. pp. 12123–12130. AAAI Press (2021), <https://ojs.aaai.org/index.php/AAAI/article/view/17439>
- [45] Kalyanpur, A., Breloff, T., Ferrucci, D.A.: Braid: Weaving symbolic and neural knowledge into coherent logical explanations. In: Thirty-Sixth AAAI Conference on Artificial Intelligence, AAAI 2022, Thirty-Fourth Conference on Innovative Applications of Artificial Intelligence, IAAI 2022, The Twelveth Symposium on Educational Advances in Artificial Intelligence, EAAI 2022 Virtual Event, February 22 - March 1, 2022. pp. 10867–10874. AAAI Press (2022), <https://ojs.aaai.org/index.php/AAAI/article/view/21333>
- [46] Kambhampati, S., Sreedharan, S., Verma, M., Zha, Y., Guan, L.: Symbols as a lingua franca for bridging human-ai chasm for explainable and advisable AI systems. In: Thirty-Sixth AAAI Conference on Artificial Intelligence, AAAI 2022, Thirty-Fourth Conference on Innovative Applications of Artificial Intelligence, IAAI 2022, The Twelveth Symposium on Educational Advances in Artificial Intelligence, EAAI 2022 Virtual Event, February 22 - March 1, 2022. pp. 12262–12267. AAAI Press (2022), <https://ojs.aaai.org/index.php/AAAI/article/view/21488>
- [47] Kasioumis, T., Townsend, J., Inakoshi, H.: Elite backprop: Training sparse interpretable neurons. In: d’Avila Garcez, A.S., Jiménez-Ruiz, E. (eds.) Proceedings of the 15th International Workshop on Neural-Symbolic Learning and Reasoning as part of the 1st International Joint Conference on Learning & Reasoning (IJCLR 2021), Virtual conference, October 25-27, 2021. CEUR Workshop Proceedings, vol. 2986, pp. 82–93. CEUR-WS.org (2021), <https://ceur-ws.org/Vol-2986/paper6.pdf>
- [48] Kasirzadeh, A.: Reasons, values, stakeholders: A philosophical framework for explainable artificial intelligence. In: Elish, M.C., Isaac, W., Zemel, R.S. (eds.) FAccT ’21: 2021 ACM Conference on Fairness, Accountability, and Transparency, Virtual Event / Toronto, Canada, March 3-10, 2021. p. 14. ACM (2021). , <https://doi.org/10.1145/3442188.3445866>
- [49] Kazemi, S.M., Kim, N., Bhatia, D., Xu, X., Ramachandran, D.: LAMBADA: backward chaining for automated reasoning in natural language. CoRR **abs/2212.13894** (2022). , <https://doi.org/10.48550/arXiv.2212.13894>
- [50] Kusters, R., Kim, Y., Collery, M., de Sainte Marie, C., Gupta, S.: Differentiable rule induction with learned relational features. In: d’Avila Garcez, A.S., Jiménez-Ruiz, E. (eds.) Proceedings of the 16th International Workshop on Neural-Symbolic Learning and Reasoning as part of the 2nd International Joint Conference on Learning & Reasoning (IJCLR 2022), Cumberland Lodge, Windsor Great Park, UK, September 28-30, 2022. CEUR Workshop Proceedings, vol. 3212, pp. 30–44. CEUR-WS.org (2022), <https://ceur-ws.org/Vol-3212/paper3.pdf>
- [51] Lamb, L.C., d’Avila Garcez, A.S., Gori, M., Prates, M.O.R., Avelar, P.H.C., Vardi, M.Y.: Graph neural networks meet neural-symbolic computing: A survey and perspective. In: Bessiere, C. (ed.) Proceedings of the Twenty-Ninth International Joint Conference on Artificial Intelligence, IJCAI 2020. pp. 4877–4884. ijcai.org (2020). , <https://doi.org/10.24963/ijcai.2020/679>
- [52] Landajuela, M., Petersen, B.K., Kim, S., Santiago, C.P., Glatt, R., Mundhenk, T.N., Pettit, J.F., Faissol, D.M.: Discovering symbolic policies with deep reinforcement learning. In: Meila, M., Zhang, T. (eds.) Proceedings of the 38th International Conference on Machine Learning, ICML 2021, 18-24 July 2021, Virtual Event. Proceedings of Machine Learning Research, vol. 139, pp. 5979–5989. PMLR (2021), <http://proceedings.mlr.press/v139/landajuela21a.html>

- [53] Langer, M., Oster, D., Speith, T., Hermanns, H., Kästner, L., Schmidt, E., Sesing, A., Baum, K.: What do we want from explainable artificial intelligence (xai)? – a stakeholder perspective on xai and a conceptual model guiding interdisciplinary xai research. *Artificial Intelligence* **296**, 103473 (Jul 2021). , <http://dx.doi.org/10.1016/j.artint.2021.103473>
- [54] Lee, S., Wang, X., Han, S., Yi, X., Xie, X., Cha, M.: Self-explaining deep models with logic rule reasoning. In: *NeurIPS* (2022), [http://papers.nips.cc/paper\\_files/paper/2022/hash/1548d98b62d3a4382a31ba77d89186cd-Abstract-Conference.html](http://papers.nips.cc/paper_files/paper/2022/hash/1548d98b62d3a4382a31ba77d89186cd-Abstract-Conference.html)
- [55] Li, B., Qi, P., Liu, B., Di, S., Liu, J., Pei, J., Yi, J., Zhou, B.: Trustworthy AI: from principles to practices. *CoRR* **abs/2110.01167** (2021), <https://arxiv.org/abs/2110.01167>
- [56] Li, S., Feng, M., Wang, L., Essofi, A., Cao, Y., Yan, J., Song, L.: Explaining point processes by learning interpretable temporal logic rules. In: *The Tenth International Conference on Learning Representations, ICLR 2022, Virtual Event, April 25-29, 2022*. OpenReview.net (2022), <https://openreview.net/forum?id=P07dq7iSAGr>
- [57] Lipton, Z.C.: The mythos of model interpretability. *Commun. ACM* **61**(10), 36–43 (2018). , <https://doi.org/10.1145/3233231>
- [58] Liu, D., Lian, J., Liu, Z., Wang, X., Sun, G., Xie, X.: Reinforced anchor knowledge graph generation for news recommendation reasoning. In: *Zhu, F., Ooi, B.C., Miao, C. (eds.) KDD '21: The 27th ACM SIGKDD Conference on Knowledge Discovery and Data Mining, Virtual Event, Singapore, August 14-18, 2021*. pp. 1055–1065. ACM (2021). , <https://doi.org/10.1145/3447548.3467315>
- [59] Liu, H., Wang, Y., Fan, W., Liu, X., Li, Y., Jain, S., Liu, Y., Jain, A., Tang, J.: Trustworthy ai: A computational perspective. *ACM Trans. Intell. Syst. Technol.* **14**(1) (nov 2022). , <https://doi.org/10.1145/3546872>
- [60] Liu, W., Cheng, Y., Wang, H., Tang, J., Liu, Y., Zhao, R., Li, W., Zheng, Y., Liang, X.: "my nose is running." "are you also coughing?": Building A medical diagnosis agent with interpretable inquiry logics. In: *Raedt, L.D. (ed.) Proceedings of the Thirty-First International Joint Conference on Artificial Intelligence, IJCAI 2022, Vienna, Austria, 23-29 July 2022*. pp. 4266–4272. ijcai.org (2022). , <https://doi.org/10.24963/ijcai.2022/592>
- [61] Liu, X., Lei, W., Lv, J., Zhou, J.: Abstract rule learning for paraphrase generation. In: *Raedt, L.D. (ed.) Proceedings of the Thirty-First International Joint Conference on Artificial Intelligence, IJCAI 2022, Vienna, Austria, 23-29 July 2022*. pp. 4273–4279. ijcai.org (2022). , <https://doi.org/10.24963/ijcai.2022/593>
- [62] McCulloch, W.S., Pitts, W.: A logical calculus of the ideas immanent in nervous activity. *The Bulletin of Mathematical Biophysics* **5**(4), 115–133 (Dec 1943). , <https://doi.org/10.1007/bf02478259>
- [63] Paleja, R.R., Ghuy, M., Arachchige, N.R., Jensen, R., Gombolay, M.C.: The utility of explainable AI in ad hoc human-machine teaming. In: *Ranzato, M., Beygelzimer, A., Dauphin, Y.N., Liang, P., Vaughan, J.W. (eds.) Advances in Neural Information Processing Systems 34: Annual Conference on Neural Information Processing Systems 2021, NeurIPS 2021, December 6-14, 2021, virtual*. pp. 610–623 (2021), <https://proceedings.neurips.cc/paper/2021/hash/05d74c48b5b30514d8e9bd60320fc8f6-Abstract.html>
- [64] Papernot, N.: What does it mean for ml to be trustworthy? *ICML Workshop on Participatory Approaches to Machine Learning* (2020), <https://youtu.be/UpGgIqLhaqo>
- [65] Peng, S., Fu, D., Cao, Y., Liang, Y., Xu, G., Gao, L., Tang, Z.: Compute like humans: Interpretable step-by-step symbolic computation with deep neural network. In: *Zhang, A., Rangwala, H. (eds.) KDD '22: The 28th ACM SIGKDD Conference on Knowledge Discovery and Data Mining, Washington, DC, USA, August 14 - 18, 2022*. pp. 1348–1357. ACM (2022). , <https://doi.org/10.1145/3534678.3539276>
- [66] Peng, X., Riedl, M.O., Ammanabrolu, P.: Inherently explainable reinforcement learning in natural language. In: *NeurIPS* (2022), [http://papers.nips.cc/paper\\_files/paper/2022/hash/672e44a114a41d5f34b97459877c083d-Abstract-Conference.html](http://papers.nips.cc/paper_files/paper/2022/hash/672e44a114a41d5f34b97459877c083d-Abstract-Conference.html)
- [67] Qu, M., Chen, J., Xhonneux, L.A.C., Bengio, Y., Tang, J.: Rnnlogic: Learning logic rules for reasoning on knowledge graphs. In: *9th International Conference on Learning Representations, ICLR 2021, Virtual Event, Austria, May 3-7, 2021*. OpenReview.net (2021), <https://openreview.net/forum?id=tGZu6DlbreV>
- [68] Raedt, L.D., Manhaeve, R., Dumancic, S., Demeester, T., Kimmig, A.: Neuro-symbolic = neural + logical + probabilistic. In: *Doran, D., d'Avila Garcez, A.S., Lécué, F. (eds.) Proceedings of the 2019 International Workshop on Neural-Symbolic Learning and Reasoning (NeSy 2019), Annual workshop of the Neural-Symbolic Learning and Reasoning Association, Macao, China, August 12, 2019* (2019)
- [69] Rajapaksha, D., Bergmeir, C.: LIMREF: local interpretable model agnostic rule-based explanations for forecasting, with an application to electricity smart meter data. In: *Thirty-Sixth AAAI Conference on Artificial Intelligence, AAAI 2022, Thirty-Fourth Conference on Innovative Applications of Artificial Intelligence, IAAI 2022, The Twelveth Symposium on Educational Advances in Artificial Intelligence, EAAI 2022 Virtual Event, February 22 - March 1, 2022*. pp. 12098–12107. AAAI Press (2022), <https://ojs.aaai.org/index.php/AAAI/article/view/21469>
- [70] Ribeiro, M.T., Singh, S., Guestrin, C.: "why should I trust you?": Explaining the predictions of any classifier. In: *Krishnapuram, B., Shah, M., Smola, A.J., Aggarwal, C.C., Shen, D., Rastogi, R. (eds.) Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, San Francisco, CA, USA, August 13-17, 2016*. pp. 1135–1144. ACM (2016). , <https://doi.org/10.1145/2939672.2939778>
- [71] Rudin, C.: Stop explaining black box machine learning models for high stakes decisions and use interpretable models instead. *Nat. Mach. Intell.* **1**(5), 206–215 (2019). , <https://doi.org/10.1038/s42256-019-0048-x>
- [72] Sarker, M.K., Zhou, L., Eberhart, A., Hitzler, P.: Neuro-symbolic artificial intelligence: Current trends. *CoRR* **abs/2105.05330** (2021), <https://arxiv.org/abs/2105.05330>
- [73] Schoeffer, J., Kuehl, N., Machowski, Y.: "there is not enough information": On the effects of explanations on perceptions of informational fairness and trustworthiness in automated decision-making. In: *2022 ACM Conference on Fairness, Accountability, and Transparency*. ACM (Jun 2022). , <https://doi.org/10.1145/3531146.3533218>

- [74] Sen, P., de Carvalho, B.W.S.R., Riegel, R., Gray, A.G.: Neuro-symbolic inductive logic programming with logical neural networks. In: Thirty-Sixth AAAI Conference on Artificial Intelligence, AAAI 2022, Thirty-Fourth Conference on Innovative Applications of Artificial Intelligence, IAAI 2022, The Twelveth Symposium on Educational Advances in Artificial Intelligence, EAAI 2022 Virtual Event, February 22 - March 1, 2022. pp. 8212–8219. AAAI Press (2022), <https://ojs.aaai.org/index.php/AAAI/article/view/20795>
- [75] Serban, A., van der Blom, K., Hoos, H.H., Visser, J.: Practices for engineering trustworthy machine learning applications. In: 1st IEEE/ACM Workshop on AI Engineering - Software Engineering for AI, WAIN@ICSE 2021, Madrid, Spain, May 30-31, 2021. pp. 97–100. IEEE (2021), <https://doi.org/10.1109/WAIN52551.2021.00021>
- [76] Sharan, S.P., Zheng, W., Hsu, K., Xing, J., Chen, A., Wang, Z.: Symbolic distillation for learned TCP congestion control. In: NeurIPS (2022), [http://papers.nips.cc/paper\\_files/paper/2022/hash/4574ac9854d4defe3bf119d07b817084-Abstract-Conference.html](http://papers.nips.cc/paper_files/paper/2022/hash/4574ac9854d4defe3bf119d07b817084-Abstract-Conference.html)
- [77] Shi, S., Chen, H., Ma, W., Mao, J., Zhang, M., Zhang, Y.: Neural logic reasoning. In: d'Aquin, M., Dietze, S., Hauff, C., Curry, E., Cudré-Mauroux, P. (eds.) CIKM '20: The 29th ACM International Conference on Information and Knowledge Management, Virtual Event, Ireland, October 19-23, 2020. pp. 1365–1374. ACM (2020), <https://doi.org/10.1145/3340531.3411949>
- [78] Shvo, M., Li, A.C., Icarte, R.T., McIlraith, S.A.: Interpretable sequence classification via discrete optimization. In: Thirty-Fifth AAAI Conference on Artificial Intelligence, AAAI 2021, Thirty-Third Conference on Innovative Applications of Artificial Intelligence, IAAI 2021, The Eleventh Symposium on Educational Advances in Artificial Intelligence, EAAI 2021, Virtual Event, February 2-9, 2021. pp. 9647–9656. AAAI Press (2021), <https://ojs.aaai.org/index.php/AAAI/article/view/17161>
- [79] Soares, E., Angelov, P.: Fair-by-design explainable models for prediction of recidivism (2019), <https://arxiv.org/abs/1910.02043>
- [80] Sokol, K., Flach, P.A.: Explainability fact sheets: A framework for systematic assessment of explainable approaches. CoRR **abs/1912.05100** (2019), <http://arxiv.org/abs/1912.05100>
- [81] de Sousa Ribeiro, M., Leite, J.: Aligning artificial neural networks and ontologies towards explainable AI. In: Thirty-Fifth AAAI Conference on Artificial Intelligence, AAAI 2021, Thirty-Third Conference on Innovative Applications of Artificial Intelligence, IAAI 2021, The Eleventh Symposium on Educational Advances in Artificial Intelligence, EAAI 2021, Virtual Event, February 2-9, 2021. pp. 4932–4940. AAAI Press (2021), <https://ojs.aaai.org/index.php/AAAI/article/view/16626>
- [82] Speith, T.: A review of taxonomies of explainable artificial intelligence (XAI) methods. In: 2022 ACM Conference on Fairness, Accountability, and Transparency. ACM (Jun 2022), <https://doi.org/10.1145/3531146.3534639>
- [83] Sreedharan, S., Soni, U., Verma, M., Srivastava, S., Kambhampati, S.: Bridging the gap: Providing post-hoc symbolic explanations for sequential decision-making problems with inscrutable representations. In: The Tenth International Conference on Learning Representations, ICLR 2022, Virtual Event, April 25-29, 2022. OpenReview.net (2022), <https://openreview.net/forum?id=o-1v9hdSult>
- [84] Stein, G.J.: Generating high-quality explanations for navigation in partially-revealed environments. In: Ranzato, M., Beygelzimer, A., Dauphin, Y.N., Liang, P., Vaughan, J.W. (eds.) Advances in Neural Information Processing Systems 34: Annual Conference on Neural Information Processing Systems 2021, NeurIPS 2021, December 6-14, 2021, virtual. pp. 17493–17506 (2021), <https://proceedings.neurips.cc/paper/2021/hash/926ec030f29f83ce5318754fdb631a33-Abstract.html>
- [85] Susskind, Z., Arden, B., John, L.K., Stockton, P., John, E.B.: Neuro-symbolic AI: an emerging class of AI workloads and their characterization. CoRR **abs/2109.06133** (2021), <https://arxiv.org/abs/2109.06133>
- [86] Topin, N., Milani, S., Fang, F., Veloso, M.: Iterative bounding mdps: Learning interpretable policies via non-interpretable methods. In: Thirty-Fifth AAAI Conference on Artificial Intelligence, AAAI 2021, Thirty-Third Conference on Innovative Applications of Artificial Intelligence, IAAI 2021, The Eleventh Symposium on Educational Advances in Artificial Intelligence, EAAI 2021, Virtual Event, February 2-9, 2021. pp. 9923–9931. AAAI Press (2021), <https://ojs.aaai.org/index.php/AAAI/article/view/17192>
- [87] Varshney, K.R.: Trustworthy machine learning and artificial intelligence. XRDS **25**(3), 26–29 (2019), <https://doi.org/10.1145/3313109>
- [88] Verma, P., Marpally, S.R., Srivastava, S.: Asking the right questions: Learning interpretable action models through query answering. In: Thirty-Fifth AAAI Conference on Artificial Intelligence, AAAI 2021, Thirty-Third Conference on Innovative Applications of Artificial Intelligence, IAAI 2021, The Eleventh Symposium on Educational Advances in Artificial Intelligence, EAAI 2021, Virtual Event, February 2-9, 2021. pp. 12024–12033. AAAI Press (2021), <https://ojs.aaai.org/index.php/AAAI/article/view/17428>
- [89] Vilone, G., Longo, L.: Classification of explainable artificial intelligence methods through their output formats. Mach. Learn. Knowl. Extr. **3**(3), 615–661 (2021), <https://doi.org/10.3390/make3030032>
- [90] Wang, N., Nie, S., Wang, Q., Wang, Y., Sanjabi, M., Liu, J., Firooz, H., Wang, H.: COFFEE: counterfactual fairness for personalized text generation in explainable recommendation. CoRR **abs/2210.15500** (2022), <https://doi.org/10.48550/arXiv.2210.15500>
- [91] Wang, W., Yang, Y., Wu, F.: Towards data-and knowledge-driven artificial intelligence: A survey on neuro-symbolic computing (2022), <https://arxiv.org/abs/2210.15889>
- [92] Wang, Z., Zhang, W., Liu, N., Wang, J.: Scalable rule-based representation learning for interpretable classification. In: Ranzato, M., Beygelzimer, A., Dauphin, Y.N., Liang, P., Vaughan, J.W. (eds.) Advances in Neural Information Processing Systems 34: Annual Conference on Neural Information Processing Systems 2021, NeurIPS 2021, December 6-14, 2021, virtual. pp. 30479–30491 (2021), <https://proceedings.neurips.cc/paper/2021/hash/fbdc6cbb019a1413183c8d08f2929307-Abstract.html>
- [93] Weller, A.: Transparency: Motivations and challenges. In: Samek, W., Montavon, G., Vedaldi, A., Hansen, L.K., Müller, K. (eds.) Explainable AI: Interpreting, Explaining and Visualizing Deep Learning, Lecture Notes in Computer Science, vol. 11700, pp. 23–40. Springer (2019), [https://doi.org/10.1007/978-3-030-28954-6\\_2](https://doi.org/10.1007/978-3-030-28954-6_2)
- [94] Wu, J., Yin, F., Zhang, Y., Zhang, X., Liu, C.: Graph-to-graph: Towards accurate and interpretable online handwritten mathematical expression recognition. In: Thirty-Fifth AAAI Conference on Artificial Intelligence, AAAI 2021, Thirty-Third Conference on Innovative Applications of Artificial Intelligence, IAAI 2021, The Eleventh Symposium on Educational Advances in Artificial Intelligence, EAAI 2021, Virtual Event, February 2-9, 2021. pp. 2925–2933. AAAI Press (2021), <https://ojs.aaai.org/index.php/AAAI/article/view/16399>

- [95] Yadav, R.K., Jiao, L., Granmo, O., Goodwin, M.: Human-level interpretable learning for aspect-based sentiment analysis. In: Thirty-Fifth AAAI Conference on Artificial Intelligence, AAAI 2021, Thirty-Third Conference on Innovative Applications of Artificial Intelligence, IAAI 2021, The Eleventh Symposium on Educational Advances in Artificial Intelligence, EAAI 2021, Virtual Event, February 2-9, 2021. pp. 14203–14212. AAAI Press (2021), <https://ojs.aaai.org/index.php/AAAI/article/view/17671>
- [96] Yadav, R.K., Jiao, L., Granmo, O., Goodwin, M.: Robust interpretable text classification against spurious correlations using and-rules with negation. In: Raedt, L.D. (ed.) Proceedings of the Thirty-First International Joint Conference on Artificial Intelligence, IJCAI 2022, Vienna, Austria, 23-29 July 2022. pp. 4439–4446. ijcai.org (2022). , <https://doi.org/10.24963/ijcai.2022/616>
- [97] Yang, C., Chaudhuri, S.: Safe neurosymbolic learning with differentiable symbolic execution. In: The Tenth International Conference on Learning Representations, ICLR 2022, Virtual Event, April 25-29, 2022. OpenReview.net (2022), <https://openreview.net/forum?id=NYBmJN4MyZ>
- [98] Yang, F., He, K., Yang, L., Du, H., Yang, J., Yang, B., Sun, L.: Learning interpretable decision rule sets: A submodular optimization approach. In: Ranzato, M., Beygelzimer, A., Dauphin, Y.N., Liang, P., Vaughan, J.W. (eds.) Advances in Neural Information Processing Systems 34: Annual Conference on Neural Information Processing Systems 2021, NeurIPS 2021, December 6-14, 2021, virtual. pp. 27890–27902 (2021), <https://proceedings.neurips.cc/paper/2021/hash/ea32c96f620053cf442ad32258076b9-Abstract.html>
- [99] Yang, R., Wang, X., Jin, Y., Li, C., Lian, J., Xie, X.: Reinforcement subgraph reasoning for fake news detection. In: Zhang, A., Rangwala, H. (eds.) KDD '22: The 28th ACM SIGKDD Conference on Knowledge Discovery and Data Mining, Washington, DC, USA, August 14 - 18, 2022. pp. 2253–2262. ACM (2022). , <https://doi.org/10.1145/3534678.3539277>
- [100] Yang, Y., Kerce, J.C., Fekri, F.: LOGICDEF: an interpretable defense framework against adversarial examples via inductive scene graph reasoning. In: Thirty-Sixth AAAI Conference on Artificial Intelligence, AAAI 2022, Thirty-Fourth Conference on Innovative Applications of Artificial Intelligence, IAAI 2022, The Twelveth Symposium on Educational Advances in Artificial Intelligence, EAAI 2022 Virtual Event, February 22 - March 1, 2022. pp. 8840–8848. AAAI Press (2022), <https://ojs.aaai.org/index.php/AAAI/article/view/20865>
- [101] Yao, H., Chen, Y., Ye, Q., Jin, X., Ren, X.: Refining language models with compositional explanations. In: Ranzato, M., Beygelzimer, A., Dauphin, Y.N., Liang, P., Vaughan, J.W. (eds.) Advances in Neural Information Processing Systems 34: Annual Conference on Neural Information Processing Systems 2021, NeurIPS 2021, December 6-14, 2021, virtual. pp. 8954–8967 (2021), <https://proceedings.neurips.cc/paper/2021/hash/4b26dc4663ccf960c8538d595d0a1d3a-Abstract.html>
- [102] Zha, H., Chen, Z., Yan, X.: Inductive relation prediction by BERT. In: Thirty-Sixth AAAI Conference on Artificial Intelligence, AAAI 2022, Thirty-Fourth Conference on Innovative Applications of Artificial Intelligence, IAAI 2022, The Twelveth Symposium on Educational Advances in Artificial Intelligence, EAAI 2022 Virtual Event, February 22 - March 1, 2022. pp. 5923–5931. AAAI Press (2022). , <https://doi.org/10.1609/aaai.v36i5.20537>
- [103] Zhang, D., Zhang, H., Zhou, H., Bao, X., Huo, D., Chen, R., Cheng, X., Wu, M., Zhang, Q.: Building interpretable interaction trees for deep NLP models. In: Thirty-Fifth AAAI Conference on Artificial Intelligence, AAAI 2021, Thirty-Third Conference on Innovative Applications of Artificial Intelligence, IAAI 2021, The Eleventh Symposium on Educational Advances in Artificial Intelligence, EAAI 2021, Virtual Event, February 2-9, 2021. pp. 14328–14337. AAAI Press (2021), <https://ojs.aaai.org/index.php/AAAI/article/view/17685>
- [104] Zhong, W., Huang, J., Liu, Q., Zhou, M., Wang, J., Yin, J., Duan, N.: Reasoning over hybrid chain for table-and-text open domain question answering. In: Raedt, L.D. (ed.) Proceedings of the Thirty-First International Joint Conference on Artificial Intelligence, IJCAI 2022, Vienna, Austria, 23-29 July 2022. pp. 4531–4537. ijcai.org (2022). , <https://doi.org/10.24963/ijcai.2022/629>
- [105] Zhu, Z., Galkin, M., Zhang, Z., Tang, J.: Neural-symbolic models for logical queries on knowledge graphs. In: Chaudhuri, K., Jegelka, S., Song, L., Szepesvári, C., Niu, G., Sabato, S. (eds.) International Conference on Machine Learning, ICML 2022, 17-23 July 2022, Baltimore, Maryland, USA. Proceedings of Machine Learning Research, vol. 162, pp. 27454–27478. PMLR (2022), <https://proceedings.mlr.press/v162/zhu22c.html>