

Symbol Grounding in Neuro-Symbolic AI: A Gentle Introduction to Reasoning Shortcuts

EMANUELE MARCONATO^{*†}, University of Trento, Italy

SAMUELE BORTOLOTTI[†], University of Trento, Italy

EMILE VAN KRIEKEN[†], Vrije Universiteit Amsterdam, The Netherlands

PAOLO MORETTIN, University of Trento, Italy

ELENA UMILI, Sapienza University of Rome, Italy

ANTONIO VERGARI, University of Edinburgh, United Kingdom

EFTHYMIA TSAMOURA, Huawei Labs, United Kingdom

ANDREA PASSERINI, University of Trento, Italy

STEFANO TESO, CIMEC & DISI, University of Trento, Italy

Neuro-symbolic (NeSy) AI aims to develop deep neural networks whose predictions comply with *prior knowledge* encoding, e.g., safety or structural constraints. As such, it represents one of the most promising avenues for reliable and trustworthy AI. The core idea behind NeSy AI is to combine neural and symbolic steps: neural networks are typically responsible for mapping low-level inputs into high-level symbolic concepts, while symbolic reasoning infers predictions compatible with the extracted concepts and the prior knowledge. Despite their promise, it was recently shown that – whenever the concepts are not supervised directly – NeSy models can be affected by **Reasoning Shortcuts (RSs)**. That is, they can achieve high label accuracy by grounding the concepts incorrectly. RSs can compromise the interpretability of the model’s explanations, performance in out-of-distribution scenarios, and therefore reliability. At the same time, RSs are difficult to detect and prevent unless concept supervision is available, which is typically not the case. However, the literature on RSs is scattered, making it difficult for researchers and practitioners to understand and tackle this challenging problem. This overview addresses this issue by providing a gentle introduction to RSs, discussing their causes and consequences in intuitive terms. It also reviews and elucidates existing theoretical characterizations of this phenomenon. Finally, it details methods for dealing with RSs, including mitigation and awareness strategies, and maps their benefits and limitations. By reformulating advanced material in a digestible form, this overview aims to provide a unifying perspective on RSs to lower the bar to entry for tackling them. Ultimately, we hope this overview contributes to the development of reliable NeSy and trustworthy AI models.

JAIR Track: Integration of Logical Constraints in Deep Learning

JAIR Associate Editor: Matteo Zavatteri

^{*}Corresponding Author.

[†]Equal Contribution.

Authors’ Contact Information: Emanuele Marconato, ORCID: [0000-0002-7407-5465](https://orcid.org/0000-0002-7407-5465), emanuele.marconato@unitn.it, University of Trento, Trento, Italy; Samuele Bortolotti, ORCID: [0009-0009-4028-4806](https://orcid.org/0009-0009-4028-4806), samuele.bortolotti@unitn.it, University of Trento, Trento, Italy; Emile van Krieken, ORCID: [0000-0001-5502-4817](https://orcid.org/0000-0001-5502-4817), e.van.krieken@vu.nl, Vrije Universiteit Amsterdam, Amsterdam, The Netherlands; Paolo Morettin, ORCID: [0000-0003-4321-5215](https://orcid.org/0000-0003-4321-5215), paolo.morettin@unitn.it, University of Trento, Trento, Italy; Elena Umili, ORCID: [0000-0002-5639-6038](https://orcid.org/0000-0002-5639-6038), umili@diag.uniroma1.it, Sapienza University of Rome, Rome, Italy; Antonio Vergari, ORCID: [0000-0003-0036-5678](https://orcid.org/0000-0003-0036-5678), avergari@ed.ac.uk, University of Edinburgh, Edinburgh, United Kingdom; Efthymia Tsamoura, ORCID: [0009-0008-8302-1902](https://orcid.org/0009-0008-8302-1902), efthymia.tsamoura@gmail.com, Huawei Labs, Cambridge, United Kingdom; Andrea Passerini, ORCID: [0000-0002-2765-5395](https://orcid.org/0000-0002-2765-5395), andrea.passerini@unitn.it, University of Trento, Trento, Italy; Stefano Teso, ORCID: [0000-0002-2340-9461](https://orcid.org/0000-0002-2340-9461), stefano.teso@unitn.it, CIMEC & DISI, University of Trento, Trento, Italy.



This work is licensed under a Creative Commons Attribution International 4.0 License.

© 2026 Copyright held by the owner/author(s).

DOI: [10.1613/jair.1.22389](https://doi.org/10.1613/jair.1.22389)

JAIR Reference Format:

Emanuele Marconato, Samuele Bortolotti, Emile van Krieken, Paolo Morettin, Elena Umili, Antonio Vergari, Efthymia Tsamoura, Andrea Passerini, and Stefano Teso. 2026. Symbol Grounding in Neuro-Symbolic AI: A Gentle Introduction to Reasoning Shortcuts. *Journal of Artificial Intelligence Research* 86, Article 40 (July 2026), 63 pages. DOI: [10.1613/jair.1.22389](https://doi.org/10.1613/jair.1.22389)

1 Introduction

In cognitive science and psychology, *symbols*—or more generally, *concepts*—are the compositional building blocks of human thought (Mandler, 2004; Spelke and Kinzler, 2007). They enable the abstraction and structuring of perception, supporting high-level cognitive functions (Whitehead, 1927; Johnson-Laird, 1994) such as language, reasoning, and planning. In recent years, there has been growing interest in using concepts in neural networks (Smolensky, 1987; Sun, 1992; Greff et al., 2020), which bring not only improved generalization but also greater interpretability to their decision-making processes. While the emergence of abstract concepts in the human brain is well-documented (Mandler, 2004), how—and even whether—such entities arise in neural networks remains fundamentally unclear (Jo and Bengio, 2017; Lake and Baroni, 2018; Park et al., 2024).

Neuro-symbolic (NeSy) models aim to bridge this gap by integrating neural learning with symbolic reasoning (Harnad, 1990; McMillan et al., 1991; Sun, 1992), thereby enabling explicit manipulation of concepts within the inference process. Some approaches embed prior symbolic structure—such as logical rules (De Raedt et al., 2021)—while others attempt to learn such structure through specialized, trainable components (Lake and Baroni, 2018; Ellis et al., 2021). A key promise of this integration is **reliability** and **trustworthiness**: by making all decisions traceable to interpretable, well-defined concepts, decisions become more transparent (Kambhampati et al., 2022). Moreover, thanks to the modularity of symbolic components, learned concepts can be reused in novel scenarios without significant performance degradation. Yet, the effectiveness of this integration hinges on solving a core challenge in NeSy AI: the problem of **symbol grounding** (Harnad, 1990)—that is, how to connect low-level perceptual data with high-level, abstract concepts.

Example 1.1. Consider a model trained on autonomous driving scenarios, where the car must decide whether to stop or go based on the visual input, as per Fig. 1. The model leverages prior knowledge K_1 which imposes that whenever pedestrian or red_light are detected, the car must stop. Attaining accurate action predictions requires assigning valid binary values to the concepts pedestrian and red_light (among others) from the visual input. Therefore, concepts are grounded by the model through learning to give correct driving predictions. Because pedestrian and red_light are treated symmetrically by the knowledge (due to the disjunction)^a the self-driving car may confuse the two while still achieving correct predictions (e.g., in Fig. 1 (Left), red_light fires when either “pedestrians” or “red lights” are detected).

^aWe will precisely define the symmetries of the knowledge in later sections.

This example highlights a fundamental and underappreciated incompatibility between how NeSy models are trained and what is expected of them: on the one hand, the symbolic knowledge explicitly assigns *names* (like pedestrian and red_light) to the symbols with the expectation that these reflect their meaning, on the other, the model might assign them completely different semantics. In our example, training the autonomous car to give correct predictions does not guarantee that concepts activate when they should: provided other objects apart from red lights can be perceived by the sensors, the autonomous car’s red_light may not solely contain information about the presence of “red lights” in the raw sensory data but also information about the other objects, if the knowledge is ambiguous enough to allow so.

Example 1.1 illustrates the central challenge of symbol grounding in Neuro-Symbolic AI: ensuring that abstract predicted concepts inside the model maintain a consistent and semantically correct link to the real-world entities

they are meant to represent. Unfortunately, in many learning contexts NeSy models—even when they perform nearly optimally on a task—may fail to assign the right meaning to learned concepts, leading to concepts with *unintended semantics* (Marconato et al., 2023a,b). This issue originates from *reasoning shortcuts* (RSs): unwanted concept assignments that allow models to reach correct label predictions. These have been shown to affect many NeSy models (Marconato et al., 2023b; Yang et al., 2024). While the decision-making appears correct on the surface, the underlying concepts are flawed, undermining the reliability and trustworthiness of the NeSy model. For example, in the context of autonomous driving of Example 1.1 and of Fig. 1, the correct “stop” prediction can be achieved by mistaking “pedestrians” with “red_lights”, as both choices entail the same decision.

RSs essentially imply that concept learning is ill-posed in the context of NeSy AI: in many tasks, there is simply not enough information for a NeSy model to learn the correct concepts. Still, RSs are non-trivial, for several reasons. Again, they cannot be detected based on training (and possibly test) accuracy alone. Moreover, researchers, practitioners and stakeholders are likely to assume that the names and semantics of the symbols in the knowledge have to be aligned *by construction* – precisely because this is what happens in logic programming. For instance, imagine having to write a logic program for the addition of two digits. The first step would be to choose how many predicates to use and how to name them. The most intuitive choice is to introduce ten binary predicates to model the (one-hot encoding of the) left digit and ten for the right one. These symbols will then be used in the constraints assuming their name matches their meaning. However, as we will see, there is no guarantee that a NeSy model trained to solve said task will use them as expected: it might well use the first ten predicates for the right digit and the others for the left one. This is at best surprising, and at worst deleterious.

While this misalignment does not induce a drop in prediction accuracy on data *in-distribution*, it becomes critical in situations where such distinctions matter. In fact, ***poorly grounded concepts may not transfer to out-of-distribution scenarios***. This has consequences in applications where the data differs from the training distribution, e.g., in continual learning (Marconato et al., 2023a): despite being able to ensure predictions comply with the prior knowledge, NeSy models may behave erratically due to incorrectly ground concepts, similarly to Fig. 1 (Right), where the wrong activation of model concepts leads to a faulty prediction. This also ***impacts the model’s interpretability***. When asked why it decided to “stop” in Fig. 1 (Left), the model would explain that it stopped because there is a red_light on the road, while clearly this is not the case, and similarly for Fig. 1 (Right). Finally, incorrect grounding also ***affects down-stream applications*** like Neuro-Symbolic verification (Xie et al., 2019), which enables designers to formally verify whether learned models satisfy predetermined constraints defined in terms of high-level concepts, like “does this autonomous vehicle always stop when there is a pedestrian on the road?”. Clearly, incorrect grounding of pedestrian would impair the trustworthiness of these tools. In a nutshell, RSs can compromise some of those desirable properties – among which reliability and interpretability – that NeSy architectures are designed to possess, making them a primary target for mitigation.

RSs were observed in early NeSy work (Manhaeve et al., 2018; Chang et al., 2020; Topan et al., 2021) and were shown to affect a variety of NeSy tasks (Bortolotti et al., 2024), including high-stakes ones like autonomous driving. Yet, they only recently received a formal treatment, revealing that RSs are hard to tackle (Marconato et al., 2023b; Umili et al., 2023; Yang et al., 2024; DeLong et al., 2024), despite arising naturally from certain symmetries in the symbolic component of NeSy systems. In turn, theoretical studies (Marconato et al., 2023b; Wang et al., 2023; Yang et al., 2024; Bortolotti et al., 2025) have begun to identify under which conditions RSs can be provably avoided and when they cannot. Intuitively, these conditions depend on four factors: the structure of the prior knowledge, the coverage of the training set, the hypothesis class of learnable neural networks, and the choice of loss function. Robustness to RSs often requires pairing standard NeSy training methods with ***mitigation strategies*** that encourage better concept grounding by targeting the four aforementioned factors. Building on these theoretical insights, we summarize and categorize existing mitigation strategies – including knowledge editing, multi-task learning, prototypes, entropy maximization, and contrastive learning – and analyse their

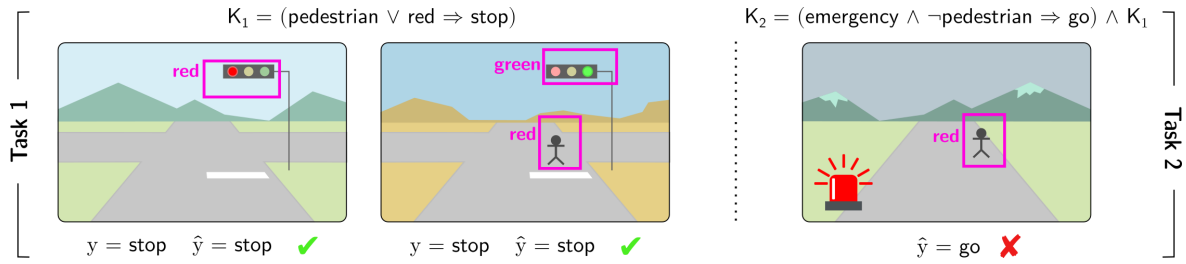


Fig. 1. **Reasoning shortcuts are failures of correct symbol grounding.** Consider an autonomous driving car that must drive in compliance with traffic rules. (Left) The car learns to make **correct predictions** by leveraging background knowledge K_1 in Task 1, i.e., “if there is a pedestrian or a red light the car has to stop”. However, in the process, it may incorrectly associate the concept of red_light with the presence of either *pedestrians* or *red lights* in the dashcam, resulting in what is known as a **reasoning shortcut**. (Right) Consequently, the learned concepts can lead to incorrect (and potentially catastrophic) predictions when the background knowledge changes. For example, K_2 in Task 2 introduces an exception in the presence of an emergency situation, i.e., “if there is an emergency and no pedestrians are on the street, the car may proceed”.¹

effectiveness, scope and overall cost. We conclude by noting how designing effective mitigation strategies that do not require a significant annotation cost remains an open problem.

Overall, we aim to consolidate the scattered literature on RSs in NeSy AI and offer a comprehensive overview of different approaches to mitigate them.

Contributions. This article provides a gentle introduction to reasoning shortcuts, unifies the existing literature on the topic, and connects it to the well-known symbol grounding problem. We provide a general perspective on RSs, highlighting that they cannot be avoided simply by designing different neuro-symbolic architectures. To this end, we compile all relevant theory on RSs through the perspectives of identifiability and statistical learning. We then organize known mitigation strategies to prevent RSs into a taxonomy, and explain which strategies can effectively reduce the likelihood of RSs. Finally, we identify open problems and future directions, showing that RSs – and thus correct symbol grounding – extend beyond the NeSy models studied so far.

Outline. The remainder of the article is organized as follows. [Section 2](#) introduces the preliminaries necessary for understanding the theoretical material. [Section 3](#) presents the problem of reasoning shortcuts at a high level, providing examples and discussing their causes and impacts. [Section 4](#) delves into the details by introducing the mathematical tools required to analyze reasoning shortcuts, exploring the issue from two perspectives: identification and statistical learning. [Section 5](#) adopts a more applied perspective by examining the root causes of RSs, together with practical strategies for diagnosing their presence and for mitigating or addressing their effects, and discusses the benefits of each strategy. [Section 6](#) explores various extensions of the reasoning shortcut problem across different domains and discusses open problems and future directions. Finally, [Section 7](#) reviews relevant related research, and [Section 8](#) provides concluding remarks.

¹To be precise, the prior knowledge for Task 2 would be $K = (\text{alarm} \Rightarrow K_1) \wedge (\neg \text{alarm} \Rightarrow K_2)$ where K_1 is the knowledge used for Task 1, and $K_2 = (\text{emergency} \wedge \neg \text{pedestrian} \Rightarrow \text{go})$.

2 Preliminaries

Notation. Throughout, we denote scalar constants by lowercase letters x , scalar (random) variables by uppercase letters X , vectors of constants $\mathbf{x}$ and (random) variables $\mathbf{X}$ in bold typeface, and sets (e.g., $\mathcal{X}$) by calligraphic letters. Throughout, we use $p(X)$ to indicate the distribution of X and $p(x)$ to indicate the probability of the event $X = x$. If φ is a logical formula over variables $\mathbf{X}$, we say that the constants $\mathbf{x}$ entail the formula φ ($\mathbf{x} \models \varphi$) if and only if replacing the variables $\mathbf{X}$ in the formula φ with $\mathbf{x}$ makes the formula true. In this case, we say that $\mathbf{x}$ satisfies or is consistent with the formula; otherwise, it violates or is inconsistent with the formula. See Table 1 for a glossary.

Table 1. Glossary of used symbols.

Symbol	Meaning	Symbol	Meaning
$x, y, z,$	Scalar constants	X, Y, Z	Scalar variables
$\mathbf{x}, \mathbf{X}$	Vectors	$\mathcal{X}, \mathcal{Y}, \mathcal{C}$	Sets
$f(\mathbf{x}), p(\mathbf{c} \mid \mathbf{x})$	Concept extractor	$\beta(\mathbf{c}), p(\mathbf{y} \mid \mathbf{c}; \mathbf{K})$	Inference layer
$\models$	Logical entailment	$\mathbf{K}$	Prior knowledge
θ	Network parameters	$\mathcal{D}$	Dataset

2.1 From Neuro-Symbolic AI to Neuro-Symbolic Predictors

Neuro-symbolic models aim to solve the long-standing problem of integrating learning and reasoning. Over time, many radically different NeSy architectures have emerged (Garcez et al., 2022; De Raedt et al., 2021; Feldstein et al., 2024) that differ not only in what kind of logic they implement reasoning with (e.g., propositional (Hoernle et al., 2022; Ahmed et al., 2022; Buffelli and Tsamoura, 2023) vs. first-order (Lippi and Frasconi, 2009; Diligenti et al., 2012; Manhaeve et al., 2018) and classical (Zhou, 2019) vs. fuzzy (Donadello et al., 2017; van Krieken et al., 2020) vs. probabilistic (Manhaeve et al., 2018; Ahmed et al., 2022; Feldstein et al., 2023)), but also in how they interpret logic reasoning itself (e.g., model vs. proof-based inference (De Raedt et al., 2021)) and in how they implement the computations (e.g., fully neural (Rocktäschel and Riedel, 2016) vs. hybrid neural-symbolic (Manhaeve et al., 2018)). This variety is reflected by the diversity of tasks that NeSy architectures can tackle, which range from hierarchical classification (Giunchiglia and Lukasiewicz, 2020; Hoernle et al., 2022; Ahmed et al., 2022) and knowledge base completion (Rocktäschel and Riedel, 2016), to open-ended reasoning (Manhaeve et al., 2018; Badreddine et al., 2022) and learning knowledge graph embeddings (Maene and Tsamoura, 2025).

Reasoning shortcuts have been studied mainly in the context of **NeSy predictors**, a class of NeSy architectures specialized for constrained prediction tasks (Giunchiglia et al., 2022; Dash et al., 2022; Marconato et al., 2023a), denoted here as **NeSy tasks**; a discussion of RS in other NeSy architectures is left to Section 6.1. A NeSy task requires predicting labels that are consistent with known constraints. Each such task is defined by four elements: a description of the **input** $\mathbf{X}$ with domain $\mathcal{X}$, a description of the k **concepts** $\mathbf{C} = (C_1, \dots, C_k)$ with domain $\mathcal{C} = C_1 \times \dots \times C_k$, a description of m discrete **labels** $\mathbf{Y} = (Y_1, \dots, Y_m)$ with domain $\mathcal{Y} = \mathcal{Y}_1 \times \dots \times \mathcal{Y}_m$, and **prior knowledge** $\mathbf{K}$. Here, $\mathcal{X}$ is the input domain (e.g., a vector space), $\mathcal{C}$ is a finite set of **concept vectors** $\mathbf{c}$, and so are the label sets $\mathcal{Y}_i$.

The input $\mathbf{X}$ is usually high-dimensional and sub-symbolic (e.g., an image), while the concepts $\mathbf{C}$ are high-level and possibly human readable properties of the input (e.g., the objects appearing in the image). The prior knowledge $\mathbf{K}$ encodes the constraints that we wish the model to satisfy—e.g., safety constraints or relevant regulations—as a single propositional logical formula.² Both the concepts $\mathbf{C}$ and the labels $\mathbf{Y}$ appear as logic variables in $\mathbf{K}$. We ground these notions using a running example:

²Although some NeSy predictor architectures support first-order logical formulas, we restrict our description to propositional formulas, for ease of exposition. We postpone a discussion of how RSs transfer to the first-order case to Section 6.1.

Example 2.1. Consider the simplified BDD-OIA autonomous driving task (Xu et al., 2020). Here, the inputs $\mathbf{x}$ are dashcam images and the (binary) label y indicates whether the vehicle is allowed to proceed. The task involves three binary concepts C_{grn} , C_{red} , C_{ped} encoding the presence of green lights, red lights, and pedestrians in the input image, respectively, and the prior knowledge K encodes the constraint that if either a pedestrian or a red light is visible in the image, the vehicle must stop: $K = (C_{\text{ped}} \vee C_{\text{red}} \Leftrightarrow \neg Y)$. In this scenario, a NeSy predictor has to predict vehicle actions that are both accurate and comply with traffic regulations.

Given a (noiseless) training set $\mathcal{D} = \{(\mathbf{x}, y)\}$, a NeSy predictor is tasked with learning a mapping from inputs $\mathbf{x}$ to labels y that comply with the knowledge, that is, $(c, y) \models K$.³ The key challenge is that, just like in Example 2.1, the prior knowledge K is not specified at the input level, but rather over the concepts. These, in turn, are complex functions of the input that, in practice, cannot be manually specified. For this reason, NeSy predictors adopt modular architectures that comprise a learnable **concept extractor** responsible for predicting the concepts C from the inputs X , and an **inference layer** responsible for inferring the labels Y from C compatibly with the prior knowledge K . The former is typically implemented as a feed-forward neural network,⁴ and the latter using some form of (differentiable) symbolic reasoning. Given an input $\mathbf{x}$, a NeSy predictor first applies the concept extractor to obtain a distribution $p(C | \mathbf{x})$ over concepts, and then applies the inference layer to obtain a distribution $p(Y | C; K)$ over outputs y , given the concepts. Overall, the predictor defines a predictive distribution $p(Y | X; K)$ (see Section 2.2 for details).

Example 2.2. Consider the BDD-OIA (Xu et al., 2020) dataset illustrated in Example 2.1. Here, the input $\mathbf{x} \in X$ is a dashcam image. The concept extractor is a neural network that takes $\mathbf{x}$ and predicts the probabilities of three independent binary concepts (C_{grn} , C_{red} , C_{ped}). For example, given an image $\mathbf{x}$, the extractor may output $p(C_{\text{red}} = 1 | \mathbf{x}) = 0.9$, $p(C_{\text{ped}} = 1 | \mathbf{x}) = 0.2$, $p(C_{\text{grn}} = 1 | \mathbf{x}) = 0.1$. These concept predictions are then passed to an inference layer, which combines them with prior knowledge $K = (C_{\text{ped}} \vee C_{\text{red}}) \Leftrightarrow \text{stop}$ to obtain the final decision y . Depending on the chosen implementation, this reasoning step may enforce the rule strictly, approximate it through fuzzy logic, or learn it implicitly during training. In this example, since the concept extractor assigns high probability to $C_{\text{red}} = 1$, the system will likely infer $Y = \text{stop}$, indicating that the vehicle should stop.

In most cases, NeSy predictors can be, and in fact are, trained using label supervision only. This is not surprising if we consider that per-concept annotations can be expensive to collect in practice. However—and this is critical for our discussion of RSs—this also means that *the concepts themselves are often not supervised, and therefore act as latent variables*.

2.2 The Variety of NeSy Predictor Architectures

Different NeSy predictor architectures differ in how they implement and use the concept extractor and the inference layer. Next, we introduce the four architectures that have been most extensively studied in the RS literature, namely probabilistic NeSy predictors, the Semantic Loss, Logic Tensor Networks, and Abductive Learning. These are illustrated in Fig. 2.

³Here, (c, y) is to be read as the concatenation of c and y .

⁴Although other architectures can be used (Diligenti et al., 2012).

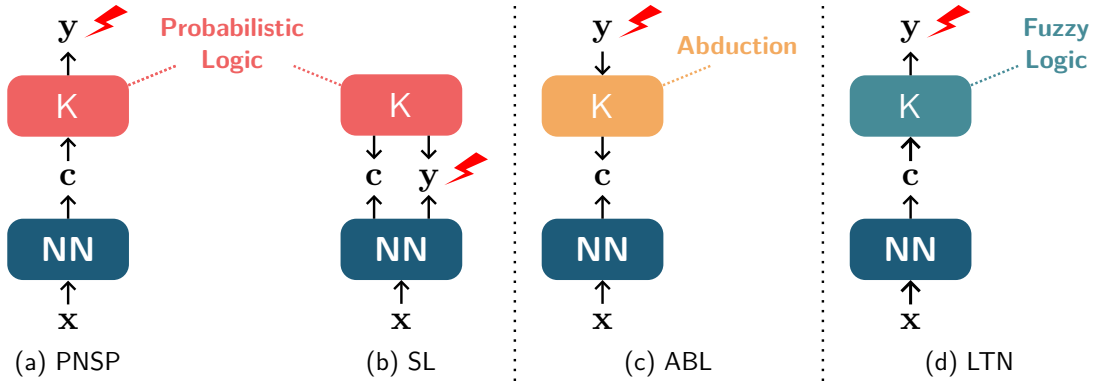


Fig. 2. **Schematic illustration of NeSy predictors.** All of them employ a neural network for extracting concepts and prior knowledge K . PNSPs (a) map inputs x to (a distribution over) concepts c and then uses probabilistic logic to infer labels y consistent with K . The SL (b) maps inputs to both concepts c and labels y and uses K to define a loss term penalizing inconsistent predictions. During training, ABL (c) abduces concepts from the ground-truth label y using K and uses them as pseudo-labels to train the concept extractor. LTN (d) works similarly to PNSPs but uses *fuzzy logic* to make K differentiable. Lightning strikes indicate supervision.

Probabilistic NeSy Predictors. A common method for designing NeSy architectures is the family of **Probabilistic NeSy Predictors** (PNSPs) (Manhaeve et al., 2018, 2021a; Yang et al., 2020; Huang et al., 2021b). PNSPs can be viewed as combining a neural concept extractor with a probabilistic-logic reasoning layer (De Raedt and Kimmig, 2015), as shown in Fig. 2 (a). Perhaps the best-known member of this family is DeepProbLog (DPL) (Manhaeve et al., 2018, 2021a), a fully-fledged neuro-symbolic programming language based on Prolog and its probabilistic extension, ProbLog (De Raedt et al., 2007; Kimmig et al., 2011). From a simplified point of view, PNSPs define the following predictive distribution for any x and y :

$$p(y | x; K) = \sum_c p(y, c | x; K) = \sum_c p(y | c; K) p(c | x) = \frac{1}{Z_x} \sum_c \mathbb{1}\{(c, y) \models K\} p(c | x) \quad (1)$$

Here, $p(C | x)$ is the predictive distribution over concepts computed by the concept extractor, $\mathbb{1}\{\cdot\}$ is the indicator function, and Z_x is a normalization constant ensuring that the result is indeed a conditional distribution. In words, the probability of a label y is the sum of the probabilities of all concept vectors c that are consistent with that label y according to the knowledge K . Inference in PNSPs amounts to computing a most likely label, which usually requires solving the Maximum A Posteriori (MAP) problem $\arg\max_y p(y | x; K)$ (Koller and Friedman, 2009). Importantly, if a label y , together with the predicted concepts, would violate the knowledge K , then *its probability is exactly zero and it will not be predicted*. Learning in PNSPs is implemented via maximum likelihood estimation. That is, given a training set $\mathcal{D} = \{(x, y)\}$, learning amounts to maximizing the log-likelihood of the data, namely

$$\mathcal{L}(p, \mathcal{D}, K) = \frac{1}{|\mathcal{D}|} \sum_{(x, y) \in \mathcal{D}} \log p(y | x; K). \quad (2)$$

by tuning the parameters of the concept extractor. In practice, this is done via gradient descent, as the predictive distribution $p(y | x; K)$ is differentiable (Manhaeve et al., 2018).

PNSPs require evaluating the predictive probability $p(y | x; K)$, which involves summing over a potentially exponential number (in k , the cardinality of the concept vector) of vectors c , as per Eq. (1). Hence, both inference

and learning are worst-case intractable.⁵ DPL works around this issue by exploiting *knowledge compilation* techniques (Darwiche and Marquis, 2002) that leverage symmetries in the summation to rewrite it into a data structure – a *probabilistic circuit* (Choi et al., 2020; Vergari et al., 2021; Maene et al., 2025; Derkinderen et al., 2025) – that is potentially much more compact and supports efficient evaluation. This allows DPL to scale to practical NeSy tasks. Other PNSPs refine and approximate these ideas (specifically, Eq. (1)) to further improve scalability (Manhaeve et al., 2021b; Huang et al., 2021b; Winters et al., 2022; De Smet et al., 2023a; van Krieken et al., 2023; Choi et al., 2026; Chen et al., 2025).

Semantic Loss. The **Semantic Loss** (SL) (Xu et al., 2018) also relies on probabilistic logic and knowledge compilation, but uses them to convert the prior knowledge K into a differentiable penalty term. This can be used to steer any neural network classifier toward allocating low or even zero probability to inconsistent outputs. While the SL is a general purpose strategy, here we discuss how it can be exploited for defining a NeSy predictor (Marconato et al., 2023a). The idea is to apply the SL to a neural network that acts *both* as concept extractor and as classifier, *i.e.*, that outputs both $p(C | X)$ and $p(Y | X)$. This encourages the concepts and labels predicted by the network to be logically consistent with each other according to K .⁶ The SL grows proportionally to how much probability mass the concept extractor associates to invalid configurations, or more precisely:

$$SL(p_\theta, (x, y), K) = -\log \sum_c \mathbb{1}\{(c, y) \models K\} p_\theta(c | x) \quad (3)$$

As in DPL, efficient computation of this sum relies on knowledge compilation. During training, the SL is paired with a standard supervised loss ℓ (*e.g.*, the cross-entropy loss), resulting in a joint objective of the form:

$$\mathcal{L}(p_\theta, \mathcal{D}, K) = \frac{1}{|\mathcal{D}|} \sum_{(x, y) \in \mathcal{D}} \ell(p_\theta, (x, y)) + \mu SL(p_\theta, (x, y), K) \quad (4)$$

with $\mu > 0$ a hyperparameter. Learning amounts to minimizing this joint objective on training data. At test time, predictions are computed directly by the neural network classifier $p_\theta(y | x)$ rather than by a symbolic layer as in PNSPs.

Logic Tensor Networks. Another well-known approach is **Logic Tensor Networks** (LTNs) (Donadello et al., 2017; Badreddine et al., 2022), which share elements with both PNSPs and SL. As with the SL, they transform the symbolic knowledge K into a differentiable real-valued function $\mathcal{T}_K$ that evaluates how well predictions conform to the logical constraints using *fuzzy logic* (van Krieken et al., 2020). This transformation is a proper *relaxation*. Specifically, it converts all the original Boolean variables in K (namely, the concepts and the labels) into real-valued variables in the range $[0, 1]$ encoding *degrees of truth*. At the same time, it converts all logic connectives (conjunctions, disjunctions, and negations) into real-valued operators capable of handling soft degrees of truth. The translation yields a function $\mathcal{T}_K$ that takes as input a distribution over the concepts C and a distribution over the labels Y , and outputs the degree (in $[0, 1]$) to which these satisfy the knowledge K , the higher the better. By default, LTNs employ a transformation based on the *product real logic* (Donadello et al., 2017) T-norm to ensure that $\mathcal{T}_K$ is fully differentiable, although other options are available (van Krieken et al., 2020). During training, LTNs penalize the concept extractor $p(C | X)$ for violating the prior knowledge by minimizing the following:

$$\mathcal{L}(p, \mathcal{D}, K) = 1 - \frac{1}{|\mathcal{D}|} \sum_{(x, y) \in \mathcal{D}} \mathcal{T}_K(p(C | x), \mathbb{1}\{Y = y\}) \quad (5)$$

⁵Specifically, Eq. (1) can be understood as an instance of (weighted) model counting, the problem of computing (the total weight of all) models, or valid assignments, of a propositional logic formula. This problem is well known to be #P-complete in general, see *e.g.*, (Roth, 1996).

⁶Whether this is a single neural network or two specialized networks makes no difference from a modeling perspective.

Here, $\mathbb{1}\{Y = y\}$ is a (deterministic) distribution that assigns all probability mass to the ground-truth label y . The purpose of this objective is to encourage $p(C | X)$ to concentrate its mass on concepts that satisfy K . Similarly to PNSPs, LTNs employ the prior knowledge at inference time. They proceed in two steps: they first compute the most probable concept vector $\hat{c} = \operatorname{argmax}_c p_\theta(c | x)$ in a forward pass, then select the label $\hat{y}$ that maximizes the satisfaction of the knowledge $\mathcal{T}_K(\hat{c}, \mathbb{1}\{Y = \hat{y}\})$.

Abductive Learning. Another well-studied approach is **Abductive Learning** (ABL) (Zhou, 2019). Compared to the above NeSy predictors, it works backwards: rather than inferring a prediction from the concepts, it constrains the predicted concepts using the ground-truth label. To this end, during training, ABL finds the concept vectors $\hat{c}$ that are both close to those predicted by the concept extractor (according to some pre-defined distance metric) and that according to K entail the ground-truth label y . It then uses them as pseudo-labels to supervise the concept extractor. The pseudo-labels are obtained by solving the following optimization problem:

$$\hat{c} = \operatorname{argmin}_{c' \in C} d(\bar{c}, c') \quad \text{s.t.} \quad (c', y) \models K \quad (6)$$

where $\bar{c} = \operatorname{argmax}_{c \in C} p(c | x)$, and d is a suitable distance metric. The constraint ensures that $\hat{c}$ entails the ground-truth label y . The choice of the distance metric influences the type of weak supervision obtained, thereby intuitively biasing learning toward certain solutions. The training objective of the concept extractor is to maximize the (log)-likelihood of the pseudo-labels using $\hat{c}$ from Eq. (6):

$$\mathcal{L}(p, \mathcal{D}, K) = \frac{1}{|\mathcal{D}|} \sum_{(x, \hat{c})} \log p(\hat{c} | x) \quad (7)$$

At inference time, ABL first predicts the concepts and then uses the symbolic knowledge to obtain the final prediction.

2.3 NeSy Predictor Architectures: Differences and Similarities

These different NeSy predictor architectures all share a key property: **if the concept extractor allocates all probability mass to a single concept vector, then all architectures will output a label compatible with the rules of propositional logic**. Consider MNIST-Add (Manhaeve et al., 2018), where the task is to predict the sum (e.g., $y = 9$) of the digits appearing in two MNIST (LeCun, 1998) images (e.g., $x = \begin{bmatrix} 4 & 5 \end{bmatrix}$). The concepts $C = (C_1, C_2)$ encode the individual digits (e.g., $c = (4, 5)$).⁷ Prior knowledge enforces the prediction to be their arithmetic sum: $K = (Y = C_1 + C_2)$. In this example, if the concepts $C_1 = 2$ and $C_2 = 3$ are predicted with certainty, then all architectures will predict the label $Y = 5$ with certainty—thus matching what any logical encoding of the arithmetic sum would do. However, for concept predictions that are not certain, different architectures can output different label distributions. This distinction is particularly important during optimization, as gradients are shaped by how probability mass over concepts is aggregated and filtered by the inference mechanism (van Krieken et al., 2020).

At the same time, these four NeSy predictor architectures differ in several subtle but significant ways. The first one concerns **efficiency**. LTN and ABL can have an advantage over probabilistic logic approaches (like PNSPs and the SL) in that inference—and therefore training—does not require summing over all possible concept configurations, making it potentially more efficient, although knowledge compilation and approximation strategies help bridge the gap (Huang et al., 2021b). This, however, often comes at a cost, in that LTN and ABL tend to be more susceptible to local minima (Badreddine et al., 2022). On the other hand, fuzzy relaxations (as used by LTN) may not be entirely accurate to the original prior knowledge, in the sense that, for certain problems, the optima

⁷Here and elsewhere, we simplify the presentation by using numerical variables. It is always possible to encode these and the corresponding constraints into propositional logic.

of the satisfaction function $\mathcal{T}_K$ may not correspond to models (that is, 0–1 solutions) of K (Giannini et al., 2018; van Krieken et al., 2022). Moreover, the fuzzy transformation may lead to mathematically different relaxations for logically equivalent constraints (Di Liello et al., 2020; van Krieken et al., 2022). Probabilistic logic is not affected by these issues (Xu et al., 2018). ABL is special in that it does not require the reasoning step to be differentiable, and as such it does not need to relax or extend the semantics of the prior knowledge at all.

A second difference concerns *validity guarantees*, which set apart layer-based approaches from penalty-based ones. The SL “bakes” the prior knowledge directly into the neural network, meaning that finding a high-quality output amounts to a simple forward pass over the network: K plays no role during inference. This also results in a lower memory footprint compared to PNSPs, as the probabilistic circuit can be dropped after training (Di Liello et al., 2020). The downside is that, while PNSPs and ABL ensure invalid outputs will not be predicted, this is not the case for the SL, unless the neural network attains exactly zero Semantic Loss during training, and even then this property might not carry over to test and out-of-distribution samples. Regardless, it has been shown that PNSPs and the SL have the same effect on the underlying neural network, that is, when learned to optimality, both models yield concept extractors that comply with the prior knowledge K (Marconato et al., 2023a). LTN sits in-between these alternatives, as it uses the prior knowledge at inference time, but may output inconsistent predictions if the most likely inputs violate the knowledge. To resolve this issue, fuzzy logic layers such as CCN+ (Giunchiglia et al., 2024) ensure all outputs comply with the knowledge.

2.4 NeSy Predictors: Benefits

Despite these differences, all NeSy predictors offer a number of key benefits:

- B1 Performance:** Just like regular neural networks, they are fully-fledged deep learning architectures that can handle complex low-level inputs via end-to-end training and latent representations.
- B2 Validity:** Unlike regular neural networks, their predictions *comply with the prior knowledge* K , possibly also with guarantees and out-of-distribution, cf. Section 2.3. This is essential in high-stakes applications where predictions have to comply with safety or structural constraints.
- B3 Reusability:** It is straightforward to *reuse the learned concept extractors in downstream tasks*, e.g., out-of-distribution tasks (Marconato et al., 2023b), continual learning scenarios (Marconato et al., 2023a), model verification (Xie et al., 2022; Zaid et al., 2023; Morettin et al., 2024), and shielding (Yang et al., 2023a).
- B4 Interpretability:** Users can not only inspect the concept-level predictions and the symbolic inference steps to make sense of the model’s predictions, they can also *trace back the model’s prediction to the underlying concepts* using gradient-based (Sundararajan et al., 2017) or formal (Huang et al., 2021a) explainability techniques. This supports stakeholders in assessing the reliability of the predictor and potentially enables them to supply corrective feedback (Teso et al., 2023).

Compared to regular neural networks, which primarily target **B1**, benefits **B2–B4** make NeSy predictors an ideal choice for high-stakes applications (Di Liello et al., 2020; Hoernle et al., 2022; Marconato et al., 2023b; Yang et al., 2023a). However, reusability (**B3**) and interpretability (**B4**) *hinge on the concepts being grounded appropriately*. In fact, unless the learned concepts possess reasonable semantics, their meaning might be opaque to stakeholders, making it difficult to properly explain the model’s predictions with them. Furthermore, as exemplified by Fig. 1, reusing poorly grounded concepts in downstream applications may lead to unintended consequences. This is precisely where reasoning shortcuts enter the picture, as discussed next.

3 A Gentle Introduction to Reasoning Shortcuts

When does a NeSy predictor ground the concepts incorrectly? It may be tempting to assume that—as long as the training data $\mathcal{D} = \{(\mathbf{x}, \mathbf{y})\}$ is abundant and noiseless, and the prior knowledge is complete and correct⁸—NeSy predictors that minimize the training loss will predict the concepts correctly, just like in standard supervised learning. However, this is not the case. This is due to **reasoning shortcuts** (RSs), which we informally define as follows:

Definition 3.1 (Reasoning Shortcut, Informal). *A reasoning shortcut is a situation in which a NeSy predictor attains accurate label predictions that comply with the prior knowledge by grounding concepts incorrectly.*

To ground intuition, consider the following example:

Example 3.2 (RSs in MNIST-Add (Manhaeve et al., 2018)). Consider a toy instance of MNIST-Add where the training set consists of just two examples:

$$(\mathbf{4} \ \mathbf{5}) \mapsto 9 \quad \text{and} \quad (\mathbf{3} \ \mathbf{2}) \mapsto 5. \quad (8)$$

Assume that the concept extractor processes the MNIST images separately. Ideally, it would ground the concept correctly, that is, learn the intended image-concept mapping – that is, $\{\mathbf{2} \mapsto 2, \mathbf{3} \mapsto 3, \mathbf{4} \mapsto 4, \mathbf{5} \mapsto 5\}$ – as it adheres to the constraints and attains high label accuracy. But it can alternatively ground them incorrectly, that is, learn the less intuitive input-concept mapping $\{\mathbf{2} \mapsto 4, \mathbf{3} \mapsto 1, \mathbf{4} \mapsto 3, \mathbf{5} \mapsto 6\}$, which also complies with the constraints and achieves high label accuracy.

While this is not a realistic application, the same issue affects also high-stakes tasks:

Example 3.3 (RSs in BDD-OIA). Consider Example 2.1. A NeSy predictor that grounds the concepts of green light, red light, and pedestrian correctly achieves high accuracy, as it will also correctly infer that it has to stop whenever a pedestrian or a red light appear in the input image. However, a different NeSy predictor that systematically confuses pedestrians with red lights achieves the same accuracy, since, according to K, both lead to the correct stop action (Marconato et al., 2023b), as shown in Fig. 3 below.

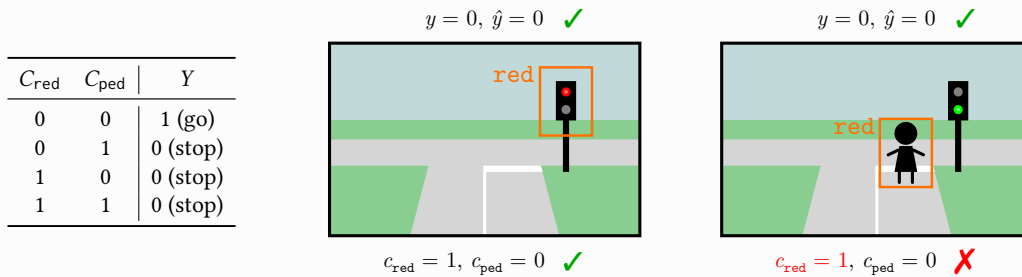


Fig. 3. Left: truth table of the simplified prior knowledge K used in Example 2.1: the two concepts are interchangeable, in that as soon as one of them fires, the predictor infers the same (correct) action. Right: illustration of this RS: the model can confuse pedestrians and red lights with no drop in training loss.

⁸That is, for all inputs, it correctly and unambiguously allows to infer the ground-truth label from the ground-truth concepts, as in MNIST-Add.

These examples show that correct and faulty NeSy predictors cannot be distinguished by label accuracy alone, and therefore **there is no reason why, during training, NeSy predictors should favor one concept mapping over the other**. They also indicate that RSs can occur even when the prior knowledge is accurate and complete, and the training data is noiseless. Moreover, RSs can persist even if the training data is exhaustive (e.g., if the BDD-OIA training set encompasses examples containing all possible combinations of green lights, red lights, and pedestrians), as the correct and faulty models yield the same predictions on all examples.

3.1 Causes, Frequency and Impact

Why do RSs arise? Roughly speaking, RSs stem from two related issues: on the one hand, the prior knowledge K may allow inferring the correct labels y from improperly grounded concept vectors c ; on the other, the concept extractor is expressive enough (e.g., has enough layers) to acquire any such incorrect input-concept mapping. Our two examples satisfy both conditions. Their combination **introduces ambiguity in the learning problem**, meaning that NeSy predictors are free to learn incorrect concept mappings and still achieve high accuracy. Properly understanding the root causes of RSs, however, requires more technical background, which we provide in [Section 4](#), hence we postpone a more comprehensive discussion to [Section 5.1](#). Importantly, clarifying the causes of RSs allows us to discriminate between risky and safe NeSy tasks, and to design RS mitigation strategies. We explore these topics in [Section 5](#).

Do RSs occur in practice? The above discussion entails that RSs can occur whenever the prior knowledge is not “strict” enough, but it does not prove that they *must* occur. Arguably, NeSy predictors *can* learn the intended concepts even in this case. The question is then whether RSs occur in practice. While it is difficult to gauge their frequency in real-world applications, the literature abounds with situations in which NeSy predictors trained in standard conditions fall prey to RSs ([Manhaeve et al., 2018](#); [Wang et al., 2019](#); [Chang et al., 2020](#); [Manhaeve et al., 2021a](#); [Marconato et al., 2023b, 2024](#); [Bortolotti et al., 2024](#); [Yang et al., 2024](#); [DeLong et al., 2024](#)), suggesting that **RSs occur naturally in NeSy benchmarks**.

Do RSs make NeSy predictors unusable? The short answer is *no*. More precisely, in applications where the only goal is to obtain accurate (B1) and valid (B2) predictions, RSs are not an issue. In fact, [Definition 3.1](#) makes it clear that RSs are a symbol grounding issue that leaves label accuracy unchanged. Hence, **the label predictions of affected and unaffected NeSy predictors can be just as accurate**. We also remark that RSs **have no impact on validity (B2)**. As mentioned in [Section 2.3](#), predictors like PNSPs, ABL, and LTN ensure their output is *always* consistent with the prior knowledge by explicitly searching for a label y that, paired with the predicted concepts vectors $\hat{c}$ (or their predictive distribution $p(C | x)$), satisfies the knowledge K . Validity is built-in, *i.e.*, it does not depend on the concepts predicted by the model.

A tricky aspect of RSs is that we cannot detect them by monitoring label predictions alone. We will discuss more reliable diagnostic techniques in [Section 5.2](#) and strategies for rendering predictors aware of their own RSs in [Section 5.4](#). Furthermore, **RS mitigation strategies can greatly reduce the presence of RSs or even entirely remove them** to ensure NeSy predictors are also reusable (B3) and interpretable (B4). We overview these strategies in [Section 5.3](#).

What consequences do RSs have? RSs can seriously compromise reusability (B3) and interpretability (B4). Let us begin from reusability. RSs may exploit concept groundings that do not work in other application settings, meaning that affected concepts suffer from **poor generalization beyond the specific data and prior knowledge used during training** ([Marconato et al., 2023b](#); [Li et al., 2023](#); [Bortolotti et al., 2024](#)). To see this, consider the following example, taken from ([Marconato et al., 2023b](#)):

Example 3.4. Consider the NeSy predictor confusing pedestrians with red lights in [Example 3.3](#). Imagine pairing it with prior knowledge intended for an autonomous ambulance use-case, encoding that, in case of an emergency, it is allowed to cross red lights. When emergencies arise, the resulting NeSy predictor would decide to $y = \text{go}$ when there are pedestrians on the road.

RSs also affect concept reuse in downstream applications where the semantics of the concepts matters. For instance, in **continual learning** one is concerned with learning models that generalize across a sequence of learning tasks. NeSy predictors were shown to struggle with this precisely because of RSs ([Marconato et al., 2023a](#)):

Example 3.5. Consider the shortcut predictor from [Example 3.3](#). Imagine fine-tuning it on a new NeSy task that also includes the concept of “stop sign” and a new rule based on it, and specifically that whenever a stop sign is detected, the vehicle must stop. When fine-tuned on this new task, the concept extractor may mistakenly learn to equate stop signs and red lights, whose semantics are corrupted.

Another affected application is **neuro-symbolic verification** ([Liu et al., 2021](#); [Xie et al., 2022](#); [Zaid et al., 2023](#); [Moretting et al., 2024](#); [Manginas et al., 2025](#)). In model verification, the goal is to formally verify whether a neural network satisfies some property of interest, such as robustness to adversarial attacks or fairness. This involves writing a formal specification of the property, translating the model into a logical formula (or a similar representation), and then using tools from automated reasoning to check whether the model’s formula is compatible with the specification. NeSy verification is similar, except that the property of interest is specified in terms of high-level concepts. Following the neuro-symbolic setup, their definitions are provided by a trained concept extractor – provided by a third party – that also gets translated into a logical formula for the purpose of verification. Clearly, if the concepts are incorrect, this will affect the meaning of the property to be verified.

Finally, **interpretability (B4)** also relies on the concepts being grounded appropriately ([Marconato et al., 2023b](#); [Koh et al., 2020](#); [Poeta et al., 2023](#); [Tapli et al., 2026](#)). Unless these possess human-aligned semantics, their meaning might be opaque to stakeholders, impairing understanding:

Example 3.6. In [Example 3.3](#), imagine the input scene portrays a pedestrian. An affected NeSy predictor would (correctly) predict stop, but its explanation would indicate that the decision depends on the presence of a red light. This explanation is misleading, and in fact it does not reflect the input at all.

4 Theory of Reasoning Shortcuts

In this section, we outline the two mainstream formalizations of RSs. [Section 4.2](#) discusses RSs from an *identifiability* perspective, investigating under what conditions models achieving high likelihood can in fact learn correct concepts in the infinite sample case. [Section 4.3](#) instead is concerned with guarantees – in terms of empirical risk minimization – for learning correct concepts from finite data. The so-far vague notion of “correct” concept will be clarified in the following sections.

Additional notation. Going forward, we will often treat distributions over a finite set of values as vectors in a simplex. We will write Δ_C to indicate the simplex defined by the values contained in C :

$$\Delta_C = \{\mathbf{p} \in [0, 1]^{|C|} \mid \sum_{i=1}^{|C|} p_i = 1\}$$

Furthermore, we will write $\text{Vert}(\Delta_C) = \{\mathbf{p} \in \{0, 1\}^{|C|} \mid \sum_{i=1}^{|C|} p_i = 1\}$ to indicate the *vertices* of the simplex Δ_C . This is simply a collection of one-hot vectors, one for each value in C .

4.1 Setting and Assumptions

In order to formally describe RSs, we have to clarify the link between the observed data, the ground-truth concepts, and the concepts learned by the predictor. This link is provided by the **data generation process**, described next, which forms a solid basis for both theoretical approaches.

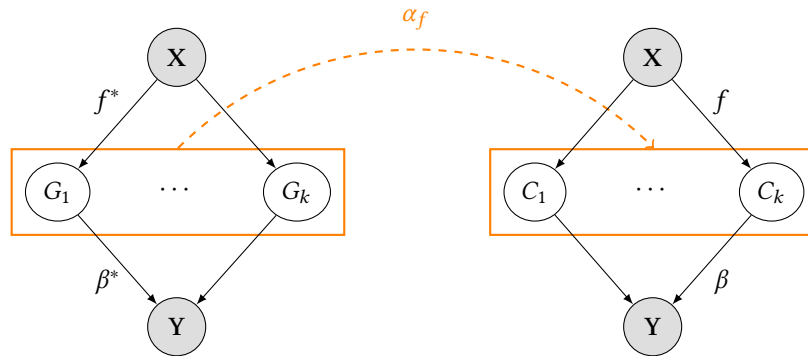


Fig. 4. **Left.** The data generation process: f^* maps inputs X onto (distributions over) ground-truth concepts G , and β^* maps these to (distributions over) labels Y . In a standard learning setting, only X and Y are observed (in light gray), whereas G are latent. **Right.** A NeSy predictor maps inputs X to concepts C via a learned f and infers labels Y via an inference layer β . By [Eq. \(19\)](#), doing so entails a map α_f (in orange) from ground-truth concepts G to learned concepts C .

The data generation process. First we distinguish between **ground-truth concepts** $\mathbf{g} = (g_1, \dots, g_k) \in C$ underlying the data and **learned concepts** $\mathbf{c} = (c_1, \dots, c_k) \in C$. For instance, in [Example 3.3](#) there are two binary concepts encoding the presence of pedestrians and red traffic lights: g_{ped} and g_{red} are the true values, while c_{ped} and c_{red} are the predicted values.

The *training examples* $(\mathbf{x}, \mathbf{y})$ are sampled from a ground-truth joint distribution $p^*(\mathbf{X}, \mathbf{Y}) = p^*(\mathbf{Y} \mid \mathbf{X})p^*(\mathbf{X})$, which is determined by the ground-truth concepts $\mathbf{g} \in C$ associated with the input $\mathbf{x}$. Specifically, the values of the ground-truth concepts are sampled from a ground-truth conditional distribution $p^*(\mathbf{G} \mid \mathbf{X})$, while the ground-truth labels $\mathbf{y}$ are sampled from $p^*(\mathbf{Y} \mid \mathbf{G}; K)$. It is assumed that this distribution implements the background

knowledge K , that is, it assigns zero or low probability to labels incompatible with the constraints.⁹ It is useful to think of these two distributions as *functions*: the former as a function $f^* : \mathcal{X} \rightarrow \Delta_C$ mapping each input to a *distribution* over ground-truth concepts, and the latter as a function $\beta^* : \Delta_C \rightarrow \Delta_Y$ mapping each distribution over concepts to a distribution over labels. Elements $f(\mathbf{x}) \in \Delta_C$ are probability distributions $p(C | \mathbf{x})$ over the possible concept vectors $\mathbf{c} \in C$. Whenever $f(\mathbf{x})$ allocates all probability mass to one specific $\mathbf{c} \in C$, *i.e.*, it is one-hot, it is a vertex of the simplex and as such it belongs to $\text{Vert}(\Delta_C) \subset \Delta_C$. Moreover, abusing notation, we write

$$\beta^*(\mathbf{c}) := \beta^*(\mathbb{1}\{C = \mathbf{c}\}) \quad \forall \mathbf{c} \in C \quad (9)$$

to indicate the map $\beta^* : C \rightarrow \Delta_Y$ from the set of *concepts* to distribution over labels.

Overall, $p^*(Y | X; K)$ results from composing f^* and β^* . A visualization of the data generation process is reported in Fig. 4 (Left). With this in mind, we present the first assumption capturing how data is generated, which will be used throughout the theoretical material.

Assumption 4.1 (Completeness). *Let $f^* : \mathcal{X} \rightarrow \Delta_C$ be the ground-truth concept extractor, and $\beta^* : \Delta_C \rightarrow \Delta_Y$ the inference function induced by the prior knowledge K . Let $p^*(X)$ be the marginal distribution over input variables. We assume that the joint data distribution is generated as:*

$$p^*(X, G, Y) = p^*(Y | G; K) p^*(G | X) p^*(X) \quad (10)$$

where $p^(G | X) := f^*(X)$ and $p^*(Y | G; K) = \beta^*(G)$. In particular, we have $p^*(Y | X; K) := (\beta^* \circ f^*)(X)$.*

Under [Assumption 4.1](#),¹⁰ the ground-truth concept extractor becomes a *sufficient statistic* of the input X for inferring the ground-truth labels Y , that is, the ground-truth concept distribution output by f^* retains all information necessary to find the correct labels with β^* . This is possible in all tasks from rsbench ([Bortolotti et al., 2024](#)), where ground-truth concepts are fully predictive of final labels. Notice that, however, this is not necessarily true in all practical cases, *e.g.*, when the concept space is “too restrictive” compared to the label space. For example, in MNIST-Add, if there are only two ground-truth digits, then sums above 18 cannot be generated. The above assumption excludes such cases from the training data. Furthermore, [Assumption 4.1](#) ensures the labels are independent from the input given the ground-truth concepts, that is, $X \perp\!\!\!\perp Y | G$, and that all pairs $(G, Y) \sim p^*(G, Y)$ are consistent with the knowledge, *i.e.*, $(G, Y) \models K$.

Simplifying assumptions on the data generation process. While [Assumption 4.1](#) underlies all theoretical works on RSs ([Yang et al., 2024](#); [Wang et al., 2023](#); [Marconato et al., 2023b](#); [He and Li, 2025](#)), we introduce two additional assumptions that make the theoretical analysis more manageable. The first assumption restricts the ground-truth concept extractor f^* : It should assign all probability mass to a single vector of concepts $\mathbf{g}$ for each input $\mathbf{x}$.

Assumption 4.2 (Extrapolability). *The ground-truth concept extractor is a function $f^* : \mathcal{X} \rightarrow \text{Vert}(\Delta_C)$, *i.e.*, for all inputs $\mathbf{x} \in \mathcal{X}$, the ground-truth concept distribution $f^*(\mathbf{x})$ is a one-hot (deterministic) distribution.*

[Assumption 4.2](#) allows us to retrieve ground-truth concepts G from inputs X via f^* . The next assumption ensures that, with the prior knowledge K , each vector of concepts $\mathbf{g}$ is associated with a single label $\mathbf{y}$.

⁹More specifically, for technical reasons, we assume the inference layer of the predictor behaves as in [Eq. \(11\)](#).

¹⁰[Assumption 4.1](#) can be viewed as the distributional analogue of the invertibility assumption adopted by ([Marconato et al., 2023b](#)). In that case, inputs are generated through an invertible map $f^* : (G, S) \mapsto X$, where G denotes the ground-truth concepts and S the style variables. The invertibility of f^* guarantees that the concepts can be recovered from the input (as in [Assumption 4.2](#)) and therefore constitute a sufficient statistic of X for predicting Y . Our formulation encodes the same requirement directly at the level of probability distributions, without requiring concepts to be deterministically associated to the input X .

Assumption 4.3 (Determinism). *The ground-truth inference layer defines a function $\beta^* : C \rightarrow \text{Vert}(\Delta_Y)$ that maps each concept vector $\mathbf{c} \in C$ to a one-hot (deterministic) distribution over labels $\beta^*(\mathbf{c})$.*

In words, [Assumption 4.3](#) ensures that β^* assigns a single label $\mathbf{y}$ to each concept vector $\mathbf{c}$. [Assumptions 4.2](#) and [4.3](#) often apply in experimental regimes. For MNIST-Add and BDD-OIA, the concepts in the input image can be unambiguously determined ([Assumption 4.2](#) is satisfied), and all concept vectors $\mathbf{c}$ determine a single output label ([Assumption 4.3](#) is satisfied), e.g., in MNIST-Add we use $Y = C_1 + C_2$. Together, [Assumptions 4.2](#) and [4.3](#) ensure that each input $\mathbf{x} \in X$ is mapped to a single label $\mathbf{y} \in Y$. The same holds for many common NeSy datasets ([Bortolotti et al., 2024](#)).

Before proceeding, we have to define the support of the input and ground-truth concept distributions. We consider the marginal distributions $p^*(X)$ and $p^*(G) := \mathbb{E}_{\mathbf{x} \sim p^*(X)} [f^*(\mathbf{x})]$. We indicate the support of input data with $\text{supp}(X) \subseteq X$, which contains the inputs with non-zero probability according to $p^*(X)$. Similarly, we denote the support of ground-truth concepts by $\text{supp}(G) \subseteq C$.

4.2 Perspective: Identification

In this section, we ask whether a NeSy predictor trained in a standard supervised fashion is guaranteed to learn the ground-truth concept extractor. Specifically, we want to determine whether—depending on the choice of knowledge, training data, and architectural bias—the ground-truth concept extractor is the *unique maximum* of the likelihood of the training data. If so, then any NeSy predictor that attains maximum likelihood will necessarily learn the ground-truth concept extractor, and thus ground all concepts correctly. This question puts us firmly in the context of *identifiability theory*; for additional background, please see [Section 7](#).

Below we provide a revised perspective on the identifiability of RSs ([Marconato et al., 2023a,b](#); [Bortolotti et al., 2025](#)) by clarifying the setting, presenting a first non-identifiability result from ([Marconato et al., 2023a](#)) in [Theorem 4.4](#), and reviewing the characterization of RSs from ([Marconato et al., 2023b](#); [Bortolotti et al., 2025](#)) in [Theorems 4.10](#) and [4.13](#). All results assume the data follows the generating process described in [Fig. 4](#) and that we have access to possibly infinite amounts of data.

4.2.1 Model class of NeSy predictors. Like in the data generation process, NeSy predictors can be understood as a pair of functions: a concept extractor $f : X \rightarrow \Delta_C$ and an inference layer $\beta : \Delta_C \rightarrow \Delta_Y$. We indicate with $\mathcal{F}$ the space of learnable concept extractors f . Most theoretical accounts on RSs work in a *non-parametric* setting, whereby $\mathcal{F}$ contains all possible functions from X to Δ_C . In words, they assume the neural network implementing the concept extractor can express *any* function $f : X \rightarrow \Delta_C$ ([Marconato et al., 2023b](#); [Yang et al., 2024](#); [Bortolotti et al., 2025](#)). We do the same, but we will relax this assumption when discussing mitigation and awareness strategies in [Section 5.3](#) and [Section 5.4](#). Furthermore, we also assume that, when provided with the prior knowledge K , the inference layer β of a NeSy predictor architecture behaves exactly like the ground-truth inference layer β^* , i.e.,

$$\beta(\mathbf{c}) = \beta^*(\mathbf{c}), \forall \mathbf{c} \in C. \quad (11)$$

We denote with $\mathcal{B}$ the space of inference layers that satisfy [Eq. \(11\)](#). The choice of the NeSy predictor architecture (e.g., PNSPs or LTN) determines how the inference layer β behaves, as discussed in [Section 2.2](#). Notice that, for LTN, [Eq. \(11\)](#) is respected if all fuzzy logic operators are created from a t-norm ([van Krieken et al., 2022](#)).

A specific NeSy predictor thus amounts to a pair $(f, \beta) \in \mathcal{F} \times \mathcal{B}$,¹² its label distribution is given by:

$$p_f(Y | X; K) := (\beta \circ f)(X), \quad (12)$$

¹¹Recall that, for all $\beta \in \mathcal{B}$, we use the notation $\beta(\mathbf{c}) := \beta(\mathbb{1}\{\mathbf{C} = \mathbf{c}\})$.

¹²Unless mentioned otherwise, we present the results in the non-parametric setting where $\mathcal{F}$ contains all possible concept extractors that map inputs to concept conditional probabilities ([Bortolotti et al., 2025](#)).

For this class of models,¹³ we consider the maximum log-likelihood objective over infinite data:

$$\arg \max_{f \in \mathcal{F}} \mathbb{E}_{(x,y) \sim p^*(X,Y)} [\log p_f(y \mid x; K)] . \quad (13)$$

It turns out that this problem may not admit a unique solution, as explained next.

4.2.2 Non-identifiability of the concept extractor. We start by presenting the necessary and sufficient conditions for a NeSy predictor to achieve maximum likelihood on data and then show that this can lead to non-identifiability. Before proceeding, for a probability measure $p^*(X)$ and any two measurable functions $r(x)$ and $s(x)$ with respect to p^* , we will use the shorthand

$$r(X) = s(X) \quad (14)$$

to mean that r takes values equal to s for p^* -almost every $x \in \text{supp}(X)$.

THEOREM 4.4 (REVISITED FROM (MARCONATO ET AL., 2023A)). *Under Assumption 4.1, a NeSy predictor $(f, \beta) \in \mathcal{F} \times \mathcal{B}$ attains maximum likelihood with respect to the distribution $p^*(X, Y)$ if and only if*

$$(\beta \circ f)(X) = (\beta^* \circ f^*)(X) . \quad (15)$$

Unless otherwise stated, proofs for all results can be found in the original papers. This immediately yields the following important corollary: when the same label can be predicted using different concept vectors, **we cannot identify the concept extractor**, i.e., $f^* \in \mathcal{F}$ is not the only choice that leads to obtaining maximum likelihood for the ground-truth inference layer β^* , but there are alternatives $f \neq f^*$.

Corollary 4.5 (Non-identifiability). *Let $\Delta_Y^* \subseteq \Delta_Y$ be the subspace obtained from mapping inputs to labels, which is given by $\Delta_Y^* := (\beta^* \circ f^*)(\text{supp}(X))$, and let $\Delta_C^* \subseteq \Delta_C$ be the preimage of Δ_Y^* over the ground-truth inference layer β^* , defined as $\Delta_C^* := \{p \in \Delta_C \mid \beta^*(p) \in \Delta_Y^*\}$. Under Assumption 4.1, if β^* is not injective from Δ_C^* to Δ_Y^* , then*

$$(\beta^* \circ f)(X) = (\beta^* \circ f^*)(X) \not\Rightarrow f(X) = f^*(X) . \quad (16)$$

PROOF. By construction. If $\beta^* : \Delta_C^* \rightarrow \Delta_Y^*$ is not injective, we can find $p, p' \in \Delta_C^*$ such that $p \neq p'$ and

$$\beta^*(p) = \beta^*(p') . \quad (17)$$

Hence, for all $x \in \text{supp}(X)$ such that $f^*(x) = p$ we can construct another function $f \in \mathcal{F}$ such that $f(x) = p'$. By Eq. (17), the NeSy predictor (f, β^*) will obtain maximum likelihood (Theorem 4.4), but $f(X) \neq f^*(X)$. $\square$

Corollary 4.5 formalizes the intuition in (Marconato et al., 2023a) that, when ground-truth inference layer β^* maps multiple concept vectors to the same label, an optimal NeSy predictor can learn an RS. **The important consequence is that the concept extractor f attaining maximum likelihood is not unique**, even when using the ground-truth inference layer β^* and an infinite amount of data. In other words, **we do not always correctly ground concepts**. Motivated by these failures modes, we provide a formal definition of **shortcut NeSy predictors** that can be encountered in training.¹⁴

¹³Eq. (12) captures PNSPs, (some versions of) LTN and ABL but not the SL, as the latter does not have separate concept extraction and reasoning stages. Nevertheless, (deterministic) RSs behave essentially the same in all these models, also in practice, so we employ this equation as a useful abstraction throughout.

¹⁴This was the original definition of RSs in (Marconato et al., 2023b) with unlimited support $\text{supp}(X)$. Here, we change the original name to specialize it to NeSy predictors and distinguish it from the abstract notion of RSs, as explained in Definition 4.7.

Definition 4.6 (Shortcut NeSy predictor (Marconato et al., 2023b)). Consider a data generation process that satisfies Assumption 4.1 and let $(f^*, \beta^*) \in \mathcal{F} \times \mathcal{B}$ be the ground-truth concept extractor and inference layer. We say that $(f, \beta) \in \mathcal{F} \times \mathcal{B}$ is a **shortcut NeSy predictor** if (f, β) attains maximum likelihood on data (Theorem 4.4) and we have

$$f(\mathbf{X}) \neq f^*(\mathbf{X}). \quad (18)$$

From the definition, the set of shortcut NeSy predictors both depends on the choice of admitted concept extractors $\mathcal{F}$ and the possible inference layers $\mathcal{B}$. We will show in Section 5.3 how, by constraining the space $\mathcal{F}$ through mitigation strategies, certain shortcut NeSy predictors can vanish. From Section 2.2, each NeSy architecture specifies a particular $\beta \in \mathcal{B}$, which limits shortcut NeSy predictors to concept extractors $f \in \mathcal{F}$ for which (f, β) is faulty. This also implies that some concept extractors $f \neq f^*$ may constitute a shortcut NeSy predictor when paired with the inference layer of one model class, e.g., PNSPs, but not of another, e.g., LTN.

While Corollary 4.5 establishes non-identifiability in general, it does not provide insight on what characterizes RSs, and in particular shortcut NeSy predictors. Without this characterization, it is unclear when the ground-truth concept extractor can be identified and RSs consistently avoided. We next describe results that characterize RSs in the problem, allowing us to identify the ground-truth concept extractor f^* .

4.2.3 Concept remapping distributions. Following (Marconato et al., 2023b; Bortolotti et al., 2025), we shift the perspective from describing which concept extractors $f \in \mathcal{F}$ are valid optima of the likelihood, to studying how ground-truth concept vectors can be remapped by NeSy models while retaining optimal likelihood. Formally, under Assumption 4.1, we let $p^*(\mathbf{X} \mid \mathbf{G})$ be the posterior distribution induced by the ground-truth concept extractor f^* . Then, to describe how concepts learned by a NeSy predictor $(f, \beta) \in \mathcal{F} \times \mathcal{B}$ relate to ground-truth concepts, we define

$$\alpha_f(\mathbf{g}) := \mathbb{E}_{\mathbf{x} \sim p^*(\mathbf{X} \mid \mathbf{g})} [f(\mathbf{x})], \quad (19)$$

the *concept remapping distribution* $\alpha_f : \mathcal{C} \rightarrow \Delta_{\mathcal{C}}$ induced by f (van Krieken et al., 2025). These distributions α_f 's describe how ground-truth concept vectors are mapped to learned concept vectors by f . Hereafter, we use $\mathcal{A}$ to denote the space of these α_f 's, which is a simplex (Bortolotti et al., 2025). Furthermore, we call $\text{Vert}(\mathcal{A})$ the set of (deterministic) concept remappings: for each ground-truth vector $\mathbf{g} \in \mathcal{C}$, the output of a concept remapping $\alpha_f(\mathbf{g})$ is one-hot. That is, for such (deterministic) concept remappings α_f , there is a single concept vector $\mathbf{c} \in \mathcal{C}$ with probability one:

$$\max_{\mathbf{c} \in \mathcal{C}} \alpha_f(\mathbf{g})_{\mathbf{c}} = 1. \quad (20)$$

Any element $\alpha_f \in \mathcal{A}$ can be expressed using convex combinations of concept remappings, $\mathbf{a} \in \text{Vert}(\mathcal{A})$, that is:

$$\alpha_f(\mathbf{g}) = \sum_{\mathbf{a} \in \text{Vert}(\mathcal{A})} \lambda_{\mathbf{a}}^f \mathbf{a}(\mathbf{g}), \quad (21)$$

where we have $\lambda_{\mathbf{a}}^f \geq 0$ for all $\mathbf{a} \in \text{Vert}(\mathcal{A})$ and $\sum_{\mathbf{a} \in \text{Vert}(\mathcal{A})} \lambda_{\mathbf{a}}^f = 1$. Hence, a NeSy predictor can be seen both as a pair $(f, \beta) \in \mathcal{F} \times \mathcal{B}$ and is associated to a pair $(\alpha_f, \beta) \in \mathcal{A} \times \mathcal{B}$. Notice that, because the learnable α 's depend on the space of learnable concept extractors $\mathcal{F}$, constraining $\mathcal{F}$ also reduces the size of $\mathcal{A}$. This fact is also explicit in some mitigation strategies (see Section 5.3.10).

We depict the precise relation between the learnable concept extractors $\mathcal{F}$ and concept remapping distributions $\mathcal{A}$ in Fig. 5, where a one-to-one correspondence exists only between concept remappings $\alpha \in \text{Vert}(\mathcal{A})$ and functions $f \in \mathcal{F}$ restricted to the support of $\mathbf{X}$.¹⁵ Studying concept remapping distributions is especially helpful

¹⁵For the precise relation, refer to (Marconato et al., 2023b, Lemma 1).

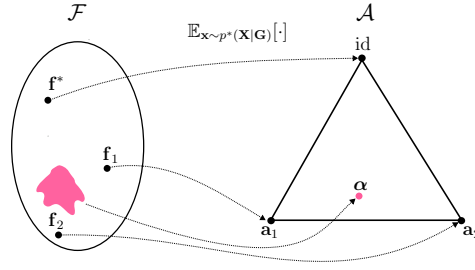


Fig. 5. We connect concept extractors $f \in \mathcal{F}$ and maps $\alpha \in \mathcal{A}$ through Eq. (19). Here, we represent the space of learnable concept extractors $\mathcal{F}$ and the space $\mathcal{A}$ as a two-dimensional simplex. The three vertices correspond to the concept remappings $\text{id}, a_1, a_2 \in \text{Vert}(\mathcal{A})$. id is the intended solution, and the remaining mappings are deterministic RSs. From (Marconato et al., 2023b, Lemma 1), any concept remapping $a \in \text{Vert}(\mathcal{A})$ is in one-to-one correspondence with a concept extractor $f \in \mathcal{F}$ with domain restricted to $\text{supp}(\mathbf{X})$ and, under Assumption 4.2, f^* is in one-to-one correspondence with the identity function. However, non-deterministic concept remapping distributions $\alpha \in \mathcal{A} \setminus \text{Vert}(\mathcal{A})$ can arise from several concept extractors in $\mathcal{F}$, here represented by the magenta colored subset mapping to one point in the simplex $\mathcal{A}$.

under Assumption 4.2. Then, the concept remapping for the ground-truth f^* is the identity function $\text{id}(\cdot)$ (Marconato et al., 2023b), i.e.,

$$\alpha_{f^*}(\mathbf{G}) = \text{id}(\mathbf{G}). \quad (22)$$

4.2.4 Reasoning shortcuts as unintended, optimal concept remapping distributions. Based on this construction, we can give a formal definition of what a reasoning shortcut is and illustrate it with an example:

Definition 4.7 (Reasoning Shortcut for β (Formal)). Let $\beta \in \mathcal{B}$ be the inference layer of a NeSy predictor (e.g., PNSPs or LTN) and $\mathcal{F}$ the space of learnable concept extractors. We say that a concept remapping distribution $\alpha \in \mathcal{A}$ is a **reasoning shortcut for the inference layer** β if (i) $\alpha \neq \text{id}$, (ii) there exists a concept extractor $f \in \mathcal{F}$ such that $\mathbb{E}_{\mathbf{x} \sim p^*(\mathbf{X}|\mathbf{G})}[f(\mathbf{x})] = \alpha(\mathbf{G})$, and (iii) we have

$$(\beta \circ \alpha)(\mathbf{g}) = \beta^*(\mathbf{g}), \forall \mathbf{g} \in \text{supp}(\mathbf{G}). \quad (23)$$

With this definition, we formalize reasoning shortcuts as concept remapping distributions that differ from the identity on at least one ground truth concept $\mathbf{g} \in \mathcal{C}$.¹⁶ There are two reasons why α might not be the identity. The first and most obvious is that we are dealing with a shortcut NeSy predictor (f, β) whose concept extractor f induces a RS $\alpha_f \neq \text{id}$ in distribution. The second and more subtle reason is that, even if the NeSy predictor is not faulty, in the sense that it grounds all concepts correctly *in distribution*, the resulting map α may still behave differently from the identity outside of the support $\text{supp}(\mathbf{G})$, that is, *out-of-distribution*. This may happen when the model fails to generalize combinatorially (Montero et al., 2021; Hwang et al., 2023; Duan et al., 2025). When the training data includes all combinations of ground-truth concepts, i.e., $\text{supp}(\mathbf{G}) = \mathcal{C}$, RSs originate exclusively from shortcut NeSy predictors. The definition of RSs allows us to focus on concept remapping distributions $\alpha \in \mathcal{A}$ instead of concept extractors, that is, on the symbolic level rather than the sub-symbolic level.

Notice that, likewise for shortcut NeSy predictors, Definition 4.7 specializes RSs to the specific inference layer β at hand. This, in turn, implies that some concept remapping distributions α 's can be an RS for one NeSy model class, say PNSPs, but not for another, say LTN. This may happen when two inference layers $\beta, \beta' \in \mathcal{B}$

¹⁶In contrast to (Marconato et al., 2023b, Definition 1), we refer to RSs as the maps $\alpha \in \mathcal{A}$ for a given $\beta \in \mathcal{B}$, not as the concept extractors $f \in \mathcal{F}$ that induce shortcut NeSy predictors.

differ in behavior when mapping non-deterministic distributions $p(C | x) \in \Delta_C \setminus \text{Vert}(\Delta_C)$. This distinction is particularly evident between probabilistic logic approaches (which essentially implement integration through marginalization, like PNSPs and SL) and fuzzy logic methods (where the inference layer depends on the fuzzy logic chosen). However, all *deterministic* concept remappings $\alpha \in \text{Vert}(\mathcal{A})$ that are RSs, are RSs for *all* inference layers. This is because, by Eq. (11), they output the same values on one-hot concept distributions $\text{Vert}(\Delta_C)$ (see also (Marconato et al., 2023b, Appendix A)). We will refer to these RSs as **deterministic RSs**.

Definition 4.8 (Deterministic Reasoning Shortcut). *We say that $\alpha \in \text{Vert}(\mathcal{A})$ is a **deterministic reasoning shortcut** if α is a reasoning shortcut for $\beta^* \in \mathcal{B}$ (Definition 4.7).*

Example 4.9. Consider a simplified version of BDD-OIA (Xu et al., 2020) with label $y \in \{\text{go}, \text{stop}\} = \mathcal{Y}$, three binary concepts ($C_{\text{green}}, C_{\text{red}}, C_{\text{ped}}$) with domain $\mathcal{C}$, and prior knowledge $K = (C_{\text{ped}} \vee C_{\text{red}} \Leftrightarrow (Y = \text{stop})) \wedge (C_{\text{green}} \wedge \neg(Y = \text{stop}) \Leftrightarrow (Y = \text{go}))$. In Fig. 6, we visualize the four concept remapping distributions appearing in Fig. 5. The concept remapping distributions $\alpha \in \mathcal{A}$ determine how ground-truth concepts are mapped to conditional probability distributions on model concepts. First, we can find an RS (a_1 in Fig. 6 left-center) that predicts $C_{\text{ped}} = 1$ with probability one whenever $G_{\text{red}} = 1$ or $G_{\text{ped}} = 1$. Similarly, we have an RS (a_2 in Fig. 6 right-center) that predicts $C_{\text{red}} = 1$ with probability one whenever $G_{\text{red}} = 1$ or $G_{\text{ped}} = 1$. A nondeterministic RS (α in the rightmost part of Fig. 6) assigns non-zero probability to both $C_{\text{red}} = 1$ and $C_{\text{ped}} = 1$.

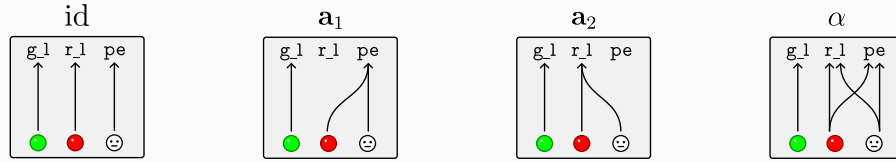


Fig. 6. The four concept remapping distributions appearing in Fig. 5 according to Example 4.9. To aid visualization, we display the concept remapping distributions as directed graphs where arrows denote how ground-truth concepts are mapped to the predicted concepts. On the left, id corresponds to the ground-truth map α_{f^*} from the data generation process. At the center, $a_1, a_2 \in \text{Vert}(\mathcal{A})$ represent two deterministic RSs. On the right, $\alpha \in \mathcal{A} \setminus \text{Vert}(\mathcal{A})$ is a non-deterministic RS. Among these RSs, C_{green} is always predicted correctly, whereas C_{red} and C_{ped} are incorrectly predicted.

4.2.5 Characterization of deterministic reasoning shortcuts. The description using maps $\alpha \in \mathcal{A}$ is useful because we can count the number of deterministic RSs explicitly. As noted above, under Assumption 4.2, the ground-truth concept remapping α_{f^*} is the identity $\text{id}(\cdot)$. Intuitively, any learnable $\alpha \in \text{Vert}(\mathcal{A})$ that differs from the identity function is a deterministic RS. The following result makes the number of deterministic RSs explicit:

THEOREM 4.10 (NUMBER OF DETERMINISTIC REASONING SHORTCUTS (MARCONATO ET AL., 2023B)). *Let $\mathcal{F}$ be a space of concept extractors. Under Assumptions 4.1 to 4.3, the number of deterministic RSs (Definition 4.7) is*

$$\sum_{\alpha \in \text{Vert}(\mathcal{A}_{\mathcal{F}})} \mathbb{1} \left\{ \bigwedge_{g \in \text{supp}(G)} (\beta^* \circ \alpha)(g) = \beta^*(g) \right\} - 1, \quad (24)$$

where $\mathcal{A}_{\mathcal{F}} := \{\alpha \in \mathcal{A} : \exists f \in \mathcal{F} \text{ s.t. } \alpha = \alpha_f\}$.

When this count is greater than 1, the NeSy task admits (deterministic) RSs. A NeSy predictor can attain high label accuracy by learning any one of them or any (non-deterministic) mixture thereof. Because the count does not depend on the particular choice of $\beta \in \mathcal{B}$, **all NeSy architectures have the same deterministic reasoning shortcuts** (Marconato et al., 2023b). This happens because all NeSy inference layers $\beta \in \mathcal{B}$ integrating the prior knowledge K satisfy Eq. (11).

When $\text{Vert}(\mathcal{A}_{\mathcal{F}})$ contains all maps, i.e., $\text{Vert}(\mathcal{A}_{\mathcal{F}}) = \text{Vert}(\mathcal{A})$, and $\text{supp}(\mathbf{G})$ is complete, i.e., $\text{supp}(\mathbf{G}) = \mathcal{C}$, it is possible to explicitly count the number of deterministic reasoning shortcuts in the task, see (Marconato et al., 2023b, Appendix C). The number of deterministic RSs rapidly explodes based on how many concept vectors correspond to a single label through β . For example, if we can predict a label y_0 from M different concept vectors, the number multiplies by a factor of M^M . Accounting for practical restrictions in the set $\text{Vert}(\mathcal{A}_{\mathcal{F}}) \subseteq \text{Vert}(\mathcal{A})$ due to, e.g., the choice of specific neural backbones, complicates the explicit counting of RSs, but this can be done using model counting, see Section 5.2.1. We illustrate next how Eq. (24) depends on the quantities inside the sum, but defer the in-depth treatment of causes and mitigations of RSs to Section 5.3.

Example 4.11 (Deterministic Reasoning Shortcuts in XOR (Marconato et al., 2023b; van Krieken et al., 2025)). Consider two binary concepts $G_1, G_2 \in \{0, 1\}$ that both generate an input $\mathbf{x} \sim p(\mathbf{X} \mid G_1, G_2)$, which can be thought as the image visualization of the values (e.g., $\mathbf{g} = (1, 0)$ corresponds to $\mathbf{x} = \begin{bmatrix} 1 \\ 0 \end{bmatrix}$). The label Y results from the XOR operation between the binary concepts: $p^*(Y \mid \mathbf{G}; K) = \mathbb{1}\{Y = G_1 \oplus G_2\}$. The deterministic maps in $\text{Vert}(\mathcal{A})$ can be seen as all maps $\alpha : \{0, 1\}^2 \rightarrow \{0, 1\}^2$, and so $|\text{Vert}(\mathcal{A})| = 4^4$. We analyze how the RS count in Eq. (24) varies when: (1) the support of ground-truth concepts is complete and all deterministic α 's are allowed; (2) we impose a constraint on deterministic α 's.

Case 1. When $\mathcal{A}_{\mathcal{F}} = \mathcal{A}$ and the support of ground-truth concepts in the training data is exhaustive, i.e., $\text{supp}(\mathbf{G}) = \{0, 1\}^2$, we obtain a total of $2^2 \cdot 2^2 - 1$ deterministic RSs. To show this, we construct one where

$$\alpha(0, 0) = (0, 0), \quad \alpha(0, 1) = (0, 1), \quad \alpha(1, 0) = (0, 1), \quad \alpha(1, 1) = (0, 0). \quad (25)$$

This shows that, for $b_0, b_1 \in \{0, 1\}$, we can choose between mapping both $\{(0, 0), (1, 1)\} \mapsto (b_0, b_0)$ and $\{(0, 1), (1, 0)\} \mapsto (1 - b_1, b_1)$ so that we get valid label predictions $(\beta^* \circ \alpha)(G_1, G_2) = \beta(G_1, G_2)$. Another deterministic RS flips the values of the bits, resulting in

$$\alpha(0, 0) = (1, 1), \quad \alpha(0, 1) = (1, 0), \quad \alpha(1, 0) = (0, 1), \quad \alpha(1, 1) = (0, 0). \quad (26)$$

Denote with $C_0 = \{(0, 0), (1, 1)\}$ and $C_1 = \{(0, 1), (1, 0)\}$. We consider maps α where for all $\mathbf{c} \in C_0$, $\alpha(\mathbf{c}) \in C_0$, so there are two valid options for each element of C_0 . The same holds for C_1 . Now each map α that satisfies these constraints is a valid map. There are then $|C_0|^{|C_0|} \cdot |C_1|^{|C_1|} = 2^2 \cdot 2^2 = 16$ valid α 's, and all of them are deterministic RSs except the identity.

Case 2. We consider full support of ground-truth concepts in the training data, i.e., $\text{supp}(\mathbf{G}) = \{0, 1\}^2$, but limit the maps α 's to be factorized, that is $\alpha(G_1, G_2) = (\hat{\alpha}(G_1), \hat{\alpha}(G_2))$. This means that $\text{Vert}(\mathcal{A}_{\mathcal{F}}) \subset \text{Vert}(\mathcal{A})$, and in total there are $2^2 \cdot 2^2$ deterministic α 's in $\text{Vert}(\mathcal{A}_{\mathcal{F}})$.^a Because each $\hat{\alpha}$ depends only on the binary value of the i -th concept, i.e., $\hat{\alpha} : \{0, 1\} \rightarrow \{0, 1\}$, we only have one deterministic RS resulting from $\hat{\alpha}(b) = 1 - b$ for $b \in \{0, 1\}$. This gives the map

$$\begin{aligned} \alpha(0, 0) &= (\hat{\alpha}(0), \hat{\alpha}(0)) = (1, 1), & \alpha(0, 1) &= (\hat{\alpha}(0), \hat{\alpha}(1)) = (1, 0), \\ \alpha(1, 0) &= (\hat{\alpha}(1), \hat{\alpha}(0)) = (0, 1), & \alpha(1, 1) &= (\hat{\alpha}(1), \hat{\alpha}(1)) = (0, 0). \end{aligned} \quad (27)$$

In fact, one can check that other combinations with $\hat{\alpha}(G_i) = b$ with b a binary constant, are not valid α 's.

^aWe later discuss in Section 4.3 and in Section 5.3.10 that this can result from concept extractors being architecturally disentangled.

Theorem 4.10 only counts *deterministic* RSs. If the count is reduced to zero, then NeSy predictors may still suffer from *non-deterministic* RSs (those α 's in $\mathcal{A} \setminus \text{Vert}(\mathcal{A}_{\mathcal{F}})$). To ensure that we can deal with *all* RSs by avoiding deterministic RSs, we need to introduce an additional assumption on the inference layer β :

Assumption 4.12 (Extremality (Bortolotti et al., 2025)). *The inference layer $\beta \in \mathcal{B}$ satisfies extremality if, for all $\lambda \in (0, 1)$ and for all $\mathbf{c} \neq \mathbf{c}' \in C$ such that $\arg\max_{y \in \mathcal{Y}} \beta(\mathbf{c})_y \neq \arg\max_{y \in \mathcal{Y}} \beta(\mathbf{c}')_y$, where $\beta(\cdot)_y$ is the y -th component of β , we have*

$$\max_{y \in \mathcal{Y}} \beta(\lambda \mathbb{1}\{C = \mathbf{c}\} + (1 - \lambda) \mathbb{1}\{C = \mathbf{c}'\})_y < \max \left(\max_{y \in \mathcal{Y}} \beta(\mathbf{c})_y, \max_{y \in \mathcal{Y}} \beta(\mathbf{c}')_y \right).$$

A NeSy predictor satisfies **Assumption 4.12** if its inference layer β is “peaked”: for any two concept vectors $\mathbf{c}$ and $\mathbf{c}'$ with different predictions, the label probability given by β for a *mixture* of these predictions is no larger than the label probability it associates to $\mathbf{c}$ and $\mathbf{c}'$. This happens in PNSPs on MNIST-Add, where the knowledge specifies that $\beta : (C = (0, 1)^\top) \mapsto (Y = 1)$ and $\beta : (C = (0, 2)^\top) \mapsto (Y = 2)$, both with probability one. Any convex combination crossing $\mathbf{c}' = (0, 1)^\top$ to $\mathbf{c} = (0, 2)^\top$ would not result in any label with probability one, see Fig. 7. E.g., PNSPs will assign a higher probability to the sum $Y = 1$ when $\lambda < 0.5$. In essence, distributing probability mass across different sets of concepts does not increase label likelihood. Extremality holds for PNSPs and SL, and holds for LTN when using fuzzy logic operators based on t-norms and t-conorms (van Krieken et al., 2022). An investigation for which models **Assumption 4.12** holds is reported in (Bortolotti et al., 2025, Appendix E).

Under **Assumption 4.12**, the next result holds:

THEOREM 4.13 (IDENTIFIABILITY (BORTOLOTTI ET AL., 2025)). *Under Assumptions 4.1 to 4.3 and for all $\beta \in \mathcal{B}$ satisfying Assumption 4.12, if the count of deterministic RSs is zero (Eq. (24)), then any NeSy predictor $(f, \beta) \in \mathcal{F} \times \mathcal{B}$ attaining maximum likelihood (Theorem 4.4) has $\alpha_f(\mathbf{G}) = \text{id}(\mathbf{G})$ and so finds the ground-truth concept extractor, i.e., $f(\mathbf{X}) = f^*(\mathbf{X})$.*

This result highlights that *extremality reduces the space of equivalent optima* by preventing different concept assignments from being indistinguishable at the label level. It indicates we can mitigate RSs for all NeSy predictors that obey **Assumption 4.12**. In fact, zeroing the count on deterministic RSs (**Theorem 4.10**) implies avoiding all RSs in the learning task and, in particular, avoiding shortcut NeSy predictors. For NeSy predictors model classes obeying **Assumption 4.12**, the identifiability result gives that, in the absence of deterministic RSs, **maximizing the likelihood on labels permits the correct grounding of the concepts**. This, in turn, presupposes finding additional mitigations that, paired with maximum-likelihood Eq. (2), guarantee the *updated count* of deterministic reasoning shortcuts equals zero. We will discuss next (Section 5.3) how different mitigations act on the count reduction.

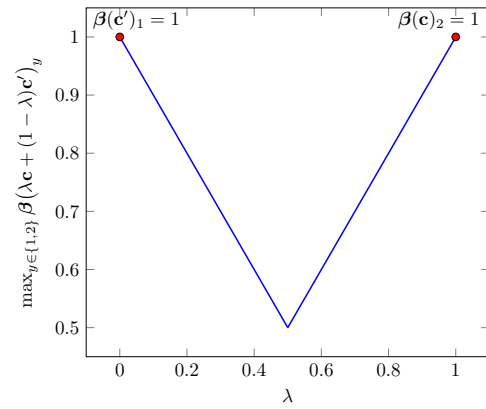


Fig. 7. The extremality condition holds for PNSPs and SL.

Remark 4.14. As a corollary of [Theorem 4.13](#), for probabilistic logic methods (including PNSPs like DPL ([Manhaeve et al., 2018](#)), SL ([Xu et al., 2018](#)), and SPL ([Ahmed et al., 2022](#))), all non-deterministic RSs can be constructed from convex combinations of deterministic RSs ([Marconato et al., 2023b](#), Proposition 3). This guarantees that if there are no deterministic RSs, then there are also no other RSs. Moreover, as will be discussed in [Section 5.4](#), it is possible to leverage this geometrical aspect to construct non-deterministic RSs from convex combinations of deterministic ones, whose uncertainty on (some of) the learned concepts reveals which of them are affected by RSs.

4.3 Perspective: Statistical Learning

Learning a NeSy model amounts to learning (the parameters of) a concept extractor that—once combined with the inference layer—attains high label likelihood. From a statistical learning perspective, this is equivalent to finding a model with minimal empirical *label* risk. Statistical learning theory allows one to bound the true risk of a model—that is, to assess how well it generalizes—based on its observed empirical risk and other factors, such as training set size and model complexity ([Valiant, 1984](#); [Vapnik, 2013](#); [Shalev-Shwartz and Ben-David, 2014](#)). Since the inference layer of a NeSy model is fixed and supplied upfront via the background knowledge K , it is natural to ask when and how minimizing the *label* risk causes a reduction of the *concept* risk. Or similarly, when does low risk on label predictions bound the mismatch between the predicted and the ground-truth concepts?

We expect that if RSs are present, such a bound could be vacuous for some concepts. In fact, from the results in the previous section, we know that the concepts predicted by a shortcut NeSy predictor can drastically differ from ground-truth ones, even when the model is trained with an infinite amount of data. To make this explicit, we follow the analysis of ([Yang et al., 2024](#)), which indicates from a statistical learning perspective the relation between label and concept risks. The subsequent study of [He and Li \(2025\)](#) follows similar steps.

4.3.1 Knowledge complexity. The statistical learning perspective on RSs follows steps similar to identifiability analysis ([Section 4.2](#)). First, data are assumed to be generated as per [Assumption 4.1](#), whereby the ground-truth concepts and labels are sampled from the ground-truth concept extractor $f^* : \mathcal{X} \rightarrow \Delta_C$ and inference function $\beta^* : \Delta_C \rightarrow \Delta_Y$. From the joint distribution $p^*(\mathbf{X}, \mathbf{G}, \mathbf{Y})$ of [Eq. \(10\)](#), it is possible to determine the marginal distribution over labels $p^*(\mathbf{Y})$. This marginal distribution is used together with the knowledge K to determine the complexity of the knowledge:

Definition 4.15 (Knowledge complexity ([Yang et al., 2024](#))). *The knowledge complexity of K for labels distributed according to the label distribution $p^*(\mathbf{Y})$ ([Assumption 4.1](#)) is defined as*

$$\text{KC}(K; p^*) := \mathbb{E}_{\mathbf{Y} \sim p^*(\mathbf{Y})} \left[\sum_{\mathbf{c} \in C} \mathbb{1}\{(\mathbf{c}, \mathbf{y}) \not\models K\} \right]. \quad (28)$$

This quantity gives a global description of how many concepts are not compatible with the knowledge, based on the labels $\mathbf{y} \in \text{supp}(\mathbf{Y})$ observed during training. Intuitively, when only few concepts allow inferring any given label vector, the knowledge complexity increases, indicating that the knowledge is more explicit about what concept vectors lead to that label vector. On the other hand, a low knowledge complexity indicates that many concept vectors are mapped to same label vector.

For instance, in the BDD-OIA example ([Example 2.1](#)) with three binary concepts ($|C| = 8$), go can only be inferred from $C_{\text{green}} = 1$, $C_{\text{ped}} = 0$, and $C_{\text{red}} = 0$, meaning that 1 out of 8 concept vectors $\mathbf{c}$ entail go and 7 out of 8 entail stop. Hence, whenever both go and stop instances are observed, we get $\text{KC}(K; p^*) < |C| - 1 = 7$. Notice that, by construction, the complexity of any knowledge K on any distribution $p^*(\mathbf{Y})$ over labels is upper-bounded

by

$$\text{KC}(\mathbf{K}; p^*) \leq |\mathcal{C}| - 1.^{17} \quad (29)$$

Next, we show how the complexity of the knowledge is intimately connected to the emergence of RSs.

4.3.2 Impossibility result in bounding the reasoning shortcut risk. We start by recalling the relevant notation introduced in [Section 4.2](#). For a NeSy predictor $(f, \beta) \in \mathcal{F} \times \mathcal{B}$, the corresponding concept and label distributions are given by

$$p_f(\mathbf{C} \mid \mathbf{X}) := f(\mathbf{X}), \quad p_f(\mathbf{Y} \mid \mathbf{X}; \mathbf{K}) := (\beta \circ f)(\mathbf{X}). \quad (30)$$

We assume, as before, that the space of learnable concept extractors $\mathcal{F}$ is unrestricted and that the inference layer β matches the ground-truth inference layer at the vertices ([Eq. \(11\)](#)). The expected label risk and expected concept risk are given by

$$\mathcal{R}_Y(f) := -\mathbb{E}_{(\mathbf{x}, \mathbf{y}) \sim p^*(\mathbf{X}, \mathbf{Y})} [\log p_f(\mathbf{y} \mid \mathbf{x}; \mathbf{K})], \quad \mathcal{R}_C(f) := -\mathbb{E}_{(\mathbf{x}, \mathbf{g}) \sim p^*(\mathbf{X}, \mathbf{G})} [\log p_f(\mathbf{C} = \mathbf{g} \mid \mathbf{x})], \quad (31)$$

respectively. Their difference gives the expected *reasoning shortcut risk* ([Yang et al., 2024](#)):

$$\mathcal{R}_{RS}(f) := \mathcal{R}_C(f) - \mathcal{R}_Y(f). \quad (32)$$

Notice that when the RS risk is positive, concept risk exceeds label risk; hence, despite more accurate label predictions, concepts are less accurately predicted. Conversely, a concept risk less than or equal to zero indicates that concept predictions are at least as accurate as label predictions. The next result connects RSs to the *knowledge complexity* showing that, when it is not high enough, an analogous result to [Corollary 4.5](#) holds:

THEOREM 4.16 (UNBOUNDED RSs RISK ([YANG ET AL., 2024](#))). *Under [Assumption 4.1](#), if the knowledge complexity is bounded by the size of the concept space as $\text{KC}(\mathbf{K}; p^*) < |\mathcal{C}| - 1$, then there is a concept extractor f such that $\mathcal{R}_Y = 0$ and $\mathcal{R}_C \rightarrow \infty$, so that the RS risk approaches infinity, i.e.,*

$$\mathcal{R}_{RS} \rightarrow \infty. \quad (33)$$

[Theorem 4.16](#) shows that minimizing the label loss alone does not appropriately guide concept learning. If the inference layer implementing knowledge $\mathbf{K}$ maps many concept vectors to the same label, a model can achieve perfect label accuracy while predicting completely wrong concepts. In this case the concept risk can grow arbitrarily large even though the label risk is zero. Notice that the main causes for this behavior are that: (1) the knowledge complexity is not maximal, and therefore each label $\mathbf{y} \in \text{supp}(\mathbf{Y})$ is associated to multiple concept vectors $\mathbf{c} \in \mathcal{C}$, (2) the fact that the space of learnable concept extractors is unrestricted and can express all possible conditional concept distributions. Both are natural targets for mitigation. For instance, one can increase knowledge complexity by introducing more specialized constraints into the knowledge, whenever feasible. Alternatively, one can shrink the space of learnable concept extractors $\mathcal{F}$ by constraining the extractor's architecture. We will discuss these and strategies further in [Section 5.3](#).

Relation to ([He and Li, 2025](#)). [He and Li \(2025\)](#) provide a complementary learnability analysis of neuro-symbolic learning through the notion of a derived constraint satisfaction problem (DCSP). Their identification result (ref. to [He and Li, 2025](#), Lemma 4.7) shows that a NeSy task is learnable if and only if the associated DCSP admits a unique solution; when multiple solutions satisfy the constraints, the ground-truth concepts cannot be identified from label supervision alone, even with infinite data. This procedure resembles the RSs count in [Eq. \(24\)](#), although their enumeration focuses only on those RSs that affect in-distribution concept risk. Moreover, for general NeSy tasks they show that the asymptotic concept error is controlled by the disagreement among DCSP solutions. From

¹⁷ Notice that, for corner cases in which observed labels in $\text{supp}(\mathbf{Y})$ cannot be predicted from the concepts, i.e., $(\mathbf{c}, \mathbf{y}) \not\models \mathbf{K}$ for all $\mathbf{c} \in \mathcal{C}$, we get $\text{KC}(\mathbf{K}; p^*) = |\mathcal{C}|$. We systematically avoid it through [Assumption 4.1](#), which ensures that concepts are predictive of the observed labels.

the perspective of reasoning shortcuts, while Yang et al. (2024) emphasize the mismatch between label risk and concept risk, He and Li (2025) characterize this phenomenon in terms of the solutions to the DCSP brought by the knowledge and derive corresponding learnability guarantees. Notice that the statistical view allows to derive bounds on the reasoning shortcut risk, which has been done for the case of concept smoothing and pretraining of the concept extractor. Overall, the goal is to provide guarantees about avoiding learning shortcut NeSy predictors. We refer the reader to (Yang et al., 2024; He and Li, 2025) for more details on these bounds. A closely related constraint-based perspective is developed by Takemura et al. (2026), who formalize reasoning shortcuts as a constraint satisfaction problem and ask when the constraints uniquely determine the intended concept mapping. They identify a *discrimination property* — no valid mapping can be turned into another by swapping two concept values — which is necessary but not sufficient for shortcut-freeness.

4.3.3 PAC learnability with factorized concept extractors. Note that Theorem 4.16 implicitly assumes the hypothesis space $\mathcal{F}$ is very expressive. A natural question is whether this negative result can be worked around if we restrict $\mathcal{F}$. In many NeSy tasks, inputs naturally decompose into independent units (e.g., digits in MNIST-Add (Manhaeve et al., 2018)). This motivates studying factorized concept extractors, where each input component is processed independently.

The work by Wang et al. (2023) studies one special case which is often available in NeSy tasks: when the input can be naturally decomposed into a tuple of independent units, e.g., the digits in MNIST-Add, and the same concept extractor can be employed to predict the concepts of each input instance, e.g., the value of each digit is obtained by mapping the input alone, what conditions are required to avoid RSs and guarantee correct learning of the concept extractor? This setting, which builds on factorized concept extractors of the form $f = \hat{f}_1 \times \dots \times \hat{f}_k$, has been studied via the *probably approximate correct* (PAC) framework (Valiant, 1984; Shalev-Shwartz and Ben-David, 2014), and connects the *learnability* of NeSy predictors to that of weakly-supervised models (Steinhardt and Liang, 2015; Raghunathan et al., 2016).

4.3.4 Setting. Following (Wang et al., 2023), we consider training samples $(\mathbf{x}_i, y_i) \in \mathcal{X} \times \mathcal{Y}$ with a single categorical label y (instead of a vector label $\mathbf{y}$ used in the previous sections) and where each input $\mathbf{x}$ is a sequence of k vectors $\mathbf{x} = (\bar{\mathbf{x}}_1, \dots, \bar{\mathbf{x}}_k) \in \mathcal{X} = \bar{\mathcal{X}}^k$, with shared domain $\bar{\mathcal{X}} \subseteq \mathbb{R}^n$. Similarly, the concept vectors $\mathbf{c} = (c_1, \dots, c_k) \in \mathcal{C}$ concatenates k copies of the same concept $\bar{\mathcal{C}}$, one for each element of the input, with overall domain $\mathcal{C} = \bar{\mathcal{C}}^k$. The simplest way to visualize this setup is to think of MNIST-Add (Example 3.2): the elements $\bar{\mathbf{x}}_i$ would be the individual MNIST digits, $\bar{\mathcal{C}}$ any digit, and y their sum.

Next, we have to fix and construct the family of learnable concept extractors $f \in \mathcal{F}$. We do so by imposing that f extracts the k concepts by applying the same function $\bar{f} : \bar{\mathcal{X}} \rightarrow \Delta_{\bar{\mathcal{C}}} \in \bar{\mathcal{F}}$ to each of them in turn, that is,

$$f(\mathbf{x}) = (\bar{f}(\bar{\mathbf{x}}_1), \dots, \bar{f}(\bar{\mathbf{x}}_k)), \quad (34)$$

In MNIST-Add, this would equate to processing each digit *independently from the others* with the same function $\bar{f}$. Overall, this construction implies that $\mathcal{F}$ *factorizes* as $\mathcal{F} = \bar{\mathcal{F}}^k$ or, equivalently, that f is *architecturally disentangled*, as we will discuss in Section 5.3.10. This is in contrast with Corollary 4.5 and Theorem 4.16, which neither assume the concept extractor is architecturally disentangled nor that the predicted concepts are independent.

Recall that f and $\bar{f}$ output *distributions* over concept values, e.g., distributions over digits. We need an easy way to refer to their *most likely values*, e.g., the predicted digits themselves. For this purpose, we introduce the shorthand $[\bar{f}](\bar{\mathbf{x}}) := \operatorname{argmax}_{c \in \bar{\mathcal{C}}} \bar{f}(\bar{\mathbf{x}})_c$ to denote their predictions. By construction, we have that $[\bar{f}] : \bar{\mathcal{X}} \rightarrow \bar{\mathcal{C}}$, and we denote with $[\bar{\mathcal{F}}]$ the space of these functions.

As customary in statistical learning, we are interested in bounding the *risk* of the model. Wang et al. (2023) do so for the risk induced by the zero-one losses over labels and concepts, which we introduce next. For the concepts,

the *zero-one loss* is $\ell_C^{01}(c, c') := \mathbb{1}\{c \neq c'\}$ for any two values $c, c' \in \bar{C}$ and therefore the expected concept risk is:

$$\mathcal{R}_C^{01}(f) := \mathbb{E}_{(\mathbf{x}, \mathbf{g}) \sim p^*(\mathbf{X}, \mathbf{G})} \left[\frac{1}{k} \sum_{j=1}^k \ell_C^{01}([\bar{f}](\bar{\mathbf{x}}_j), g_j) \right]. \quad (35)$$

Here, $p^*(\mathbf{X}, \mathbf{G})$ is the ground-truth distribution of inputs and concept vectors. For the labels, the zero-one loss (also known as *zero-one partial loss* (Wang et al., 2023)) is simply $\ell_{\beta^*}^{01}(\mathbf{c}, y) := \mathbb{1}\{\beta^*(\mathbf{c}) \neq y\}$, where $\beta^* \in \mathcal{B}$ is the ground-truth inference layer. It follows that the expected label risk is:

$$\mathcal{R}_Y^{01}(f, \beta^*) := \mathbb{E}_{(\mathbf{x}, y) \sim p^*(\mathbf{X}, Y)} \left[\ell_{\beta^*}^{01}([\bar{f}](\bar{\mathbf{x}}_1), \dots, [\bar{f}](\bar{\mathbf{x}}_k), y) \right]. \quad (36)$$

Here, $p^*(\mathbf{X}, Y)$ is the ground-truth distribution that the training data $(\mathbf{x}_i, y_i) \in \mathcal{X} \times \mathcal{Y}$ is sampled from. (This setup is reminiscent of Section 4.3.)

We say that $\mathcal{F}$ is *realizable* when there exists a concept extractor $f^* \in \mathcal{F}$ that achieves zero risk, that is, $\mathcal{R}_Y^{01}(f^*; \beta^*) = 0$. As customary in *empirical risk minimization*, we train the concept extractor using a training set $\mathcal{D}_{p^*}$ with m_{p^*} many examples $(\mathbf{x}, y)$ sampled from $p^*(\mathbf{X}, Y)$. Hence, the associated *empirical risk* is:

$$\hat{\mathcal{R}}_Y^{01}(f, \beta^*, \mathcal{D}_{p^*}) := \frac{1}{m_{p^*}} \sum_{(\mathbf{x}, y) \in \mathcal{D}_{p^*}} \ell_{\beta^*}^{01}([\bar{f}](\bar{\mathbf{x}}_1), \dots, [\bar{f}](\bar{\mathbf{x}}_k), y). \quad (37)$$

PAC learnability of this NeSy task (Shalev-Shwartz and Ben-David, 2014; Wang et al., 2023) asks whether there exists a large enough training set that allows one to learn a concept extractor with low enough concept risk with high enough probability. Formally, we ask if, for any distribution $p^*(\mathbf{X}, Y)$ and any ϵ and $\delta \in (0, 1)$, the *zero-one concept risk* of the empirical risk minimizer f is less than or equal to ϵ (i.e., $\mathcal{R}_C^{01}(f) \leq \epsilon$) with probability at least $1 - \delta$ when the number of training examples m_{p^*} is at least $m_{\epsilon, \delta}$. Next, we present how Wang et al. (2023) derive such a m_{p^*} for this learning setting.

4.3.5 PAC learnability with k -unambiguous inference layers. To prove the PAC learnability of NeSy tasks, we must bound the concept risk $\mathcal{R}_C^{01}(f)$ by the label risk $\mathcal{R}_Y^{01}(f; \beta^*)$ for any distribution of the training data. Below, we report the conditions that β^* has to satisfy to ensure that this is the case. We already know that if the β^* allows to infer the same label from multiple concept vectors, then label supervision is not enough to recover the ground-truth concepts because the model cannot distinguish between correct and incorrect concepts based on labels alone. This holds even if the concept extractor is factorized (Eq. (34)). The next condition formalizes when the inference layer avoids this ambiguity by focusing on diagonal vectors $\mathbf{c} = (c, \dots, c)$:

Definition 4.17 (k -unambiguity (Wang et al., 2023)). For k -dimensional concept vectors $\mathbf{c} \in C$, an inference layer β^* is **k -unambiguous** if for any two diagonal concept vectors $\mathbf{c} = (c, \dots, c)$ and $\mathbf{c}' = (c', \dots, c') \in C$, such that $c \neq c'$, we have that $\beta^*(\mathbf{c}') \neq \beta^*(\mathbf{c})$. Otherwise, the inference layer β^* is **k -ambiguous**.

This definition essentially guarantees that passing diagonal vectors of the form $\mathbf{c} = c(1, \dots, 1)^\top$ and $\mathbf{c}' = c'(1, \dots, 1)^\top$ to β^* never returns the same label predictions for $c \neq c'$, guaranteeing the injectivity of β^* for these pairs. Although Definition 4.17 might seem strict, NeSy tasks like MNIST-Add satisfy this condition, see Example 3.2. Moreover, it is possible to extend k -unambiguity to the case where β^* is non-deterministic and to establish the relationship with small ambiguity degree from partial label learning (Liu and Dietterich, 2014). With this condition, the following result holds:

Lemma 4.18 (Risk bounds under k -unambiguity (Wang et al., 2023)). *If the space of learnable concept extractors $\mathcal{F}$ is constructed like in Eq. (34) and β^* is k -unambiguous, then we have:*

$$\mathcal{R}_C^{01}(f) \leq O(\mathcal{R}_Y^{01}(f; \beta^*)^{1/k}) \quad \text{as} \quad \mathcal{R}_Y^{01}(f; \beta^*) \rightarrow 0 \quad (38)$$

Moreover, if β^ is k -ambiguous, then a concept extractor f attaining label risk $\mathcal{R}_Y^{01}(f; \beta^*) = 0$ can have a concept risk of $\mathcal{R}_C^{01}(f) = 1$.*

Lemma 4.18 shows that under k -unambiguity, minimizing label risk also drives the concept extractor toward correct concept predictions. Hence, k -unambiguity is central to avoiding RSs for the specific concept extractors modeled in $\mathcal{F}$. In fact, whenever β^* is k -ambiguous, concept risk can be non-zero, resulting in RSs. This happens naturally for the XOR of two bits, but not for the sum of the two, as detailed next.

Example 4.19. *Consider a setting where the input $\mathbf{x} = (\bar{x}_1, \bar{x}_2)$ comprises two bits, e.g., $\mathbf{x} = (\text{0}, \text{1})$. Consider the knowledge for the XOR between the two, i.e., $K = (C_1 \oplus C_2 \Leftrightarrow Y)$. The function β^* constructed from K is k -ambiguous ($k = 2$), as there are only two vectors with the same concept value that return the label. That is, $\beta^*((0, 0)^\top) = \beta^*((1, 1)^\top)$, since $0 \oplus 0 = 1 \oplus 1 = 0$. We can see that, in this setting, there is a deterministic RS that confuses the zero for one and vice versa, that is, $\text{0} \mapsto 1$ and $\text{1} \mapsto 0$. Consider now, instead, the sum of the two digits, with the knowledge given by $K = (Y = C_1 + C_2)$. Here, the inference layer β^* is k -unambiguous ($k = 2$), and there are no concept remapping distributions other than the identity that give the correct label prediction. k -unambiguity is also satisfied for the complete MNIST-Add, which is a NeSy task that does not have any RS for the disentangled concept extractors in $\mathcal{F}$.*

Based on this result, Wang et al. (2023) show that it is possible to guarantee PAC learnability of the NeSy task. The result incorporates the Natarajan dimension of the space of classifiers $[\tilde{f}]$, denoted as $d_{[\tilde{f}]}$. We refer readers to (Shalev-Shwartz and Ben-David, 2014) for the explicit definition of the Natarajan dimension.

THEOREM 4.20 (ERM LEARNABILITY UNDER k -UNAMBIGUITY (WANG ET AL., 2023)). *Let $\mathcal{F}$ be the space of learnable concept extractors $\mathcal{F}$ constructed like in Eq. (34). Suppose that $\mathcal{F}$ is realizable under ℓ_Y^{01} and $[\tilde{f}]$ has a finite Natarajan dimension $d_{[\tilde{f}]}$. Then, for any $\epsilon, \delta \in (0, 1)$, there exists a universal constant $C_0 > 0$, such that with probability at least $1 - \delta$, the empirical label risk minimizer f on the dataset $\mathcal{T}_{p^*}$, such that $\hat{\mathcal{R}}_Y^{01}(f; \beta^*; \mathcal{T}_{p^*}) = 0$, has a concept risk $\mathcal{R}_C^{01}(f) < \epsilon$, if*

$$m_{p^*} \geq C_0 \frac{\eta_1}{\epsilon^k} \left(\eta_2 \log \frac{1}{\epsilon^k} + \log \frac{1}{\delta} + \eta_3 \right), \quad (39)$$

where $\eta_1, \eta_2, \eta_3 \in \mathbb{R}^+$ are constants that depend on k , the $d_{[\tilde{f}]}$, and the number of labels $|\mathcal{Y}|$.

This shows that if the inference layer is k -unambiguous and the concept extractor class has finite complexity (as determined by the Natarajan dimension), then minimizing empirical label risk is sufficient to reduce the concept risk with high probability. We refer the reader to (Wang et al., 2023) for the explicit values of the constants. Furthermore, the work presents other results regarding the PAC learnability of the NeSy task under unknown β^* 's and error bounds when considering the semantic loss (Xu et al., 2018). As a final note, we remark that the bound in Eq. (39) holds only in the setting of factorized concept extractor and generalizing it to wider hypothesis spaces $\mathcal{F}$ is open.

4.4 Relationship between Theories

A common trait between the two perspectives can be found in [Corollary 4.5](#) and [Theorem 4.16](#): both results indicate that, even in the limit of infinite data, when conditions on the knowledge are not met (either because of non-injectivity or low knowledge complexity), RSs can arise during training. Both results highlight a general problem in NeSy predictors: when we have a large hypothesis space of learnable concept extractors $\mathcal{F}$ and if the knowledge admits multiple solutions for label predictions, *there is no unique and correct grounding of the concepts*. Conversely, if a one-to-one mapping between concepts and labels is implemented by the knowledge (either because the ground-truth inference layer β^* is injective in [Corollary 4.5](#) and $KC(K, p^*) = |C| - 1$ in [Theorem 4.16](#)), achieving optimal likelihood on data implies learning well-grounded concepts. Achieving this condition, however, may be difficult in practice whenever the number of possible concept vectors exceeds that of labels, see e.g. BDD-OIA in [Example 4.9](#). Constrained hypothesis spaces for the concept extractor are more likely to meet the conditions for avoiding RSs with a less restricted K , as shown in [Lemma 4.18](#), and lead to statistical learning bounds for recovering ground-truth concepts with finite-sized datasets ([Theorem 4.20](#)).

We notice that the two perspectives differ in their main object of study. While the identifiability picture treats concept remapping distributions and studies identification of f^* (e.g., [Theorem 4.13](#)), the statistical learning picture focuses on shortcut NeSy predictors and on bounding the concept risk (e.g., [Theorem 4.20](#)). Because the statistical learning theory perspective ensures the avoidance of shortcut NeSy predictors, avoiding RSs is guaranteed only when $C = \text{supp}(G)$, i.e., when RSs in the problem are caused by shortcut NeSy predictors. In other cases where models are tested on unobserved concept combinations as in combinatorial generalization ([Duan et al., 2025](#)), i.e., when training and test distributions differ, ensuring concept risk minimization might not ensure avoiding all RSs.

Moreover, the two perspectives differ in how they specialize in treating mitigations. The identifiability perspective ([Marconato et al., 2023b, 2024](#); [Bortolotti et al., 2025](#); [van Krieken et al., 2025](#)) has so far investigated the impact of strategies on count reduction for deterministic α 's, studying how different mitigations change the quantities in the count ([Eq. \(24\)](#)). This also includes evaluating how the count reduces when different strategies are used together, but it is limited to studying optima of the learning problem. The statistical learning perspective ([Wang et al., 2023](#); [Yang et al., 2024](#); [He and Li, 2025](#)) has mainly explored how different mitigations that constrain the space of learnable concept extractors $\mathcal{F}$ can be used to bound the empirical RSs risk. These results do not explicitly require studying the optima of the learning problem, in the sense that, constraints on $\mathcal{F}$ can be directly translated to bounds on the RS risk ([Eq. \(32\)](#)). In fact, strategies like concept smoothing change the optima of the learning problem, and may not allow reaching maximum likelihood. Despite the different treatments, we note that the effect of some mitigation strategies can be described from both perspectives: a mitigation strategy can both reduce the count in [Eq. \(24\)](#) and reduce the risk in [Eq. \(32\)](#). In this sense, we view both perspectives as complementary: through identifiability theory, one can analyze how mitigations act to reduce deterministic RSs at optimal likelihood, whereas through statistical learning theory, bounds on the empirical concept risk can be found even when likelihood is not optimal. Formulating a theory that takes both perspectives into account remains open.

5 Handling Reasoning Shortcuts

We proceed by detailing the root causes of RSs (in Section 5.1), and outlining existing diagnostic tools (Section 5.2), mitigation strategies (Section 5.3), and awareness strategies (Section 5.4).

5.1 Root Causes

As anticipated in Section 3.1, RSs arise whenever the NeSy learning process does not penalize models that learn incorrect concepts, *i.e.*, concepts that do not match the ground-truth ones. Moreover, the theory in Section 4 provides additional details. Specifically, Theorem 4.16 suggests that RSs can occur when the knowledge complexity KC (Definition 4.15) is less than $|C| - 1$ and the space of learnable concept extractors $\mathcal{F}$ is broad enough. Corollary 4.5 supports the same conclusion. Following (Marconato et al., 2023b), additional insights can be gleaned by studying the structure of Eq. (24), reported below for reference, which counts the number of deterministic RSs (Definition 4.8) affecting a NeSy task. Since this is useful for understanding the mitigation strategies in Section 5.3, we analyze it here:

$$\sum_{\alpha \in \text{Vert}(\mathcal{A}_{\mathcal{F}})} \mathbb{1} \left\{ \bigwedge_{g \in \text{supp}(G)} (\beta^* \circ \alpha)(g) = \beta^*(g) \right\} - 1. \quad (40)$$

learnable maps α prior knowledge K support

The count is controlled by the four colored elements, which we discuss below.

5.1.1 The knowledge. The **prior knowledge K** determines the inference layer β^* . To understand its role, imagine two NeSy tasks with two binary concepts that are identical except for the prior knowledge K: one uses $K_1 = (Y_1 = 2C_1 + C_2)$ and the other $K_2 = (Y_2 = C_1 \vee C_2)$, see Fig. 8. Assuming the training set covers all classes (see Section 5.1.2), the task using K_1 does not admit any RSs: each label y can only be inferred by a unique choice of c , so high label accuracy entails high concept accuracy. The task using K_2 , however, is affected by RSs: the positive label can be inferred from $c = (0, 1)$, $c = (1, 0)$ and $c = (1, 1)$, hence NeSy predictors can confuse them with no impact on accuracy. From the perspective of Eq. (40), this happens because, all else being equal, swapping K_1 with K_2 softens the equality condition, increasing the count.

C_1	C_2	Y_1	Y_2
0	0	0	0
0	1	1	1
1	0	2	1
1	1	3	1

Fig. 8. K_1 is immune to RSs, K_2 is not.

5.1.2 The training distribution. Another important factor is the set of configurations of ground-truth concepts g that underlie the training data, that is, the **support** $\text{supp}(G)$ of the training distribution. To see this, consider the MNIST-Add example in Example 3.2. Suppose we add two new examples, (**55**, 10) and (**53**, 8), to the training set. Intuitively, the concept extractor *has* to map **5** to the digit 5, otherwise it would mispredict the (**55**, 10) example; likewise, now that **5** has only one correct grounding, it has to map **4** to 4 otherwise it would mispredict the (**54**, 9) example. The same reasoning applies for the remaining examples, namely (**53**, 8) and (**32**, 5). In summary, this small change completely removes all RSs, as the only way to match all training examples is to assign each MNIST digit to its correct numeric value. Intuitively, the larger the support – *i.e.*, the more combinations of ground-truth concepts g the training set represents – the more restrictive the conjunction in Eq. (40), lowering the count.

5.1.3 The optimality condition. The **optimality condition** $\mathcal{L}$ is responsible for filtering out those maps α that do not fit the ground-truth labels. Consider Example 3.2 again. One possible RS is $\{\mathbf{2} \mapsto 4, \mathbf{3} \mapsto 1, \mathbf{4} \mapsto 5, \mathbf{5} \mapsto 4\}$, which collapses **5** and **2** into the same concept value, 4. Now, imagine augmenting the training objective, by

default the cross-entropy, with a reconstruction loss. This mapping no longer achieves low loss, since the value 4 decodes to both **2** and **5**, which clearly hinders reconstruction. In terms of Eq. (40), modifying the training objective changes the optimality condition, potentially decreasing the number of RSs.

5.1.4 The family of learnable maps. Finally, the sum in Eq. (24) runs over (the vertices of) *the space of learnable concept remappings* α , denoted by $\text{Vert}(\mathcal{A}_{\mathcal{F}})$. This, in turn, depends on the space of learnable concept extractors $\mathcal{F}$ (Theorem 4.10). By introducing architectural bias in the architecture of the concept extractor (e.g., smoothing concept extractors by tuning a temperature parameter), we can constrain the space of learnable concept remappings α 's and therefore decrease the number of learnable RSs $\text{Vert}(\mathcal{A}_{\mathcal{F}})$. We will see in Section 5.3 how to effectively control this aspect of the learning problem.

5.2 How to Diagnose Reasoning Shortcuts

As mentioned in Section 3, in-distribution label accuracy is not affected by RSs and is therefore insufficient for diagnosing them. In this section, we briefly outline more effective diagnostic techniques.

5.2.1 Task-level diagnosis. Even before we commit to an architecture and train the model, we can *quantify how many deterministic RSs can affect a given NeSy task*. We accomplish this by evaluating Eq. (24), which is defined on known properties, such as the number of concepts, their cardinality and the given annotations. However, the count is very hard to compute manually even for relatively small problems, as the number of maps α increases doubly exponentially with the number of concepts k .

The countrss tool (Bortolotti et al., 2024) works around this by building on the following observation: all deterministic maps α in $\text{Vert}(\mathcal{A})$ can be viewed as binary matrices that indicate which predicted concepts C_j correspond to which ground-truth concepts G_i . The assumptions that constrain the number of mappings can be encoded into a propositional logical formula, whose solutions (that is, admissible binary matrices) are in one-to-one correspondence to the deterministic RSs α . Then, the count in Eq. (24) can be reduced to *model counting* (#SAT) (Gomes et al., 2021), i.e., the task of counting the solutions of a logical formula.

	C_{red}	$\neg C_{\text{red}}$	C_{ped}	$\neg C_{\text{ped}}$
G_{red}	0	0	1	0
$\neg G_{\text{red}}$	0	0	0	1
G_{ped}	1	0	0	0
$\neg G_{\text{ped}}$	0	1	0	0

(a)

	C_{red}	$\neg C_{\text{red}}$	C_{ped}	$\neg C_{\text{ped}}$
G_{red}	1	0	0	0
$\neg G_{\text{red}}$	0	1	0	0
G_{ped}	0	0	0	1
$\neg G_{\text{ped}}$	0	0	0	1

(b)

Fig. 9. Binary matrices encoding two possible mappings α from Example 2.1. (a) is a solution, resulting in a mapping that mistakes red lights with pedestrians and vice versa but is otherwise consistent with K for any combination of ground truth concepts. (b) is not a solution, since the resulting mapping makes the predictor unable to detect pedestrians and give correct predictions on some, but not all, combinations of ground truth concepts.

Although one could use any #SAT procedure, the complexity of the encoding for any task of practical interest warrants the use of approximate #SAT solvers (Chakraborty et al., 2021). Under the hood, countrss employs the state-of-the-art hashing-based counter ApproxMC¹⁸, which outputs an ϵ approximation of the exact count with probability δ . countrss inherits these statistical (PAC-style) guarantees, providing close approximations to Eq. (24) with high probability¹⁹. The estimates returned by countrss can provide an initial indication of the

¹⁸<https://github.com/meelgroup/approxmc>

¹⁹Regardless of the solver used, the reduction assumes disentanglement. Otherwise, the resulting number of logical variables in the encoding would be orders of magnitude above the capabilities of current solvers.

hardness of the learning problem and help make informed architectural choices according to the task at hand. `countrss` is included in the `rsbench` library (Bortolotti et al., 2024).

An alternative to count-based diagnosis is to certify shortcut-freeness directly. Takemura et al. (2026) cast the problem as a constraint satisfaction problem and propose a sound Answer Set Programming (Lifschitz, 2019) procedure that checks whether a given constraint set uniquely determines the intended concept mapping. Whereas `countrss` estimates *how many* shortcuts a task admits, this procedure answers *whether* any exist, making the two complementary.

5.2.2 Model metrics. Whenever concept annotations are available, the most effective way to evaluate the quality of the learned concepts is to use standard metrics such as accuracy, F_1 score, and concept-level confusion matrices. Another useful metric is *concept collapse* (Bortolotti et al., 2024). Roughly speaking, it quantifies to what extent the learned concept extractor compresses distinct ground-truth concepts into the same learned concept. More formally, let $C \in [0, 1]^{m \times m}$ be a concept confusion matrix, where m denotes its size (e.g., $m = 2^k$ when all ground-truth concepts are observed). We define $\text{Cls} = 1 - \frac{p}{m}$, where $p = \sum_{j=1}^m \mathbb{1}\{\exists i : C_{ij} > 0\}$. We remark that, however, collapse cannot detect all RSs. For instance, imagine a NeSy predictor trained on some variant of MNIST-Add affected by an RS mapping  to 1 and  to 0. This RS is completely invisible to the collapse metric, as the model does not collapse ground-truth concepts together, cf. Section 5.3.8. Still, it can be useful to gauge the presence of RSs that do. Finally, a major downside of these metrics is that they all require concept-level annotations, which are not always available.

With insufficient or missing concept annotations, another approach is to train multiple NeSy predictors to high label accuracy, and check their concept extractors align on concept predictions. When (deterministic) RSs are many in a NeSy task, it is likely that each predictor finds a different RS. Furthermore, if there are no RSs, concept extractors giving optimal likelihood must predict the same concepts, although the reverse is not necessarily true. The alignment between different concept extractors of the ensemble can be indirectly measured via the overall conditional entropy of the averaged concept predictions of the ensemble. Precisely, for $r \in \mathbb{N}$ members of the ensemble, we can account the mean conditional Shannon entropy $H(\frac{1}{r} \sum_j p_j(C | \mathbf{x}))$. This quantity is central for unsupervised RS-aware methods such as bears (Marconato et al., 2024) and NeSyDM (van Krieken et al., 2026), see Section 5.4.

5.3 A Tour of Mitigation Strategies

Several strategies relevant to mitigating RSs have been proposed. We list them in Table 2 and discuss them next.

5.3.1 Concept supervision. The most direct approach for encouraging correct grounding is to exploit **concept annotations** during training, thus altering the **optimality condition**. This requires introducing an additional loss term penalizing mispredicted concepts. Doing so effectively constrains the concept extractor f to recover the annotations, dramatically reducing the number of learnable RSs, and potentially avoiding them altogether.

This solution is, however, neither cheap nor perfect. It can be costly because concept-level annotations are frequently unavailable in the wild, and collecting them involves running potentially expensive annotation campaigns. Following recent trends (Oikarinen et al., 2022; Yang et al., 2023b; Rao et al., 2024; Srivastava et al., 2024; Yuksekgonul et al., 2023), one could obtain these annotations with foundation models. This solution, however, risks injecting substantial amounts of noise into the learning loop (Debole et al., 2025): despite being trained on vast amounts of data, foundation models can mispredict even basic concepts (Wüst et al., 2025).

It is also imperfect because, in the worst case, ruling out all RSs requires exhaustively annotating *all possible combinations of concepts*. Unless this is the case, a sufficiently expressive concept extractor might fit all annotations on the supervised combinations and fluke the remaining ones. Naturally, annotating an exponential (in the number of concepts) number of combinations can be infeasible. A possible workaround is to restrict the expressivity of

the concept extractor—e.g., by processing inputs separately, as typically done in MNIST-Add—so as to effectively transfer concept annotations across concept combinations. Another option is to intelligently select instances or concepts to be annotated; see Section 5.5.

Table 2. **Overview of mitigation strategies.** The columns indicate, for each strategy, what component(s) of Eq. (40) it targets (**Target**), whether it requires extra annotations (**Annot**) and guarantees reducing the count in Eq. (24) (**Count**), and how well it worked on different NeSy Tasks. Specifically, “✓” means that it prevents all RSs, “✗” that it fails to do so, and “?” that it has not been evaluated. Each **Target** value corresponds to a component of Eq. (40): **K** reshapes the knowledge constraint $(\beta^* \circ \alpha)(g) = \beta^*(g)$, while **F** and **L** both act through $\text{Vert}(\mathcal{A}_{\mathcal{F}})$: **F** as the family of representable maps α , and **L** as the optimality condition filtering these to the optimal ones. **Note:** With MNIST-Add * we refer to the family of datasets based on the MNIST-Add benchmark (Manhaeve et al., 2018).

Strategy	Properties			Arithmetic	Logic		High-Stakes
	Target	Annot	Count	MNIST-Add *	Kandinsky	Clevr	BDD-OIA
Concept Supervision (Koh et al., 2020)	L	✓	✓	✓	✓	✓	✗
Knowledge editing	K	✗	✗	?	?	?	?
Multi-task Learning (Caruana, 1997)	K	✓	✓	✓	?	?	?
Prototype-based grounding (Snell et al., 2017)	L	✓	✓	✓	✓	✓	✗
Weak Supervision (Yang et al., 2024)	L	✗	✗	✓	?	?	✗
Entropy Maximization (Manhaeve et al., 2021a)	L	✗	✗	?	?	?	✗
Smoothing (Szegedy et al., 2016)	L	✗	✗	?	?	?	?
Reconstruction (Kingma, 2013)	L	✗	✓	✓	?	?	?
Contrastive Learning (Chen et al., 2020)	L	✗	✗	?	?	?	?
Disentanglement (Suter et al., 2019)	F	✗	✓	?	✗	?	?

5.3.2 Knowledge editing. An alternative is to edit the *prior knowledge* **K** by introducing additional constraints, so as to remove ambiguities between concepts. For instance, in Example 2.1, one can break the symmetry between pedestrian and red_light by introducing a mutual exclusivity constraint between red_light and green_light. This way, the model can no longer freely confuse them, provided they appear together in the data, as in Fig. 3 (Right). While this strategy can be very effective, it requires figuring out what to change in the knowledge and sufficient expertise to implement the changes appropriately. Moreover, it is not always applicable. To see this, consider MNIST-Add. Here, the knowledge specifies the rule of addition: there is no way to meaningfully edit it without altering the semantics of this operation. This means that knowledge editing cannot be used to avoid RSs in any variants of this task.

5.3.3 Multi-task learning. Another way to improve the *prior knowledge* **K** is to train a NeSy predictor to solve *multiple NeSy tasks* (sharing some of the same concepts) jointly (Caruana, 1997). Doing so equates to training on a single NeSy task whose prior knowledge is the *conjunction* of the prior knowledge of the different tasks (Marconato et al., 2023b). As per Eq. (40), doing so reduces the number of potential RSs. Intuitively, this constraints the map β^* to distinguish among more different concept vectors than with a single task and excludes concept remapping distributions that are valid for tasks in isolation.

However, not all combinations of NeSy tasks are equally effective. To see this, imagine training a NeSy predictor to output both the sum of two MNIST digits and whether the first one is greater than the second one. Observations take the form $((\boxed{4} \boxed{5}), (9, 0))$ and $((\boxed{3} \boxed{2}), (5, 1))$, where the second label indicates whether the inequality holds. The RS in Example 3.2, i.e., $\{\boxed{2} \mapsto 4, \boxed{3} \mapsto 1, \boxed{4} \mapsto 3, \boxed{5} \mapsto 6\}$, is no longer valid, as it incorrectly predicts $(5, 0)$ for the second example. However, an alternative shortcut exists, namely $\{\boxed{2} \mapsto 1, \boxed{3} \mapsto 4, \boxed{4} \mapsto 3, \boxed{5} \mapsto 6\}$. A better alternative is to pair MNIST addition and multiplication. This is identical to Example 3.2, except now

there are two labels, e.g., $((4, 5), (9, 20))$ and $((3, 2), (5, 6))$. It can be verified that the only concept mapping that correctly predicts both the sum and the product is the intended one, i.e., the identity, ruling out all RSs.

Naturally, multi-task learning requires gathering label annotations for all involved tasks, which can be non-trivial. Another downside is that it may not be obvious how to construct a “natural” NeSy task that guarantees the removal of all RSs. We will discuss possible solutions in [Section 6.6](#).

5.3.4 Abductive weak supervision. This strategy requires training the concept extractor using *procedurally generated* pseudo-labels obtained with logical abduction (Yang et al., 2024). ABL embodies this approach: in ABL, the model uses logical abduction to infer the most plausible concept vector c that logically explains the ground-truth label y according to the prior knowledge. In practice, ABL uses a pre-defined distance function $d(c, c')$ to select this concept vector from alternatives. Then, ABL uses this new concept vector as the learning target for the loss function. Choosing an appropriate distance function can provide ABL with concept risk reduction; specifically, when the distance metric is always zero, Yang et al. (2024) showed that the risk of learning a shortcut solution reduces compared to other approaches. The downside is that there is no guarantee that the abduction process will recover the ground-truth concepts.

5.3.5 Prototype-based grounding. Another way to mitigate RSs is to rely on **prototypes**, i.e., anchoring each concept to a small *support set* of representative examples (Snell et al., 2017; Bontempelli et al., 2023; Andolfi and Giunchiglia, 2025). Following Andolfi and Giunchiglia (2025), prototype-based grounding addresses RSs by constraining the concept extractor to assign concepts c to an input x based on their *similarity* to the prototypes in the embedding space. For each concept c , a prototype is computed as the centroid of the embeddings of a (possibly very small) set of labeled examples. Compared to a standard NeSy predictor, during training Andolfi and Giunchiglia (2025) include an extra objective in which the model has to correctly classify the elements in the *support set* with respect to their corresponding prototypes based on distances (e.g., L2 distance). At inference time, instead, the probability of a predicted concept depends on the distance between the input representation and the corresponding prototype, typically via a softmax over distances. From the perspective of Eq. (40), prototype-based grounding targets the **optimality condition**: by enforcing concept assignments to respect distances to their prototypes, it imposes a specific geometry on the embedding space, thereby reducing the number of possible solutions the model can learn. The main advantage of this approach is that it requires relatively few annotations, in contrast to dense concept supervision (Marconato et al., 2023a), as only a small set of samples per concept is needed to compute the prototypes. A limitation, however, is that this approach assumes concepts are reasonably clusterable in the learned embedding space and that representative prototypes can be found. In domains where concepts are highly entangled or lack clear visual or semantic representations, defining meaningful prototypes can be challenging. Andolfi and Giunchiglia (2025) provide guarantees on the number of RSs when prototype-based grounding is employed. The main intuition is that, under the assumption that concepts form separable clusters in the embedding space, anchoring each concept with at least one labeled prototype reduces the number of deterministic RSs to the same level as full concept supervision (Marconato et al., 2023b).

5.3.6 Entropy maximization. Perhaps the simplest *unsupervised* mitigation strategy is **entropy maximization** (Manhaeve et al., 2021a). A simple idea here is to introduce an entropy-based loss term encouraging the model to evenly distribute probability mass across *all* concept combinations. This can be made practical by averaging over all predictions in a batch, and maximizing the Shannon entropy $H(p(C))$ (Manhaeve et al., 2021a).

Entropy maximization is easy to implement and computationally lightweight. However, it only encourages the predictor to spread concept predictions, but does not guarantee that the learned concepts will match the ground-truth ones. Furthermore, there are no formal results that show this method reduces the number of RSs.

5.3.7 Smoothing. The idea behind **concept smoothing** is to encourage the probabilities of the most and least probable concept vectors c to be “close”. The main idea is to never allow NeSy predictors to learn “peaked” concept

distributions, *i.e.*, $p(C | X)$ of the model is never one-hot. Smoothing can be enforced either architecturally, *e.g.*, by introducing a temperature parameter that reduces the magnitude of the concept activations, or during training, by penalizing high activations on input samples. This strategy permits reducing the RSs risk (Eq. (32)), the full derivation can be found in (Yang et al., 2024). This strategy, however, only circumvents *some* RSs: while smoothing prevents deterministic RSs, it does not affect non-deterministic ones. *E.g.*, imagine that K allows inferring the same label vector y from both c' and c'' (that is, $\beta^*(c') = \beta^*(c'') = y$, and that a NeSy predictor predicts a *mixture* of the two, *e.g.*, because it has learned the map $\alpha(g) = 0.5\mathbb{1}\{C = c'\} + 0.5\mathbb{1}\{C = c''\}$, where $g \in \{c', c''\}$. The mapping α is smooth – the highest probability it can output is 0.5 – but it is also a (non-deterministic) reasoning shortcut for PNSPs and SL, as it produces a correct prediction: $0.5\beta^*(c') + 0.5\beta^*(c'') = \mathbb{1}\{Y = y\}$.

5.3.8 Reconstruction. Another option is to introduce a **reconstruction penalty** into the learning process (Cottrell et al., 1987). Doing so encourages learning concepts that allow reconstructing the input x , or part of it. This rules out RSs that conflate unrelated concepts together and thus prevent reconstruction. To see how it works, consider MNIST-Add (Example 3.2) again: NeSy predictors achieving low reconstruction loss cannot map different MNIST digits to the same concept, as doing so prevents them from faithfully reconstructing the original inputs (see Section 5.1.3). For instance, the mapping $\mathbf{2} \mapsto 4$, $\mathbf{3} \mapsto 1$, $\mathbf{4} \mapsto 5$, $\mathbf{5} \mapsto 4$ implies that a single concept, 2, has to decode into two different digits, $\mathbf{2}$ and $\mathbf{5}$, leading to subpar reconstruction loss.

While reconstruction can effectively rule out such RSs, it cannot prevent all of them. In fact, a mapping like $\mathbf{2} \mapsto 4$, $\mathbf{3} \mapsto 1$, $\mathbf{4} \mapsto 3$, $\mathbf{5} \mapsto 6$ satisfies the knowledge *and* does not hinder reconstruction. It is easy to see that this mitigation targets the **optimality condition** in Eq. (40), and that as such it can decrease the number of RSs. Moreover, it comes at a cost: reconstruction requires adding a decoder module, increasing model size, computational overhead, and training complexity in practice. Crucially, for reconstruction to be effective, the learned concepts must correspond to distinct and separable visual features, as in MNIST-Add. This is not always the case. In real-world settings, *e.g.*, BDD-OIA (Xu et al., 2020), reconstructing the input image can be very difficult, requiring additional information such as object bounding boxes and segmentation masks. *E.g.*, reconstructing a pedestrian requires isolating it from the background, which is not possible from raw pixel input alone. This makes reconstruction impractical in NeSy tasks involving complex, real-world inputs.

5.3.9 Contrastive learning. Another effective unsupervised strategy targeting the **optimality condition** is **contrastive learning** (Chen et al., 2020). The core idea is to encourage similar inputs to produce the same concepts. This is typically achieved by forcing augmented views of the same input (positive pairs) to predict the same concepts, while ensuring that different inputs (negative pairs) yield distinct ones. For instance, given a MNIST digit, slight rotations or shifts can be applied to create positive pairs, while the remaining images in the batch serve as negatives. This encourages the model to map together different visual appearances of the same concept (*e.g.*, $\mathbf{3}$) while keeping the others separate (*e.g.*, $\mathbf{5}$).

Although contrastive learning provides no formal guarantees for reducing RSs, in practice it has an effect similar to reconstruction, as it addresses RSs that collapse semantically distinct inputs into a single concept. For instance, in Example 3.2 both strategies can prevent the model from mapping visually distinct digits such as $\mathbf{3}$ and $\mathbf{5}$ to the same concept, *e.g.*, 4. Contrastive learning is, however, easier to implement (*e.g.*, using stock implementations of the infNCE loss (Oord et al., 2018)) and significantly easier to optimize than reconstruction-based approaches. Nevertheless, care must be taken when designing augmentations to ensure they preserve the semantic meaning of the underlying concepts. For example, rotating a digit such as $\mathbf{4}$ in MNIST can transform it into a $\mathbf{9}$, which completely changes the semantics of the image, introducing noise (Chang et al., 2020).

5.3.10 Architectural disentanglement. The last unsupervised strategy we consider is **architectural disentanglement**. While the notion of disentanglement has different meanings in the literature (Higgins et al., 2018; Suter et al., 2019), here it simply means that the concept extractor processes independent objects in the input separately.

For instance, in visual recognition tasks in which multiple objects appear in the scene and each object is associated with a separate set of concepts (e.g., shape, size and color), architectural disentanglement amounts to processing each object separately. To see why this makes sense, consider MNIST-Add again. Here, the input comprises two digits. When using a single concept extractor that processes *both* digits at once, there is a chance that it will swap the order of the two digits, as doing so leaves their sum unchanged. Separate processing resolves this ambiguity. More generally, this strategy prevents the model from confusing the concepts of one object with those of a different object, thus reducing the number of RSs.

A nice benefit of this architectural bias is that it dramatically reduces the space of learnable concept extractors $\mathcal{F}$, which shrinks the space of learnable models (i.e., $\text{Vert}(\mathcal{A}_{\mathcal{F}})$ in Eq. (40)) and therefore the number of learnable RSs. Moreover, it enables parameter sharing across objects, reducing sample complexity. It is, however, not always applicable or easy to implement. On the one hand, it only makes sense for NeSy tasks in which the inputs contain multiple independent objects, such as MNIST-Add. On the other hand, it requires being able to isolate what part of the input encodes what object, which can be non-trivial depending on the complexity of the input. For instance, applying disentanglement to tasks like BDD-OIA requires using either ground-truth bounding boxes, which are seldom available, or predicted segmentations from tools such as Yolo (Redmon et al., 2016; Shindo et al., 2023) and FastRCNN (Girshick, 2015; Marconato et al., 2023b). These, however, can introduce noise in the learning process.

5.3.11 Which mitigation strategy is best? As shown in Table 2, different strategies offer different efficacy-cost trade-offs. The main takeaway is twofold: on the one hand, **supervised strategies** like concept supervision Section 5.3.1 and multi-task learning Section 5.3.3 can be very effective but also expensive; on the other, **unsupervised strategies** are cheaper to implement but also less reliable. Specifically, while some of them (like reconstruction Section 5.3.8, abduction-based weak supervision Section 5.3.4 and smoothing Section 5.3.7) provably avoid *certain* kinds of RSs, they are ineffective for others.

There is no one-size-fits-all strategy and the best option is application-specific, in that it depends on factors like the cost of obtaining annotations and the kind of RSs one wishes to eliminate. One option is to **mix multiple strategies**, yet not all combinations are useful. For instance, high smoothness is at odds with both concept supervision and reconstruction: the former aims to spread probability mass across concepts, increasing their uncertainty and uninformative-ness, whereas the latter attempt to fit the annotations with high certainty and to inject as much information as possible into the concepts, respectively. No model can fully satisfy both objectives. Finally, combining multiple strategies can yield overly complex loss functions that are challenging to optimize in practice, e.g., because they require careful hyperparameter tuning.

5.3.12 Can one simply recover learned concepts? Concepts learned without supervision risk encoding incorrect, misaligned semantics. A natural workaround, that has been sometimes used in the literature (Daniele et al., 2023; Tang and Ellis, 2023), is to figure out what concepts the model has learned in a *post-hoc* fashion, by matching concept predictions with ground-truth concept annotations.

However, this strategy may not work: it only works insofar as the concept extractor *identifies* (in a technical sense, see (Bortolotti et al., 2025)) the ground-truth concepts modulo a permutation of their indices and element-wise invertible transformations (Bortolotti et al., 2025). Given how difficult it is to ensure latent variables can be identified (Hyvärinen et al., 2024), this is unlikely to happen by chance. For instance, in MNIST-Add, this strategy only works if the concept extractor has learned to predict the ground-truth digits, except shuffled—e.g., it predicts 4 with 8, but predicts them perfectly when reordered. On the flip side, this strategy cannot work whenever the concept extractor collapses together multiple ground-truth concepts, e.g., it predicts 6 whenever it sees a 5 and an 6.

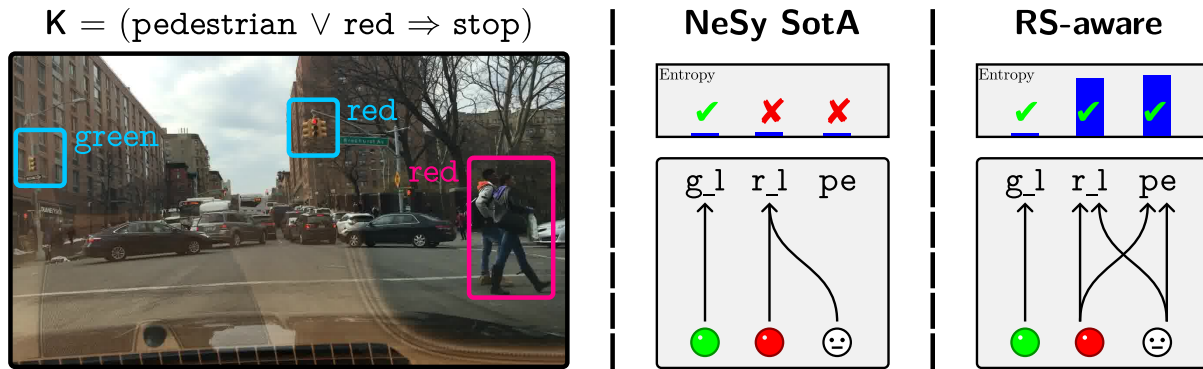


Fig. 10. **NeSy predictors are not RS-aware** (Marconato et al., 2024). Left: in the BDD-OIA autonomous driving task (Xu et al., 2020), NeSy predictors can achieve high accuracy and satisfy the prior knowledge even when they confuse pedestrians (ped) and red lights (red) (Marconato et al., 2023b). Right: An RS-aware predictor assigns high confidence to shortcut concepts, while giving low confidence to concepts not affected by RS.

5.4 Awareness Strategies

Another strategy for dealing with RSs is to *make NeSy predictors aware of the RSs that affect them*. We say that a NeSy predictor is **RS-aware** if it picks a non-deterministic RS that exhibits high confidence for those predicted concepts that are *not* affected by RSs and low confidence for the others (Marconato et al., 2024). As customary, confidence can be assessed using the *entropy* of the predictive concept distribution $p(C | \mathbf{x})$.

While awareness does *not* reduce RSs, RS-aware models supply stakeholders with valuable insights into the quality of the learned concepts, helping them to identify and distrust predictions obtained with poorly grounded concepts. To see this, consider the following example adapted from (Marconato et al., 2024):

Example 5.1. The left NeSy predictor in Fig. 10 is not RS-aware: despite confusing pedestrians with red lights, it is highly certain about the concepts it predicts. The one on the right is equally confused, but it is also RS-aware, in that it is more uncertain about low-quality concepts (pedestrian and red lights) than about high-quality ones (green lights). This enables users to distinguish between these cases.

Moreover, meaningful uncertainty estimates are also a powerful tool for steering interactive acquisition of annotations (Section 5.5), which can substantially reduce the cost of supervised mitigation (see Section 5.3.1). RS-awareness is complementary to mitigation strategies: when all realistic mitigation strategies are exhausted, some RSs may still be present. In that case, RS-aware models will improve uncertainty quantification. Unfortunately, in practice, NeSy predictors are *not* RS-aware: they tend to be very confident even about poorly grounded concepts (Marconato et al., 2024; van Krieken et al., 2026; Bortolotti et al., 2026).

What sets awareness apart from mere mitigation is that it is possible to make NeSy predictors RS-aware *in a principled manner*, but *without* concept supervision, as we will see. As mentioned in Section 5.3.6, a simple heuristic is to **maximize the entropy** of the predictive concept distribution. Is this always the best we can do? To answer this question, we will need to define RS-awareness.

5.4.1 RS-awareness as mixtures of deterministic RSs. RS awareness can be understood as mixing over the set of deterministic RSs that the NeSy predictor is affected by (van Krieken et al., 2025). We denote this set as $\text{Vert}(\mathcal{A})_{\text{RS}} \subseteq \text{Vert}(\mathcal{A})$, and it can be seen as the set of “irreducible” deterministic RSs that are still present after

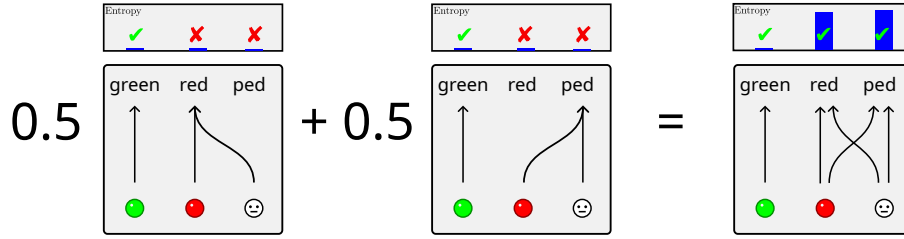


Fig. 11. **Intuition behind bears.** bears (Marconato et al., 2024) constructs ensembles of concept predictors that employ different sets of concepts to solve the task. Concepts that need to be learned correctly have to match across models, while those that cannot will have different groundings. As a result, averaging their predictions increases uncertainty only where disagreement exists.

exhausting the feasible mitigation strategies. Formally, we consider RS-aware NeSy predictors as distributions that “mix” over the concept remapping distributions (Eq. (19)) induced by each irreducible deterministic RS $a \in \text{Vert}(\mathcal{A})_{\text{RS}}$. We achieve this using the convex combination of deterministic RSs from Eq. (21), where each weight $\lambda_a > 0$. The convex combination ensures each RS contributes to the final concept remapping distribution. The resulting distribution is itself a (nondeterministic) RS (Marconato et al., 2024, Lemma 1), and hence has perfect label accuracy.

5.4.2 Building RS-aware NeSy predictors. How do we achieve such nondeterministic RSs in NeSy predictors? This usually requires that the space of learnable concept extractors $\mathcal{F}$ (Section 4.2.1) is powerful enough to model dependencies between different concepts (van Krieken et al., 2025). However, many implementations of NeSy predictors use concept extractors that assume (when conditioned on the input $\mathbf{x}$) the different concepts $C_1, \dots, C_k$ are independent, i.e., $p(C_1, \dots, C_k \mid \mathbf{x}) = \prod_{i=1}^k p(C_i \mid \mathbf{x})$ (van Krieken et al., 2024b). NeSy predictors exploit the independence assumption for efficient inference (Darwiche and Marquis, 2002; van Krieken et al., 2023; Choi et al., 2026; De Smet et al., 2023a). However, it also greatly restricts the space of learnable concept extractors $\mathcal{F}$.

Example 5.2 (Example in (van Krieken et al., 2025)). We consider the MNIST-XOR task where the NeSy predictor should predict the output of the XOR function applied to two MNIST digit. The XOR function returns 1 if the two input digits are different, and 0 otherwise. E.g., $\mathbf{x} = \begin{bmatrix} 1 & 0 \end{bmatrix}$ gives $y = 1$, while $\mathbf{x} = \begin{bmatrix} 1 & 1 \end{bmatrix}$ gives $y = 0$. This problem contains the shortcut $\begin{bmatrix} 1 \\ 0 \end{bmatrix} \mapsto 0, \begin{bmatrix} 0 \\ 0 \end{bmatrix} \mapsto 1$ (under disentanglement, Section 5.3.10). Given input $\mathbf{x} = \begin{bmatrix} 1 & 0 \end{bmatrix}$, an RS-aware NeSy predictor should assign 0.5 probability to both the ground-truth concept vector $(1, 0)$ and the RS concept vector $(0, 1)$. If we would try to model this distribution with the independence assumption, we find that $p(C_1 = 1 \mid \mathbf{x}) = p(C_2 = 1 \mid \mathbf{x}) = 0.5$. But then $p(C_1 = 1, C_2 = 1 \mid \mathbf{x}) = p(C_1 = 1 \mid \mathbf{x}) \cdot p(C_2 = 1 \mid \mathbf{x}) = 0.25$, i.e., it assigns probability to a concept vector that does not give the correct output label. Hence, the resulting distribution is not an RS, and might make incorrect label predictions.

In general, van Krieken et al. (2025) proves that the independence assumption prevents models from being aware of RSs for most inference layers. Motivated by these results, we discuss two recent methods that go beyond the independence assumption to build RS-aware NeSy predictors.

5.4.3 RS-Awareness via Ensembles. bears (Marconato et al., 2024) is a model-agnostic technique for making NeSy predictors RS-aware. bears replaces the concept extractor with an *ensemble* of concept extractors, each using the independence assumption. bears constructs this ensemble specifically such that each member captures a different deterministic RS in $\text{Vert}(\mathcal{A})_{\text{RS}}$. This construction has ties with Bayesian deep learning techniques (Wang and

Yeung, 2020; Daxberger et al., 2021; Osawa et al., 2019; Lakshminarayanan et al., 2017; Gal and Ghahramani, 2016) but it is also backed by a separate theoretical setup (Marconato et al., 2024). For instance, in BDD-OIA (Example 3.3) one extractor might conflate pedestrians with red lights, while another might do the opposite, as in Fig. 11 (left). bears computes the predictive concept distribution by averaging over all extractors in the ensemble. Intuitively, this works for two reasons. First, if all extractors have high label accuracy, their average will as well (Eq. (21)). Second, the concept entropy of the ensemble hinges on the semantics associated by the different extractors to the concepts themselves. Then, the ensemble has high entropy for concepts that are affected by RSs—like pedestrian and red light—as different extractors each learn different semantics. For unaffected concepts—like green light—the extractors agree on their semantics and the ensemble has low entropy. This intuition is illustrated in Fig. 11 (right).

In practice, bears extends deep ensembles (Lakshminarayanan et al., 2017) with a loss term that encourages ensemble members to model different RSs, and an entropy term (reminiscent of Section 5.3.6) that spreads concept predictions. Compared to other Bayesian deep learning methods, such as MC dropout (Gal and Ghahramani, 2016) and the Laplace approximation (Daxberger et al., 2021), deep ensembles allow learning highly multi-modal distributions, making them an excellent choice for learning multiple, diverse RSs. This is supported by experiments, which show that bears significantly improves RS-awareness of NeSy architectures on different tasks and that it does so better than other Bayesian alternatives (Marconato et al., 2024). One downside of bears is that it requires learning an ensemble of concept extractors. Fortunately, small ensembles of 5 to 10 members were shown to be sufficient in practice.

5.4.4 RS-Awareness via Diffusion. Motivated by going beyond the independence assumption, NeSyDM (van Krieken et al., 2026) uses expressive discrete diffusion models as concept extractors. Unlike bears, it trains a single concept extractor $p(C \mid \mathbf{x})$ via masked diffusion models (MDMs) (Austin et al., 2021; Sahoo et al., 2024). MDMs start with a concept vector that consists entirely of masks. Masks essentially indicate we do not yet know what the value of the concept should be. Then, MDMs are tasked with iteratively unmasking the various concepts. The main motivation for using MDMs to go beyond independence is scalability: each unmasking step can be seen as a concept extractor with the independence assumption. Hence, NeSyDM can directly reuse existing inference methods that exploit the independence assumption.

NeSyDM consists of three loss components. The first loss aims to unmask partially masked concept, the second unmask partially masked labels, and the final loss encourages high entropy among the concepts. One benefit compared to bears is that NeSyDM trains only a single model, eliminating the need to choose the ensemble size. Experimentally, NeSyDM also matches or improves on calibration and out-of-distribution concept accuracy compared to bears.

5.5 Awareness Helps Mitigation

To support mitigation strategies, specifically concept supervision (discussed in Section 5.3.1), one can leverage human input to request targeted annotations. This idea is well established in the field of Interactive Machine Learning (Ware et al., 2001; Fails and Olsen, 2003; Amershi et al., 2014; Teso and Kersting, 2019). In the context of RSs, the most relevant method explored so far is Active Learning (Settles, 2012), which we briefly describe below.

Active Learning. Rather than relying on costly full supervision, Active Learning (Settles, 2012) allows a model to interactively query a human (or oracle) for labels on the most informative or uncertain examples. The core idea is to focus annotation efforts where they are most impactful, thereby improving model performance while minimizing human effort. As demonstrated in (Marconato et al., 2024), uncertainty estimates over concepts – provided by bears (see Section 5.4) or NeSyDM (van Krieken et al., 2026) (see Section 5.4.4) – can be used to identify and query the most uncertain concepts, effectively addressing RSs with fewer annotations on concepts.

This targeted strategy outperforms traditional active learning methods, such as querying concept uncertainty in PNSPs, where the uncertainty signal is often uninformative (Marconato et al., 2024).

5.6 How to choose between mitigations?

Below, we provide a short guide for choosing among mitigation strategies. Recall that RSs can be avoided simply by obtaining concept annotations, with the caveat that many may be required in order to avoid all possible RSs. If doing so is not feasible, the choice of mitigation strategy largely depends on the application. We summarize common practical issues in Table 3.

- (1) Not all RSs affect downstream performance equally. This typically occurs in tasks where distinguishing between certain concepts is unnecessary for the final prediction (e.g., the type of obstacle on the road). In contrast, other (high-stakes) concepts must be clearly distinguished to ensure correct generalization to new tasks (e.g., pedestrian vs red_light). To enforce such distinctions, one approach is to directly supervise the targeted concepts (Section 5.3.1). This can be combined with prototypical representations (Andolfi and Giunchiglia, 2025) to reduce annotation costs (Section 5.3.5). Alternatively, designing opportune multi-task objectives can promote learning separate concept representations (Section 5.3.3).
- (2) In practice, NeSy models are often prone to learning RSs that collapse concepts into a limited set of combinations, thus impairing generalization and post-hoc identification of the learned concepts (Bortolotti et al., 2025). Augmenting the learning objective with entropy regularization (Section 5.3.6) encourages more diverse concept mappings, mitigating collapse. More resource-intensive alternatives include contrastive learning (Section 5.3.9) and reconstruction penalties (Section 5.3.8). Preventing collapse also enables post-hoc renaming procedures (Fokkema et al., 2025), where learned concepts are reassigned meaningful labels. This allows both concept representations and inference layers to be aligned with ground truth (Bortolotti et al., 2025).
- (3) Collapsed representations often lead to overconfident predictions. This issue can be addressed using bears (Section 5.4.3) and NeSyDM (Section 5.4.4), which explicitly identify concepts affected by RSs. However, these approaches can be computationally expensive. More lightweight alternatives with weaker guarantees include label smoothing (Section 5.3.7) and entropy regularization (Section 5.3.6). Label smoothing operates at the output level, while entropy-based methods act at the concept level.
- (4) RSs can also create entanglement between concepts. The most effective solution is to use model backbones that enforce disentangled concept representations by design (Section 5.3.10). When such architectures are not available, weak supervision through multiple views (Locatello et al., 2020a) can help promote disentanglement, although formal guarantees may not be met.
- (5) Eliminating deterministic RSs (i.e., zeroing their count) does not necessarily remove all RSs in practice. This is particularly relevant for NeSy predictors that do not satisfy the extremality condition (Assumption 4.12). In such cases, it is preferable to adopt inference layers that meet these requirements, such as PNSPs, or to rely on suitable fuzzy logic fragments (Section 2.2). Additionally, weak supervision through ABL (Section 5.3.4) can further help reduce concept-level risk.

Table 3. A list of issues that might arise due to RSs. We outline possible mitigations and the rationale behind each of them.

Issue	Mitigations	Rationale
Lack of generalization in new tasks	Concept Supervision + Prototypes (<i>if concept embeddings are clusterizable</i>), Multi-task learning (<i>if available</i>)	Distinguish between high-stakes concepts supervising on them, if available integrate new tasks that can be relevant in novel scenarios. Prototype-based grounding is effective when concept embeddings are clusterizable.
Collapse onto few concept combinations	Entropy, Contrastive Learning, Reconstruction	Use additional regularizations which penalize collapse, then apply a post-hoc renaming phase.
Overconfident predictions	bears and NeSyDM (<i>higher computational cost</i>), Smoothing and Entropy (<i>lesser computational cost</i>)	Model uncertainty on concepts affected by RSs, which informs about uncertain predictions.
Entangled concept predictions	Disentanglement (<i>if possible</i>), Contrastive Learning with Concept Supervision	Adopt architectural biases of the network that can help separate between concept predictions. Contrast different concept values by supervising them.
Non-deterministic concepts	Inference layers that satisfy Assumption 4.12 , Weak-supervision with ABL	Avoid methods which might not guarantee full RSs removal when zeroing the count. Use methods which allow for stronger concept risk reduction.

6 Extensions and Open Problems

In this section, we discuss RSs in contexts beyond those studied in this paper, and outline potential extensions and open research problems.

6.1 Reasoning Shortcuts in NeSy AI Beyond Predictors

RSs stem from the intrinsic ambiguity of logic inference when the prior knowledge K admits inferring the ground-truth label from unintended concepts. This suggests that RSs arise *regardless of how the reasoning step is implemented*, that is, also for NeSy architectures beyond the four predictors we considered here.²⁰ *E.g.*, RSs plausibly affect extensions of these models, such as (Huang et al., 2021b; Winters et al., 2022; De Smet et al., 2023b; Badreddine et al., 2023), as well as predictors that combine probabilistic *and* fuzzy semantics, like NeuPSL (Pryor et al., 2023), models that embed symbols into continuous vector spaces, such as neural theorem provers with fuzzy semantics (Rocktäschel and Riedel, 2016) and probabilistic semantics (Maene and Raedt, 2023), models that implement hard constraints with fuzzy logic, such as (Giunchiglia and Lukasiewicz, 2020; Giunchiglia et al., 2024), and diffusion-based models (van Krieken et al., 2026). Moreover, RSs have also been identified in SatNet (Wang et al., 2019; Simard et al., 1991), an architecture that performs perception and reasoning *entirely in the embedding space*. This leads us to postulate that—barring implicit mitigation due to the use of reconstruction penalties, see [Section 5.3.8](#)—RSs might also affect *generative* NeSy models, such as (Misino et al., 2022; Ferber et al., 2024).

While it is likely that RSs affect a broad class of NeSy architectures, more work is needed to extend the existing theoretical frameworks and mitigation strategies to these models.

6.2 Reasoning Shortcuts in Concept-based Models

So far, we only considered NeSy predictors where the prior knowledge K , and therefore the inference layer, is supplied externally and fixed. Recently, RSs have also been studied in two related families of models whose *inference layer is learned* (Bortolotti et al., 2025): *i*) NeSy modular predictors that acquire the knowledge from data (Daniele et al., 2023; Tang and Ellis, 2023; Shindo et al., 2023, 2024; Wüst et al., 2024; Debot et al., 2024),

²⁰Models like DPL and LTN are full-fledged first-order programming languages, *i.e.*, they can handle first-order logic specifications and work with a variable number of logic entities and concepts. Given that RSs affect them in the propositional case and that FOL is a strict generalization of propositional logic, we can safely assume that RSs exist also in the FOL case.

and ii) Concept-based Models (CBMs), which pair a concept extractor with an interpretable (typically linear) inference layer (Koh et al., 2020; Espinosa Zarlenga et al., 2022; Marconato et al., 2022; Schwalbe, 2022; Poeta et al., 2023; Pugnana et al., 2026; Knab et al., 2026). It was shown that, without the aid of concept supervision, these models suffer from *joint reasoning shortcuts* (JRSs): failure modes involving extracting wrong concepts and/or wrong logical assignments that lead to correct label predictions (see Example 6.1 for an example).

Example 6.1 (MNIST-SumParity). *To provide an example, imagine a variant of MNIST-Add where the goal is to predict whether the sum of two digits is odd or even, that is, $Y = (C_1 + C_2) \bmod 2$. Suppose that the training set includes only examples for which C_2 is odd. A possible JRS amounts to mapping all even digits to 0, i.e., $\{0, 2, 4, 6, 8\} \mapsto 0$, and all odd digits to 1, i.e., $\{1, 3, 5, 7, 9\} \mapsto 1$. Consistently, the learned knowledge is the difference among the predicted digits $\hat{Y} = \hat{C}_1 - \hat{C}_2$, as it would return the very same labels as using $Y = (C_1 + C_2) \bmod 2$. In contrast to regular RSs, here both the concepts and the learned knowledge are “wrong”, with the result that in the OOD setting where C_2 includes even digits, the learned model can predict $\hat{Y} = -1$ (see Fig. 12 for a visual representation).*

All RSs can also be viewed as JRSs, but the opposite is not true: the set of learnable JRSs can vastly outnumber that of regular RSs (Bortolotti et al., 2025). For example, permuting concept values does not alter the concept semantics and the learned rules, leading to equivalent solutions upon renaming the extracted concepts.²¹ Counting RSs can be adapted to joint-RSs, and an equivalent result to Theorem 4.13 also holds when K is learned, showing that dealing with deterministic joint-RSs is, in many cases, sufficient to rule out all joint-RSs. However, the benefits of many mitigation strategies for RSs discussed in Section 5.3 do not transfer to joint-RSs. Supervised strategies, like concept supervision and multi-task learning, offer some improvement but are insufficient to address all joint-RSs. In particular, concept supervision can lead to high concept accuracy, yet it does not necessarily recover the correct knowledge. It is unknown whether more powerful strategies for combating JRSs exist, and a theoretical characterization of JRSs from a statistical learning perspective is currently missing.

In the context of RSs, investigating the scenario where learned concept vectors ($c \in C$) and ground-truth concept vectors ($g \in \mathcal{G}$) belong to different spaces, i.e., $C \neq \mathcal{G}$, is open and so far unexplored. Considering two different spaces C and $\mathcal{G}$ results in the NeSy predictor inference layer β being structurally different from β^* , because of the different mapping domain. This setting naturally draws a connection between RSs and JRSs, where in the latter a concept-based model may learn to use fewer concepts in the bottleneck and fine-tune a $\beta \neq \beta^*$ that still matches label predictions. In this sense, because NeSy predictors (f, β) have $\beta \neq \beta^*$, one has to check whether learned concepts and the inference layer possess (a variation of) the intended semantics (Bortolotti et al., 2025, Definition 3.3). We expect that, in this scenario, RSs can likely increase due to a mismatch between the NeSy predictor’s and the ground-truth knowledge.

6.3 Foundation and Large Language Models

An interesting direction for future research is to investigate the connection between RSs and foundation models (Bommasani et al., 2021; Radford et al., 2021). Several recent approaches already leverage information *elicited from foundation models* to prime NeSy (Stammer et al., 2024; Helff et al., 2025; Wüst et al., 2026) and concept-based architectures (Oikarinen et al., 2022; Yang et al., 2023b; Rao et al., 2024; Srivastava et al., 2024; Yuksekgonul et al., 2023), e.g., using pre-trained representations or weak concept annotations obtained by prompting visual-language models. Clearly, the benefit is that foundation models are trained on vast amounts of data, which allows covering a large support of (potentially new) inputs, and are multimodal—allowing the use of textual concept

²¹Specifically, since concepts are *anonymous* during training, they can be easily aligned to ground-truth concepts at test time by leveraging annotations. This, however, is not possible if a joint-RS is learned by the model (Bortolotti et al., 2025).

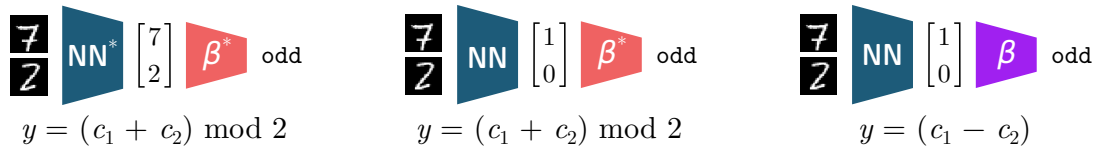


Fig. 12. **Example of Joint Reasoning Shortcuts.** The task is to predict whether the sum of two MNIST digits is odd, using training data that cover all (even, even), (odd, odd), and (odd, even) pairs. Red denotes fixed knowledge, purple denotes learned knowledge. In all the presented cases, the models achieve optimal performance. *Left:* ground-truth concepts with the correct inference layer. *Middle:* NeSy model with prior knowledge may still learn shortcuts by assigning unintended semantics to concepts. *Right:* CBMs, both concepts and the inference layer are misaligned, resulting in *joint reasoning shortcuts*. **Note:** In both of the previous cases, the OOD setting (even, odd) produces a consistent prediction, whereas the joint reasoning shortcut yields -1 as the predicted value.

descriptions to estimate the concepts in the data. Moreover, we expect that pre-training implicitly implements some of the RS mitigation strategies we outlined: concept supervision (Section 5.3.1) and multi-task learning (Section 5.3.3). In fact, it is quite likely that some elementary concepts, like the color or shape of objects, need to be separated in model representations to allow them to match a huge amount of image-caption pairs (e.g., from the Web Image Text dataset (Srinivasan et al., 2021)). For these reasons, it is plausible that, depending on the domain, foundation models may produce useful, high-quality concepts.

Empirical evidence, however, calls for caution. Recent findings (Debole et al., 2025; Wüst et al., 2025) in the context of VLM-augmented CBMs indicate that foundation models can fail to identify some basic visual concepts. In addition, they may fall short due to issues such as hallucinations (Huang et al., 2024), shortcuts (Yuan et al., 2024), and limited logical (Calanzone et al., 2025) and conceptual consistency (Sahu et al., 2022). In a direction closely related to RSs, Stein et al. (2025) have shown that language models are affected by **symbol hallucinations**, a phenomenon where models prompted to generate concepts and executable programs to solve a task end up hallucinating poorly grounded symbols that still manage to achieve high accuracy. Intuitively, symbol hallucinations can be thought of as the analog of RSs for prompting, although they originate not because of ambiguities related to a fixed prior knowledge in NeSy. This, in turn, complicate their characterization. An important direction is to determine whether foundation models themselves are affected by (joint) RSs proper, but we expect that the current theory will need to be extended to also treat models that cannot be described as NeSy predictors. Attempts to provide identifiability guarantees for concept learning in model representations (among which those of foundation models) have recently surfaced (Zheng et al., 2025; Liu et al., 2025; Rajendran et al., 2024), but it remains to be understood how this relates to shortcuts when these guarantees are not given.

Reasoning shortcuts have also been studied in language modelling, though with a different scope from the one of this paper. In question answering, multi-hop reasoning is the model’s ability to answer a question by combining information from multiple pieces of evidence rather than relying on a single fact (Yang et al., 2018). As an example of two-hop reasoning, we would expect that to answer the question “In which city of the Orange Free State was the author of *The Lord of the Rings* born?”, the model first has to identify J. R. R. Tolkien as the author of *The Lord of the Rings*, and second has to determine that Tolkien was born in Bloemfontein, which is in the Orange Free State. Jiang and Bansal (2019) found that when models are trained on certain multi-hop QA datasets, such as HotpotQA (Yang et al., 2018), they often exploit shortcuts that allow them to return correct answers by matching keywords or surface patterns in the question and context, *without performing the intended multi-hop reasoning steps*. In this context, the authors refer to it as a *reasoning shortcut* for the model. For instance, in the example above, an incorrect reasoning step consists of predicting Bloemfontein just because of the presence of “the Orange Free State”. Although these unintended solutions differ from the reasoning shortcuts discussed in our manuscript, they can be related to joint reasoning shortcuts, whereby correct concept assignments are given but

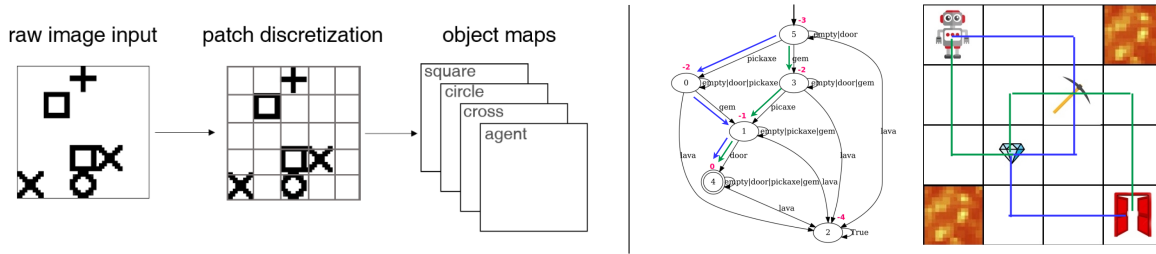


Fig. 13. (Left) Environment introduced in (Garnelo et al., 2016) and mainly used in (Badreddine and Spranger, 2019; Amador and Gierasimczuk, 2025). The agent (+ sign) has to maximize the sum of rewards by collecting shapes, each shape gives a different immediate reward. In both works, they use a sort of soft knowledge injection that is biased for this specific environment: they divide the input image into 25 patches (one per grid cell), and the shape recognizer classifies each patch singularly. (Right) Environment introduced in (Umili et al., 2024) for temporally extended visual tasks. The agent has to maximize the reward on the automaton. It changes automaton state by reaching specific items in the environment. On the left side, the automaton for the task "reach the pickaxe and the gem before the door, while always avoiding the "lava", is displayed. On the rightmost side, a visualization of the environment is portrayed. In green and blue two optimal trajectories causing the indistinguishability of symbols "pickaxe" and "gem" due to the structure of the task.

the learned knowledge fails to capture the ground truth. Making this link concrete is open and is an important research direction.

6.4 Reasoning Shortcuts in Reinforcement Learning

NeSy reinforcement learning (RL) (Acharya et al., 2024; Veronese et al., 2025) is a new field in which logical knowledge is integrated into the reinforcement learning workflow for various purposes, including guiding exploration (Anderson et al., 2020), ensuring safe trajectories (Yang et al., 2023a), improving explainability (Baugh et al., 2025; Deane and Ray, 2025), performance (Umili et al., 2024; Mitchener et al., 2022), or generalization to new environments (Garnelo et al., 2016; Badreddine and Spranger, 2019). Most of the works consider environments with symbolic observations, where grounding symbols or concepts is not required. A smaller set tackles more realistic environments by discretizing or clustering the neural features extracted from high-dimensional inputs (Umili et al., 2021; Garnelo et al., 2016; Hafner et al., 2021). In such cases, the extracted concepts lack a predefined meaning, potentially limiting the human interpretability of the agent's behavior. Only a few works have investigated embedding prior symbolic knowledge into RL policy learning (Badreddine and Spranger, 2019; Umili et al., 2024; Amador and Gierasimczuk, 2025). While reasoning shortcuts are still relatively underexplored in NeSy RL, first studies in this direction have already proven valuable for: (i) certifying the quality of learned concepts and its reusability in single-task vs. multitask settings (Umili et al., 2024), and (ii) improving the explainability of NeSy-RL agents when knowledge is learned rather than injected (Umili and Capobianco, 2025).

Badreddine and Spranger (2019) and Amador and Gierasimczuk (2025) focus on solving the RL environment introduced in (Garnelo et al., 2016) and shown in Fig. 13 (Left). They both employ a NeSy shape recognizer exploiting knowledge encoded in LTN. In these works, RSs do not appear because they rely on mitigation strategies specifically tailored to that particular environment, such as the use of pretrained shape detectors in (Badreddine and Spranger, 2019), the additional knowledge of the environment model in (Amador and Gierasimczuk, 2025), and patch discretization in both works.

A more recent work in NeSy-RL (Umili et al., 2024) explicitly investigates reasoning shortcuts through neural reward machines (NRMs). There, the RL agent learns to solve temporally extended tasks expressed in linear temporal logic over finite traces (LTLf) (Giacomo and Vardi, 2013) formulas. Fig. 13 (Right) illustrates an example

of such a task. NRMs encode temporal knowledge through a probabilistic relaxation of the deterministic finite automaton (DFA) equivalent to LTLf task. Here, the only bias is that rewards are shaped according to the automaton: reward is maximum at the automaton’s accepting states (reachable only by satisfying trajectories) and reduced elsewhere. This is a minimal bias that allows applying the framework to diverse environments and tasks without any change in the framework. In the context of the example of Fig. 13, the green and blue trajectories correspond to equivalent solutions for the agent, which may induce the agent to swap the concepts “gem” and “pickaxe”, resulting in an RS. Moreover, in (Umili et al., 2024) an efficient algorithm to identify all reasoning shortcuts that depend on the temporal logical specification is introduced, allowing to count RSs over three order of magnitude faster than a brute-force baseline.²² Applying this algorithm to the LTLf formulas commonly used in non-Markovian RL showed an extremely high number of RSs. Interestingly, in single-task RL, these RSs do not harm learning performance, since they involve only concept renamings and do not induce any concept collapse. Thus, all task-relevant concepts remain distinguished, but they may be swapped. However, this swap can harm the transfer of learned concepts to new multitask settings (Kuo et al., 2021). A follow-up work (Dewidar and Umili, 2025) extends NRMs to the cases where the task logical knowledge is not provided a priori: both the knowledge and the concepts are learned directly from exploration data (Umili and Capobianco, 2025). It investigates an algorithm to jointly minimize the automaton states and concepts learned via RS detection, yielding knowledge that is significantly more compact and explainable.

Future research can likely explore RSs in multitask generalization. In this context, prior knowledge is increasingly used to automatically generate new environments and tasks (Team et al., 2021; Bauer et al., 2023), making it timely to exploit (parts of) this knowledge during training to boost the RL agent’s performance. Another important open question concerns the role of temporal dependencies in NeSy RL. Unlike static tasks (e.g., classification or regression), RL data consists of exploration traces, where each observation depends on the previous ones, raising the challenge of how best to exploit this interdependence for better RS detection (van Krieken et al., 2025).

6.5 Imbalanced Learning

Tsamoura et al. (2025) identified another phenomenon that is inherent to NeSy: learning imbalances, which are major differences in errors that occur when classifying instances of different classes (aka *class-specific risks*) and can lead to poor learning of concept extractors. Existing research on machine learning (Menon et al., 2021; Cao et al., 2019; Wang et al., 2022) has studied imbalances, but this was only under the prism of *long-tailed* (aka *imbalanced*) data, where instances from different classes occur with very different frequencies, (He and Garcia, 2009; Horn and Perona, 2017). However, these results cannot fully characterize learning imbalances in all NeSy tasks. This is because *the background knowledge may cause learning imbalances in the labels even when the concept data is uniformly distributed, or vice versa*. We illustrate the above phenomenon by the following example:

²²While the brute-force baseline constructs all possible α ’s in $\text{Vert}(\mathcal{A})$, the algorithm introduced in (Umili et al., 2024) exploits common properties of RL tasks, such as terminal states and self-loop transitions in automata, to reduce the space of the search.

Example 6.2 (Learning Imbalances in MNIST-Max (Tsamoura et al., 2025)). *Let us consider a variant of the classical MNIST-Add scenario, referred to as MNIST-Max, in which instead of predicting the sum, we predict the maximum of two MNIST digits, e.g., $(\text{5 } 9) \mapsto 9$. We will adopt the notation in Wang et al. (2023) and denote each training sample by $(\mathbf{x}_1, \mathbf{x}_2, \mathbf{y})$, where $\mathbf{y}$ corresponds to the target maximum. We distinguish the following two cases:*

- (1) *The marginal of $\mathbf{Y}$ is uniform. In other words, the number of training samples for which the maximum is 0 ($\mathbf{y} = 0$) is equal to the number of training samples where the maximum $\mathbf{y}$ takes any of the remaining nine digits. In this settings, we expect that learning the digit zero (concept $\mathbf{c} = 0$) to be easier than learning the digit nine (concept $\mathbf{c} = 9$), i.e., the classifier will incur fewer errors in recognizing the digit zero than when recognizing any other digit. This is because training samples in which the maximum is zero, i.e., $(\mathbf{x}_1, \mathbf{x}_2, 0)$, can only occur when both digits are zero and the label $\mathbf{y} = 0$ appears uniformly in the training set. As a result, such samples force the NeSy model to learn the concept zero, since it is the only one satisfying the prior knowledge.*
- (2) *The marginal of the ground-truth concepts $\mathbf{G}$, i.e., the values of MNIST digits, is uniform. However, when we create training samples $(\mathbf{x}_1, \mathbf{x}_2, \mathbf{y})$ by sampling two MNIST digits independently, the resulting label $\mathbf{y} = \max(\mathbf{x}_1, \mathbf{x}_2)$ is not uniform. As a result, learning the digit nine (concept $\mathbf{c} = 9$) is easier than learning the digit zero (concept $\mathbf{c} = 0$) as $\mathbf{y} = 9$ appears much more frequently than $\mathbf{y} = 0$, even though samples of the form $(\mathbf{x}_1, \mathbf{x}_2, 0)$ provide a direct and unambiguous way to learn the concept zero.*

Tsamoura et al. (2025) theoretically characterized this phenomenon in NeSy, extending previous results (Cour et al., 2011). The characterization involved solving a non-linear program whose solutions are upper bounds to class-specific risks of the concepts. In addition, the authors proposed algorithms to mitigate imbalances during training and testing time. The former algorithm assigns pseudolabels to training data based on a novel formulation of NeSy as an integer linear program (Srikumar and Roth, 2023). The latter algorithm constrains the model's predictions on test data using robust semi-constrained optimal transport (Le et al., 2021). Both mitigation algorithms rely on a common idea in long-tailed learning: enforcing the class priors to the predictions of a concept extractor. The intuition is that the concept extractor will tend to predict the labels that appear more often in the training data. Hence, enforcing the priors gives more importance to the minority concepts at training time and encourages the model to predict minority concepts at testing time.

6.6 Additional Open problems

Here, we list a few more directions which were not treated in previous subsections:

- (1) Although several mitigation strategies have been proposed (some of them quite powerful, such as relying on concept supervision (Koh et al., 2020)), the most effective ones are costly in terms of human annotations. Related to this, multitask learning (Section 5.3.3) can be employed to obtain such annotations by designing combinations of NeSy tasks that promote the removal of reasoning shortcuts. However, it remains unclear how to construct such tasks effectively. A naive approach might involve combining learned programs (Muggleton and De Raedt, 1994; De Raedt and Kersting, 2008; Wüst et al., 2024) with rs-count (Bortolotti et al., 2024) to evaluate how the newly generated programs affect the number of reasoning shortcuts. This, however, requires access to ground-truth concepts, which is typically infeasible in real-world applications, and can also be highly inefficient in terms of computational cost. Developing efficient methods for artificially constructing such NeSy tasks therefore remains an open research problem. On the other hand, inexpensive methods such as entropy regularization (Manhaeve et al., 2021a) do not offer guarantees regarding their effectiveness in reducing the reasoning shortcuts count. There is thus a need for mitigation strategies that are both theoretically grounded

and cost-efficient, requiring minimal human annotation effort while still providing provable guarantees of effectiveness.

- (2) The mitigation strategies discussed in [Section 5.3](#) focus on avoiding *deterministic* RSs [Eq. \(40\)](#) and as such provide explicit reduction of their number, see [Eq. \(24\)](#). Implicitly, doing so equates to assuming that all RSs are equally impactful. However, this is not necessarily true: mistaking red lights for pedestrians is often worse than confusing red lights for stop signs. An alternative is to design mitigation strategies that take this into consideration. By prioritizing more costly RSs, this could reduce the need for expensive concept supervision and, more generally, reduce the down-stream cost of RSs. One option is to design mitigation strategies by building on the statistical learning perspective we discussed in [Section 4.3](#), which natively enables modelling the risk [Eq. \(32\)](#) of RSs in a cost-sensitive fashion. How to properly do so remains an open question.
- (3) Another open direction concerns the development of stronger concept detection pipelines, which could improve the overall identification process ([Locatello et al., 2020b](#)). For example, in computer vision, RSs are closely related to visual grounding ([Xiao et al., 2025](#)), where much work has focused, for instance, on object detection; however, a concrete link between object detection and RSs has yet to be established. Moreover, incorporating temporal dependencies (e.g., in reinforcement learning or in textual data) may further improve early detection and mitigation of RSs ([Lippe et al., 2022, 2023](#)).
- (4) To better advance the study of RSs, it is also extremely useful to extend the available experimental resources. While `rsbench` ([Bortolotti et al., 2024](#)) represents the first benchmark suite for a standardized evaluation of the impact of RSs in NeSy predictors and the efficacy of mitigation strategies, the field lacks benchmarks to examine how RSs affect model performance in highly challenging settings. An unexplored direction is to evaluate RSs on large, real-world NeSy datasets such as ROAD-R ([Giunchiglia et al., 2023](#)). Furthermore, to facilitate deploying NeSy predictors, it is beneficial to develop modular implementations of NeSy models that are both easy to use and freely available. Recent initiatives, like ULLER ([van Krieken et al., 2024a](#)) and DeepLog ([Derkinderen et al., 2025; Manhaeve et al., 2026](#)), are promising steps towards these implementations.

7 Related Work

Symbol Grounding. The symbol grounding problem concerns how symbols acquire meaning independently of human interpretation ([Harnad, 1990](#)). Reasoning shortcuts correspond to a faulty attribution of meaning to symbols (*w.r.t.* their human-interpretable semantics), which can hinder interpretability and out-of-distribution generalization. [Harnad \(1990\)](#) pioneered that symbol grounding should chiefly depend on the *identification* of the symbols. He argues that *discrimination* is the ability of systems to assign inputs representing different concepts with different symbols (e.g., images of horses and cats are associated to different symbols), but that this is not sufficient: *identification* further requires that these symbols are exactly the concepts we have in mind for them (e.g., for an image of a horse, the symbol must be exactly the concept of horse). Identification is precisely what underlies the grounding of the symbols. Similarly, an RS ([Section 4](#)) is when we *discriminate* for the purposes of the task, but fail to completely *identify* concepts.

The symbol grounding problem has inspired challenging datasets for visual reasoning like CLEVR ([Johnson et al., 2017](#)) and the visual abstractions benchmark ([Hsu et al., 2025](#)). The former requires scene understanding based on multiple 3D objects annotated with textual descriptions, whereas the latter focuses on grounding highly varied instantiations of the same concept (e.g., mazes). They argue that *schemas*, collections of connected lower-level concepts, together form the basis for grounding higher-level abstract concepts like a maze.

Symbol grounding is also closely related to binding ([Greff et al., 2020](#)) in neural networks, that is, how models derive a compositional understanding of the world in terms of symbolic entities, like objects or concepts. Recently, researchers asked to what degree foundation models require — or even have already achieved — symbol grounding ([Hsu et al., 2023; Jiang et al., 2024](#)). Some authors argue that the problem remains for models that are trained purely

on language (Pavlick, 2023; Levine, 2025), which are symbolic in nature. Others argue that forming meaningful concepts requires embodied experiences (Barsalou, 2020, 1999), which current models fundamentally lack (Dove, 2024; Levine, 2025). Steels et al. (2008) argues the focus should be on AI systems that autonomously create and maintain the grounding of concepts, and that supervised machine learning is insufficient as the semantics still comes from humans. Part of the focus has shifted to what forms of grounding —if linguistic, multimodal, or embodied— are necessary for genuine conceptual understanding (Barsalou, 2020; Gubelmann, 2024) and to whether this should be more framed specifically for distributed representations (Mollo and Millière, 2023). While NeSy predictors go beyond supervised machine learning and reduce reliance on direct concept supervision, they still require human-provided knowledge to be effective and currently do not support autonomous concept acquisition.

Inductive Logic Programming. Inductive logic programming (ILP) (De Raedt and Kersting, 2008; Muggleton and De Raedt, 1994) is a branch of machine learning that induces logic programs from examples and background knowledge. Hypotheses are typically represented as sets of first-order Horn clauses, and the goal is to find a symbolic program that entails all positive examples while excluding negative ones. One important technique in ILP is *predicate invention*, where the system autonomously generates new predicates to capture latent relationships in the data (Stahl, 1993). By creating intermediate concepts, predicate invention allows ILP models to abstract complex relational patterns, simplify rule sets, and improve generalization. For example, in an animal classification task, an ILP system might invent a `mammalian(X)` predicate to encode the mammal concept. However, the grounding of invented predicates can misalign with the one intended by the human, which makes predicate invention an interesting research direction to study connections with reasoning shortcuts and human interpretability.

A recurring challenge in ILP, both with and without predicate invention, is the presence of *symmetries* in the hypothesis space, which often leads to redundant search (Tarzariol et al., 2022). For instance, two clauses such as `parent(X,Y) :- mother(X,Y).` and `parent(A,B) :- mother(A,B).` are equivalent as they differ only in variable naming. Similarly, two clauses that differ only in the order of their literals are identical: for example, `sibling(X,Y) :- parent(Z,X), parent(Z,Y).` is equivalent to `sibling(X,Y) :- parent(Z,Y), parent(Z,X).` Exploring the hypothesis space in ILP without addressing symmetries can lead to the generation of many redundant, equivalent clauses. To mitigate this issue, various ILP systems employ *symmetry-breaking* techniques that constrain the search space and avoid unnecessary symmetries. For example, FOIL (Quinlan, 1990) addressed this by enforcing a consistent left-to-right ordering of literals, while Progol (Muggleton, 1995) and Aleph (Srinivasan, 2001) used mode declarations, which restrict the possible forms that clauses can take, and Metagol (Cropper and Muggleton, 2016), instead, provided metarules, which define general clause schemata that implicitly prevent the generation of equivalent variants.

In ILP, symmetries arise from structural redundancies in the hypothesis space, such as clauses that are equivalent up to variable renaming, literal order, or clause permutation. These symmetries may not permit identification: there can be multiple programs that fit the data equally well, yet differ in how they capture the intended abstraction. In particular, recurring relational patterns in the optimal hypothesis space can create multiple equivalent representations, each providing the same explanatory power. An analogous problem also occurs in satisfiability (SAT), where variables or clauses that play equivalent roles introduce redundancy in the search space (Sakallah, 2021). Just as in ILP, SAT solvers employ symmetry-breaking techniques to eliminate redundant solutions (Anders et al., 2024; Ulyantsev et al., 2016; Bogaerts et al., 2022). Studying the connection between symmetries in the search space and reasoning shortcuts can be of interest, as it may offer insights into both improving learning efficiency and understanding how learned hypotheses relate to intended abstractions.

Shortcuts in Machine Learning. Machine learning models often rely on shortcuts, spurious correlations between input features and target labels that are not causally related, leading to poor performance, particularly in OOD settings (Geirhos et al., 2020; Teso et al., 2023; Ye et al., 2024; Steinmann et al., 2024). Prior work related to

concepts has examined how spurious correlations affect both CBMs and NeSy models (Margeloiu et al., 2021; Mahinpei et al., 2021; Havasi et al., 2022; Raman et al., 2023; Stammer et al., 2021). These studies primarily focus on understanding and correcting concept quality in the presence of spurious confounders in the data. To address this problem, strategies such as Explanatory Interactive Learning (Teso and Kersting, 2019; Teso et al., 2023; Bontempelli et al., 2021; Teso et al., 2021; Bontempelli et al., 2023) have been developed, aiming to identify such spurious correlation and allow humans to correct them during the model debugging process. However, this line of work does not consider the high-level abstraction required to connect symbols to input features, a challenge central to the symbol grounding problem. In contrast, this paper focuses on reasoning shortcuts (Li et al., 2023; Marconato et al., 2023b; Wang et al., 2023; Umili et al., 2024; Marconato et al., 2024; Bortolotti et al., 2024; DeLong et al., 2024; Bortolotti et al., 2025), which arise not from low-level feature correlations but from the misuse of concepts during decision-making, as their semantics differ from the intended ones.

Identifiability. Identifiability has been largely studied in representation learning, especially within independent component analysis (ICA) (Hyvärinen et al., 2009, 2024) and causal representation learning (CRL) (Schölkopf et al., 2021). In ICA, the goal is to determine the independent underlying components from observed data. The central question of identifiability is: under what conditions can the same independent components (up to tolerable ambiguities) be consistently recovered (Buchholz et al., 2022; Gresele, 2023)? If components are not identifiable, re-estimating them at different times may yield conflicting results, meaning there is no unique way to determine the same independent factors. The same lack of uniqueness also appears in reasoning shortcuts. Seminal works have explored conditions under which non-linear independent components can be uniquely recovered, including using auxiliary variables (Hyvärinen et al., 2019; Khemakhem et al., 2020), sparsity (Lachapelle et al., 2022), anchor points (Moran et al., 2021), weak supervision (Locatello et al., 2020a), multiple views (Gresele et al., 2020), and constraints to the model family (Gresele et al., 2021). These ideas have also been extended to CRL, where identifiability is necessary not only to discover the underlying components but also to reconstruct the causal relationships between them (Von Kügelgen et al., 2021; Lippe et al., 2022; Buchholz et al., 2023; Ahuja et al., 2023; Lippe et al., 2023; von Kügelgen et al., 2024; Fokkema et al., 2025). We foresee that many techniques and learning setups that provide guarantees for identification can also help in mitigating RSs, but so far the connection to ICA and CRL has not been investigated.

Identifiability is also studied in supervised classification (Reizinger et al., 2025), multi-task learning (Lachapelle et al., 2023; Fumero et al., 2023), contrastive learning (Von Kügelgen et al., 2021; Zimmermann et al., 2021), and language models (Roeder et al., 2021; Marconato et al., 2025). In these settings, models may achieve close predictive performance while learning dissimilar internal representations (Nielsen et al., 2024, 2025). This situation is analogous to reasoning shortcuts, where correct outputs are obtained through misaligned or unintended representations. Revealing the precise connection between these models and RSs may help extend beyond the current paradigm of NeSy predictors.

8 Conclusion

Symbol grounding aims to link humans' and machines' high-level abstractions of the world. Well-grounded symbols, or concepts, contribute to generalization and interpretability. NeSy AI holds the great promise of facilitating this alignment in hybrid systems that are human interpretable and understandable. Our work characterizes the past, present and future challenges of symbol grounding in NeSy AI systems through the lens of RSs. The presence of RSs prevents identification of the correct concepts, and thus possibly compromises the benefits of NeSy AI. Our theoretical analysis in Section 4 indicates that, in general, all NeSy predictors we consider—and likely most others—are naturally susceptible to RSs. That is, concepts are non-identifiable in general. However, the theory also characterizes the mechanisms behind the occurrence of RSs, and therefore paves the way for designing a new generation of NeSy AI systems that are more reliable and easier to align with humans' expectations through

effective mitigation strategies. More work is, however, needed to articulate what conditions guarantee provable retrieval of the ground-truth concepts without requiring dense concept annotations. We expect that this can be done in a similar spirit and with similar strategies as those used to guarantee identifiability of representations in independent component analysis and causal representation learning (see [Section 7](#)).

Acknowledgments

We thank Samy Badreddine for useful input while writing this paper. Part of this work was initiated before Efthymia Tsamoura joined Huawei Labs. Funded by the European Union. The views and opinions expressed are however those of the author(s) only and do not necessarily reflect those of the European Union, the European Health and Digital Executive Agency (HaDEA) or the European Research Executive Agency. Neither the European Union nor the granting authority can be held responsible for them. Grant Agreement no. 101120763 - TANGO. Emile van Krieken is funded by the NWO AiNed project “Human-Centric AI Agents with Common Sense” under contract number NGF.1607.22.044. Paolo Morettin is supported by the MSCA project “Probabilistic Formal Verification for Provably Trustworthy AI - PFV-4-PTAI” under GA no. 101110960. Elena Umili is supported by the PNRR MUR project PE0000013-FAIR. AV is supported by the “UNREAL: Unified Reasoning Layer for Trustworthy ML” project (EP/Y023838/1) selected by the ERC and funded by UKRI EPSRC. Stefano Teso was partially supported by the Flemish Government under the Flemish research foundation (FWO) project “Neurosymbolic AI for Constraint Learning” (Project G047124N).

References

- Kamal Acharya, Waleed Raza, Carlos M. J. M. Dourado Júnior, Alvaro Velasquez, and Houbing Herbert Song. 2024. Neurosymbolic Reinforcement Learning and Planning: A Survey. *IEEE Trans. Artif. Intell.* 5, 5 (2024), 1939–1953. <https://doi.org/10.1109/TAI.2023.3311428>
- Kareem Ahmed, Stefano Teso, Kai-Wei Chang, Guy Van den Broeck, and Antonio Vergari. 2022. Semantic Probabilistic Layers for Neuro-Symbolic Learning. In *Advances in Neural Information Processing Systems*, S. Koyejo, S. Mohamed, A. Agarwal, D. Belgrave, K. Cho, and A. Oh (Eds.), Vol. 35. Curran Associates, Inc., 29944–29959. https://proceedings.neurips.cc/paper_files/paper/2022/file/c182ec594f38926b7fcb827635b9a8f4-Paper-Conference.pdf
- Kartik Ahuja, Divyat Mahajan, Yixin Wang, and Yoshua Bengio. 2023. Interventional causal representation learning. In *International conference on machine learning*. PMLR, 372–407.
- Ivo Amador and Nina Gierasimczuk. 2025. SymDQN: Symbolic Knowledge and Reasoning in Neural Network-based Reinforcement Learning. In *19th International Conference on Neurosymbolic Learning and Reasoning*. <https://openreview.net/forum?id=ncEGGRYska>
- Saleema Amershi, Maya Cakmak, William Bradley Knox, and Todd Kulesza. 2014. Power to the People: The Role of Humans in Interactive Machine Learning. *AI Magazine* 35, 4 (Dec. 2014), 105–120. <https://doi.org/10.1609/aimag.v35i4.2513>
- Markus Anders, Sofia Brenner, and Gaurav Rattan. 2024. Satsuma: Structure-based symmetry breaking in SAT. *arXiv preprint arXiv:2406.13557* (2024).
- Greg Anderson, Abhinav Verma, Isil Dillig, and Swarat Chaudhuri. 2020. Neurosymbolic Reinforcement Learning with Formally Verified Exploration. In *Advances in Neural Information Processing Systems 33: Annual Conference on Neural Information Processing Systems 2020, NeurIPS 2020, December 6-12, 2020, virtual*. <https://proceedings.neurips.cc/paper/2020/hash/448d5eda79895153938a8431919f4c9f-Abstract.html>
- Luca Andolfi and Eleonora Giunchiglia. 2025. Right for the Right Reasons: Avoiding Reasoning Shortcuts via Prototypical Neurosymbolic AI. In *The Thirty-ninth Annual Conference on Neural Information Processing Systems*. <https://openreview.net/forum?id=eb71SNTjux>

- Jacob Austin, Daniel D Johnson, Jonathan Ho, Daniel Tarlow, and Rianne Van Den Berg. 2021. Structured denoising diffusion models in discrete state-spaces. *Advances in neural information processing systems* 34 (2021), 17981–17993.
- Samy Badreddine, Artur d’Avila Garcez, Luciano Serafini, and Michael Spranger. 2022. Logic tensor networks. *Artificial Intelligence* 303 (2022), 103649.
- Samy Badreddine, Luciano Serafini, and Michael Spranger. 2023. logLTN: Differentiable Fuzzy Logic in the Logarithm Space. arXiv:2306.14546 [cs.AI] <https://arxiv.org/abs/2306.14546>
- Samy Badreddine and Michael Spranger. 2019. Injecting Prior Knowledge for Transfer Learning into Reinforcement Learning Algorithms using Logic Tensor Networks. In *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*.
- Lawrence W. Barsalou. 1999. Perceptual symbol systems. *Behavioral and Brain Sciences* 22, 4 (1999), 577–660. <https://doi.org/10.1017/S0140525X99002149>
- Lawrence W Barsalou. 2020. Challenges and opportunities for grounding cognition. *Journal of Cognition* 3, 1 (2020), 31.
- Jakob Bauer, Kate Baumli, Feryal Behbahani, Avishkar Bhoopchand, Nathalie Bradley-Schmieg, Michael Chang, Natalie Clay, Adrian Collister, Vibhavari Dasagi, Lucy Gonzalez, Karol Gregor, Edward Hughes, Sheleem Kashem, Maria Loks-Thompson, Hannah Openshaw, Jack Parker-Holder, Shreya Pathak, Nicolas Perez-Nieves, Nemanja Rakicevic, Tim Rocktäschel, Yannick Schroecker, Satinder Singh, Jakub Sygnowski, Karl Tuyls, Sarah York, Alexander Zacherl, and Lei M Zhang. 2023. Human-Timescale Adaptation in an Open-Ended Task Space. In *Proceedings of the 40th International Conference on Machine Learning (Proceedings of Machine Learning Research, Vol. 202)*, Andreas Krause, Emma Brunskill, Kyunghyun Cho, Barbara Engelhardt, Sivan Sabato, and Jonathan Scarlett (Eds.). PMLR, 1887–1935. <https://proceedings.mlr.press/v202/bauer23a.html>
- Kexin Gu Baugh, Luke Dickens, and Alessandra Russo. 2025. Neural DNF-MT: A Neuro-symbolic Approach for Learning Interpretable and Editable Policies. In *Proceedings of the 24th International Conference on Autonomous Agents and Multiagent Systems, AAMAS 2025, Detroit, MI, USA, May 19-23, 2025*. 252–260. <https://doi.org/10.5555/3709347.3743538>
- Bart Bogaerts, Stephan Gocht, Ciaran McCreesh, and Jakob Nordström. 2022. Certified symmetry and dominance breaking for combinatorial optimisation. In *Proceedings of the AAAI Conference on Artificial Intelligence*, Vol. 36. 3698–3707.
- Rishi Bommasani, Drew A. Hudson, Ehsan Adeli, Russ Altman, Simran Arora, Sydney von Arx, Michael S. Bernstein, Jeannette Bohg, Antoine Bosselut, Emma Brunskill, Erik Brynjolfsson, Shyamal Buch, Dallas Card, Rodrigo Castellon, Niladri Chatterji, Annie Chen, Kathleen Creel, Jared Quincy Davis, Dora Demszky, Chris Donahue, Moussa Doumbouya, Esin Durmus, Stefano Ermon, John Etchemendy, Kawin Ethayarajh, Li Fei-Fei, Chelsea Finn, Trevor Gale, Lauren Gillespie, Karan Goel, Noah Goodman, Shelby Grossman, Neel Guha, Tatsunori Hashimoto, Peter Henderson, John Hewitt, Daniel E. Ho, Jenny Hong, Kyle Hsu, Jing Huang, Thomas Icard, Saahil Jain, Dan Jurafsky, Pratyusha Kalluri, Siddharth Karamcheti, Geoff Keeling, Fereshte Khani, Omar Khattab, Pang Wei Koh, Mark Krass, Ranjay Krishna, Rohith Kuditipudi, Ananya Kumar, Faisal Ladhak, Mina Lee, Tony Lee, Jure Leskovec, Isabelle Levent, Xiang Lisa Li, Xuechen Li, Tengyu Ma, Ali Malik, Christopher D. Manning, Suvir Mirchandani, Eric Mitchell, Zanele Munyikwa, Suraj Nair, Avanika Narayan, Deepak Narayanan, Ben Newman, Allen Nie, Juan Carlos Niebles, Hamed Nilforoshan, Julian Nyarko, Giray Ogut, Laurel Orr, Isabel Papadimitriou, Joon Sung Park, Chris Piech, Eva Portelance, Christopher Potts, Aditi Raghunathan, Rob Reich, Hongyu Ren, Frieda Rong, Yusuf Roohani, Camilo Ruiz, Jack Ryan, Christopher Ré, Dorsa Sadigh, Shiori Sagawa, Keshav Santhanam, Andy Shih, Krishnan Srinivasan, Alex Tamkin, Rohan Taori, Armin W. Thomas, Florian Tramèr, Rose E. Wang, William Wang, Bohan Wu, Jiajun Wu, Yuhuai Wu, Sang Michael Xie, Michihiro Yasunaga, Jiaxuan You, Matei Zaharia, Michael Zhang, Tianyi Zhang, Xikun Zhang, Yuhui Zhang, Lucia Zheng,

- Kaitlyn Zhou, and Percy Liang. 2021. On the Opportunities and Risks of Foundation Models. *arXiv preprint arXiv:2108.07258* (July 2021).
- Andrea Bontempelli, Fausto Giunchiglia, Andrea Passerini, and Stefano Teso. 2021. Toward a Unified Framework for Debugging Gray-box Models. In *The AAAI-22 Workshop on Interactive Machine Learning*.
- Andrea Bontempelli, Stefano Teso, Fausto Giunchiglia, and Andrea Passerini. 2023. Concept-level debugging of part-prototype networks. In *International Conference on Learning Representations*.
- Samuele Bortolotti, Emanuele Marconato, Tommaso Carraro, Paolo Morettin, Emile van Krieken, Antonio Vergari, Stefano Teso, and Andrea Passerini. 2024. A Neuro-Symbolic Benchmark Suite for Concept Quality and Reasoning Shortcuts. In *Advances in Neural Information Processing Systems*, A. Globerson, L. Mackey, D. Belgrave, A. Fan, U. Paquet, J. Tomczak, and C. Zhang (Eds.), Vol. 37. Curran Associates, Inc., 115861–115905. https://proceedings.neurips.cc/paper_files/paper/2024/file/d1d11bf8299334d354949ba8738e8301-Paper-Datasets_and_Benchmarks_Track.pdf
- Samuele Bortolotti, Emanuele Marconato, Paolo Morettin, Andrea Passerini, and Stefano Teso. 2025. Shortcuts and Identifiability in Concept-based Models from a Neuro-Symbolic Lens. <https://openreview.net/forum?id=rdp1dLxyMI>
- Samuele Bortolotti, Emanuele Marconato, Andrea Pugnana, Andrea Passerini, and Stefano Teso. 2026. Concise and Logically Consistent Conformal Sets for Neuro-Symbolic Concept-Based Models. *arXiv:2605.18202* [cs.LG] <https://arxiv.org/abs/2605.18202>
- Simon Buchholz, Michel Besserve, and Bernhard Schölkopf. 2022. Function classes for identifiable nonlinear independent component analysis. *Advances in Neural Information Processing Systems* 35 (2022), 16946–16961.
- Simon Buchholz, Goutham Rajendran, Elan Rosenfeld, Bryon Aragam, Bernhard Schölkopf, and Pradeep Ravikumar. 2023. Learning linear causal representations from interventions under general nonlinear mixing. *Advances in Neural Information Processing Systems* 36 (2023), 45419–45462.
- Davide Buffelli and Efthymia Tsamoura. 2023. Scalable Theory-Driven Regularization of Scene Graph Generation Models. In *Thirty-Seventh AAAI Conference on Artificial Intelligence, Washington, DC, USA, February 7-14, 2023*. 6850–6859. <https://doi.org/10.1609/AAAI.V37I6.25839>
- Diego Calanzone, Stefano Teso, and Antonio Vergari. 2025. Logically Consistent Language Models via Neuro-Symbolic Integration. In *The Thirteenth International Conference on Learning Representations*.
- Kaidi Cao, Colin Wei, Adrien Gaidon, Nikos Arechiga, and Tengyu Ma. 2019. Learning Imbalanced Datasets with Label-Distribution-Aware Margin Loss. In *NeurIPS*. 1567–1578.
- Rich Caruana. 1997. Multitask learning. *Machine learning* (1997).
- Supratik Chakraborty, Kuldeep S. Meel, and Moshe Y. Vardi. 2021. Approximate Model Counting. In *Handbook of Satisfiability - Second Edition*, Armin Biere, Marijn Heule, Hans van Maaren, and Toby Walsh (Eds.). Frontiers in Artificial Intelligence and Applications, Vol. 336. IOS Press, 1015–1045. <https://doi.org/10.3233/FAIA201010>
- Oscar Chang, Lampros Flokas, Hod Lipson, and Michael Spranger. 2020. Assessing SATNet’s ability to solve the symbol grounding problem. *Advances in Neural Information Processing Systems* 33 (2020), 1428–1439.
- Ting Chen, Simon Kornblith, Mohammad Norouzi, and Geoffrey Everest Hinton. 2020. A Simple Framework for Contrastive Learning of Visual Representations.
- Weixin Chen, Simon Yu, Huajie Shao, Lui Sha, and Han Zhao. 2025. Neural Probabilistic Circuits: Enabling Compositional and Interpretable Predictions through Logical Reasoning. *arXiv preprint arXiv:2501.07021* (2025).
- Seewon Choi, Alaia Solko-Breslin, Rajeev Alur, and Eric Wong. 2026. Ctsketch: Compositional tensor sketching for scalable neurosymbolic learning. *Advances in Neural Information Processing Systems* 38 (2026), 75588–75612.
- Y Choi, Antonio Vergari, and Guy Van den Broeck. 2020. Probabilistic circuits: A unifying framework for tractable probabilistic models. *UCLA*. URL: <http://starai.cs.ucla.edu/papers/ProbCirc20.pdf> (2020), 6.
- Garrison W Cottrell, Paul Munro, and David Zipser. 1987. Learning internal representation from gray-scale images: An example of extensional programming. In *Proceedings of the Annual Meeting of the Cognitive Science*

- Society*, Vol. 9.
- Timothee Cour, Ben Sapp, and Ben Taskar. 2011. Learning from Partial Labels. *Journal of Machine Learning Research* 12 (2011), 1501–1536.
- Andrew Cropper and Stephen H Muggleton. 2016. Metagol system.
- Alessandro Daniele, Tommaso Campari, Sagar Malhotra, and Luciano Serafini. 2023. Deep symbolic learning: discovering symbols and rules from perceptions. In *IJCAI*.
- Adnan Darwiche and Pierre Marquis. 2002. A knowledge compilation map. *Journal of Artificial Intelligence Research* 17 (2002), 229–264.
- Tirtharaj Dash, Sharad Chitlangia, Aditya Ahuja, and Ashwin Srinivasan. 2022. A review of some techniques for inclusion of domain-knowledge into deep neural networks. *Scientific Reports* (2022).
- Erik Daxberger, Agustinus Kristiadi, Alexander Immer, Runa Eschenhagen, Matthias Bauer, and Philipp Hennig. 2021. Laplace redux-effortless bayesian deep learning. *Advances in neural information processing systems* 34 (2021), 20089–20103.
- Luc De Raedt, Sebastijan Dumančić, Robin Manhaeve, and Giuseppe Marra. 2021. From statistical relational to neural-symbolic artificial intelligence. In *IJCAI*.
- Luc De Raedt and Kristian Kersting. 2008. Probabilistic inductive logic programming. In *Probabilistic inductive logic programming: theory and applications*. Springer, 1–27.
- Luc De Raedt and Angelika Kimmig. 2015. Probabilistic (logic) programming concepts. *Machine Learning* 100 (2015), 5–47.
- Luc De Raedt, Angelika Kimmig, and Hannu Toivonen. 2007. ProbLog: A probabilistic Prolog and its application in link discovery. In *IJCAI*.
- Lennert De Smet, Emanuele Sansone, and Pedro Zuidberg Dos Martires. 2023a. Differentiable sampling of categorical distributions using the catlog-derivative trick. *Advances in Neural Information Processing Systems* 36 (2023), 30416–30428.
- Lennert De Smet, Pedro Zuidberg Dos Martires, Robin Manhaeve, Giuseppe Marra, Angelika Kimmig, and Luc De Raedt. 2023b. Neural probabilistic logic programming in discrete-continuous domains. In *Proceedings of the Thirty-Ninth Conference on Uncertainty in Artificial Intelligence (Proceedings of Machine Learning Research, Vol. 216)*, Robin J. Evans and Ilya Shpitser (Eds.). PMLR, 529–538. <https://proceedings.mlr.press/v216/de-smet23a.html>
- Oliver Deane and Oliver Ray. 2025. Neuro-Symbolic Inverse Constrained Reinforcement Learning. In *19th International Conference on Neurosymbolic Learning and Reasoning*. <https://openreview.net/forum?id=oVb3sJAnfx>
- Nicola Debole, Pietro Barbiero, Francesco Giannini, Andrea Passerini, Stefano Teso, and Emanuele Marconato. 2025. If Concept Bottlenecks are the Question, are Foundation Models the Answer? arXiv:2504.19774 [cs.LG] <https://arxiv.org/abs/2504.19774>
- David Debot, Pietro Barbiero, Francesco Giannini, Gabriele Ciravegna, Michelangelo Diligenti, and Giuseppe Marra. 2024. Interpretable Concept-Based Memory Reasoning. In *NeurIPS*.
- Lauren Nicole DeLong, Yojana Gadiya, Paola Galdi, Jacques D Fleuriot, and Daniel Domingo-Fernández. 2024. MARS: A neurosymbolic approach for interpretable drug discovery. *arXiv preprint arXiv:2410.05289* (2024).
- Vincent Derkinderen, Robin Manhaeve, Rik Adriaensen, Lucas Van Praet, Lennert De Smet, Giuseppe Marra, and Luc De Raedt. 2025. The DeepLog Neurosymbolic Machine. arXiv:2508.13697 [cs.AI] <https://arxiv.org/abs/2508.13697>
- Hazem Dewidar and Elena Umili. 2025. Fully Learnable Neural Reward Machines. In *Proceedings of the 7th International Workshop on Artificial Intelligence and Formal Verification, Logic, Automata, and Synthesis (OVERLAY 2025) co-located with 28th European Conference on Artificial Intelligence (ECAI 2025), Bologna, Italy, October 26, 2025 (CEUR Workshop Proceedings, Vol. 4142)*, Angelo Montanari, AndreA Orlandini, Nicola Saccomanno, and Stefano Tonetta (Eds.). CEUR-WS.org, 59–66. <https://ceur-ws.org/Vol-4142/paper7.pdf>

- Luca Di Liello, Pierfrancesco Ardino, Jacopo Gobbi, Paolo Morettin, Stefano Teso, and Andrea Passerini. 2020. Efficient generation of structured objects with constrained adversarial networks. *Advances in neural information processing systems* 33 (2020), 14663–14674.
- Michelangelo Diligenti, Marco Gori, Marco Maggini, and Leonardo Rigutini. 2012. Bridging logic and kernel machines. *Machine learning* 86, 1 (2012), 57–88.
- Ivan Donadello, Luciano Serafini, and Artur D’Avila Garcez. 2017. Logic tensor networks for semantic image interpretation. In *IJCAL*.
- Guy Dove. 2024. Symbol ungrounding: what the successes (and failures) of large language models reveal about human cognition. *Philosophical Transactions B* 379, 1911 (2024), 20230149.
- Xintong Duan, Yutong He, Fahim Tajwar, Wentse Chen, Ruslan Salakhutdinov, and Jeff Schneider. 2025. State Combinatorial Generalization In Decision Making With Conditional Diffusion Models. *Transactions on Machine Learning Research* (2025). <https://openreview.net/forum?id=XB1dd01Ozz>
- Kevin Ellis, Catherine Wong, Maxwell Nye, Mathias Sablé-Meyer, Lucas Morales, Luke Hewitt, Luc Cary, Armando Solar-Lezama, and Joshua B Tenenbaum. 2021. Dreamcoder: Bootstrapping inductive program synthesis with wake-sleep library learning. In *Proceedings of the 42nd acm sigplan international conference on programming language design and implementation*. 835–850.
- Mateo Espinosa Zarlenga, Pietro Barbiero, Gabriele Ciravegna, Giuseppe Marra, Francesco Giannini, Michelangelo Diligenti, Frederic Precioso, Stefano Melacci, Adrian Weller, Pietro Lio, et al. 2022. Concept embedding models. In *NeurIPS 2022-36th Conference on Neural Information Processing Systems*.
- Jerry Alan Fails and Dan R. Olsen. 2003. Interactive machine learning. In *Proceedings of the 8th International Conference on Intelligent User Interfaces (Miami, Florida, USA) (IUI '03)*. Association for Computing Machinery, New York, NY, USA, 39–45. <https://doi.org/10.1145/604045.604056>
- Jonathan Feldstein, Paulius Dilkas, Vaishak Belle, and Efthymia Tsamoura. 2024. Mapping the Neuro-Symbolic AI Landscape by Architectures: A Handbook on Augmenting Deep Learning Through Symbolic Reasoning. arXiv:2410.22077 [cs.AI] <https://arxiv.org/abs/2410.22077>
- Jonathan Feldstein, Modestas Jurcius, and Efthymia Tsamoura. 2023. Parallel Neurosymbolic Integration with Concordia. In *International Conference on Machine Learning (ICML), 23-29 July 2023, Honolulu, Hawaii, USA (Proceedings of Machine Learning Research, Vol. 202)*. 9870–9885.
- Aaron M Ferber, Arman Zharmagambetov, Taoan Huang, Bistra Dilkina, and Yuandong Tian. 2024. GenCO: Generating Diverse Designs with Combinatorial Constraints. In *Proceedings of the 41st International Conference on Machine Learning (Proceedings of Machine Learning Research, Vol. 235)*, Ruslan Salakhutdinov, Zico Kolter, Katherine Heller, Adrian Weller, Nuria Oliver, Jonathan Scarlett, and Felix Berkenkamp (Eds.). PMLR, 13445–13459. <https://proceedings.mlr.press/v235/ferber24a.html>
- Hidde Fokkema, Tim van Erven, and Sara Magliacane. 2025. Sample-efficient learning of concepts with theoretical guarantees: from data to concepts without interventions. *arXiv preprint arXiv:2502.06536* (2025).
- Marco Fumero, Florian Wenzel, Luca Zancato, Alessandro Achille, Emanuele Rodolà, Stefano Soatto, Bernhard Schölkopf, and Francesco Locatello. 2023. Leveraging sparse and shared feature activations for disentangled representation learning. *Advances in Neural Information Processing Systems* 36 (2023), 27682–27698.
- Yarin Gal and Zoubin Ghahramani. 2016. Dropout as a bayesian approximation: Representing model uncertainty in deep learning. In *international conference on machine learning*. PMLR, 1050–1059.
- Artur d’Avila Garcez, Sebastian Bader, Howard Bowman, Luis C Lamb, Leo de Penning, BV Illuminoo, Hoi-fung Poon, and COPPE Gerson Zaverucha. 2022. Neural-symbolic learning and reasoning: A survey and interpretation. *Neuro-Symbolic Artificial Intelligence: The State of the Art* 342 (2022), 1.
- Marta Garnelo, Kai Arulkumaran, and Murray Shanahan. 2016. Towards Deep Symbolic Reinforcement Learning. *CoRR abs/1609.05518* (2016). arXiv:1609.05518 <http://arxiv.org/abs/1609.05518>

- Robert Geirhos, Jörn-Henrik Jacobsen, Claudio Michaelis, Richard Zemel, Wieland Brendel, Matthias Bethge, and Felix A Wichmann. 2020. Shortcut learning in deep neural networks. *Nature Machine Intelligence* 2, 11 (2020), 665–673.
- Giuseppe De Giacomo and Moshe Y. Vardi. 2013. Linear Temporal Logic and Linear Dynamic Logic on Finite Traces. In *IJCAI 2013, Proceedings of the 23rd International Joint Conference on Artificial Intelligence, Beijing, China, August 3-9, 2013*. 854–860. <http://www.aaai.org/ocs/index.php/IJCAI/IJCAI13/paper/view/6997>
- Francesco Giannini, Michelangelo Diligenti, Marco Gori, and Marco Maggini. 2018. On a convex logic fragment for learning and reasoning. *IEEE Transactions on Fuzzy Systems* 27, 7 (2018), 1407–1416.
- Ross Girshick. 2015. Fast r-cnn. In *Proceedings of the IEEE international conference on computer vision*. 1440–1448.
- Eleonora Giunchiglia and Thomas Lukasiewicz. 2020. Coherent Hierarchical Multi-label Classification Networks. *NeurIPS* (2020).
- Eleonora Giunchiglia, Mihaela Cătălina Stoian, Salman Khan, Fabio Cuzzolin, and Thomas Lukasiewicz. 2023. ROAD-R: the autonomous driving dataset with logical requirements. *Machine Learning* 112, 9 (2023), 3261–3291.
- Eleonora Giunchiglia, Mihaela Catalina Stoian, and Thomas Lukasiewicz. 2022. Deep Learning with Logical Constraints. *arXiv preprint arXiv:2205.00523* (2022).
- Eleonora Giunchiglia, Alex Tatomir, Mihaela Cătălina Stoian, and Thomas Lukasiewicz. 2024. CCN+: A Neuro-symbolic Framework for Deep Learning with Requirements. *International Journal of Approximate Reasoning* (2024), 109124.
- Carla P. Gomes, Ashish Sabharwal, and Bart Selman. 2021. Model Counting. In *Handbook of Satisfiability - Second Edition*, Armin Biere, Marijn Heule, Hans van Maaren, and Toby Walsh (Eds.). Frontiers in Artificial Intelligence and Applications, Vol. 336. IOS Press, 993–1014. <https://doi.org/10.3233/FAIA201009>
- Klaus Greff, Sjoerd Van Steenkiste, and Jürgen Schmidhuber. 2020. On the binding problem in artificial neural networks. *arXiv preprint arXiv:2012.05208* (2020).
- Luigi Gresele. 2023. *Learning Identifiable Representations: Independent Influences and Multiple Views*. Ph.D. Dissertation. Eberhard Karls Universität Tübingen Tübingen.
- Luigi Gresele, Paul K Rubenstein, Arash Mehrjou, Francesco Locatello, and Bernhard Schölkopf. 2020. The incomplete rosetta stone problem: Identifiability results for multi-view nonlinear ica. In *Uncertainty in Artificial Intelligence*. PMLR, 217–227.
- Luigi Gresele, Julius Von Kügelgen, Vincent Stimper, Bernhard Schölkopf, and Michel Besserve. 2021. Independent mechanism analysis, a new concept? *Advances in neural information processing systems* 34 (2021), 28233–28248.
- Reto Gubelmann. 2024. Pragmatic Norms Are All You Need – Why The Symbol Grounding Problem Does Not Apply to LLMs. In *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing*, Yaser Al-Onaizan, Mohit Bansal, and Yun-Nung Chen (Eds.). Association for Computational Linguistics, Miami, Florida, USA, 11663–11678. <https://doi.org/10.18653/v1/2024.emnlp-main.651>
- Danijar Hafner, Timothy P Lillicrap, Mohammad Norouzi, and Jimmy Ba. 2021. Mastering Atari with Discrete World Models. In *International Conference on Learning Representations*. <https://openreview.net/forum?id=0oabwyZbOu>
- Stevan Harnad. 1990. The symbol grounding problem. *Physica D: Nonlinear Phenomena* 42, 1–3 (June 1990), 335–346. [https://doi.org/10.1016/0167-2789\(90\)90087-6](https://doi.org/10.1016/0167-2789(90)90087-6)
- Marton Havasi, Sonali Parbhoo, and Finale Doshi-Velez. 2022. Addressing leakage in concept bottleneck models. *NeurIPS*.
- Haibo He and Edwardo A. Garcia. 2009. Learning from Imbalanced Data. *IEEE Transactions on Knowledge and Data Engineering* 21, 9 (2009), 1263–1284.
- Hao-Yuan He and Ming Li. 2025. A Learnability Analysis on Neuro-Symbolic Learning. *arXiv preprint arXiv:2503.16797* (2025).

- Lukas Helff, Ruben Härle, Wolfgang Stammer, Felix Friedrich, Manuel Brack, Antonia Wüst, Hikaru Shindo, Patrick Schramowski, and Kristian Kersting. 2025. ActivationReasoning: Logical Reasoning in Latent Activation Spaces. *arXiv preprint arXiv:2510.18184* (2025).
- Irina Higgins, David Amos, David Pfau, Sebastien Racaniere, Loic Matthey, Danilo Rezende, and Alexander Lerchner. 2018. Towards a definition of disentangled representations. *arXiv preprint arXiv:1812.02230* (2018).
- Nick Hoernle, Rafael Michael Karampatsis, Vaishak Belle, and Kobi Gal. 2022. Multiplexnet: Towards fully satisfied logical constraints in neural networks. In *AAAI*.
- Grant Van Horn and Pietro Perona. 2017. The Devil is in the Tails: Fine-grained Classification in the Wild. *CoRR abs/1709.01450* (2017).
- Joy Hsu, Jiayuan Mao, Josh Tenenbaum, and Jiajun Wu. 2023. What’s left? concept grounding with logic-enhanced foundation models. *Advances in Neural Information Processing Systems* 36 (2023), 38798–38814.
- Joy Hsu, Jiayuan Mao, Joshua B Tenenbaum, Noah D Goodman, and Jiajun Wu. 2025. What Makes a Maze Look Like a Maze? International Conference on Learning Representations (ICLR).
- Jiani Huang, Ziyang Li, Binghong Chen, Karan Samel, Mayur Naik, Le Song, and Xujie Si. 2021b. Scallop: From Probabilistic Deductive Databases to Scalable Differentiable Reasoning. *NeurIPS* (2021).
- Lei Huang, Weijiang Yu, Weitao Ma, Weihong Zhong, Zhangyin Feng, Haotian Wang, Qianglong Chen, Weihua Peng, Xiaocheng Feng, Bing Qin, et al. 2024. A survey on hallucination in large language models: Principles, taxonomy, challenges, and open questions. *ACM Transactions on Information Systems* (2024).
- Xuanxiang Huang, Yacine Izza, Alexey Ignatiev, Martin C Cooper, Nicholas Asher, and Joao Marques-Silva. 2021a. Efficient explanations for knowledge compilation languages. *arXiv preprint arXiv:2107.01654* (2021).
- Geonho Hwang, Jaewoong Choi, Hyunsoo Cho, and Myungjoo Kang. 2023. MAGANet: Achieving Combinatorial Generalization by Modeling a Group Action. In *Proceedings of the 40th International Conference on Machine Learning (Proceedings of Machine Learning Research, Vol. 202)*, Andreas Krause, Emma Brunskill, Kyunghyun Cho, Barbara Engelhardt, Sivan Sabato, and Jonathan Scarlett (Eds.). PMLR, 14237–14248.
- Aapo Hyvärinen, Jarmo Hurri, and Patrik O. Hoyer. 2009. *Independent Component Analysis*. Springer London, London, 151–175. https://doi.org/10.1007/978-1-84882-491-1_7
- Aapo Hyvärinen, Ilyes Khemakhem, and Ricardo Monti. 2024. Identifiability of latent-variable and structural-equation models: from linear to nonlinear. *Annals of the Institute of Statistical Mathematics* 76, 1 (2024), 1–33.
- Aapo Hyvärinen, Hiroaki Sasaki, and Richard Turner. 2019. Nonlinear ICA using auxiliary variables and generalized contrastive learning. In *The 22nd international conference on artificial intelligence and statistics*. PMLR, 859–868.
- Bowen Jiang, Yangxinyu Xie, Xiaomeng Wang, Weijie J Su, Camillo Jose Taylor, and Tanwi Mallick. 2024. Multi-Modal and Multi-Agent Systems Meet Rationality: A Survey. In *ICML 2024 Workshop on LLMs and Cognition*. <https://openreview.net/forum?id=9Rtm2gAVjo>
- Yichen Jiang and Mohit Bansal. 2019. Avoiding Reasoning Shortcuts: Adversarial Evaluation, Training, and Model Development for Multi-Hop QA. In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*, Anna Korhonen, David Traum, and Lluís Màrquez (Eds.). Association for Computational Linguistics, Florence, Italy, 2726–2736. <https://doi.org/10.18653/v1/P19-1262>
- Jason Jo and Yoshua Bengio. 2017. Measuring the tendency of cnns to learn surface statistical regularities. *arXiv preprint arXiv:1711.11561* (2017).
- Justin Johnson, Bharath Hariharan, Laurens Van Der Maaten, Li Fei-Fei, C Lawrence Zitnick, and Ross Girshick. 2017. Clevr: A diagnostic dataset for compositional language and elementary visual reasoning. In *CVPR*.
- Philip N Johnson-Laird. 1994. Mental models and probabilistic thinking. *Cognition* 50, 1-3 (1994), 189–209.
- Subbarao Kambhampati, Sarath Sreedharan, Mudit Verma, Yantian Zha, and Lin Guan. 2022. Symbols as a lingua franca for bridging human-ai chasm for explainable and advisable ai systems. In *Proceedings of the AAAI*

- Conference on Artificial Intelligence*, Vol. 36. 12262–12267.
- Ilyes Khemakhem, Diederik Kingma, Ricardo Monti, and Aapo Hyvarinen. 2020. Variational autoencoders and nonlinear ica: A unifying framework. In *AISTATS*.
- Angelika Kimmig, Bart Demoen, Luc De Raedt, Vitor Santos Costa, and Ricardo Rocha. 2011. On the implementation of the probabilistic logic programming language ProbLog. *Theory and Practice of Logic Programming* (2011).
- Diederik P Kingma. 2013. Auto-encoding variational bayes. *arXiv preprint arXiv:1312.6114* (2013).
- Patrick Knab, David Steinmann, Christian Bartelt, Kristian Kersting, Bernt Schiele, Thomas Seidl, Udo Schlegel, and Wolfgang Stammer. 2026. What’s in the Bottle? A Survey and Roadmap of Concept Bottleneck Models. *Transactions on Machine Learning Research* (2026). <https://openreview.net/forum?id=IF5vnqxBEW>
- Pang Wei Koh, Thao Nguyen, Yew Siang Tang, Stephen Mussmann, Emma Pierson, Been Kim, and Percy Liang. 2020. Concept bottleneck models. In *International Conference on Machine Learning*. PMLR, 5338–5348.
- Daphne Koller and Nir Friedman. 2009. *Probabilistic graphical models: principles and techniques*. MIT press.
- Yen-Ling Kuo, Boris Katz, and Andrei Barbu. 2021. Compositional RL Agents That Follow Language Commands in Temporal Logic. *Frontiers Robotics AI* 8 (2021), 689550. <https://doi.org/10.3389/FROBT.2021.689550>
- Sébastien Lachapelle, Tristan Deleu, Divyat Mahajan, Ioannis Mitliagkas, Yoshua Bengio, Simon Lacoste-Julien, and Quentin Bertrand. 2023. Synergies between disentanglement and sparsity: Generalization and identifiability in multi-task learning. In *International Conference on Machine Learning*. PMLR, 18171–18206.
- Sébastien Lachapelle, Pau Rodriguez, Yash Sharma, Katie E Everett, Rémi Le Priol, Alexandre Lacoste, and Simon Lacoste-Julien. 2022. Disentanglement via mechanism sparsity regularization: A new principle for nonlinear ICA. In *Conference on Causal Learning and Reasoning*. PMLR, 428–484.
- Brenden Lake and Marco Baroni. 2018. Generalization without systematicity: On the compositional skills of sequence-to-sequence recurrent networks. In *International conference on machine learning*. PMLR, 2873–2882.
- Balaji Lakshminarayanan, Alexander Pritzel, and Charles Blundell. 2017. Simple and scalable predictive uncertainty estimation using deep ensembles. *Advances in neural information processing systems* 30 (2017).
- Khang Le, Huy Nguyen, Quang M Nguyen, Tung Pham, Hung Bui, and Nhat Ho. 2021. On Robust Optimal Transport: Computational Complexity and Barycenter Computation. In *Advances in Neural Information Processing Systems*. 21947–21959.
- Yann LeCun. 1998. The MNIST database of handwritten digits. <http://yann.lecun.com/exdb/mnist/> (1998).
- Sergey Levine. 2025. Language Models in Plato’s Cave. <https://sergeylevine.substack.com/p/language-models-in-platos-cave> Accessed: 2025-10-02.
- Zenan Li, Zehua Liu, Yuan Yao, Jingwei Xu, Taolue Chen, Xiaoxing Ma, L Jian, et al. 2023. Learning with Logical Constraints but without Shortcut Satisfaction. In *ICLR*.
- Vladimir Lifschitz. 2019. *Answer set programming*. Vol. 3. Springer Cham.
- Phillip Lippe, Sara Magliacane, Sindy Löwe, Yuki M Asano, Taco Cohen, and Efstratios Gavves. 2023. Biscuit: Causal representation learning from binary interactions. In *Uncertainty in Artificial Intelligence*. PMLR, 1263–1273.
- Phillip Lippe, Sara Magliacane, Sindy Löwe, Yuki M Asano, Taco Cohen, and Stratis Gavves. 2022. Citris: Causal identifiability from temporal intervened sequences. In *International Conference on Machine Learning*. PMLR, 13557–13603.
- Marco Lippi and Paolo Frasconi. 2009. Prediction of protein β -residue contacts by Markov logic networks with grounding-specific weights. *Bioinformatics* (2009).
- Changliu Liu, Tomer Arnon, Chris Lazarus, Christopher Strong, Clark Barrett, and Mykel J Kochenderfer. 2021. Algorithms for verifying deep neural networks. *Foundations and Trends in Optimization* 4, 3-4 (2021), 244–404.
- Li-Ping Liu and Thomas G. Dietterich. 2014. Learnability of the Superset Label Learning Problem. In *ICML*. 1629–1637.

- Yuhang Liu, Dong Gong, Yichao Cai, Erdun Gao, Zhen Zhang, Biwei Huang, Mingming Gong, Anton van den Hengel, and Javen Qinfeng Shi. 2025. I Predict Therefore I Am: Is Next Token Prediction Enough to Learn Human-Interpretable Concepts from Data? *arXiv preprint arXiv:2503.08980* (2025).
- Francesco Locatello, Ben Poole, Gunnar Rätsch, Bernhard Schölkopf, Olivier Bachem, and Michael Tschannen. 2020a. Weakly-supervised disentanglement without compromises. In *International conference on machine learning*. PMLR, 6348–6359.
- Francesco Locatello, Dirk Weissenborn, Thomas Unterthiner, Aravindh Mahendran, Georg Heigold, Jakob Uszkoreit, Alexey Dosovitskiy, and Thomas Kipf. 2020b. Object-centric learning with slot attention. *Advances in neural information processing systems* 33 (2020), 11525–11538.
- Jaron Maene, Vincent Derkinderen, and Pedro Zuidberg Dos Martires. 2025. Klay: Accelerating arithmetic circuits for neurosymbolic ai. In *ICLR*.
- Jaron Maene and Luc De Raedt. 2023. Soft-Unification in Deep Probabilistic Logic. In *Thirty-seventh Conference on Neural Information Processing Systems*.
- Jaron Maene and Efthymia Tsamoura. 2025. Embeddings as Probabilistic Equivalence in Logic Programs. In *Proceedings of the Thirty-Ninth Conference on Neural Information Processing Systems (NeurIPS)*.
- Anita Mahinpei, Justin Clark, Isaac Lage, Finale Doshi-Velez, and Weiwei Pan. 2021. Promises and pitfalls of black-box concept learning models. In *International Conference on Machine Learning: Workshop on Theoretic Foundation, Criticism, and Application Trend of Explainable AI*, Vol. 1. 1–13.
- Jean Matter Mandler. 2004. *The foundations of mind: Origins of conceptual thought*. Oxford University Press.
- Vasileios Manginas, Nikolaos Manginas, Edward Stevinson, Sherwin Varghese, Nikos Katzouris, Georgios Paliouras, and Alessio Lomuscio. 2025. A Scalable Approach to Probabilistic Neuro-Symbolic Robustness Verification. In *19th International Conference on Neurosymbolic Learning and Reasoning*. <https://openreview.net/forum?id=DAP8WCTGVj>
- Robin Manhaeve, Stefano Colamonaco, Vincent Derkinderen, Rik Adriaensen, Lucas Van Praet, Luc De Raedt, and Giuseppe Marra. 2026. Deeplog: A software framework for modular neurosymbolic ai. *arXiv preprint arXiv:2605.10279* (2026).
- Robin Manhaeve, Sebastijan Dumancic, Angelika Kimmig, Thomas Demeester, and Luc De Raedt. 2018. Deep-ProbLog: Neural Probabilistic Logic Programming. *NeurIPS* (2018).
- Robin Manhaeve, Sebastijan Dumančić, Angelika Kimmig, Thomas Demeester, and Luc De Raedt. 2021a. Neural probabilistic logic programming in DeepProbLog. *Artificial Intelligence* 298 (2021), 103504.
- Robin Manhaeve, Giuseppe Marra, and Luc De Raedt. 2021b. Approximate Inference for Neural Probabilistic Logic Programming. In *KR*.
- Emanuele Marconato, Gianpaolo Bontempo, Elisa Ficarra, Simone Calderara, Andrea Passerini, and Stefano Teso. 2023a. Neuro Symbolic Continual Learning: Knowledge, Reasoning Shortcuts and Concept Rehearsal. In *ICML*.
- Emanuele Marconato, Samuele Bortolotti, Emile van Krieken, Antonio Vergari, Andrea Passerini, and Stefano Teso. 2024. BEARS Make Neuro-Symbolic Models Aware of their Reasoning Shortcuts. *Uncertainty in AI* (2024).
- Emanuele Marconato, Sebastien Lachapelle, Sebastian Weichwald, and Luigi Gresele. 2025. All or None: Identifiable Linear Properties of Next-Token Predictors in Language Modeling. In *The 28th International Conference on Artificial Intelligence and Statistics*. <https://openreview.net/forum?id=XCmIlemQP5>
- Emanuele Marconato, Andrea Passerini, and Stefano Teso. 2022. Glancenets: Interpretable, leak-proof concept-based models. *Advances in Neural Information Processing Systems* 35 (2022), 21212–21227.
- Emanuele Marconato, Stefano Teso, Antonio Vergari, and Andrea Passerini. 2023b. Not All Neuro-Symbolic Concepts Are Created Equal: Analysis and Mitigation of Reasoning Shortcuts. In *NeurIPS*.
- Andrei Margeloiu, Matthew Ashman, Umang Bhatt, Yanzhi Chen, Mateja Jamnik, and Adrian Weller. 2021. Do Concept Bottleneck Models Learn as Intended? *arXiv preprint arXiv:2105.04289* (2021).

- Clayton McMillan, Michael C Mozer, and Paul Smolensky. 1991. Rule induction through integrated symbolic and subsymbolic processing. *Advances in neural information processing systems* 4 (1991).
- Aditya Krishna Menon, Sadeep Jayasumana, Ankit Singh Rawat, Himanshu Jain, Andreas Veit, and Sanjiv Kumar. 2021. Long-tail learning via logit adjustment. In *ICLR*.
- Eleonora Misino, Giuseppe Marra, and Emanuele Sansone. 2022. VAEI: Bridging Variational Autoencoders and Probabilistic Logic Programming. *NeurIPS* (2022).
- Ludovico Mitchener, David Tuckey, Matthew Crosby, and Alessandra Russo. 2022. Detect, Understand, Act: A Neuro-symbolic Hierarchical Reinforcement Learning Framework. *Mach. Learn.* 111, 4 (2022), 1523–1549. <https://doi.org/10.1007/S10994-022-06142-7>
- Dimitri Coelho Mollo and Raphaël Millière. 2023. The vector grounding problem. *arXiv preprint arXiv:2304.01481* (2023).
- Milton Llera Montero, Casimir JH Ludwig, Rui Ponte Costa, Gaurav Malhotra, and Jeffrey Bowers. 2021. The role of disentanglement in generalisation. In *International Conference on Learning Representations*.
- Gemma E Moran, Dhanya Sridhar, Yixin Wang, and David M Blei. 2021. Identifiable deep generative models via sparse decoding. *arXiv preprint arXiv:2110.10804* (2021).
- Paolo Morettin, Andrea Passerini, and Roberto Sebastiani. 2024. A Unified Framework for Probabilistic Verification of AI Systems via Weighted Model Integration. *arXiv preprint arXiv:2402.04892* (2024).
- Stephen Muggleton. 1995. Inverse entailment and Progol. *New generation computing* 13, 3 (1995), 245–286.
- Stephen Muggleton and Luc De Raedt. 1994. Inductive logic programming: Theory and methods. *The Journal of Logic Programming* 19 (1994), 629–679.
- Beatrix Miranda Ginn Nielsen, Luigi Gresele, and Andrea Dittadi. 2024. Challenges in explaining representational similarity through identifiability. In *UniReps: 2nd Edition of the Workshop on Unifying Representations in Neural Models*.
- Beatrix M. G. Nielsen, Emanuele Marconato, Andrea Dittadi, and Luigi Gresele. 2025. When Does Closeness in Distribution Imply Representational Similarity? An Identifiability Perspective. *arXiv:2506.03784* [cs.LG] <https://arxiv.org/abs/2506.03784>
- Tuomas Oikarinen, Subhro Das, Lam M Nguyen, and Tsui-Wei Weng. 2022. Label-free Concept Bottleneck Models. In *ICLR*.
- Aaron van den Oord, Yazhe Li, and Oriol Vinyals. 2018. Representation learning with contrastive predictive coding. *arXiv preprint arXiv:1807.03748* (2018).
- Kazuki Osawa, Siddharth Swaroop, Mohammad Emtiyaz E Khan, Anirudh Jain, Runa Eschenhagen, Richard E Turner, and Rio Yokota. 2019. Practical deep learning with Bayesian principles. *Advances in neural information processing systems* 32 (2019).
- Kiho Park, Yo Joong Choe, and Victor Veitch. 2024. The Linear Representation Hypothesis and the Geometry of Large Language Models. In *International Conference on Machine Learning*. PMLR, 39643–39666.
- Ellie Pavlick. 2023. Symbols and grounding in large language models. *Philosophical Transactions of the Royal Society A* 381, 2251 (2023), 20220041.
- Eleonora Poeta, Gabriele Ciravegna, Eliana Pastor, Tania Cerquitelli, and Elena Baralis. 2023. Concept-based explainable artificial intelligence: A survey. *Comput. Surveys* (2023).
- Connor Pryor, Charles Dickens, Eriq Augustine, Alon Albalak, William Yang Wang, and Lise Getoor. 2023. NeuPSL: neural probabilistic soft logic. In *Proceedings of the Thirty-Second International Joint Conference on Artificial Intelligence*. 4145–4153.
- Andrea Pugnana, Riccardo Massidda, Francesco Giannini, Pietro Barbiero, Mateo Espinosa Zarlenga, Roberto Pellungrini, Gabriele Dominici, Fosca Giannotti, and Davide Bacciu. 2026. Deferring Concept Bottleneck Models: Learning to Defer Interventions to Inaccurate Experts. In *The Thirty-ninth Annual Conference on Neural Information Processing Systems*. <https://openreview.net/forum?id=QdfdwsoOE>

- J. Ross Quinlan. 1990. Learning logical definitions from relations. *Machine learning* 5, 3 (1990), 239–266.
- Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sastry, Amanda Askell, Pamela Mishkin, Jack Clark, et al. 2021. Learning transferable visual models from natural language supervision. In *International conference on machine learning*. PmLR, 8748–8763.
- Aditi Raghunathan, Roy Frostig, John Duchi, and Percy Liang. 2016. Estimation from Indirect Supervision with Linear Moments. In *ICML*, Vol. 48. 2568–2577.
- Goutham Rajendran, Simon Buchholz, Bryon Aragam, Bernhard Schölkopf, and Pradeep Kumar Ravikumar. 2024. From Causal to Concept-Based Representation Learning. In *The Thirty-eighth Annual Conference on Neural Information Processing Systems*.
- Naveen Raman, Mateo Espinosa Zarlenga, Juyeon Heo, and Mateja Jamnik. 2023. Do Concept Bottleneck Models Obey Locality?. In *XAI in Action: Past, Present, and Future Applications*.
- Sukrut Rao, Sweta Mahajan, Moritz Böhle, and Bernt Schiele. 2024. Discover-then-Name: Task-Agnostic Concept Bottlenecks via Automated Concept Discovery. In *European Conference on Computer Vision*.
- Joseph Redmon, Santosh Divvala, Ross Girshick, and Ali Farhadi. 2016. You only look once: Unified, real-time object detection. In *Proceedings of the IEEE conference on computer vision and pattern recognition*. 779–788.
- Patrik Reizinger, Alice Bizeul, Attila Juhos, Julia E Vogt, Randall Balestriero, Wieland Brendel, and David Klindt. 2025. Cross-Entropy Is All You Need To Invert the Data Generating Process. In *The Thirteenth International Conference on Learning Representations*.
- Tim Rocktäschel and Sebastian Riedel. 2016. Learning knowledge base inference with neural theorem provers. In *Proceedings of the 5th workshop on automated knowledge base construction*. 45–50.
- Geoffrey Roeder, Luke Metz, and Durk Kingma. 2021. On linear identifiability of learned representations. In *International Conference on Machine Learning*. PMLR, 9030–9039.
- Dan Roth. 1996. On the hardness of approximate reasoning. *Artificial intelligence* 82, 1-2 (1996), 273–302.
- Subham Sahoo, Marianne Arriola, Yair Schiff, Aaron Gokaslan, Edgar Marroquin, Justin Chiu, Alexander Rush, and Volodymyr Kuleshov. 2024. Simple and effective masked diffusion language models. *Advances in Neural Information Processing Systems* 37 (2024), 130136–130184.
- Pritish Sahu, Michael Cogswell, Yunye Gong, and Ajay Divakaran. 2022. Unpacking large language models with conceptual consistency. *arXiv preprint arXiv:2209.15093* (2022).
- Karem A Sakallah. 2021. Symmetry and satisfiability. In *Handbook of Satisfiability*. IOS press, 509–570.
- Bernhard Schölkopf, Francesco Locatello, Stefan Bauer, Nan Rosemary Ke, Nal Kalchbrenner, Anirudh Goyal, and Yoshua Bengio. 2021. Toward causal representation learning. *Proc. IEEE* 109, 5 (2021), 612–634.
- Gesina Schwalbe. 2022. Concept embedding analysis: A review. *arXiv preprint arXiv:2203.13909* (2022).
- Burr Settles. 2012. *Active Learning*. Vol. 6. <https://doi.org/10.2200/S00429ED1V01Y201207AIM018>
- Shai Shalev-Shwartz and Shai Ben-David. 2014. *Understanding Machine Learning: From Theory to Algorithms*. Cambridge University Press, USA.
- Hikaru Shindo, Viktor Pfanschilling, Devendra Singh Dhami, and Kristian Kersting. 2023. (alpha)ILP: thinking visual scenes as differentiable logic programs. *Machine Learning* 112, 5 (2023), 1465–1497.
- Hikaru Shindo, Viktor Pfanschilling, Devendra Singh Dhami, and Kristian Kersting. 2024. Learning differentiable logic programs for abstract visual reasoning. *Machine Learning* 113, 11 (2024), 8533–8584.
- Patrice Simard, Bernard Victorri, Yann LeCun, and John S Denker. 1991. Tangent prop-a formalism for specifying selected invariances in an adaptive network. In *NIPS*, Vol. 91. 895–903.
- Paul Smolensky. 1987. Analysis of distributed representation of constituent structure in connectionist systems. In *Neural Information Processing Systems*.
- Jake Snell, Kevin Swersky, and Richard Zemel. 2017. Prototypical networks for few-shot learning. *Advances in neural information processing systems* 30 (2017).
- Elizabeth S Spelke and Katherine D Kinzler. 2007. Core knowledge. *Developmental science* 10, 1 (2007), 89–96.

- Vivek Srikumar and Dan Roth. 2023. The Integer Linear Programming Inference Cookbook. *ArXiv abs/2307.00171* (2023). <https://api.semanticscholar.org/CorpusID:259316294>
- Ashwin Srinivasan. 2001. The aleph manual. (2001).
- Krishna Srinivasan, Karthik Raman, Jiecao Chen, Michael Bendersky, and Marc Najork. 2021. WIT: Wikipedia-Based Image Text Dataset for Multimodal Multilingual Machine Learning. In *Proceedings of the 44th International ACM SIGIR Conference on Research and Development in Information Retrieval* (Virtual Event, Canada) (SIGIR '21). Association for Computing Machinery, New York, NY, USA, 2443–2449. <https://doi.org/10.1145/3404835.3463257>
- Divyansh Srivastava, Ge Yan, and Tsui-Wei Weng. 2024. VLG-CBM: Training Concept Bottleneck Models with Vision-Language Guidance. In *Advances in Neural Information Processing Systems*, A. Globerson, L. Mackey, D. Belgrave, A. Fan, U. Paquet, J. Tomczak, and C. Zhang (Eds.), Vol. 37. Curran Associates, Inc., 79057–79094. https://proceedings.neurips.cc/paper_files/paper/2024/file/90043ebd68500f9efe84fedf860a64f3-Paper-Conference.pdf
- Irene Stahl. 1993. Predicate invention in ILP—an overview. In *European conference on machine learning*. Springer, 311–322.
- Wolfgang Stammer, Patrick Schramowski, and Kristian Kersting. 2021. Right for the Right Concept: Revising Neuro-Symbolic Concepts by Interacting With Their Explanations. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*. 3619–3629.
- Wolfgang Stammer, Antonia Wüst, David Steinmann, and Kristian Kersting. 2024. Neural Concept Binder. In *The Thirty-eighth Annual Conference on Neural Information Processing Systems*.
- Luc Steels et al. 2008. The symbol grounding problem has been solved. so what’s next. *Symbols and embodiment: Debates on meaning and cognition* (2008), 223–244.
- Adam Stein, Aaditya Naik, Neelay Velingker, Mayur Naik, and Eric Wong. 2025. Neuro-Symbolic Programming in the Age of Foundation Models: Pitfalls and Opportunities. *Proceedings of Machine Learning Research vol vvv 1* (2025), 18.
- Jacob Steinhardt and Percy S Liang. 2015. Learning with Relaxed Supervision. In *NeurIPS*, Vol. 28.
- David Steinmann, Felix Divo, Maurice Kraus, Antonia Wüst, Lukas Struppek, Felix Friedrich, and Kristian Kersting. 2024. Navigating Shortcuts, Spurious Correlations, and Confounders: From Origins via Detection to Mitigation. *arXiv preprint arXiv:2412.05152* (2024).
- Ron Sun. 1992. On variable binding in connectionist networks. *Connection Science* 4, 2 (1992), 93–124.
- Mukund Sundararajan, Ankur Taly, and Qiqi Yan. 2017. Axiomatic attribution for deep networks. In *International conference on machine learning*. PMLR, 3319–3328.
- Raphael Suter, Djordje Miladinovic, Bernhard Schölkopf, and Stefan Bauer. 2019. Robustly disentangled causal mechanisms: Validating deep representations for interventional robustness. In *ICML*.
- Christian Szegedy, Vincent Vanhoucke, Sergey Ioffe, Jon Shlens, and Zbigniew Wojna. 2016. Rethinking the inception architecture for computer vision. In *Proceedings of the IEEE conference on computer vision and pattern recognition*. 2818–2826.
- Akihiro Takemura, Katsumi Inoue, and Masaaki Nishino. 2026. Constraint-Based Analysis of Reasoning Shortcuts in Neurosymbolic Learning. *arXiv preprint arXiv:2604.23377* (2026).
- Hao Tang and Kevin Ellis. 2023. From perception to programs: regularize, overparameterize, and amortize. In *ICML*.
- Merve Tapli, Quentin Bouniot, Wolfgang Stammer, Zeynep Akata, and Emre Akbas. 2026. Rethinking Concept Bottleneck Models: From Pitfalls to Solutions. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*. 9901–9910.
- Alice Tarzariol, Martin Gebser, and Konstantin Schekotihin. 2022. Lifting symmetry breaking constraints with inductive logic programming. *Machine Learning* 111, 4 (2022), 1303–1326.

- Open Ended Learning Team, Adam Stooke, Anuj Mahajan, Catarina Barros, Charlie Deck, Jakob Bauer, Jakub Sygnowski, Maja Trebacz, Max Jaderberg, Michaël Mathieu, Nat McAleese, Nathalie Bradley-Schmieg, Nathaniel Wong, Nicolas Porcel, Roberta Raileanu, Steph Hughes-Fitt, Valentin Dalibard, and Wojciech Marian Czarnecki. 2021. Open-Ended Learning Leads to Generally Capable Agents. *CoRR* abs/2107.12808 (2021). arXiv:[2107.12808](https://arxiv.org/abs/2107.12808) <https://arxiv.org/abs/2107.12808>
- Stefano Teso, Öznur Alkan, Wolfgang Stammer, and Elizabeth Daly. 2023. Leveraging Explanations in Interactive Machine Learning: An Overview. *Frontiers in Artificial Intelligence* (2023).
- Stefano Teso, Andrea Bontempelli, Fausto Giunchiglia, and Andrea Passerini. 2021. Interactive Label Cleaning with Example-based Explanations. In *Proceedings of the 35th International Conference on Neural Information Processing Systems*.
- Stefano Teso and Kristian Kersting. 2019. Explanatory interactive machine learning. In *Proceedings of the 2019 AAAI/ACM Conference on AI, Ethics, and Society*. 239–245.
- Sever Topan, David Rolnick, and Xujie Si. 2021. Techniques for symbol grounding with satnet. *Advances in Neural Information Processing Systems* 34 (2021), 20733–20744.
- Efthymia Tsamoura, Kaifu Wang, and Dan Roth. 2025. Imbalances in Neurosymbolic Learning: Characterization and Mitigating Strategies. In *Proceedings of the Thirty-Ninth Conference on Neural Information Processing Systems (NeurIPS)*.
- Vladimir Ulyantsev, Ilya Zakirzyanov, and Anatoly Shalyto. 2016. Symmetry breaking predicates for sat-based DFA identification. *arXiv preprint arXiv:1602.05028* (2016).
- Elena Umili, Emanuele Antonioni, Francesco Riccio, Roberto Capobianco, Daniele Nardi, and Giuseppe De Giacomo. 2021. Learning a Symbolic Planning Domain through the Interaction with Continuous Environments. In *Workshop on Bridging the Gap Between AI Planning and Reinforcement Learning (PRL)*.
- Elena Umili, Francesco Argenziano, and Roberto Capobianco. 2024. Neural Reward Machines. In *ECAI 2024 - 27th European Conference on Artificial Intelligence, 19-24 October 2024, Santiago de Compostela, Spain - Including 13th Conference on Prestigious Applications of Intelligent Systems (PAIS 2024)*. 3055–3062. <https://doi.org/10.3233/FAIA240847>
- Elena Umili and Roberto Capobianco. 2025. Learning Minimal Symbolic Representations and Temporal Rules from Visual Sequences. (04 2025). <https://doi.org/10.13140/RG.2.2.27713.26729>
- Elena Umili, Roberto Capobianco, and Giuseppe De Giacomo. 2023. Grounding LTLf Specifications in Image Sequences. In *Proceedings of the 20th International Conference on Principles of Knowledge Representation and Reasoning, KR 2023, Rhodes, Greece, September 2-8, 2023*. 668–678. <https://doi.org/10.24963/KR.2023/65>
- Leslie G Valiant. 1984. A theory of the learnable. *Commun. ACM* 27, 11 (1984), 1134–1142.
- Emile van Krieken, Erman Acar, and Frank van Harmelen. 2020. Analyzing Differentiable Fuzzy Implications. In *Proceedings of the International Conference on Principles of Knowledge Representation and Reasoning*, Vol. 17. 893–903.
- Emile van Krieken, Erman Acar, and Frank Van Harmelen. 2022. Analyzing differentiable fuzzy logic operators. *Artificial Intelligence* 302 (2022), 103602.
- Emile van Krieken, Samy Badreddine, Robin Manhaeve, and Eleonora Giunchiglia. 2024a. ULLER: A Unified Language for Learning and Reasoning. arXiv:[2405.00532](https://arxiv.org/abs/2405.00532) [cs.AI] <https://arxiv.org/abs/2405.00532>
- Emile van Krieken, Pasquale Minervini, Edoardo Ponti, and Antonio Vergari. 2024b. On the Independence Assumption in Neurosymbolic Learning. In *Proceedings of the 41st International Conference on Machine Learning (Proceedings of Machine Learning Research, Vol. 235)*. PMLR, 49078–49097.
- Emile van Krieken, Pasquale Minervini, Edoardo Ponti, and Antonio Vergari. 2025. Neurosymbolic Reasoning Shortcuts under the Independence Assumption. In *Proceedings of the 19th International Conference on Neurosymbolic Learning and Reasoning (Proceedings of Machine Learning Research, Vol. 284)*. PMLR.

- Emile van Krieken, Pasquale Minervini, Edoardo Maria Ponti, and Antonio Vergari. 2026. Neurosymbolic diffusion models. *Advances in neural information processing systems* 38 (2026), 135889–135931.
- Emile van Krieken, Thiviyan Thanapalasingam, Jakub M Tomczak, Frank van Harmelen, and Annette ten Teije. 2023. A-NeSI: A Scalable Approximate Method for Probabilistic Neurosymbolic Inference. In *Thirty-seventh Conference on Neural Information Processing Systems*.
- Vladimir Vapnik. 2013. *The nature of statistical learning theory*. Springer science & business media.
- Antonio Vergari, YooJung Choi, Anji Liu, Stefano Teso, and Guy Van den Broeck. 2021. A compositional atlas of tractable circuit operations for probabilistic inference. *Advances in Neural Information Processing Systems* 34 (2021), 13189–13201.
- Celeste Veronese, Daniele Meli, and Alessandro Farinelli. 2025. Learning Symbolic Persistent Macro-Actions for POMDP Solving Over Time. In *Proceedings of The 19th International Conference on Neurosymbolic Learning and Reasoning (Proceedings of Machine Learning Research, Vol. 284)*, Leilani H. Gilpin, Eleonora Giunchiglia, Pascal Hitzler, and Emile van Krieken (Eds.). PMLR, 1026–1040. <https://proceedings.mlr.press/v284/veronese25a.html>
- Julius von Kügelgen, Michel Besserve, Liang Wendong, Luigi Gresele, Armin Kekić, Elias Bareinboim, David Blei, and Bernhard Schölkopf. 2024. Nonparametric identifiability of causal representations from unknown interventions. *Advances in Neural Information Processing Systems* 36 (2024).
- Julius Von Kügelgen, Yash Sharma, Luigi Gresele, Wieland Brendel, Bernhard Schölkopf, Michel Besserve, and Francesco Locatello. 2021. Self-supervised learning with data augmentations provably isolates content from style. *Advances in neural information processing systems* 34 (2021), 16451–16467.
- Haobo Wang, Mingxuan Xia, Yixuan Li, Yuren Mao, Lei Feng, Gang Chen, and Junbo Zhao. 2022. SoLar: Sinkhorn Label Refinery for Imbalanced Partial-Label Learning. In *NeurIPS*.
- Hao Wang and Dit-Yan Yeung. 2020. A survey on Bayesian deep learning. *ACM computing surveys (csur)* 53, 5 (2020), 1–37.
- Kaifu Wang, Efthymia Tsamoura, and Dan Roth. 2023. On Learning Latent Models with Multi-Instance Weak Supervision. In *NeurIPS*.
- Po-Wei Wang, Priya Donti, Bryan Wilder, and Zico Kolter. 2019. Satnet: Bridging deep learning and logical reasoning using a differentiable satisfiability solver. In *International Conference on Machine Learning*. PMLR, 6545–6554.
- Malcom Ware, Eibe Frank, Geoffrey Holmes, Mark Hall, and Ian Witten. 2001. Interactive Machine Learning: Letting Users Build Classifiers. *International Journal of Human-Computer Studies* 55 (09 2001), 281–292. <https://doi.org/10.1006/ijhc.2001.0499>
- Alfred North Whitehead. 1927. *Symbolism, its meaning and effect*. Macmillan.
- Thomas Winters, Giuseppe Marra, Robin Manhaeve, and Luc De Raedt. 2022. DeepStochLog: Neural Stochastic Logic Programming. In *AAAI*.
- Antonia Wüst, Wolfgang Stammer, Quentin Delfosse, Devendra Singh Dhami, and Kristian Kersting. 2024. Pix2Code: Learning to Compose Neural Visual Concepts as Programs. In *Proceedings of the Fortieth Conference on Uncertainty in Artificial Intelligence (Proceedings of Machine Learning Research, Vol. 244)*, Negar Kiyavash and Joris M. Mooij (Eds.). PMLR, 3829–3852. <https://proceedings.mlr.press/v244/wust24a.html>
- Antonia Wüst, Wolfgang Stammer, Hikaru Shindo, Lukas Helff, Devendra Singh Dhami, and Kristian Kersting. 2026. Synthesizing Visual Concepts as Vision-Language Programs. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*. 17346–17356.
- Antonia Wüst, Tim Tobiasch, Lukas Helff, Inga Ibs, Wolfgang Stammer, Devendra S. Dhami, Constantin A. Rothkopf, and Kristian Kersting. 2025. Bongard in Wonderland: Visual Puzzles that Still Make AI Go Mad?
- Linhui Xiao, Xiaoshan Yang, Xiangyuan Lan, Yaowei Wang, and Changsheng Xu. 2025. Towards visual grounding: A survey. *IEEE Transactions on Pattern Analysis and Machine Intelligence* (2025).

- Xuan Xie, Kristian Kersting, and Daniel Neider. 2022. Neuro-symbolic verification of deep neural networks. In *IJCAI*.
- Yaqi Xie, Ziwei Xu, Mohan S Kankanhalli, Kuldeep S Meel, and Harold Soh. 2019. Embedding symbolic knowledge into deep networks. *Advances in neural information processing systems* 32 (2019).
- Jingyi Xu, Zilu Zhang, Tal Friedman, Yitao Liang, and Guy Broeck. 2018. A Semantic Loss Function for Deep Learning with Symbolic Knowledge. In *ICML*.
- Yiran Xu, Xiaoyin Yang, Lihang Gong, Hsuan-Chu Lin, Tz-Ying Wu, Yunsheng Li, and Nuno Vasconcelos. 2020. Explainable Object-Induced Action Decision for Autonomous Vehicles. In *IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*.
- Wen-Chi Yang, Giuseppe Marra, Gavin Rens, and Luc De Raedt. 2023a. Safe reinforcement learning via probabilistic logic shields. In *Proceedings of the Thirty-Second International Joint Conference on Artificial Intelligence*. 5739–5749.
- Xiao-Wen Yang, Wen-Da Wei, Jie-Jing Shao, Yu-Feng Li, and Zhi-Hua Zhou. 2024. Analysis for Abductive Learning and Neural-Symbolic Reasoning Shortcuts. In *ICML*.
- Yue Yang, Artemis Panagopoulou, Shenghao Zhou, Daniel Jin, Chris Callison-Burch, and Mark Yatskar. 2023b. Language in a bottle: Language model guided concept bottlenecks for interpretable image classification. In *CVPR*.
- Zhun Yang, Adam Ishay, and Joohyung Lee. 2020. NeurASP: Embracing Neural Networks into Answer Set Programming. In *IJCAI 2020*.
- Zhilin Yang, Peng Qi, Saizheng Zhang, Yoshua Bengio, William Cohen, Ruslan Salakhutdinov, and Christopher D Manning. 2018. HotpotQA: A Dataset for Diverse, Explainable Multi-hop Question Answering. In *Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing*. 2369–2380.
- Wenqian Ye, Guangtao Zheng, Xu Cao, Yunsheng Ma, and Aidong Zhang. 2024. Spurious correlations in machine learning: A survey. *arXiv preprint arXiv:2402.12715* (2024).
- Yu Yuan, Lili Zhao, Kai Zhang, Guangting Zheng, and Qi Liu. 2024. Do LLMs Overcome Shortcut Learning? An Evaluation of Shortcut Challenges in Large Language Models. In *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing*. 12188–12200.
- Mert Yuksekgonul, Maggie Wang, and James Zou. 2023. Post-hoc Concept Bottleneck Models. In *The Eleventh International Conference on Learning Representations*. <https://openreview.net/forum?id=nA5AZ8CEyow>
- Faried Abu Zaid, Dennis Diekmann, and Daniel Neider. 2023. Distribution-aware neuro-symbolic verification. *AISoLA* (2023), 445–447.
- Yujia Zheng, Shaoan Xie, and Kun Zhang. 2025. Nonparametric Identification of Latent Concepts. In *Forty-second International Conference on Machine Learning*.
- Zhi-Hua Zhou. 2019. Abductive learning: towards bridging machine learning and logical reasoning. *Science China. Information Sciences* 62, 7 (2019), 76101.
- Roland S Zimmermann, Yash Sharma, Steffen Schneider, Matthias Bethge, and Wieland Brendel. 2021. Contrastive learning inverts the data generating process. In *International conference on machine learning*. PMLR, 12979–12990.

Received 31 March 2026; accepted 01 July 2026