

A Foundation Model for Zero-Shot Logical Rule Induction

Yin Jun Phua

Institute of Science Tokyo

phua@comp.isct.ac.jp

Abstract

Inductive Logic Programming (ILP) learns interpretable logical rules from data. Existing methods are transductive: their learned parameters are bound to specific predicates and require retraining for each new task. We introduce Neural Rule Inducer (NRI), a pretrained model for zero-shot rule induction. Rather than encoding literal identities, NRI represents literals using domain-agnostic statistical properties such as class-conditional rates, entropy, and co-occurrence, which generalize across variable identities and counts without retraining. The model consists of a statistical encoder and a parallel slot-based decoder. Parallel decoding preserves the permutation invariance of logical disjunction; an autoregressive decoder would instead impose an arbitrary clause order. Product T-norm relaxation makes rule execution differentiable, allowing end-to-end training on prediction accuracy alone. We evaluate NRI on rule recovery, robustness to label noise and spurious correlations, and zero-shot transfer to real-world benchmarks, and we believe this work opens up the possibility of foundation models for symbolic reasoning. Code and the reference checkpoint are available at <https://github.com/phuayj/neural-rule-inducer>. An extended version with full appendices is available at <https://arxiv.org/abs/2605.04916>.

1 Introduction

Inductive Logic Programming (ILP) systems learn interpretable logical rules purely from examples. For example, an ILP system analyzing patient records with symptoms and diagnoses might induce a rule: “ $(fever \wedge cough) \vee (chills \wedge body_aches) \rightarrow flu$.” In high-stakes domains like healthcare and finance, black-box predictions are unacceptable. Transparent logical rules like this give us insights that we can actually trust and use.

Traditional symbolic ILP methods like [Muggleton and De Raedt, 1994; Quinlan, 1990; Srinivasan, 2001; Inoue *et al.*, 2014; Muggleton, 1995; Cropper and Morel, 2021] are sensitive to noise and often have to compromise between

precision and coverage. Recently, differentiable approaches such as [Evans and Grefenstette, 2018; Gao *et al.*, 2024; Johnson *et al.*, 2025] have been proposed. These methods have proven to be more robust to noise but they are still transductive, where learned weights are tied to specific predicates. A model trained on family relationships (Parent, Grandparent) cannot be transferred to biology (Protein, Enzyme). This means that for every new dataset, we need to retrain the models from scratch.

So how do we determine which variable (boolean attributes like fever, cough) is predictive of a label (targets like flu)? We select variables by their statistical signatures rather than by name. Because these signatures are invariant to renaming and reordering and each literal is represented by a fixed-dimensional vector independent of N , they offer a plausible route to zero-shot transfer. A high class-conditional rate indicates a variable belongs in a rule, regardless of what it represents. In this paper we test this hypothesis by pretraining on synthetic DNFs and evaluating on held-out real-world tabular tasks without retraining.

This approach requires addressing three challenges. First, *variable-sized problems*: generalization from fixed-size training data to problems with fewer or more variables. Second, *inter-variable dependencies*: individual statistics miss redundancy and complementarity (e.g., XOR patterns). Third, *synthetic-to-real transfer*: learning the abstract procedure of induction rather than overfitting to synthetic data.

We present **Neural Rule Inducer (NRI)**, a foundation model [Bommasani, 2021] that consists of the following:

- **Literal Statistics Encoder** encodes literals by statistical properties such as how often the literal is true and co-occurrences with other literals.
- **Parallel Slot-Based Decoder**. This module synthesizes multiple clauses in parallel using learned slot queries. Compared to autoregressive models, this module preserves the permutation invariance of logical disjunction.
- **T-Norm Training**. We use product t-norm relaxation to execute rules differentially. This allows end-to-end training without explicit clause supervision.
- **Synthetic Data Training**. We train entirely on randomly generated boolean formulas. This trains the model to recognize the induction procedure itself rather

than domain-specific patterns, and removes the limitation on training data size.

We show that a model trained entirely on synthetic boolean formulas can perform zero-shot rule induction on diverse real-world benchmarks. Figure 1 illustrates the overall architecture of NRI. While in this work, we focus on boolean variables, the statistical encoding framework can be extended to multi-valued and continuous domains via discretization or fuzzy predicates.

2 Related Works

Differentiable ILP. [Evans and Grefenstette, 2018] proposed learning rules via gradient descent, but their method requires specifying rule templates and a fixed set of background predicates. NeuralLP [Yang *et al.*, 2017] and DRUM [Sadeghian *et al.*, 2019] extended this to knowledge base reasoning using TensorLog. Neural Theorem Provers [Rocktäschel and Riedel, 2017] introduced differentiable unification for end-to-end theorem proving. Logic Tensor Networks [Serafini and Garcez, 2016] use fuzzy semantics to integrate logical constraints into neural learning. DeepProbLog [Manhaeve *et al.*, 2018] extends probabilistic logic programming with neural predicates, enabling end-to-end learning of both neural and symbolic components. NeurASP [Yang *et al.*, 2020] integrates neural networks with answer set programming, allowing neural network outputs to serve as probabilistic inputs to logic programs. GLIDR [Johnson *et al.*, 2025] generalized this to graph-like topologies, utilizing differentiable message passing to learn recursive and cyclic dependencies.

Most differentiable ILP systems instantiate parameters against a fixed predicate or schema inventory, so transferring to a new dataset typically requires retraining. Classical symbolic systems such as FOLD-R++ [Wang and Gupta, 2022], which learns answer-set rules from mixed numerical and categorical data by top-down heuristic search, are similarly task-specific and are re-run from scratch on each new task.

Phua and Inoue [Phua and Inoue, 2024] proposed a model allowing zero-shot transfer learning. They addressed scaling issues by exploiting variable permutation symmetries. Our proposed method differs in mechanism where instead of utilizing the raw examples, we utilize statistical properties to handle missing and noisy data.

Generative Neuro-Symbolic AI. LLM-based methods like ILP-CoT [Peng *et al.*, 2025] and DeepSeek-Prover-V2 [Ren *et al.*, 2025] generate hypotheses by relying on knowledge encoded in language and human concepts. LINC [Olausson *et al.*, 2023] combines language models with first-order logic provers, using LLMs to translate natural language into formal logic for external verification. However, LLMs are not grounded in reality, they may use abstract symbols or notions that might not correspond to any measurable quantity in the real world. In contrast, our approach generates hypotheses directly from observed data and does not assume an explicit semantic model. TabPFN [Hollmann *et al.*, 2023] demonstrated that transformers trained on synthetic tabular data can perform in-context learning for classification. While TabPFN

produces black-box predictions, our approach outputs interpretable logical rules.

3 Background

We want to learn a function $f : (\mathcal{X}, \mathcal{Y}) \rightarrow \mathcal{R}$ that takes input examples $\mathcal{X}$ (boolean variables) and labels $\mathcal{Y}$ and outputs a logical hypothesis $\mathcal{R}$ in Disjunctive Normal Form (DNF). Importantly, f should work regardless of how many variables N there are and what they represent. We should not need to retrain for each new problem.

3.1 Disjunctive Normal Form (DNF)

DNF is a disjunction of conjunctions: $\mathcal{R} = C_1 \vee C_2 \vee \dots \vee C_K$, where each clause C_k is a conjunction of literals: $C_k = l_{k,1} \wedge l_{k,2} \wedge \dots$. A *literal* l_j is either a variable x_i or its negation $\neg x_i$. DNF is a canonical form that can express any boolean function. Therefore, by learning DNFs we can learn any propositional rule.

3.2 T-Norms

A triangular norm (t-norm) is a mathematical operation on $[0, 1]$ that relaxes logical conjunction to continuous values [Klement *et al.*, 2013]. The product t-norm defines the following operations:

- **Negation:** $\neg x = 1 - x$
- **Conjunction:** $x \wedge y = x \cdot y$
- **Disjunction:** $x \vee y = 1 - (1 - x)(1 - y)$ (via De Morgan)

Under product t-norm, the truth value of a clause C_k containing literals indexed by set S_k is:

$$C_k = \prod_{i \in S_k} l_i \quad (1)$$

where l_i is the truth value of literal i . A DNF formula can then be calculated differentially as follows $\mathcal{R} = 1 - \prod_k (1 - C_k)$.

3.3 Inductive Logic Programming

Inductive Logic Programming (ILP) systems learn logical rules from examples [Muggleton and De Raedt, 1994]. Given positive examples E^+ , negative examples E^- , and background knowledge B , an ILP system finds a hypothesis H such that $B \wedge H \models E^+$ (completeness) and $B \wedge H \not\models E^-$ (consistency).

Two settings exist in the ILP literature. In *learning from entailment*, examples are ground facts. The hypothesis must logically entail positives and not entail negatives. In *learning from interpretations* [Inoue *et al.*, 2014], each example is a complete state. Rules explain state transitions or classifications. Our setting mainly follows learning from interpretations, where each example is a complete boolean assignment. The output of our ILP system is a DNF rule that classifies all positive examples and none of the negative examples.

3.4 Foundation Models

A foundation model is a model trained on huge amounts of data and transfers the learning to downstream tasks without task-specific training [Bommasani, 2021]. Three properties

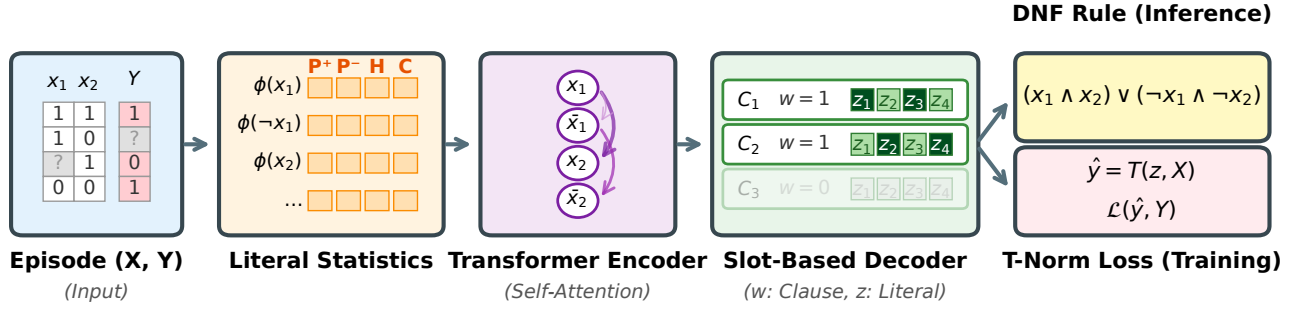


Figure 1: Neural Rule Inducer takes an episode (X, Y) as input and calculates literal statistics. For each variable we calculate $\phi(x_i)$, $\phi(\neg x_i)$ which consists of class-conditional rates (P^+ , P^-), entropy (H), and co-occurrence strength (C). We then apply cross-attention over these statistics. The slot-based decoder produces K candidate clauses in parallel using learned literal gates z and clause gates w . By evaluating the produced rule with T-norm, we can perform end-to-end training. The rules are discretized to then produce an interpretable DNF rule.

define this paradigm: training on diverse data at scale, generalization beyond the training distribution, and adaptation to new tasks via prompting or fine-tuning. GPT and CLIP are examples in language and vision.

We apply this paradigm to rule induction. Our training data consists of millions of synthetic boolean formulas with diverse rule structures. The model does not learn weights for specific predicates. Instead, it learns to recognize statistical patterns of literals that allows the model to infer that the literal belongs to a rule. At inference, the model is able to induce rules for new domains without retraining.

4 Neural Rule Inducer (NRI)

In this section we propose NRI, an end-to-end differentiable framework that does domain-agnostic rule learning. Rather than learning weights for specific predicates for a specific domain, NRI learns to select variables based on their statistical properties within the domain.

Design Rationale. NRI separates three roles: the statistical encoder provides identity-free cues about which literals matter, the example-conditioned encoder recovers which examples each literal covers, and the parallel decoder assembles clauses without imposing an artificial order on a disjunction. FiLM breaks symmetry among clause slots so different clauses can specialize. Section 4.8 and Table 2 study loss-level contributions; fuller architectural swaps are deferred because they would change these symmetry and transfer assumptions.

4.1 Problem Formulation

Our setup is structurally a meta-learning problem: the training distribution is over entire episodes, not individual literals. For a sampled N , the causal literal set is $\mathcal{L}_N = \{x_1, \neg x_1, \dots, x_N, \neg x_N\}$; we sample a bounded DNF rule R over $\mathcal{L}_N$, draw a causal matrix $X_c \in \{0, 1\}^{M \times N}$, set $Y = R(X_c)$, and then optionally add missingness, label noise, and concatenated spurious variables (Appendix A). The hypothesis space at inference is therefore bounded DNFs over the literal set of the target task. At evaluation time we

test on a fixed collection of UCI tasks rather than a parametric test distribution, with $\mathcal{L}_N$ determined by binarizing each task’s features. Within a task, generalization from the support split to the held-out split is the usual i.i.d. setting; across tasks, transfer relies on the inductive-bias assumption that many binarized tasks admit sparse bounded-DNF descriptions. We do not provide formal cross-task guarantees here; the contribution is architectural and empirical.

Given $X \in \{0, 1\}^{M \times N}$ and $Y \in \{0, 1\}^M$ where M is the number of examples and N is the number of variables (or features), we would like to find a DNF hypothesis in the following form:

$$\mathcal{R} = \bigvee_{k=1}^K \left(\bigwedge_j l_j \right) \quad (2)$$

such that $\mathcal{R}(X) = Y$.

NRI can be separated into four stages (depicted in Figure 1). First, we calculate statistical properties $\Phi \in \mathbb{R}^{2N \times D}$ for each variable. Then, we project these statistical properties via cross-attention over examples into $Z \in \mathbb{R}^{2N \times d}$. Next, we produce a DNF rule via slot-based attention. Finally, we evaluate DNF rule using differentiable T-norms for end-to-end training.

4.2 Literal Statistics Encoder

To achieve zero-shot generalization with robustness to missing and noisy data, we encode each literal by its **statistical properties** over all M examples. This differs from prior works that use raw samples directly as input.

For each literal l_j (where $j \in \{1, \dots, 2N\}$ indexes both positive literals x_i and negations $\neg x_i$), we compute a feature vector $\phi_j \in \mathbb{R}^D$ containing:

$$\phi_j = [P(l_j|y=1), P(l_j|y=0), P(l_j), \mathcal{H}(l_j), \text{sgn}_j, \bar{c}_j, \dots] \quad (3)$$

where $P(l_j|y)$ denotes literal truth rates, $\mathcal{H}(l_j)$ is the binary entropy, $\text{sgn}_j \in \{0, 1\}$ indicates literal polarity (positive/negative), and $\bar{c}_j$ is the mean absolute co-occurrence with other literals. The full feature vector includes 18 components (see Appendix B).

These statistics are fed into an MLP to produce an initial embedding:

$$h_j^{(0)} = \text{MLP}(\phi_j) \in \mathbb{R}^d \quad (4)$$

The observation-rate features in ϕ_j explicitly quantify how much of the input is observed. For example, if 30% of the values for literal j are missing among positive examples, the truth rate $P(l_j=1 \mid y=1)$ is computed only from the observed 70%, and $\text{obs}^+ = 0.7$ indicates reduced statistical support. Noisy literals manifest similarly: truth rates move toward 0.5 and entropy increases. As observation rates fall, the signal becomes correspondingly weaker because these statistics are estimated from fewer effective samples.

4.3 Example-Conditioned Encoding

The aggregate statistics ϕ_j are not claimed to be information-theoretically sufficient for arbitrary DNFs: two literals can share the same marginals yet cover different subsets of positive examples. The example-conditioned encoder is therefore used to restore this support-pattern information. Our claim is empirical adequacy of this compression rather than formal sufficiency.

Each example m is represented by a key vector e_m combining its label and literal values:

$$e_m = \text{MLP}_y([y_m, 1-y_m, \mathbf{1}_m]) + \text{MLP}_x(l^{(m)}) \quad (5)$$

where $l^{(m)} \in [0, 1]^{2N}$ is the vector of literal truth values for example m , and MLP_x uses a 64-dimensional bottleneck.

Dynamic Dimension Adaptation. MLP_x is trained with a fixed input dimension $2N_{\text{train}}$. At inference, when facing problems with $N \neq N_{\text{train}}$, we adapt the linear layer on-the-fly. If $N < N_{\text{train}}$, we zero-pad the input to match the trained dimension. If $N > N_{\text{train}}$, we expand the first layer: trained weights are copied for the first $2N_{\text{train}}$ dimensions, and the remaining $2(N - N_{\text{train}})$ dimensions are initialized randomly. This adaptation affects only the auxiliary example-conditioned branch; the main literal-statistics pathway is already dimension-free in N . The 64-dimensional bottleneck limits the influence of newly initialized weights, so this mechanism is a pragmatic approximation rather than a principled invariance guarantee.

The literal embeddings are produced by applying attention to these example keys:

$$h_j^{(1)} = h_j^{(0)} + \text{MultiHeadAttn}(h_j^{(0)}, \{e_m\}_{m=1}^M, \{e_m\}_{m=1}^M) \quad (6)$$

4.4 Clause-Conditioned Encoding via FiLM

The embeddings $h_j^{(1)}$ are shared across all clause slots. Without differentiation, clause slots will tend to converge to identical patterns during training. We therefore apply Feature-wise Linear Modulation (FiLM) [Perez *et al.*, 2018] to give each clause slot k a unique view:

$$h_{k,j} = \gamma_k \odot h_j^{(1)} + \beta_k \quad (7)$$

where $\gamma_k, \beta_k \in \mathbb{R}^d$ are learnable per-clause parameters. We initialize β_k orthogonally and $\gamma_k \sim \mathcal{N}(1, 0.5^2)$, which encourages clause slots to differentiate.

4.5 Slot-Based Decoder

We treat the DNF synthesis problem as a **slot-filling problem**. Given T clause slots, the model decides which literals to include in each clause and which clauses to activate. Unlike autoregressive generation, the design of parallel slots aims to preserve permutation invariance in the generated rule ($A \vee B \equiv B \vee A$).

The decoder is a 3-layer Transformer Decoder [Vaswani *et al.*, 2017] with 4 attention heads. Each clause slot k has a learnable query $q_k \in \mathbb{R}^d$. The decoder computes the following:

$$s_k = \text{TransformerDec}(q_k + \bar{h}, \{h_{k,j}\}_{j=1}^{2N}) \quad (8)$$

where $\bar{h} = \frac{1}{2N} \sum_j h_j^{(1)}$ is an average of all literal embeddings.

Literals Gates (“AND” Level)

For each clause k , we compute the probability a literal gets included:

$$z_{k,j} = \sigma \left(\frac{1}{\sqrt{d}} \langle W_s s_k, W_h h_{k,j} \rangle + b \right) \quad (9)$$

where $W_s, W_h \in \mathbb{R}^{d \times d}$ are learnable projections and b is a bias term. Since a clause containing both x_i and $\neg x_i$ is always false, we remove the literal with the lower score in each complementary pair.

Clause Gates (“OR” Level)

A clause gate $w_k \in [0, 1]$ determines whether slot k is active:

$$w_k = \sigma \left(\text{MLP} \left([\bar{h}, \tilde{h}_k, s_k, p_k] \right) \right) \quad (10)$$

where $\tilde{h}_k = \sum_j z_{k,j} h_{k,j} / \sum_j z_{k,j}$ is the clause’s weighted literal summary, and $p_k = 1 - \prod_j (1 - z_{k,j})$ is the probability that at least one literal is selected (non-null probability).

Clause Selection at Inference. The deployed inference rule retains every clause whose gate satisfies $w_k \geq 0.5$. As a diagnostic for clause quality we additionally compute a discrimination score:

$$d_k = \frac{1}{|Y^+|} \sum_{m: y_m=1} C_k^{(m)} - \frac{1}{|Y^-|} \sum_{m: y_m=0} C_k^{(m)} \quad (11)$$

We explored top- K' filtering by d_k ; it improves exact-match rule recovery on synthetic benchmarks but did not help UCI accuracy in our experiments, so it is not used by default.

4.6 Neuro-Symbolic Execution

The decoder produces literal gates $z_{k,j}$ and clause gates w_k that define a soft DNF rule. These gates are computed once per episode from all M examples. We then evaluate this rule on each example m using product t-norms.

For each clause k and example m , the clause truth value is:

$$C_k^{(m)} = \prod_{j=1}^{2N} \left(1 - z_{k,j} \cdot (1 - l_j^{(m)}) \right) \quad (12)$$

When $z_{k,j} \rightarrow 0$, the term becomes 1 (literal ignored). When $z_{k,j} \rightarrow 1$, the term reduces to $l_j^{(m)}$ (clause depends on literal j).

The final prediction is a combination of all clauses via the probabilistic OR:

$$\hat{y}^{(m)} = 1 - \prod_{k=1}^K (1 - w_k \cdot C_k^{(m)}) \quad (13)$$

The prediction is high when at least one active clause ($w_k \approx 1$) is satisfied ($C_k^{(m)} \approx 1$). The entire computation from Eq. 3 to Eq. 13 is differentiable, enabling end-to-end training.

4.7 Training Strategy

The combinatorial search space for ILP problems makes cold-start training difficult, particularly in a foundation model setting. To overcome this, we employ these techniques:

1. **Synthetic Pre-training:** We train the model on random boolean formulas, allowing it to learn the induction algorithm without real-world data.
2. **Spurious Environment Training:** We inject spurious features into each episode with *opposite* correlations across two environments (first and second half of examples). After shuffling examples, these features should be *marginally independent* from the label. The model must then learn to identify features that perform well on one environment but poorly on the other. See Appendix A.4 for details.
3. **Clause Dropout:** During training, we randomly drop 25% of clause slots (keeping at least 2) to prevent any single clause from dominating.

4.8 Training Objective

We utilize a multi-objective loss function to balance between accuracy, slot utilization, clause diversity, gate sharpness, margin enforcement, and counterfactual necessity:

$$\mathcal{L} = \mathcal{L}_{\text{cov}} + \lambda_b \mathcal{L}_{\text{bal}} + \lambda_r \mathcal{L}_{\text{rep}} + \lambda_e \mathcal{L}_{\text{ent}} + \lambda_m \mathcal{L}_{\text{mm}} + \lambda_{\text{cf}} \mathcal{L}_{\text{cf}} \quad (14)$$

The difficulty in training a stable, parallel slot decoder with logical consideration necessitates such a complex objective. We describe each loss component in detail in the following paragraphs.

Coverage Loss. Binary cross-entropy between predictions $\hat{y}$ and labels y .

Clause Slot Load-Balancing. Without explicit regularization, the T-norm calculation tends to lead to a few clause slots dominating while others receive vanishing gradients. To overcome this, we employ multiple complementary balancing losses from the Mixture-of-Experts literature [Fedus et al., 2022]. First, we use the auxiliary load-balancing loss from Switch Transformers: $K \cdot \sum_k u_k \cdot f_k$, where u_k is the mean routing probability to slot k and f_k is the fraction of assignments. Second, we apply a CV² (coefficient of variation squared) loss on normalized clause slot usage. Let $\bar{w}_k = \frac{1}{B} \sum_b w_k^{(b)}$ denote the mean clause gate activation across a batch:

$$\mathcal{L}_{\text{bal}} = K \cdot \sum_{k=1}^K \left(\frac{\bar{w}_k}{\frac{1}{K} \sum_{k'} \bar{w}_{k'}} - 1 \right)^2 \quad (15)$$

We additionally apply CV² on raw clause gate activations (before normalization) to prevent clause slots from going permanently inactive. These losses produce gradients that counteract the task loss gradients, reducing clause slot utilization variance from 0.35 to 0.003 in our experiments.

Max Margin Coverage. BCE rewards spreading probability across clauses (diffusion), while clause gates reward specialization. In our experiments, this conflict caused clause gate entropy to oscillate between low values (few clauses active) and high values (many clauses active).

We add a max-margin loss [Cortes and Vapnik, 1995] that only penalizes the *best* clause when it falls short of a margin. This removes the incentive for diffusion:

$$\mathcal{L}_{\text{mm}} = \mathbb{E}_{y=1} [\max(0, \tau^+ - C_{\text{max}})] + \mathbb{E}_{y=0} [\max(0, C_{\text{max}} - \tau^-)] \quad (16)$$

where $C_{\text{max}} = \max_k C_k$ is the maximum clause truth value, and $\tau^+ = 0.7$, $\tau^- = 0.3$. Unlike BCE which pushes all clauses to cover examples, max-margin allows multiple clause slots to cover the same pattern. This stabilizes training at low entropy (~ 0.26) where clauses use generalizable features rather than differentiating via rare attributes.

Counterfactual Necessity. BCE alone cannot distinguish necessary literals from spurious ones. A literal that only happens to correlate with the label may be selected even without causal relationship. To address this, we add a counterfactual test: if a literal is selected it should be causal, i.e. flipping its value should break the clause. We define the counterfactual loss as $\mathcal{L}_{\text{cf}} = \mathcal{L}_{\text{nec}} + \mathcal{L}_{\text{spur}} + \lambda_o \mathcal{L}_{\text{ovl}} + \lambda_c \mathcal{L}_{\text{cf-bal}}$, where $\mathcal{L}_{\text{ovl}}$ ($\lambda_o=0.1$) penalizes pairs of clauses that both cover the same positive example and $\mathcal{L}_{\text{cf-bal}}$ ($\lambda_c=0.01$) is a mild regularizer encouraging balanced clause usage on positives within the CF objective; full forms are in Appendix C.

Necessity: BCE tends to reward correct predictions regardless of whether selected literals are causally necessary or merely correlated. So if a literal that has no causal relationship happens to be included in a correct prediction, the model might learn a spurious correlation. To overcome this, we use this loss to provide gradient signal toward causal selection. For positive examples, we flip selected literals (gates $z_{s,j}$ above threshold) and recompute clause truth C'_k . The loss minimizes post-flip clause truth:

$$\mathcal{L}_{\text{nec}} = \mathbb{E}_{y=1} \left[\sum_k r_k \cdot C'_k \right] \quad (17)$$

where $r_k = \text{softmax}_k(C_k)$ is a weight applied to each clause to emphasize important clauses.

Spuriousness: On the other hand, if a model decided that a literal should be ignored (gate $z_{s,j}$ below threshold), the model claims it is not part of the rule. If the prediction changes when we flip an ignored literal, the model was implicitly relying on something it claimed to ignore. This loss penalizes such inconsistency. For positive examples, we flip ignored literals and penalize if the prediction changed due to the flip:

$$\mathcal{L}_{\text{spur}} = \mathbb{E}_{y=1} [|\hat{y} - \hat{y}'|] \quad (18)$$

where $\hat{y}'$ is the result of the T-norm after flipping ignored literals.

5 Experiments

We train exclusively on synthetic boolean expressions and test zero-shot on real-world UCI datasets [Asuncion *et al.*, 2007] and benchmarks probing rule complexity, noise robustness, and spurious variable robustness. For each synthetic episode we sample $N \sim \text{Unif}\{6, \dots, 12\}$ and $M \sim \text{Unif}\{24, \dots, 48\}$ independently, and DNF rules with up to $K = 6$ clauses and $L = 4$ literals per clause. Each episode includes 3 spurious environment features (Appendix A). The decoder uses $T = 8$ clause slots with 25% dropout. Training uses AdamW ($\text{lr}=6 \times 10^{-4}$, weight decay 10^{-2}), batch size 8192, and 500 steps.

5.1 Baseline Context

Table 1 compares NRI against 8 baselines on 14 UCI datasets. Continuous features are binarized using median thresholding (e.g., $\text{age}_{>\text{med}}$ is 1 if age exceeds the median). Missing values are preserved: NRI handles them through observation-rate features and by treating unknown literals as 0.5 (maximally uncertain) during rule evaluation. Baselines use median imputation.

NRI is evaluated *zero-shot* whereas other baseline methods are trained per dataset. Notably, 12 of 14 datasets have more features than the training range ($N > 12$), testing out-of-distribution generalization to larger schemas. All methods are evaluated under 5-fold stratified cross-validation: in each fold, one fold ($\sim 20\%$ of the data) acts as the support set on which the rule is induced, and accuracy is reported on the remaining $\sim 80\%$. This support-set size targets the low-data regime where zero-shot transfer is most valuable. We compare against gradient boosting methods (XGBoost, LightGBM), generalized additive models (EBM), classical rule learners (RIPPER, RuleFit, FIGS), decision trees (DT), and neural DNF (N-DNF). Full descriptions of the baselines are in Appendix E.

Zero-shot NRI achieves 69.7%, 13 points below EBM, which is unsurprising because EBM is trained separately on each dataset. Performance is strongest on the two in-distribution datasets (diabetes with $N=8$, breast-cancer with $N=9$), where NRI achieves 68.0% and 88.3% respectively.

Performance varies substantially across datasets due to two factors evident in the experimental data. Firstly, we converted car and nursery which are originally multi-class problems (4 and 5 classes respectively) to single class, the dataset is difficult to classify with just simple DNF rules. Second, kr-vs-kp requires all 8 clause slots for classification, which confirms that some datasets genuinely require more clauses than NRI’s training distribution ($K \leq 6$).

The learned rules are interpretable. For diabetes prediction (68.0% accuracy), NRI induces: $(\text{glucose}_{>\text{med}} \wedge \text{age}_{>\text{med}})$, capturing that elevated plasma glucose combined with older age predicts diabetes. For breast cancer diagnosis (88.3% accuracy): $(\text{cell_size}_{>\text{med}} \wedge \text{cell_shape}_{>\text{med}} \wedge \text{bare_nuclei}_{>\text{med}})$, identifying that larger, irregularly shaped cells with prominent nuclei indicate malignancy. These are clinically plausible summaries of the binarized benchmarks, not validated medical decision rules, but they show NRI can recover human-readable patterns without task-specific training.

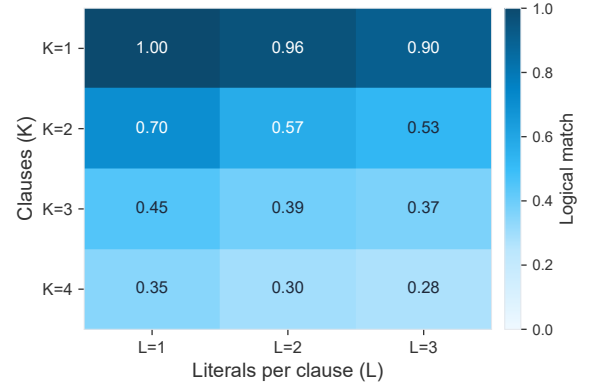


Figure 2: Heatmap of logical match rate (%) across rule complexity dimensions at $N=12$. Darker colors indicate higher accuracy.

5.2 Rule Complexity Scaling

We evaluate rule recovery on synthetic DNF formulas with varying clause count $K \in \{1, 2, 3, 4\}$ and literals per clause $L \in \{1, 2, 3\}$, testing at $N=12$ features (in-distribution, matching training range $N \in [6, 12]$). For each (K, L) configuration, we generate 200 random DNF rules per seed across 10 seeds and measure logical equivalence between the predicted and ground-truth rules.

Figure 2 shows that recovery degrades with complexity along both dimensions. For the simplest rules ($K=1, L=1$), logical match reaches 99.5%. This drops to 91.5% at $L=2$ and 81.7% at $L=3$. Increasing clause count has a larger effect: at $K=4, L=1$, recovery falls to 34.2%, and the most complex rules ($K=4, L=3$) achieve 24.0% logical match. Prediction accuracy remains high (85–100%) even when logical match is lower. The clause dimension (K) dominates difficulty because multi-clause rules require discovering multiple independent patterns, while longer clauses (L) only demand more precise literal selection within a single pattern.

5.3 Label Noise Robustness

We evaluate robustness to label noise (random flips) against two interpretable baselines, RIPPER and DT. Figure 3 shows NRI accuracy stays stable (92.3% $\rightarrow$ 87.4% at 30% noise), while RIPPER and DT degrade sharply (98.4% $\rightarrow$ 70.3% and 100% $\rightarrow$ 69.9% respectively). Symbolic methods are more accurate at low noise by fitting training data exactly, but beyond 15% noise NRI’s statistical encoding provides implicit regularization and outperforms both baselines by 17 percentage points at 30% noise.

5.4 Spurious Variable Robustness

We test whether NRI ignores spurious features that correlate with labels but are not in the true rule. We append $D \in \{0, 4, 8, 16, 32\}$ distractors per episode with $P(d=1|Y=1) = \rho$ and $P(d=1|Y=0) = 1-\rho$ for $\rho \in \{0.1, \dots, 0.9\}$, statistically predictive but causally irrelevant features. Figure 4 shows accuracy remains above 92% across all settings (97.6% at $D=32, \rho=0.9$; 96.7% at $D=32, \rho=0.1$ as the model learns to ignore the noise).

Dataset	N	XGB	LGBM	EBM	RIPPER	RuleFit	FIGS	DT	N-DNF	NRI [†]	Gap
adult	105	82.2 ± 0.2	82.7 ± 0.2	83.8 ± 0.1	81.9 ± 0.3	82.6 ± 0.2	82.8 ± 0.2	82.1 ± 0.2	83.0 ± 0.1	69.6 ± 4.4	-14.2
breast-cancer	9	92.0 ± 1.4	65.5 ± 0.0	93.0 ± 1.2	88.3 ± 1.4	90.8 ± 2.2	89.6 ± 2.3	89.3 ± 2.6	90.3 ± 1.3	88.3 ± 0.3	-4.7
car	21	80.5 ± 1.1	75.0 ± 1.6	81.2 ± 0.3	90.3 ± 0.9	94.2 ± 0.6	91.8 ± 1.3	78.7 ± 1.3	93.3 ± 1.0	51.2 ± 5.4	-43.0
credit	46	82.1 ± 1.4	55.5 ± 0.0	82.2 ± 1.6	83.1 ± 2.1	81.7 ± 1.8	81.2 ± 1.7	81.3 ± 1.7	81.4 ± 2.6	71.5 ± 7.3	-11.6
diabetes	8	67.4 ± 2.6	65.1 ± 0.0	66.9 ± 1.8	66.8 ± 2.2	65.5 ± 1.9	64.9 ± 2.0	65.3 ± 1.6	64.4 ± 2.3	68.0 ± 2.4	+0.0
german	61	66.9 ± 1.1	70.0 ± 0.0	68.1 ± 1.7	58.3 ± 3.2	65.8 ± 1.0	64.6 ± 2.0	64.6 ± 2.4	68.1 ± 1.4	59.8 ± 3.8	-10.2
hepatitis	32	79.4 ± 0.0	79.4 ± 0.0	81.2 ± 2.0	20.6 ± 0.0	20.6 ± 0.0	72.1 ± 3.2	73.8 ± 4.3	68.5 ± 5.3	55.9 ± 3.7	-25.3
ionosphere	34	66.3 ± 4.3	64.1 ± 0.0	69.1 ± 3.9	59.5 ± 3.3	69.2 ± 6.0	63.8 ± 6.3	65.4 ± 5.1	67.8 ± 4.5	62.8 ± 3.9	-6.4
kr-vs-kp	73	93.2 ± 0.6	90.7 ± 1.9	92.5 ± 0.5	89.9 ± 1.2	93.8 ± 0.7	92.7 ± 0.8	92.4 ± 0.6	92.0 ± 0.8	72.3 ± 5.3	-21.5
mushroom	116	99.0 ± 0.2	99.1 ± 0.2	99.4 ± 0.1	98.5 ± 0.4	99.5 ± 0.2	99.4 ± 0.2	99.3 ± 0.2	98.8 ± 0.3	87.8 ± 3.9	-11.7
nursery	27	90.5 ± 0.5	91.2 ± 0.4	91.1 ± 0.7	86.2 ± 1.0	88.5 ± 0.8	87.1 ± 0.5	84.8 ± 0.7	86.6 ± 0.7	71.3 ± 4.3	-19.9
spambase	57	89.7 ± 0.6	89.8 ± 0.4	91.1 ± 0.6	85.7 ± 1.2	89.8 ± 0.4	85.7 ± 0.7	85.0 ± 1.0	87.8 ± 0.7	71.9 ± 0.7	-19.2
tic-tac-toe	27	69.4 ± 2.6	65.3 ± 0.0	70.2 ± 1.2	56.0 ± 4.4	65.7 ± 1.5	64.8 ± 2.0	64.8 ± 1.3	67.8 ± 2.7	56.6 ± 2.5	-13.6
vote	32	92.9 ± 1.8	61.4 ± 0.0	87.4 ± 2.1	91.6 ± 2.2	91.1 ± 1.6	90.5 ± 1.7	90.7 ± 2.4	91.8 ± 1.6	88.3 ± 1.8	-4.6
Mean		82.2 ± 11.2	75.4 ± 13.5	82.7 ± 10.6	75.5 ± 21.0	78.5 ± 20.4	80.8 ± 12.4	79.8 ± 11.5	81.6 ± 11.8	69.7 ± 12.0	-13.0

Table 1: Baseline comparison on 14 UCI datasets (5-fold CV accuracy %). N = number of boolean features after median binarization. NRI is trained on $N \in [6, 12]$; 12/14 datasets are OOD ($N > 12$). Methods: gradient boosting (XGB, LGBM), GAM (EBM), rule-based (RIPPER, RuleFit, FIGS, DT), neural DNF (N-DNF), and our zero-shot NRI[†]. Best per row in bold.

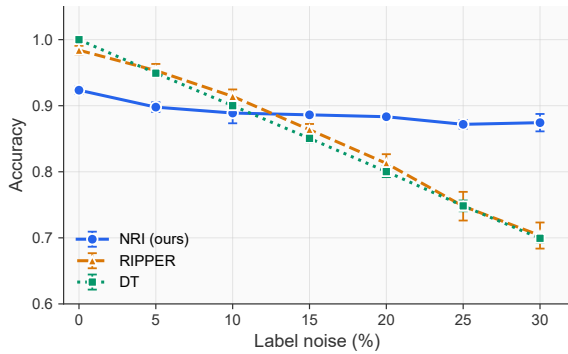


Figure 3: Noise robustness comparison. NRI (blue) maintains stable accuracy as noise increases. RIPPER (orange) and DT (green) degrade sharply. The crossover occurs around 15% noise.

Ablation	UCI Accuracy (%)	Δ Acc	Avg Clauses	Avg Literals
Full (baseline)	74.8 ± 3.8	—	7.3 ± 1.2	3.2 ± 0.2
– CF Necessity	74.1 ± 4.7	-0.6	7.8 ± 0.5	3.1 ± 0.3
– Max-Margin	72.0 ± 2.4	-2.8	7.7 ± 0.4	3.2 ± 0.2
– Slot Balance	73.3 ± 3.7	-1.5	7.0 ± 1.5	3.0 ± 0.5

Table 2: Loss function ablation study. Each row removes one loss component.

5.5 Computational Scaling

We measure inference time and memory as problem size (N) and example count (M) vary. Latency stays nearly constant (~ 7.5 ms) as M grows from 32 to 512, and increases sub-linearly with N (4.2ms $\rightarrow$ 11.8ms for a $32\times$ increase from $N=16$ to $N=512$); memory scales $O(N^2)$ due to attention. At $N=512$, inference completes in under 12ms with 593MB peak memory. Full curves are in Appendix D.

5.6 Loss Ablation Study

Each loss contributes (Table 2): removing CF Necessity, Max-Margin Coverage, or Slot Balance drops UCI accuracy by 0.7%, 2.8%, and 1.5% respectively, with characteristic failure modes (more spurious-literal selection, unstable gate

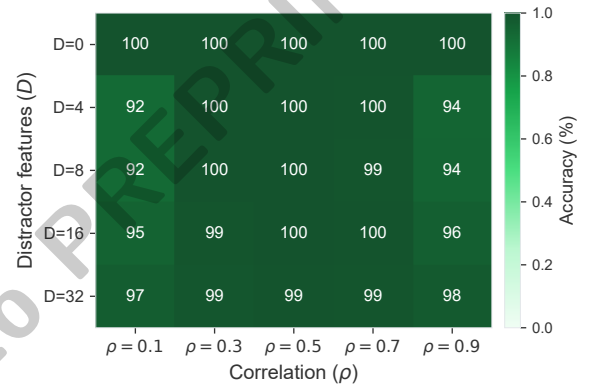


Figure 4: Accuracy heatmap across distractor count (D) and correlation strength (ρ). Accuracy remains above 92% even with many highly-correlated spurious variables.

entropy, clause monopoly with slot variance jumping from 0.003 to 0.35). Rule complexity (7–8 clauses, 3 literals) is similar across ablations, so these losses primarily affect *which* literals are selected. The full model achieves the highest accuracy (74.8%) while maintaining interpretable rule sizes.

6 Conclusion

We presented Neural Rule Inducer (NRI), a foundation model for zero-shot rule induction trained entirely on synthetic Boolean formulas. By encoding literals through identity-free statistical properties (class-conditional rates, entropy, co-occurrence) and synthesizing clauses with a parallel slot-based decoder under product T-norm semantics, NRI is trained end-to-end on prediction accuracy alone and induces interpretable DNF rules on new domains without retraining. We believe this work opens up the possibility of foundation models for symbolic reasoning. Future work extends to multi-valued and continuous variables and to first-order logic with relational predicates.

Acknowledgements

This work was supported by JSPS KAKENHI Grant Number 25K21269. This study was carried out using the TSUB-AME4.0 supercomputer at Institute of Science Tokyo. The author also thanks Professor Katsumi Inoue for a helpful discussion on this research.

References

- [Asuncion *et al.*, 2007] Arthur Asuncion, David Newman, et al. Uci machine learning repository, 2007.
- [Bommasani, 2021] Rishi Bommasani. On the opportunities and risks of foundation models. *arXiv preprint arXiv:2108.07258*, 2021.
- [Cortes and Vapnik, 1995] Corinna Cortes and Vladimir Vapnik. Support-vector networks. *Machine learning*, 20(3):273–297, 1995.
- [Cropper and Morel, 2021] Andrew Cropper and Rolf Morel. Learning programs by learning from failures. *Machine Learning*, 110(4):801–856, 2021.
- [Evans and Grefenstette, 2018] Richard Evans and Edward Grefenstette. Learning explanatory rules from noisy data. *Journal of Artificial Intelligence Research*, 61:1–64, 2018.
- [Fedus *et al.*, 2022] William Fedus, Barret Zoph, and Noam Shazeer. Switch transformers: Scaling to trillion parameter models with simple and efficient sparsity. *Journal of Machine Learning Research*, 23(120):1–39, 2022.
- [Gao *et al.*, 2024] Kun Gao, Katsumi Inoue, Yongzhi Cao, and Hanpin Wang. A differentiable first-order rule learner for inductive logic programming. *Artificial Intelligence*, 331:104108, 2024.
- [Hollmann *et al.*, 2023] Noah Hollmann, Samuel Müller, Katharina Eggenberger, and Frank Hutter. TabPFN: A transformer that solves small tabular classification problems in a second. In *International Conference on Learning Representations 2023*, 2023.
- [Inoue *et al.*, 2014] Katsumi Inoue, Tony Ribeiro, and Chikaki Sakama. Learning from interpretation transition. *Machine Learning*, 94(1):51–79, 2014.
- [Johnson *et al.*, 2025] Blair Johnson, Clayton Kerce, and Faramarz Fekri. Glidr: Graph-like inductive logic programming with differentiable reasoning. *arXiv preprint arXiv:2508.06716*, 2025.
- [Klement *et al.*, 2013] Erich Peter Klement, Radko Mesiar, and Endre Pap. *Triangular norms*, volume 8. Springer Science & Business Media, 2013.
- [Manhaeve *et al.*, 2018] Robin Manhaeve, Sebastijan Dumancic, Angelika Kimmig, Thomas Demeester, and Luc De Raedt. Deepproblog: Neural probabilistic logic programming. *Advances in Neural Information Processing Systems*, 31, 2018.
- [Muggleton and De Raedt, 1994] Stephen Muggleton and Luc De Raedt. Inductive logic programming: Theory and methods. *The Journal of Logic Programming*, 19:629–679, 1994.
- [Muggleton, 1995] Stephen Muggleton. Inverse entailment and progol. *New generation computing*, 13(3):245–286, 1995.
- [Olausson *et al.*, 2023] Theo Olausson, Alex Gu, Ben Lipkin, Cedegao Zhang, Armando Solar-Lezama, Joshua Tenenbaum, and Roger Levy. Linc: A neurosymbolic approach for logical reasoning by combining language models with first-order logic provers. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 5153–5176, 2023.
- [Peng *et al.*, 2025] Yifei Peng, Yaoli Liu, Enbo Xia, Yu Jin, Wang-Zhou Dai, Zhong Ren, Yao-Xiang Ding, and Kun Zhou. Abductive logical rule induction by bridging inductive logic programming and multimodal large language models. *arXiv preprint arXiv:2509.21874*, 2025.
- [Perez *et al.*, 2018] Ethan Perez, Florian Strub, Harm de Vries, Vincent Dumoulin, and Aaron Courville. Film: Visual reasoning with a general conditioning layer. In *AAAI Conference on Artificial Intelligence*, pages 3942–3951, 2018.
- [Phua and Inoue, 2024] Yin Jun Phua and Katsumi Inoue. Variable assignment invariant neural networks for learning logic programs. In *International Conference on Neural-Symbolic Learning and Reasoning (NeSy)*. Springer, 2024.
- [Quinlan, 1990] J. Ross Quinlan. Learning logical definitions from relations. volume 5, pages 239–266. Springer, 1990.
- [Ren *et al.*, 2025] ZZ Ren, Zhihong Shao, Junxiao Song, Huajian Xin, Haocheng Wang, Wanjia Zhao, Liye Zhang, Zhe Fu, Qihao Zhu, Dejian Yang, et al. Deepseek-prover-v2: Advancing formal mathematical reasoning via reinforcement learning for subgoal decomposition. *arXiv preprint arXiv:2504.21801*, 2025.
- [Rocktäschel and Riedel, 2017] Tim Rocktäschel and Sebastian Riedel. End-to-end differentiable proving. volume 30, 2017.
- [Sadeghian *et al.*, 2019] Ali Sadeghian, Mohammadreza Armandpour, Patrick Ding, and Daisy Zhe Wang. Drum: End-to-end differentiable rule mining on knowledge graphs. *Advances in neural information processing systems*, 32, 2019.
- [Serafini and Garcez, 2016] Luciano Serafini and Artur d’Avila Garcez. Logic tensor networks: Deep learning and logical reasoning from data and knowledge. 2016.
- [Srinivasan, 2001] Ashwin Srinivasan. The aleph manual. 2001.
- [Vaswani *et al.*, 2017] Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. Attention is all you need. volume 30, 2017.
- [Wang and Gupta, 2022] Huaduo Wang and Gopal Gupta. FOLD-R++: A scalable toolset for automated inductive learning of default theories from mixed data. In *Functional and Logic Programming - 16th International Symposium*,

FLOPS 2022, Kyoto, Japan, May 10-12, 2022, Proceedings, Lecture Notes in Computer Science, pages 224–242. Springer, 2022.

[Yang *et al.*, 2017] Fan Yang, Zhilin Yang, and William W Cohen. Differentiable learning of logical rules for knowledge base reasoning. *Advances in neural information processing systems*, 30, 2017.

[Yang *et al.*, 2020] Zhun Yang, Adam Ishay, and Joohyung Lee. Neurasp: Embracing neural networks into answer set programming. In *29th International Joint Conference on Artificial Intelligence (IJCAI 2020)*, 2020.

IJCAI-ECAI 2026 PREPRINT