
Tractable Probabilistic Neurosymbolic AI Using Boolean Tensor Factorizations

Rik Adriaensen¹

Jaron Maene¹

Luc De Raedt^{1,2}

¹KU Leuven, Belgium

²Örebro University, Sweden

Abstract

Probabilistic logic programming offers a principled way of combining deep learning and formal reasoning, but its inference is $\#P$ -hard. A simple factorization assumption introduced by TractOR makes probabilistic reasoning tractable for large relations at the cost of limiting expressivity. We reinterpret this approach using Boolean tensor factorizations to identify and resolve three practical limitations: it is restricted to symmetric relations, expresses only a limited class of distributions, and cannot reason across components.

1 INTRODUCTION

Probabilistic logic programming offers a principled paradigm for neurosymbolic AI (NeSy) [Manhaeve et al., 2021, Yang et al., 2020, Li et al., 2023], by combining deep learning and formal reasoning. Neural networks predict probabilities and probabilistic logic programs reason with them. Naturally, parameterizing the probabilities in a probabilistic logic program by neural networks results in a NeSy framework with clear semantics.

However, NeSy systems based on probabilistic logic programming have two problems. First, exact inference reduces to weighted model counting, which is $\#P$ -hard [Valiant, 1979]. Second, they model probabilistic relations using independent random variables, so all dependencies must be fully described in the logic [van Krieken et al., 2024].

TractOR [Friedman and Van den Broeck, 2020] shows that representing probabilistic relations with a simple factorization rather than using independent random variables makes probabilistic reasoning tractable and allows learning dependencies. Applying this factorization to probabilistic NeSy resolves both problems but introduces limitations of its own, as we will show.

This paper reinterprets TractOR’s factorization using Boolean tensor factorizations, identifies three practical limitations of it: (1) symmetric relations, (2) limited expressivity and (3) no interaction between components, and proposes a solution to each limitation: (1) untying the embeddings, (2) using a noisy-or instead of a mixture and (3) components with higher ranks.

2 PRELIMINARIES

Probabilistic Datalog. We assume basic familiarity with Datalog and otherwise refer to Appendix A and Abiteboul et al. [1995]. Probabilistic Datalog [Fuhr, 1995] extends Datalog with a distribution over interpretations via probabilistic facts. A *probabilistic fact* has the form $p_a :: a$, where a is a ground atom and $p_a \in [0, 1]$ its associated probability. A *probabilistic Datalog program* is a pair $(\mathcal{PF}, \mathcal{R})$ where $\mathcal{PF}$ is a set of probabilistic facts and $\mathcal{R}$ a set of rules.

Each probabilistic fact corresponds to an independent Bernoulli variable. This induces a distribution over subsets of facts w called (*possible*) *worlds*, defined as follows.

$$\Pr(w) = \prod_{a \in w} p_a \prod_{a \in \mathcal{F} \setminus w} (1 - p_a) \quad (1)$$

The *success probability* $\Pr(q)$ of a *query* atom q is defined under the *distribution semantics* [Sato, 1995] as the probability that it is entailed by the program $w \cup \mathcal{R}$, where w is a randomly sampled world. The Iverson brackets $\llbracket \cdot \rrbracket$ used below denote 1 if the condition holds and 0 otherwise.

$$\Pr(q) = \sum_{w \subseteq \mathcal{F}} \llbracket w \cup \mathcal{R} \models q \rrbracket \Pr(w) \quad (2)$$

The typical inference approach [Fierens et al., 2015] is to find a query formula over the probabilistic facts whose models are exactly the worlds w in which the atom q is entailed by $w \cup \mathcal{R}$. The success probability then equals the weighted

model count (WMC) of this formula. When the Datalog program does not contain negation or recursion, this formula can be written as a union of conjunctive queries (UCQ) [Abiteboul et al., 1995].

The complexity of computing the WMC of a formula as a function of the number of probabilistic facts is called *data complexity*, in contrast to *query complexity*, where the facts are fixed and query size varies. Dalvi and Suciu [2013] show that every UCQ formula is either safe (polynomial data complexity) or unsafe ($\#P$ -hard). They provide a lifted inference algorithm that computes the WMC in polynomial time for all safe formulas and correctly identifies (and fails on) unsafe formulas.

TractOR. TractOR [Friedman and Van den Broeck, 2020], a probabilistic relational embedding model, answers arbitrary queries with polynomial data complexity by making a simple factorization assumption. It introduces a new probabilistic fact $f_1(t)$ for every predicate t and $f_2(x)$ for every constant x . Then it assumes that every binary predicate t factorizes as

$$t(X, Y) := f_1(t), f_2(X), f_2(Y). \quad (3)$$

Tractability is then assured by the following theorem, as the UCQ representation of any query will now contain at most one variable per atom.

Theorem 1 [Friedman and Van den Broeck, 2020, Thm. 4] *Suppose that φ is a UCQ with at most one variable per atom. Then φ is safe.*

To improve expressivity, they work with a mixture of k such simple models using uniform mixing weights. The parameters of $f_1(t)$ and $f_2(x)$ across the mixture components act as embeddings.

3 BOOLEAN TENSOR FACTORIZATIONS IN DATALOG

TractOR’s factorization is closely related to Boolean tensor factorizations. Hence, we introduce the relevant concepts [Miettinen, 2011] and show how to express them in Datalog.

A *Boolean tensor* is a multidimensional array $\mathbf{T} \in \mathbb{B}^{I_1 \times \dots \times I_m}$ with all operations defined over the Boolean algebra. The *order* of the tensor equals the number of dimensions m , also called *ways* or *modes*. An m -ary relation t over a domain $\mathcal{D}$ can be represented as a Boolean tensor $\mathbf{T}$ of order m , assuming a fixed ordering of the domain.

We use lowercase symbols (e.g., t) for logical predicates and the corresponding uppercase bold symbols (e.g., $\mathbf{T}$) for their tensor representation, such that $\forall \bar{c} \in \mathcal{D}^m : t(\bar{c}) \iff \mathbf{T}(\bar{c}) = 1$. As all Boolean tensors in this paper represent

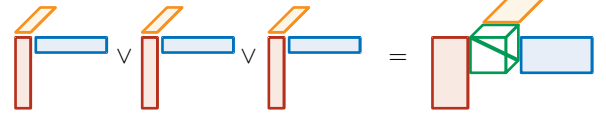


Figure 1: The Boolean CP factorization visualized.

relations, we index tensors directly with domain elements, assuming an implicit mapping to tensor indices.

A *rank-1* Boolean tensor is an outer product of Boolean vectors. This rule expresses that $\mathbf{T}$ is a rank-1 tensor.

$$t(X, Y, Z) :- f_1(X), f_2(Y), f_3(Z) \quad (4)$$

The Boolean *canonical polyadic* (CP) factorization decomposes a tensor into a disjunction of r rank-1 component tensors, where r is the (Boolean) rank of the factorization. The Datalog rules $t(X, Y, Z) :- f_1(X, 1), f_2(Y, 1), f_3(Z, 1)$ and $t(X, Y, Z) :- f_1(X, 2), f_2(Y, 2), f_3(Z, 2)$ express that $\mathbf{T}$ is a rank-2 tensor¹. The r vectors for each mode n can be collected as the columns of a *factor matrix* $\mathbf{F}_n$ corresponding to the predicate f_n in the Datalog rules. The CP factorization can then be written as $\mathbf{I}_r \times_1 \mathbf{F}_1 \times_2 \mathbf{F}_2 \times_3 \mathbf{F}_3$, with $\mathbf{I}_r$ the size r identity tensor and $\times_n$ the tensor product along the n^{th} mode (see Figure 1). The numbers in the second argument position of f_n , indicating the columns, form the *latent domain* $\mathcal{D}_{\text{lat}} = \{1, \dots, r\}$ where r is the rank of the factorization. The two rules from before can now be written as a single rule where the variable l ranges over the latent domain containing two elements.

$$t(X, Y, Z) :- f_1(X, l), f_2(Y, l), f_3(Z, l) \quad (5)$$

Example 1 Consider the third-order Boolean tensor over the domain $\{a, b, c\}$ representing the relation

$$t = \{(a, a, a), (c, a, a), (a, c, a), (c, c, a), (a, a, b), (b, a, b), (a, b, b), (b, b, b), (a, a, c), (c, a, c), (a, c, c), (c, c, c)\}.$$

Then there exists a rank-2 Boolean CP factorization of this relation shown below as a program and using tensors.

$$\begin{aligned} & f_1(a, 1).f_1(a, 2).f_1(b, 2).f_1(c, 1).f_2(a, 1).f_2(a, 2). \\ & f_2(b, 2).f_2(c, 1).f_3(a, 1).f_3(b, 2).f_3(c, 1). \\ & t(X, Y, Z) :- f_1(X, l), f_2(Y, l), f_3(Z, l). \end{aligned}$$

$$\begin{aligned} \mathbf{T} &= \begin{matrix} & \begin{matrix} a & b & c \end{matrix} \\ \begin{matrix} a \\ b \\ c \end{matrix} & \begin{bmatrix} 1 & 0 & 1 \\ 0 & 0 & 0 \\ 1 & 0 & 1 \end{bmatrix} \end{matrix} = \begin{matrix} & \begin{matrix} a & b & c \end{matrix} \\ \begin{matrix} a \\ b \\ c \end{matrix} & \begin{bmatrix} 1 & 1 & 0 \\ 1 & 1 & 0 \\ 0 & 0 & 0 \end{bmatrix} \end{matrix} = \begin{matrix} & \begin{matrix} a & b & c \end{matrix} \\ \begin{matrix} a \\ b \\ c \end{matrix} & \begin{bmatrix} 1 & 0 & 1 \\ 0 & 0 & 0 \\ 1 & 0 & 1 \end{bmatrix} \end{matrix} \\ &= \mathbf{I}_2 \times_1 \begin{bmatrix} 1 & 2 \\ 1 & 1 \\ 0 & 1 \\ 1 & 0 \end{bmatrix} \times_2 \begin{bmatrix} 1 & 2 \\ 1 & 1 \\ 0 & 1 \\ 1 & 0 \end{bmatrix} \times_3 \begin{bmatrix} 1 & 2 \\ 1 & 0 \\ 0 & 1 \\ 1 & 0 \end{bmatrix} \begin{matrix} a \\ b \\ c \end{matrix} \end{aligned}$$

¹W.l.o.g., we factorize a single relation, as the tensors of multiple relations of the same arity can always be stacked along a new mode creating a single higher-order relation (see Appendix B.1).

Each rank-1 component in the factorization covers a subset of the relation called a tile. Therefore, the CP factorization corresponds to a tiling of the reconstructed relation. See Appendix B for further discussion on Boolean tensor factorizations and their connection to tilings.

When the same factor matrix is used for different modes in the same factorization, we call these modes or embeddings *tied*. For example, in TractOR’s factorization (Equation 3) the second and third mode are tied.

4 TRACTABLE PROBABILISTIC DATALOG

We now show how modeling probabilistic relations with tensor factorizations rather than fully materializing them leads to tractable inference.

We say a program is *factorized* if all probabilistic facts are factor predicates of the same CP factorization. Hence, a factorized program contains one rule defining the factorization (Equation 5) and possibly others over the relations reconstructed by the factorization. We assume every possible probabilistic fact over the factor predicates is present, the number of which is polynomial in the domain sizes (i.e., $\mathcal{O}(m|\mathcal{D}||\mathcal{D}_{\text{lat}}|)$). Since we use the fragment of Datalog equivalent to UCQ, any formula derived from a factorized program can be written as a UCQ containing only factor predicates. Grounding all variables ranging over the latent domain results in a formula where every atom contains at most one variable. As the increase in formula size depends only on the size of the latent domain, the theorem below follows from Theorem 1.

Theorem 2 *For a fixed latent domain $\mathcal{D}_{\text{lat}}$, answering an arbitrary query over a factorized negation-free non-recursive probabilistic Datalog program is polynomial in the size of the original domain $\mathcal{D}$.*

However, answering an arbitrary query remains intractable in the latent domain size. In practice, this intractability is prohibitive, even for low-rank factorizations where this domain is small. TractOR avoids this problem by using a mixture of rank-1 factorized programs. That is, it takes k factorized programs with a rank-1 factorization rule (in which case Equation 5 reduces to Equation 4), and the probability of a query q is the average over the answers of these programs. That is, $\text{Pr}_{\text{TractOR}}(q) = \sum_{j=1}^k \frac{1}{k} \text{Pr}_j(q)$ with $\text{Pr}_j(q)$ the predicted success probability of the j^{th} rank-1 program.

Learning factorizations. Factorized programs are a tractable subset of probabilistic Datalog, so we can learn their parameters with any standard approach. However, to achieve the claimed tractability, we rewrite the query as a UCQ and use the lifted inference algorithm from Dalvi and

Suciu [2013] to compile a polynomial-sized circuit computing its WMC. The circuit is end-to-end differentiable, so we can learn the parameters using gradient descent. This is a gradient-based approach to learning the factorization.

On a task where the domain consists of subsymbolic inputs, e.g., a citation graph where every entity is a text document, the parameters of a factorized program can be predicted by neural networks. Current neurosymbolic frameworks usually predict relations over, e.g., documents directly, whereas here they would predict the latent unary properties of each document (each column of the factor matrices represents one such property), from which these relations can be reconstructed. We leave such an extension for future work.

4.1 LIMITATIONS AND SOLUTIONS

We identify three limitations of TractOR’s factorization, propose a solution for each and use experiments to illustrate them. We experiment on two tasks: a toy example containing two simple relations t_1 and t_2 (see Figure 3, ground truth) and the countries knowledge graph [Bouchard et al., 2015], which contains the *locatedIn* and *neighborOf* relations over all countries, 23 subregions, and 5 continents of the world. On both tasks, we first train a mixture or noisy-or over rank-1 programs with only the factorization rule, learning an approximation of the given relations. Then, we add various rules to the rank-1 programs to derive the symmetric closure, intersection, union, and composition of the learned relations. We report binary cross-entropy (BCE) to measure the quality of the computed success probabilities and mean average precision (mAP) to measure the correct ranking of true and false facts. For further experimental details, including the rules used, we refer to Appendix C.

1. Symmetric relations only. Tying modes restricts tiles to be symmetric in these modes. For example, TractOR’s factorization (Equation 3) can only represent symmetric relations because the final two modes are tied, as also mentioned in the paper [Friedman and Van den Broeck, 2020].

Untying the modes (i.e., a separate f_n for each argument position) allows representing non-symmetric relations, although this increases the number of parameters from $\mathcal{O}(k|\mathcal{D}|)$ to $\mathcal{O}(mk|\mathcal{D}|)$. Figure 2 (top) shows on the toy task that this has a significant impact even when factorizing very simple relations.

2. Limited expressivity. Each rank-1 program defines a distribution over possible worlds in which the reconstructed relation t is a single tile, so the success probability of $t(a, b)$ equals the marginal probability that (a, b) lies in that tile. A mixture of the predicted success probability of these rank-1 programs can only assign probability one to two facts $t(a_1, b_1)$ and $t(a_2, b_2)$ if all rank-1 programs assign only probability mass to tiles covering both of these facts. As

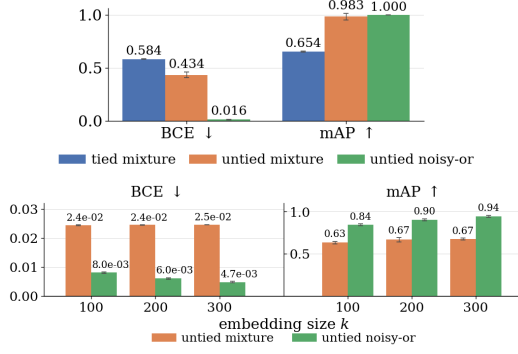


Figure 2: The results for training the rank-1 programs on the toy task (top) and countries task (bottom).

any tile covering both $t(a_1, b_1)$ and $t(a_2, b_2)$ (with $a_1 \neq a_2$ and $b_1 \neq b_2$) also covers $t(a_1, b_2)$ and $t(a_2, b_1)$ it is then impossible to *not* assign probability one to the latter two facts at the same time. These mixtures can therefore not represent arbitrary success probabilities even with an arbitrary number of components. When learning the factorization’s parameters using binary cross-entropy loss, as is common for probabilistic logic programs, this inexpressivity moreover leads to local minima that separate positive from negative queries but fail to learn the correct probabilities of these queries.

Combining the success probabilities of the rank-1 programs using a noisy-or, i.e., $\Pr(q) = 1 - \prod_{j=1}^k (1 - \Pr_j(q))$, instead of a mixture, resolves the expressivity problem. Whereas the mixture treats the programs as mutually exclusive, selecting one program’s prediction, the noisy-or treats them as independent, so a fact is derived whenever at least one program predicts it. This does allow representing arbitrary success probabilities given an arbitrary number of components, as we could define a separate tile for every true reconstructed ground fact. Figure 2 shows that using the noisy-or is better at both learning the correct marginals and correctly ranking true and false facts. Figure 3 (top) shows on the toy task that the noisy-or better approximates the ground-truth queries than the mixture.

3. No interaction between components. Finally, when adding rules to rank-1 programs, they are not applied across components. Since every program reconstructs a single tile, rule bodies only combine facts from the same tile. Whether this limits what is derivable depends on the rule. For example, applying symmetry to each tile makes the whole relation symmetric, but composing two facts from different tiles is not possible.

The third limitation would be fully resolved by a single program of sufficiently high rank to approximate the target relation, at the cost of intractability in the rank (Theorem 2). A noisy-or over low-rank components would instead trade tractability against which facts are derivable, which we leave

Table 1: The results for applying rules to learned relations for the toy task using the untied noisy-or program. Adding the rules to the factorized program is compared to a baseline which first reconstructs the learned factorization and then applies the rules.

rule	baseline		factorized	
	BCE ↓	mAP ↑	BCE ↓	mAP ↑
symmetry	0.025	1.000	0.019	1.000
intersection	0.012	1.000	1.832	1.000
union	0.029	1.000	0.581	1.000
composition	0.017	1.000	2.416	0.851

to future work. Table 1 compares our approach, which adds the rules to each rank-1 program, against a baseline that first reconstructs the relation with the noisy-or and then applies the rules. The baseline is infeasible for the larger countries knowledge graph, whereas our approach does scale to this larger task (see Appendix C.1). The factorized program perfectly applies the symmetry rule but does not consistently predict the success probabilities for the other rules, as this might require reasoning across tiles. For example, Figure 3 (bottom) shows that the composition rule is correctly applied within a single tile (composition₂, composing t_2 facts) but not across tiles (composition₁, composing t_1 facts).

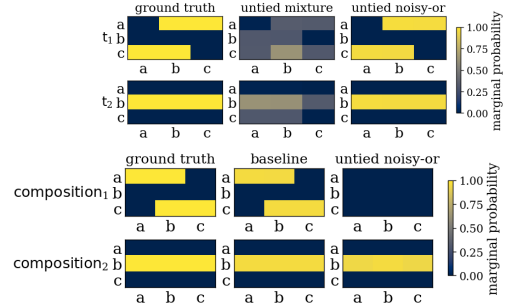


Figure 3: The success probabilities predicted by the noisy-or and mixture over rank-1 programs for the learned predicates t_1 and t_2 (top) and derived composition facts (bottom).

5 CONCLUSION

The mixture of rank-1 factorizations proposed by TractOR [Friedman and Van den Broeck, 2020] can be applied to probabilistic Datalog for tractable probabilistic neurosymbolic reasoning. However, three limitations restrict the practical use of this subset of probabilistic Datalog: it can only represent symmetric relations, expresses a limited class of distributions, and does not allow reasoning across mixture components. We proposed a solution to each: untying the embeddings, replacing the mixture with a noisy-or, and using components with a higher rank.

Acknowledgements

This research received funding from the Flemish Government under the “Flanders Artificial Intelligence Research program”. This research is co-funded by the European Union (ERC, DeepLog, 101142702). Views and opinions expressed are, however, those of the author(s) only and do not necessarily reflect those of the European Union or the European Research Council. Neither the European Union nor the granting authority can be held responsible for them. Luc De Raedt is also supported by the Wallenberg AI, Autonomous Systems and Software Program (WASP) funded by the Knut and Alice Wallenberg Foundation.

References

- Serge Abiteboul, Richard Hull, and Victor Vianu. *Foundations of Databases*. Addison-Wesley, 1995. ISBN 0-201-53771-0.
- Guillaume Bouchard, Sameer Singh, and Théo Trouillon. On approximate reasoning capabilities of low-rank vector spaces. In *Proceedings of the 2015 AAAI Spring Symposium on Knowledge Representation and Reasoning (KRR): Integrating Symbolic and Neural Approaches*, 2015.
- Nilesh Dalvi and Dan Suciu. The dichotomy of probabilistic inference for unions of conjunctive queries. *J. ACM*, 59(6), January 2013. ISSN 0004-5411. doi: 10.1145/2395116.2395119.
- Daan Fierens, Guy Van Den Broeck, Joris Renkens, Dimitar Shterionov, Bernd Gutmann, Ingo Thon, Gerda Janssens, and Luc De Raedt. Inference and learning in probabilistic logic programs using weighted Boolean formulas. *Theory and Practice of Logic Programming*, 15(3):358–401, May 2015. Publisher: Cambridge University Press.
- Tal Friedman and Guy Van den Broeck. Symbolic querying of vector spaces: Probabilistic databases meets relational embeddings. In Jonas Peters and David Sontag, editors, *Proceedings of the 36th Conference on Uncertainty in Artificial Intelligence (UAI)*, volume 124 of *Proceedings of Machine Learning Research*, pages 1268–1277. PMLR, 03–06 Aug 2020.
- Norbert Fuhr. Probabilistic datalog—a logic for powerful retrieval methods. In *Proceedings of the 18th Annual International ACM SIGIR Conference on Research and Development in Information Retrieval*, SIGIR ’95, page 282–290, New York, NY, USA, 1995. Association for Computing Machinery. ISBN 0897917146. doi: 10.1145/215206.215372.
- Floris Geerts, Bart Goethals, and Taneli Mielikäinen. Tiling databases. In Einoshin Suzuki and Setsuo Arikawa, editors, *Discovery Science*, pages 278–289, Berlin, Heidelberg, 2004. Springer Berlin Heidelberg. ISBN 978-3-540-30214-8.
- Ziyang Li, Jiani Huang, and Mayur Naik. Scallop: A language for neurosymbolic programming. *Proc. ACM Program. Lang.*, 7(PLDI), June 2023. doi: 10.1145/3591280.
- Jaron Maene, Vincent Derkinderen, and Pedro Zuidberg Dos Martires. KLayer: Accelerating arithmetic circuits for neurosymbolic AI. In *The Thirteenth International Conference on Learning Representations*, 2025.
- Robin Manhaeve, Sebastijan Dumančić, Angelika Kimmig, Thomas Demeester, and Luc De Raedt. Neural probabilistic logic programming in DeepProbLog. *Artificial Intelligence*, 298:103504, 2021. ISSN 0004-3702. doi: 10.1016/j.artint.2021.103504.
- Pauli Miettinen. Boolean tensor factorizations. In *2011 IEEE 11th International Conference on Data Mining*, pages 447–456, 2011. doi: 10.1109/ICDM.2011.28.
- Taisuke Sato. A statistical learning method for logic programs with distribution semantics. In Leon Sterling, editor, *Logic Programming, Proceedings of the Twelfth International Conference on Logic Programming, Tokyo, Japan, June 13-16, 1995*, pages 715–729. MIT Press, 1995.
- Leslie G. Valiant. The complexity of enumeration and reliability problems. *SIAM Journal on Computing*, 8(3): 410–421, 1979. doi: 10.1137/0208032. URL <https://doi.org/10.1137/0208032>.
- M. H. Van Emden and R. A. Kowalski. The semantics of predicate logic as a programming language. *J. ACM*, 23(4):733–742, October 1976. ISSN 0004-5411. doi: 10.1145/321978.321991.
- Emile van Krieken, Pasquale Minervini, Edoardo M. Ponti, and Antonio Vergari. On the independence assumption in neurosymbolic learning. In *Proceedings of the 41st International Conference on Machine Learning, ICML’24*. JMLR.org, 2024.
- Zhun Yang, Adam Ishay, and Joohyung Lee. NeurASP: Embracing neural networks into answer set programming. In Christian Bessiere, editor, *Proceedings of the Twenty-Ninth International Joint Conference on Artificial Intelligence, IJCAI-20*, pages 1755–1762. International Joint Conferences on Artificial Intelligence Organization, 7 2020. doi: 10.24963/ijcai.2020/243. Main track.

Tractable Probabilistic Neurosymbolic AI Using Boolean Tensor Factorizations (Supplementary Material)

Rik Adriaensen¹

Jaron Maene¹

Luc De Raedt^{1,2}

¹KU Leuven, Belgium

²Örebro University, Sweden

A DATALOG

A *term* is either a constant c or a variable X . An *atom* has the form $a(\bar{c})$, where a is a predicate of arity m , and $\bar{c} = c_1, \dots, c_m$ are its arguments. A *rule* is an expression of the form $h :- b_1, \dots, b_n$, where h is an atom called the *head*, and $b_1, \dots, b_n$ are atoms forming the *body*. The rule states that the head is true whenever all atoms in the body are true. A *fact* is a ground rule with an empty body representing an atom that is unconditionally true. An expression (atom, rule, query) is *ground* if it does not contain variables. A non-ground expression can be *grounded* by substituting all its variables by constants.

A *Datalog program* $\mathcal{P}$ is a pair $(\mathcal{F}, \mathcal{R})$, where $\mathcal{F}$ is a set of facts and $\mathcal{R}$ is a set of rules with non-empty bodies. The set of all constants occurring in a program is called the *domain* $\mathcal{D}$, also known as the Herbrand universe. The set of all predicates occurring in the program is denoted by $\text{Pred}(\mathcal{P})$. Throughout the paper, we will assume that the predicates occurring in $\mathcal{F}$ do not occur in the heads of the rules $\mathcal{R}$ and there is no recursion in the rule set. Rules are sometimes called conjunctive queries (CQ) and multiple rules for the same head predicate form a union of conjunctive queries (UCQ), which represents a disjunction of conjunctions.

The *Herbrand base* $\mathcal{B}$ is the set of all possible ground atoms that can be formed using predicates in $\text{Pred}(\mathcal{P})$ and constants in $\mathcal{D}$. A (*Herbrand*) *interpretation* is a set of ground atoms that are considered true; all other ground atoms in $\mathcal{B}$ are considered false. The interpretation of a predicate is called a relation, but they are often used interchangeably. A ground rule is *satisfied* by an interpretation if either the body is false or the head is true. A non-ground rule is satisfied if all its ground instances are satisfied. An interpretation is a (*Herbrand*) *model* of a program if it satisfies all rules in the program. The *least (Herbrand) model* $\mathcal{M}_{\mathcal{P}}$ is the unique model that is a subset of all other models [Van Emden and Kowalski, 1976] and it determines the intended semantics of the program. A ground query q is *entailed* by a program $\mathcal{P}$, denoted $\mathcal{P} \models q$, if q is true in the least model of $\mathcal{P}$.

Example 2 Consider the following program $\mathcal{P}$ with domain $\mathcal{D} = \{a, b, c\}$.

likes(a, a). likes(a, b). likes(b, a).
likes(b, b). likes(b, c). likes(c, b). likes(c, c).
path₃(X, Y) :- likes(X, X₁), likes(X₁, X₂), likes(X₂, Y).

The Herbrand base $\mathcal{B}$ is $\{\text{likes}(a, a), \text{likes}(a, b), \dots, \text{path}_3(c, b), \text{path}_3(c, c)\}$. The least model of $\mathcal{P}$ is $\mathcal{M}_{\mathcal{P}} = \mathcal{B} \setminus \{\text{likes}(a, c), \text{likes}(c, a)\}$.

B BOOLEAN TENSOR FACTORIZATIONS

B.1 FACTORIZING MULTIPLE RELATIONS

To factorize multiple relations of the same arity, their Boolean tensors can be stacked along a new mode, and the resulting higher-order tensor can be factorized in the usual way. This corresponds to reification in Datalog, where every atom

$t(c_1, \dots, c_m)$ containing one of the chosen predicates is replaced by $\text{stack}(t, c_1, \dots, c_m)$, with stack a fresh predicate. This preserves the semantics of the program. We allow predicates to appear as arguments in the paper, implicitly assuming reification. For example, $t(X, Y) :- f_1(t), f_2(X), f_3(Y)$ is syntactic sugar for $\text{stack}(t, X, Y) :- f_1(t), f_2(X), f_3(Y)$.

B.2 CLUSTERINGS AND TILINGS

It is common to interpret the Boolean factor matrices F_n as clusterings, where each column defines a different cluster [Miettinen, 2011]. Each column corresponds to a latent constant in $\mathcal{D}_{\text{lat}}$, so latent constants can be thought of as clusters containing all domain elements with a 1 in the corresponding column. For example, the fact $f_n(c, 1)$ can be interpreted as c being in cluster 1 under clustering n .

The corresponding interpretation of the CP decomposition is that of a tiling [Geerts et al., 2004, Miettinen, 2011]. Each rank-1 component tensor corresponds to an all-ones subtensor called a tile. Together, the r tiles must cover all ones of the reconstructed tensor and none of the zeros. A tile has the same order as the reconstructed tensor but contains only a subset of its indices along each mode. These subsets are defined by the corresponding clusters in the factor matrices. Formally, a tiling is defined as shown below.

Definition 1 (Tiling) *Given an m -ary relation $t \subseteq \mathcal{D}^m$, a tile $\tau(d_1, \dots, d_m)$ is the Cartesian product $d_1 \times \dots \times d_m$ of m subsets of the domain, that is, $d_n \subseteq \mathcal{D}$ for each $n \in \{1, \dots, m\}$. A tile covers t if all tuples in the tile are part of the relation, that is, $\forall (c_1, \dots, c_m) \in \tau(d_1, \dots, d_m) : t(c_1, \dots, c_m)$. A set of tiles forms a tiling of relation t if the union of all the tiles is equal to the relation t .*

C EXPERIMENTAL DETAILS

Both tasks consist of two steps: learn only the factorization and apply the rules. For the first step, we train the parameters of a factorized program containing only the factorization rule (e.g., $t(X, Y) :- f_1(t, l), f_2(X, l), f_3(Y, l)$ for each probabilistic relation t). Afterwards, we add a copy of the rules shown in Equation 6 (substituting t) for every probabilistic relation learned in the experiment.

$$\begin{aligned}
&\text{symmetric}(X, Y) :- t(X, Y). \\
&\text{symmetric}(X, Y) :- t(Y, X). \\
&\text{composition}(X, Z) :- t(X, Y), t(Y, Z). \\
&\text{union}(X, Y) :- t(X, Z). \\
&\text{union}(X, Y) :- t(Y, Z). \\
&\text{intersection}(X, Y) :- t(X, Z), t(Y, Z).
\end{aligned} \tag{6}$$

As a baseline for the rule applications, we first reconstruct the learned relations as the predicted success probabilities of the mixture/noisy-or over rank-1 programs and add them as explicit probabilistic relations to a probabilistic program. When learning the factorizations, we use the hyperparameters reported in Table 2. We repeat all experiments across five random seeds and report the means. For the countries experiment, we add a regularization term to the loss with regularization weight λ that ensures the L_1 norm of each column of the factor matrices is larger than one to ensure no components are unused. This is essential to make the performance improve with increasing rank.

Table 2: Hyperparameters for the toy and countries experiments, where k refers to the number of rank-1 programs.

Experiment	lr	epochs	k	loss	λ
Toy	1	1000	3	BCE	-
Countries	0.1	600	{100, 200, 300}	BCE	0.1

We train and evaluate the programs similar to Fierens et al. [2015] by compiling them once into an end-to-end differentiable circuit implemented with an extension of KLAY [Maene et al., 2025].

C.1 ADDITIONAL RESULTS

In this section, we include some additional experimental results. Table 3 shows the performance of adding different rules to the factorized programs (untied noisy-or) trained on the countries dataset for an increasing number of rank-1 programs k . Figure 4 shows the application of all four rules on the toy relations, further illustrating that these rule applications are limited to single tiles.

Table 3: The results for adding rules to the rank-1 programs using the untied noisy-or after training the factorization on the countries task.

	$k = 100$		$k = 200$		$k = 300$	
	BCE ↓	mAP ↑	BCE ↓	mAP ↑	BCE ↓	mAP ↑
symmetry	0.010	0.873	0.008	0.931	0.006	0.963
intersection	0.513	0.779	0.428	0.846	0.402	0.876
union	0.121	1.000	0.082	1.000	0.058	1.000
composition	0.074	0.398	0.087	0.400	0.094	0.439

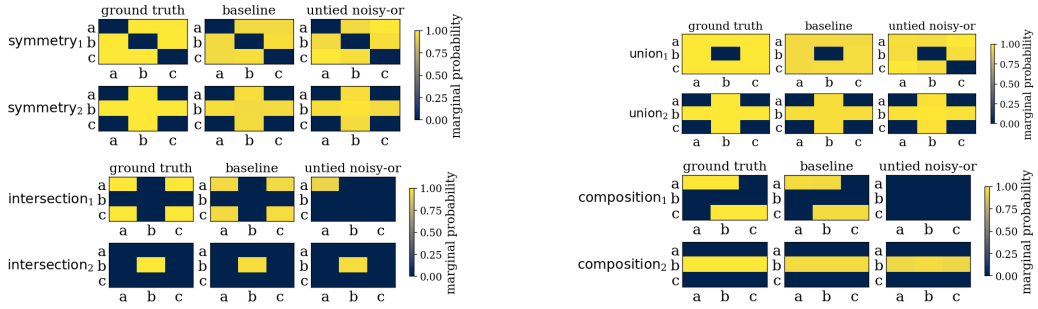


Figure 4: Success probabilities predicted by the baseline and the factorized programs for the derived symmetry, union, intersection and composition facts.