

M2-Internship: Neuro-Symbolic Framework for Explainable Inconsistency Detection in Complex Graph-Structured Systems

Duration: 6 months

Supervision: Hamida SEBA (LIRIS, Lyon 1)

Location: LIRIS, campus de la Doua Villeurbanne or IUT Lyon 1, campus de Bourg en Bresse

Application should be sent to: Hamida.seba@univ-lyon1.fr (CV+Transcripts)

Context

Modern organizational systems, comprising technical documentation, business processes, and knowledge bases, are inherently complex and evolve over time. Ensuring global consistency across these heterogeneous sources is a major challenge, as inconsistencies (e.g., a process workflow contradicting a regulatory document) can lead to significant operational risks.

While Graph Neural Networks (GNNs) are powerful for learning representations of such relational data, they operate as "black boxes" and cannot explain why a specific element is inconsistent. Conversely, symbolic reasoning provides formal logic and interpretability but lacks scalability and robustness to noise. This project aims to develop a **Neuro-Symbolic** approach that combines the perception capabilities of GNNs with the reasoning power of symbolic logic. The goal is to create a system that not only detects anomalies within complex graphs but explicitly links them to violated rules or constraints, providing an explainable bridge between data-driven detection and knowledge-driven justification.

2. Bibliography & Research Directions

The project will explore the following research pillars:

- **Relational Learning via GNNs:** Study of Graph Convolutional Networks (GCN) and Graph Attention Networks (GAT) for modeling entities and dependencies in complex systems.
- **Neuro-Symbolic Integration:** Exploration of the "Third Wave of AI," focusing on integrating logical constraints into neural loss functions or message-passing mechanisms.

- **Inconsistency vs. Anomaly Detection:** Research into the semantic difference between statistical anomalies and symbolic inconsistencies (the violation of explicit invariants).
- Tirtharaj Dash, Ashwin Srinivasan, and A. Baskar. Inclusion of domain knowledge into gnn using mode-directed inverse entailment. *Mach. Learn.*, 111(2):575–623, February 2022.
- Tirtharaj Dash, Ashwin Srinivasan, and Lovekesh Vig. Incorporating symbolic domain knowledge into graph neural networks. *Mach. Learn.*, 110(7):1609–1636, July 2021.
- Cory Davis, Patrick Stockton, Eugene B. John, Zachary Susskind, and Lizy K. John. Characterization of Neuro-Symbolic AI and Graph Convolutional Network workloads, pages 183–208. 2024.
- Vanya Bannihatti Kumar, Balaji Ganesan, Muhammed Ameen, Devbrat Sharma, and Arvind Agarwal. Automated evaluation of gnn explanations with neuro symbolic reasoning. In Douwe Kiela, Marco Ciccone, and Barbara Caputo, editors, *Proceedings of the NeurIPS 2021 Competitions and Demonstrations Track*, volume 176 of *Proceedings of Machine Learning Research*, pages 314–318. PMLR, 06–14 Dec 2022.
- Luis C. Lamb, Artur d’Avila Garcez, Marco Gori, Marcelo O.R. Prates, Pedro H.C. Avelar, and Moshe Y. Vardi. Graph neural networks meet neural-symbolic computing: a survey and perspective. In *Proceedings of the Twenty-Ninth International Joint Conference on Artificial Intelligence, IJCAI’20*, 2021.
- Ningyi Liao, Siqiang Luo, Xiaokui Xiao, and Reynold Cheng. Advances in designing scalable graph neural networks: The perspective of graph data management. In *Companion of the 2025 International Conference on Management of Data, SIGMOD/PODS ’25*, page 844–850, New York, NY, USA, 2025. Association for Computing Machinery.
- Kislay Raj. A neuro-symbolic approach to enhance interpretability of graph neural network through the integration of external knowledge. In *Proceedings of the 32nd ACM International Conference on Information and Knowledge Management, CIKM ’23*, page 5177–5180, New York, NY, USA, 2023. Association for Computing Machinery.
- Kislay Raj and Alessandra Mileo. Towards understanding graph neural networks: Functional-semantic activation mapping. In Tarek R. Besold, Artur d’Avila Garcez, Ernesto Jimenez-Ruiz, Roberto Confalonieri, Pranava Madhyastha, and Benedikt Wagner, editors, *Neural-Symbolic Learning and Reasoning*, pages 98–106, Cham, 2024. Springer Nature Switzerland.
- Ahmed Rashwan, Keith Briggs, Chris Budd, and Lisa Kreusser. Enforcing convex constraints in graph neural networks. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2026.
- Ali Sadeghian, Mohammadreza Armandpour, Patrick Ding, and Daisy Zhe Wang. DRUM: end-to-end differentiable rule mining on knowledge graphs. Curran Associates Inc., Red Hook, NY, USA, 2019.
- WenguanWang, Yi Yang, and FeiWu. Towards data-and knowledge-driven ai: A survey on neuro-symbolic computing. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 47(2):878–899, 2025.
- Zonghan Wu, Shirui Pan, Fengwen Chen, Guodong Long, Chengqi Zhang, and Philip S. Yu. A comprehensive survey on graph neural networks. *IEEE Transactions on Neural Networks and Learning Systems*, 32(1):4–24, 2021.
- Dongran Yu, Bo Yang, Dayou Liu, Hui Wang, and Shirui Pan. A survey on neural-symbolic learning systems. *Neural Networks*, 166:105–126, 2023.

3. Objectives

The main objective is to implement a pipeline that transforms unstructured system data into explainable inconsistency reports:

1. **Graph Modeling:** Design a representation layer that converts heterogeneous system components (documents, processes) into a unified graph structure where nodes are entities and edges are semantic relationships.
2. **Neuro-Symbolic Integration:** Implement a GNN architecture where symbolic domain knowledge (rules/constraints) is injected either as a regularization term in the loss function or as a constraint on node embeddings.
3. **Explainable Detection Logic:** Develop a module that maps detected "neural" anomalies back to specific violated symbolic rules, enabling the system to provide explicit justifications for each detection.
4. **Robustness Analysis:** Evaluate the framework's ability to handle noise and distribution drift using graph augmentation techniques.
5. **Validation on Use-Cases:** Test the model on real-world or synthetic datasets (e.g., business process auditing) to measure detection accuracy and explanation quality.