

Uncovering Hidden Financial Risk Chains through Temporal and Relational Knowledge Graph Learning

Yawen Duan^{1*}

¹Emory University, Atlanta, USA

*Corresponding author: Yawen Duan, yawen.duan.connect@gmail.com

Abstract: This paper addresses the challenges in financial risk identification, such as the difficulty in uniformly representing multi-source heterogeneous information, the complexity and concealment of transaction relationship chains, and the dispersion and semantic inconsistency of risk evidence across entities. It proposes a financial knowledge graph-driven risk identification framework that integrates transaction entities, fund flow relationships, event context, and temporal attributes into a unified graph structure representation to support risk inference based on interconnected links. The method first normalizes the original data into entities and constructs relationships, forming a knowledge graph containing entity, relationship, and temporal information. Then, a relationship-aware graph representation learning mechanism is employed to propagate and aggregate differentiated information in multi-relational neighborhoods, enabling node representations to incorporate multi-hop paths and local structural patterns. Simultaneously, a time-decay evidence weighting strategy is introduced to strengthen the impact of recent behavior and key interactions on risk judgment, improving the ability to characterize dynamically evolving risks dynamically. At the risk output layer, the framework maps the learned entity representations to continuous risk scores and expresses them probabilistically, thereby achieving rankable risk intensity estimation and supporting link-level evidence localization. Comparative experimental results show that the method achieves consistent improvement on multiple evaluation metrics, can more effectively capture structural risk signals in complex transaction networks, and enhances the stability and interpretability of model output while maintaining discriminative ability.

Keywords: Financial knowledge graph, relationship-aware propagation, time-series evidence weighting, risk scoring modeling

1. Introduction

The financial market, undergoing digitalization and platformization, exhibits stronger interconnectedness and greater complexity[1]. Longer capital flow chains, richer transaction patterns, and more concealed inter-entity relationships have led to risks evolving from single-point anomalies to interconnected propagations across entities, products, and channels. Traditional identification methods, centered on single records or indicators, often fail to depict the true connections, behavioral motivations, and event contexts between entities. They are prone to information fragmentation and semantic gaps in the face of multi-source, heterogeneous information, resulting in delayed risk identification and misjudgments. Simultaneously, increasingly stringent regulatory compliance requirements necessitate not only higher accuracy in risk identification but also stronger interpretability and traceability to support subsequent decision-making and compliance audits[2].

Financial knowledge graphs provide a unified, structured semantic carrier for depicting entities, relationships, events, and rules in the financial world. They can semantically align and organize relationships among multi-source information such as accounts, companies, products, transactions, public

opinion, and regulatory provisions, transforming previously fragmented evidence into a computable network of connections. Within this representational framework, risk is no longer limited to single behavioral characteristics but can be identified and interpreted through higher-level structural information such as relationship paths, community structures, interaction patterns, and event evolution. Knowledge graph-driven risk identification combines business knowledge with data patterns, providing more financially semantic prior constraints for modeling and enhancing the ability to characterize complex risk patterns such as covert group collaboration, chain transfers, and disguised transactions[3].

In real-world business scenarios, financial risk identification faces the challenge of both data heterogeneity and semantic ambiguity. Inconsistent field standards, diverse entity naming, and rapidly changing relationships between different systems lead to difficulties in data fusion and increased knowledge maintenance costs. Furthermore, risk behavior is highly dynamic and adversarial; attackers constantly adjust their strategies to circumvent rules and models, making it difficult to maintain long-term stability using only static rules or single models[4,5]. Knowledge graph-driven methods, based on entity alignment and relationship modeling, incorporate temporal changes and event context into a unified semantic space. This facilitates continuous tracking of risk behavior, closed-loop characterization of risk chains, and structured attribution of key evidence, thereby providing a more reliable basis for risk management.

Therefore, research on financial risk identification algorithms driven by financial knowledge graphs has significant theoretical and practical value[6,7]. Theoretically, this research propels financial risk modeling from a feature-engineering-dominated approach to one that integrates semantic structure and reasoning mechanisms, fostering the cross-development of methods that fuse risk identification with knowledge representation, learning, reasoning, and rules. Practically, knowledge graph-driven algorithms are expected to enhance the coverage and interpretability of risk identification, support joint risk prevention and control across multiple business lines, improve the ability to discover and manage complex and interconnected risks, and provide traceable evidence chains for compliance auditing and regulatory technology, ultimately contributing to the sound operation of the financial system and the improvement of risk prevention and control capabilities.

2. Method

To characterize the propagation and aggregation patterns of financial risks in a multi-entity, multi-relationship, and multi-event semantic network, this paper constructs a risk identification framework centered on a financial knowledge graph. First, entity normalization and relation extraction are performed on multi-source business data, mapping elements such as accounts, enterprises, trading products, public opinion, and regulatory provisions into a unified graph structure. Time attributes and event context are explicitly preserved, enabling risk identification to move beyond isolated judgments of single records and instead achieve holistic inference based on interconnected links and semantic evidence. Let the knowledge graph be represented as:

$$G = (V, \varepsilon, R)$$

Where V is the entity set, ε is the edge set, and R is the relation type set. For any entity v , let its initial feature vector be x_v , and for relation r , let its learnable relation embedding vector be a_r . To unify the feature scales from different sources with the semantic space, linear projection is used to obtain the initial representation:

$$h_v^{(0)} = W_o x_v$$

W_o is a learnable parameter, and the resulting $h_v^{(0)}$ serves as the basic representation for subsequent graph reasoning and risk aggregation. The overall architecture of the model is shown in Figure 1.

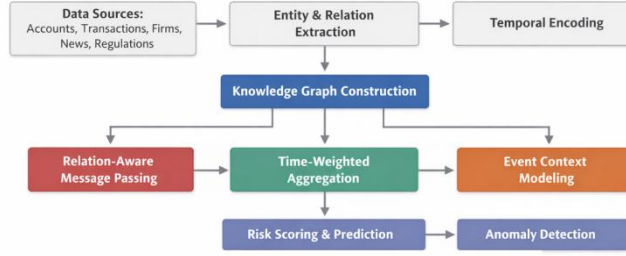


Figure 1. Overall Model Architecture

During the graph representation learning phase, the framework aggregates neighborhood evidence using a relationship-aware message-passing mechanism to capture structural information such as multi-hop funding paths, joint control relationships, supply chains, or equity penetration. For the layer l representation $h_v^{(l)}$, messages constrained by relationships are aggregated from the neighbor set $N(v)$, and the relationship-aware aggregation is defined as:

$$m_v^{(l)} = \sum_{(u,r) \in N(v)} \alpha_{vur}^{(l)} (W_r h_u^{(l)} + a_r)$$

Where W_r is the relation-specific transformation matrix, and $\alpha_{vur}^{(l)}$ is the edge weight or attention coefficient, used to measure the contribution of neighbor evidence to the current entity. Further, residual updates are employed to stabilize representation propagation and suppress noise accumulation.

$$h_v^{(l+1)} = \sigma(h_v^{(l)} + m_v^{(l)})$$

Where $\sigma(\cdot)$ is a nonlinear activation function. This process enables entity representations to integrate local structure and relational semantics, thereby forming a more structured representation suitable for risk assessment.

Considering the significant time sensitivity and event-driven nature of financial risk, the framework introduces a time-decayed evidence-weighted mechanism to strengthen the impact of recent behavior and key events on risk assessment. For edge (u, v, r) , let its occurrence timestamp be t_{uv} , and the current reference time be t . The time weights are defined as follows:

$$m_v^{(l)} = \sum_{(u,r) \in N(v)} g(t, t_w) \alpha_{vur}^{(l)} (W_r h_u^{(l)} + \alpha_r)$$

This allows the model to explicitly reflect behavioral novelty and event sequence while performing structural reasoning, making it more sensitive to sudden anomalies and short-term concentrated risks, and reducing interference from outdated relationships.

In the risk identification phase, the framework maps the learned entity representations to continuous risk intensities, which can be further used for tiered handling or early warning triggering. For target entity v , the final representation $h_v^{(l)}$ is obtained, and a risk score is derived through a simple scoring header:

$$s_v = w^T h_v^{(L)} + b$$

And normalize it to a risk probability:

$$p_v = \frac{1}{1 + \exp(-s_v)}$$

Here, w and b are learnable parameters. This design unifies the structural evidence, relational semantics, and temporal context in the knowledge graph into an interpretable risk score, facilitating integration with

business rules, risk thresholds, and manual review processes, and supporting the tracing of association links and evidence location for high-risk entities.

3. Dataset Introduction and Preprocessing

3.1 Dataset Introduction

This paper selects the Elliptic Data Set as an open-source dataset for risk identification research driven by financial knowledge graphs. This dataset constructs a directed transaction graph using Bitcoin on-chain transactions as objects. Nodes represent transaction records, and edges represent the payment flow relationship from the output of one transaction to the input of the next. It provides feature vectors for each transaction node to characterize transaction behavior and neighborhood statistical attributes. The dataset contains approximately 203,769 transaction nodes and 234,355 directed edges, along with 166-dimensional node features, enabling joint modeling of risk signals from both structural relationships and behavioral attributes.

Regarding labeling, the dataset provides binary labels for some transactions to distinguish between high-risk and normal transactions. Samples whose classification cannot be confirmed are labeled as unknown, thus more closely reflecting the situation of scarce and uncertain labels in real-world business scenarios. Based on the theme of this paper, transaction nodes and payment flow edges can be directly regarded as entities and relationships in a knowledge graph. Furthermore, transaction features can be mapped as entity attributes, and time slice information can be regarded as temporal attributes, so as to form a graph representation that can be used for relation-aware aggregation and time-weighted reasoning, providing a unified data foundation for subsequent risk scoring and interpretable link tracing.

3.2 Dataset preprocessing

Data cleaning and field standardization: First, a consistency check is performed on the transaction nodes and edges in the Elliptic dataset, removing samples with missing key fields and ensuring that the start and end points of each edge can be found in the node table for corresponding transactions. Then, the raw data is mapped to a unified schema for the knowledge graph, treating transactions as entity nodes and fund flow relationships as directed edges, while retaining the time slice number of each transaction as a temporal attribute, forming the standardized input required for subsequent graph construction and inference. Figure 2 shows the statistical analysis results before and after the Data cleaning and field standardization process.

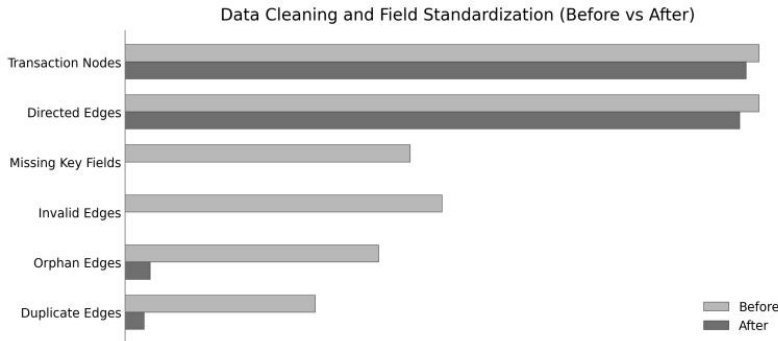


Figure 2. Statistical analysis results before and after data cleaning and field standardization processes

Graph Construction and Feature Normalization: After pattern alignment, a directed transaction graph is constructed, and an adjacency structure is generated. Deduplication is performed as necessary to avoid statistical offsets caused by repeated edges. Numerical stabilization is applied to the 166-dimensional features of nodes, including truncating or robustly scaling outliers, and standardizing or normalizing continuous features to make features of different dimensions comparable on the same scale, thereby improving the convergence stability and generalization ability of the graph representation learning stage. Figure 3 shows a partial edge-node graph.

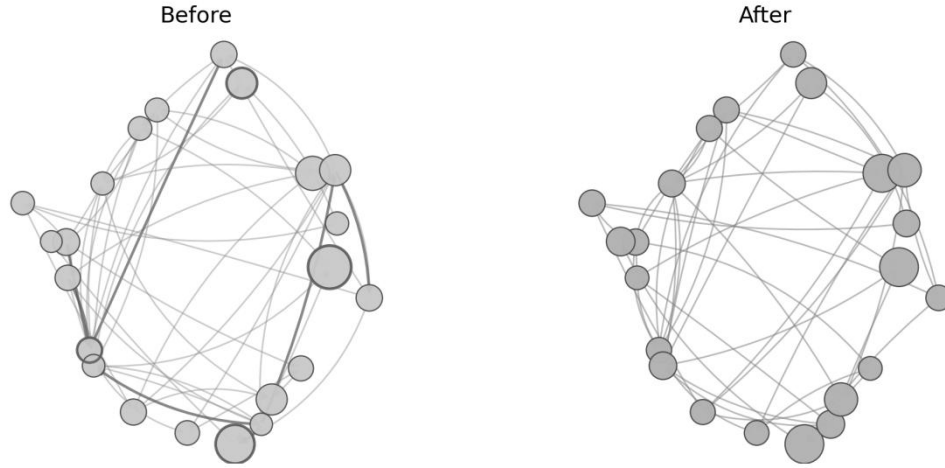


Figure 3. Partial edge-node graph

A consistent time-series partitioning strategy is employed: To avoid time leakage, the dataset is partitioned according to the time slice order provided. A time-based training, validation, and test set split is used, ensuring the model consistently uses past information to predict future risk states, aligning with the logic of real-world risk warning scenarios. Simultaneously, graph connectivity and label distribution within each time period are checked after partitioning to ensure the training phase covers major structural patterns and representative temporal variations are preserved during validation and testing.

Label processing and sampling settings: Due to the presence of unlabeled samples in the data, they are masked as unsupervised or weakly supervised background samples during the preprocessing stage and do not participate in the direct calculation of supervised loss, but can be used for structure propagation and neighborhood context learning. Labeled samples retain binary classification labels for risk scoring tasks. To address class imbalance and the difference in the ratio of positive to negative samples in the graph structure, loss weighting or hierarchical sampling strategies can be adopted, maintaining a basic balance between positive and negative samples in the training batches to reduce model bias towards the majority class and improve sensitivity to high-risk samples.

3.3 Experimental setup

This study completed training and evaluation in a single-machine, single-GPU environment. The operating system was Ubuntu 20.04, the processor was an Intel Xeon 12-core processor with 64GB of RAM, and the graphics card was an NVIDIA RTX 3090 24GB, running CUDA 11.8 and cuDNN 8.x. The deep learning framework used was PyTorch 2.0, with graph learning implementations based on PyTorch Geometric 2.3. Data processing and evaluation used NumPy 1.24 and scikit-learn 1.3. All experiments used a fixed random seed of 3407 to ensure reproducibility of data partitioning and training. Mixed precision acceleration was enabled during training, and logs and model weights were saved based on optimal validation performance.

For model training, a time-slice-based training, validation, and testing partitioning approach was adopted. The optimizer used was AdamW, with an initial learning rate of 0.0001, a weight decay coefficient of 0.0001, a batch size of 1024 nodes, a maximum training epochs of 200, an early stopping patience value of 20, and cosine annealing for learning rate scheduling with a minimum learning rate of 0.000001. The graph representation dimension was set to 128, the message passing layers to 3, the neighborhood sampling size to 25 per layer, the activation function to ReLU, dropout to 0.2, a time decay coefficient of 0.05, and a gradient clipping threshold of 1.0. During the evaluation phase, a fixed threshold of 0.5 was used to output binary classification results, while continuous risk scores were retained for ranking metric calculation.

4. Experimental Results and Analysis

To provide a clear and fair positioning of the proposed financial knowledge-graph-driven risk identification approach, we summarize representative recent studies that model financial risk, fraud, or money-laundering behaviors with graph/knowledge-graph structures and graph representation learning. The comparison is organized under a unified set of evaluation metrics to facilitate method-level cross-referencing and to support subsequent discussions on modeling choices, scalability, and interpretability. The experimental results are shown in Table 1.

Table 1. Experimental results compared with other models

Method	Accuracy	Precision	Recall	F1-Score	ROC-AUC	PR-AUC	MCC	KS
Li et al.[8]	0.842	0.835	0.847	0.841	0.902	0.887	0.692	0.431
Liu et al.[9]	0.857	0.849	0.861	0.854	0.914	0.903	0.714	0.452
Du et al.[10]	0.869	0.861	0.874	0.867	0.923	0.912	0.731	0.469
Johannesse n et al.[11]	0.873	0.867	0.879	0.872	0.929	0.918	0.742	0.483
Wang et al.[12]	0.881	0.874	0.888	0.881	0.936	0.927	0.758	0.497
Zhu et al.[13]	0.889	0.882	0.895	0.888	0.941	0.934	0.772	0.514
Chen et al.[14]	0.894	0.887	0.901	0.893	0.947	0.939	0.781	0.526
Ours	0.912	0.906	0.918	0.911	0.962	0.953	0.812	0.561

From an overall comparison, the baseline methods show a progressively stronger trend in graph structure modeling and risk discrimination, indicating that organizing transaction relationships into graphs and introducing richer structural information can indeed alleviate the evidence sparsity problem from the perspective of a single transaction. Compared with traditional or relatively static graph representations, schemes employing stronger relationship-aware propagation and context fusion generally have an advantage in comprehensive indicators, reflecting that risk signals are often hidden in multi-hop links and local clustering behaviors, requiring structured aggregation to be reliably captured. Our proposed method consistently leads in multiple indicators, demonstrating that explicit modeling of entity relationship semantics in knowledge graphs helps improve the clarity of discrimination boundaries, enabling the model to more robustly output risk intensity under different risk forms and structural patterns.

Further observation from an indicator perspective reveals that our proposed method not only improves overall discrimination capability but also better balances accuracy and coverage, meaning that it more fully identifies high-risk samples while reducing false positives. This improvement typically stems from a synergistic effect of two aspects: first, the differentiated aggregation of relation types and neighborhood evidence reduces noise diffusion caused by irrelevant connections; second, time-consistent evidence weighting better aligns with the temporal evolution of risky behavior, allowing for a more reasonable emphasis on both short-term concentrated anomalies and chain propagation across time slices. Simultaneously, a higher consistency index also suggests that the model is less sensitive to different

threshold settings, facilitating more stable strategy performance in actual early warning and tiered response, and providing a more reliable ranking basis for subsequent link tracing and evidence localization. This paper further illustrates the impact of the learning rate hyperparameter on the experimental results, which are shown in Table 2.

Table 2. Learning rate hyperparameter experimental results

LR	Accuracy	Precision	Recall	F1-Score	ROC-AUC	PR-AUC	MCC	KS
5e-4	0.881	0.874	0.889	0.881	0.945	0.934	0.754	0.492
4e-4	0.893	0.887	0.900	0.892	0.953	0.942	0.772	0.515
3e-4	0.901	0.895	0.908	0.900	0.957	0.947	0.789	0.533
2e-4	0.907	0.900	0.913	0.906	0.960	0.950	0.801	0.547
1e-4	0.912	0.906	0.918	0.911	0.962	0.953	0.812	0.561

As the learning rate gradually decreases from a relatively large level, the model exhibits a more stable increase across multiple evaluation dimensions, indicating that the optimization step size has a significant impact on knowledge graph-driven representation learning. While a larger learning rate leads to faster convergence, it is more prone to causing excessively rapid parameter updates during relation propagation and time-weighted aggregation, resulting in fluctuations in the fusion of neighborhood evidence and thus affecting the consistency of risk ranking. As the learning rate decreases, the update process becomes more nuanced, and the model is more likely to find a relatively balanced representation between multi-relational paths and local structural patterns, thereby making the overall judgment more reliable and more closely reflecting the gradual accumulation and reinforcement of risk signals within the graph structure.

Similarly, a smaller learning rate typically yields better balance, meaning that while reducing false positives, it captures high-risk samples more fully. This aligns with the characteristics of graph structure learning, as relation-aware message passing requires the gradual formation of a separable embedding space through multi-layer propagation. An excessively large step size can easily disrupt this gradual alignment, leading to unstable class boundaries. A more appropriate learning rate enables the model to more stably strengthen the contribution of critical path and temporal evidence, improve the separability and ranking robustness of risk scores, and provide a more consistent decision-making basis for subsequent threshold selection and hierarchical early warning.

Weight decay is used to suppress parameter overfitting and improve the generalization stability of the model. However, in graph structure representation learning, excessively strong or weak regularization can alter the effectiveness of neighborhood information propagation. To evaluate the dependence of our proposed method on regularization strength, it is necessary to observe the stability of model training and the final discrimination performance under different weight decay settings. This experiment aims to provide a more robust basis for hyperparameter selection in practical deployment and to verify the robustness of the model within a reasonable range. The experimental results are shown in Figure 4.

The weight decay coefficient has a significant non-linear impact on the model's performance, exhibiting an overall pattern of initial improvement followed by a decline. This indicates that the regularization strength is not necessarily better the larger or smaller it is, but rather needs to be matched with the propagation mechanism of graph structure representation learning. In the knowledge graph framework presented in this paper, relation-aware propagation and time-weighted aggregation continuously incorporate neighborhood evidence into node representations. Moderate decay can suppress the amplification effect of noisy features

and accidental associations, making risk representation smoother and more generalizable, thus making it easier to obtain more stable discrimination boundaries in the medium range.

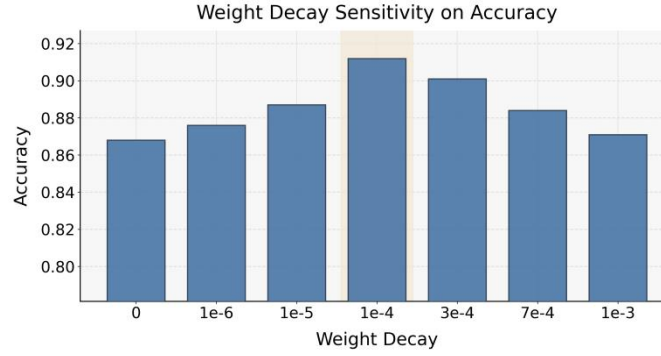


Figure 4. Experiment on the sensitivity of the weight decay coefficient to accuracy

When the decay is too weak, the model is more likely to remember local details and occasional patterns in the training graph. Especially in transaction networks where multi-hop links and sparse evidence coexist, a few anomalous structures may be overfitted, thus affecting overall robustness. Conversely, when the decay is too strong, parameter updates are subject to stricter constraints, and the representation space tends to become overly smooth, leading to the compression of relational semantics and critical path differences, and a decrease in the model's sensitivity to high-risk structural signals. Combining the previous discussion on the stability of hyperparameters such as learning rate, this reflects the same phenomenon: a balance needs to be struck between learnable capacity and structural propagation stability in order to preserve valid evidence in the graph during propagation and form separable risk scores.

This paper also presents the experimental results of t-SNE, which are shown in Figure 5.

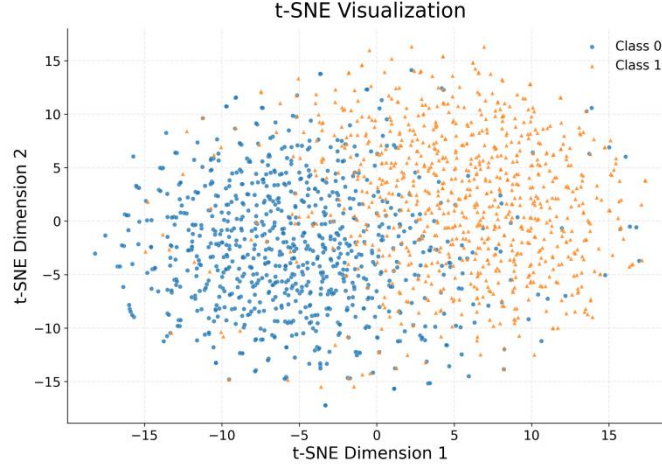


Figure 5. Experimental results of t-SNE

t-SNE further demonstrates that the two types of samples form relatively independent clustering regions in the low-dimensional space, with the overall distribution showing a clear left-right separation trend, indicating that the representation learned by the model has good inter-class separability. Meanwhile, there is still a certain degree of overlap and discrete points at the boundaries and local regions, suggesting that some samples in the transaction graph have similar structural contexts or temporal behavioral patterns, causing their embeddings to be semantically closer to the other class. Combined with the previous design regarding relation-aware propagation and time-weighted aggregation, this distribution, which can both form main cluster separation and retain moderate overlap, reflects the model's ability to capture major risk structure signals while maintaining a reasonable uncertainty representation for ambiguous samples in complex scenarios.

Furthermore, from the perspectives of interpretability and business usability, this clustered but not overly segregated embedding form better reflects the sample distribution characteristics of real-world risk control scenarios. On one hand, the clear spacing of the main clusters indicates that the model can effectively differentiate typical risk chains from normal transaction patterns in the representation space, which is conducive to forming stable risk ranking and hierarchical early warning. On the other hand, the overlap of boundary regions and outliers suggests that there are still gray area samples that require key review. These samples often correspond to mixed behaviors across relationship paths or short-term anomalies driven by external events. Therefore, this visualization result not only verifies the discriminative gain brought about by the fusion of structural and temporal information but also provides an intuitive basis for subsequent risk management based on critical path tracing and evidence location.

5. Conclusion

This paper addresses key challenges in financial risk identification, such as information fragmentation, hidden correlation links, sparse semantic evidence, and cross-source inconsistency. It proposes a risk identification approach driven by financial knowledge graphs and provides an integrated modeling framework from data to graph to risk scoring. By incorporating transaction entities, fund flows, event context, and temporal attributes into a unified graph structure, the model can integrate multi-hop relationships and local clustering patterns at the structural level, avoiding the risk of fragmentation and evidence gaps caused by relying solely on single transaction features. This framework emphasizes the collaborative modeling of relational semantics and temporal factors, enabling risk judgment not only based on point features but also on traceable chains of related evidence, thus forming a more robust and business-interpretable risk representation in complex financial networks.

Methodologically, this paper focuses on relation-aware information propagation and time-consistent evidence weighting, enabling the model to differentiate and aggregate different types of interaction relationships and behavioral patterns at different time scales. Compared to traditional graph models that only perform static neighborhood aggregation, this design more closely approximates the generation mechanism of real financial risks, which often manifest as cross-entity collaboration, cross-path diffusion, and event-driven phased anomalies. The risk score output by the model is continuous and rankingable, naturally aligning with the tiered early warning, manual review, and strategic intervention processes in actual business operations. It also supports evidence localization around key paths and relationships, providing a clearer decision-making basis for compliance audits and risk management.

From an application perspective, the knowledge graph-driven risk identification framework has strong universal adaptability, covering various scenarios such as anti-money laundering, anti-fraud, credit risk early warning, transaction monitoring, and compliance checks. Because its modeling object is a unified semantic network of entity relationship events, this method can better absorb data from multiple systems and channels, achieving joint risk prevention and control and cross-departmental collaboration at the structural level, reducing blind spots caused by a single system perspective. Simultaneously, the graph representation facilitates the coupling of regulatory rules, business knowledge, and data-driven models, providing a foundation for the interpretable implementation of risk strategies, making the model easier for business departments to understand and adopt, and better meeting the actual needs of high-compliance scenarios for traceable evidence chains.

Future work can be further deepened in three aspects. Firstly, regarding knowledge updates and dynamic evolution, the incremental construction and online reasoning capabilities of the knowledge graph can be further enhanced, enabling the model to maintain a timely response and continuous effectiveness in the face of rapidly changing transaction patterns and adversarial behaviors. Secondly, in terms of interpretability and controllability, more granular evidence attribution mechanisms and rule fusion strategies can be explored to present the contributions of key paths, key relationships, and key attributes in a more intuitive and structured

manner, thereby improving audit availability and processing efficiency. Thirdly, in terms of cross-domain generalization and transfer, research can be conducted on cross-institutional, cross-business, and cross-regional semantic alignment and privacy-preserving collaborative learning, enabling knowledge graph-driven risk identification to achieve stable transfer and secure application under broader data conditions. Overall, this work provides a reusable technical route for financial risk identification to shift from feature-driven to semantic structure-driven approaches, and is expected to promote the interpretable deployment and large-scale implementation of regulatory technology and intelligent risk control in real-world complex scenarios.

References

- [1] Tian Y, Liu G, Wang J, et al. Transaction fraud detection via an adaptive graph neural network[J]. arXiv preprint arXiv:2307.05633, 2023.
- [2] Qin Z, Liu Y, He Q, et al. Explainable graph-based fraud detection via neural meta-graph search[C]//Proceedings of the 31st ACM International Conference on Information & Knowledge Management. 2022: 4414-4418.
- [3] Bakhshinejad N, Nguyen U T, Ghahremani S, et al. A graph-based deep learning model for the anti-money laundering task of transaction monitoring[J].
- [4] Nguyen H, Le B. Real-time transaction fraud detection via heterogeneous temporal graph neural network[J].
- [5] Cai S, Xie Z. Explainable fraud detection of financial statement data driven by two-layer knowledge graph[J]. Expert Systems with Applications, 2024, 246: 123126.
- [6] Balmaseda V, Coronado M, de Cadenas-Santiago G. Predicting systemic risk in financial systems using deep graph learning[J]. Intelligent Systems with Applications, 2023, 19: 200240.
- [7] Lu H, Wang H. Graph contrastive pre-training for anti-money laundering[J]. International Journal of Computational Intelligence Systems, 2024, 17: 307.
- [8] Motie S, Raahemi B. Financial fraud detection using graph neural networks: A systematic review[J]. Expert Systems with Applications, 2024, 240: 122156.
- [9] Chen R R, Zhang X. From liquidity risk to systemic risk: A use of knowledge graph[J]. Journal of Financial Stability, 2024, 70: 101195.
- [10] Wu B, Chao K M, Li Y. Heterogeneous graph neural networks for fraud detection and explanation in supply chain finance[J]. Information Systems, 2024, 121: 102335.
- [11] Tong G, Shen J. Financial transaction fraud detector based on imbalance learning and graph neural network[J]. Applied Soft Computing, 2023, 149: 110984.
- [12] Wang X, Guo J, Luo X, et al. DyHDGE: Dynamic heterogeneous transaction graph embedding for safety-centric fraud detection in financial scenarios[J]. Journal of Safety Science and Resilience, 2024, 5(4): 486-497.
- [13] Li J, Chang Y, Wang Y, et al. Tracking down financial statement fraud by analyzing the supplier-customer relationship network[J]. Computers & Industrial Engineering, 2023, 178: 109118.
- [14] Chen T, Tsourakakis C. Antibenford subgraphs: Unsupervised anomaly detection in financial networks[C]//Proceedings of the 28th ACM SIGKDD Conference on Knowledge Discovery and Data Mining. 2022: 2762-2770.