

AI-Driven Fraud Detection in Payments Using Transformer-Based Behavioral Modeling

Abhigyan Mukherjee

abhigyan.mukherjee@yahoo.com

Abstract:

As digital payments got mainstream, fraudsters have been getting smarter. Digital payment fraud detection gets powered by AI. In this paper we propose a Transformer-based fraud detection framework based on Generative Pretrained Transformers (GPT) to encode long-term transactional behavior sequences. Based on a very large payment data, we pretrained our model to reconstruct behavior and identify anomalies in real time. Our strategy employs unsupervised pretraining, differential convolutional operations and contrastive learning for anomaly detection without requiring labeled data. Most of the work is evaluated on real transaction data from WeChat Pay, one of the largest payment platforms in China. The proposed model achieves better accuracy than traditional machine learning. According to the findings, the implementation of AI-enhanced behavioral analysis on payment transactions effectively identifies high-risk consumers or frauds while minimizing false positives and securing transactions. The paper demonstrates the advantages of deep learning in financial fraud prevention through an empirical evaluation and presents a scalable approach for real-time fraud detection in payment ecosystems.

Keywords: AI-driven fraud detection, Transformer-based modeling, Behavioral sequence analysis, Unsupervised learning, Contrastive learning, Few-shot learning, financial transaction data, Anomaly detection, Deep learning in finance, Payment risk management, Autoregressive models, Data compression in behavioral modeling, Self-supervised pre-training, Convolutional neural networks (CNN), Real-time fraud prevention.

I. INTRODUCTION

As the world of digital finance continues to grow in complexity, transaction fraud detection is no easy task. Numerous existing machine learning-based behavioral analysis models are unable to cope with the complexity and dimensionality of financial transaction data. Classical methods find it difficult to get highly accurate patterns due to the very large volumes and complexities of this type of data.

To overcome these shortcomings, we propose a novel approach that utilizes Generative Pretrained Transformers (GPT) for the modelling of financial transaction data. The models conceptualized for understanding language can successfully capture sequential dependencies and as such, are quite capable of diagnosing user payment behavior. By pretraining on large-scale transaction datasets, our framework can avoid issues arising from dependence on labelled data and inferring long-term behavior patterns.

We provide three contributions. Follow autoregressive pretraining mechanism proprietary to the financial data so that the model can learn good user behavior representations. For subtle transaction changes, we propose a differential convolutional structure in the second step for better anomaly detection. The experiments show that the model is flexible in different financial environments and that it scales well.

By extending beyond fraud detection, we aim to demonstrate the potential of unsupervised learning for financial security. Our technique paves the way for the development of safer, smarter, and more flexible fraud detection systems by lowering reliance on annotated data and securing user privacy.

II. RELATED WORK

A. Sequential Recommendation and Pretraining Strategies

Sequential recommendation systems model user preferences by analyzing their past interaction sequences. Among earlier approaches, GRU4Rec [1] employs Gated Recurrent Units (GRUs) to process session-based data. Subsequently, architectures derived from the Transformer [2] and BERT [3] paradigms have been influential, leading to models such as SASRec [4] and BERT4Rec [5] that leverage self-attention mechanisms to represent user behavior patterns. Further research has also addressed cold-start scenarios through meta-learning frameworks and enhancements to behavioral sequence data.

In contemporary machine learning, the strategy of pretraining has become widely adopted in both language and visual domains. This approach involves acquiring foundational knowledge from extensive corpora prior to specialization on narrower tasks. Within recommendation systems, methodologies like BERT4Rec [5], S3Rec [6], and PeterRec [8] implement this paradigm. The incorporation of objectives such as predicting obscured items, employing contrastive loss functions, and utilizing user metadata has been shown to enhance the fidelity of subsequent recommendations.

Contemporary research incorporates pre-trained language models within recommendation system architectures, leading to models such as UniSRec [10], P5 [11], M6-Rec, and M8Hd [12]. These methods utilize descriptive item text to enhance recommendation quality. However, their practical adoption is constrained by the requirement for auxiliary textual data and the high computational cost of large language models. To address this, we propose a versatile and generalizable architecture designed to achieve a more effective fusion of the pre-training phase with specific recommendation tasks, requiring only minimal external textual information.

B. Data Compression and Behavioral Sequence Modeling

The compression of data is very important for optimizing behavior sequence models so that information is gained from large-scale user behavior. Behavior sequences can be compressed to model the underlying motivations of a behavior. The subsequent two statements stress on behavioral modelling and data compression relationship.

- **Efficient Encodings:** Through the compression of behavioral data, such that only essential patterns remain, models will select the key components, while discarding redundancies. User interactions can be captured better in an effective way. By compressing data, dimensionality reduction on input sequences helps identify the essential behavior patterns that improve user preference and engagement modelling. Less noise in training data prevents overfitting. Hence a compact representation. It ensures that models learn relevant patterns of behavior and not unwanted information.
- **Noise Reduction:** Data about user behavior is usually very noisy which makes learning irrelevant interactions. By means of coding and compression techniques, other than behavior are removed and models focus on the behavior signal.
- **Temporal Dependency Learning:** By helping detect any anomalous events and providing relevant recommendations, it is important to capture sequential relations between user actions. The model can encode temporal relations with compressed representations.

• **Scalability:** For large-scale Recommendation and Risk detection operations, models capable of handling extensive data are necessary. Enhancing scalability through data compression reduces costs by encoding important behavior.

Integrating data compression techniques into behavioral sequence modelling enhances sequence prediction accuracy, reduces noise interference, and enables efficient processing of massive user data. Collectively, they enhance recommendations and risk detection, improving overall system performance.

III. DATA CURATION AND PREPROCESSING FOR PRETRAINING

The curation and initial transformation of training data are fundamental to the efficacy of a model trained via self-supervision. This study utilizes payment event logs from a major financial services provider in China, capturing comprehensive user activity. The data was filtered according to the following principles:

1. Prioritizing users with frequent transaction histories, as their patterns exhibit greater volume and consistency for model learning.
2. Incorporating sequences from accounts flagged for review enables the model to distinguish between normative and anomalous financial activities.
3. Segmenting transaction timelines across specific intervals facilitates modeling both transient and enduring behavioral shifts.

This curation approach ensures the training corpus mirrors genuine payment ecosystems, enhancing the model's capacity for generalization in risk assessment tasks.

The processed dataset encompasses approximately 200 million distinct users, corresponding to a sequence length of 15 billion behavioral tokens. An expansive representation of over 1.3 trillion tokens is generated by encoding features such as transaction value, temporal stamps, and payment method into a structured, multi-variable sequential format. This rich, structured input allows sequential architectures to discern intricate behavioral dynamics, leading to superior predictive performance.

IV. AUTOREGRESSIVE BEHAVIOR SEQUENCE MODELING

The temporal structure can effectively represent how user behavior runs in a sequentially natural way [13], better used than the bidirectional attention mechanisms adopted in BERT. Unlike bi-direction model that considers entire sequence at once, autoregressive model does things stepwise to predict behaviors.

In a GPT-style generative model, the probability of a behavior sequence, $B_1, B_2, \dots, B_T$, can be expressed as:

$$P(B_1, B_2, \dots, B_T) = \prod_{t=1}^T P(B_t | B_1, B_2, \dots, B_{t-1}). \quad (1)$$

The autoregressive format ensures every behavior is conditioned on what happened before (timestamped actions).

The equation denotes autoregressive (AR) or time series model. It is possible to compare the AR model mathematically to show that it is like the GPT-based model.

$$X_t = c + \sum_{i=1}^p \phi_i X_{t-i} + \epsilon_t, \quad (2)$$

where ϕ_i are autoregressive coefficients, c is a constant term, and ϵ_t represents the error term.

The two frameworks essentially employ the same principle: they make predictions based on congruence with prior observations. Assuming AR models today's values depend linearly on yesterday's values, but the GPT architecture uses very complex deep learning to model highly non-linear dependencies. We can be rest assured that autoregressive models properly capture how users behave in a robust and adaptive manner.

V. PROPOSED PRETRAINING FRAMEWORK FOR SEQUENTIAL BEHAVIOR LEARNING

To improve the modelling of sequential behaviors, we devise a very comprehensive pretraining framework that can capture temporal and contextual information in depth. This method reduces issues like a surge in tokens and rebuilding behavior without increasing costs.

A. Multivariate Behavior Representation

At each time step t , user behavior is characterized by multiple discrete attributes, denoted as:

$$a_t = [a_1, a_2, a_3, \dots]_t. \quad (3)$$

The embedding layer takes features like transaction type, device ID, location, etc., which could be categorical, and converts them to a dense vector. The concatenation happens over time to form an input matrix with dimensions afterwards (sequence length $\times$ embedding dimension).

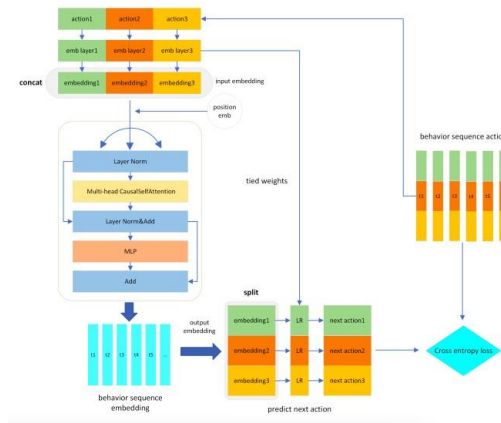


Fig. 1: Illustration of the proposed pretraining model for sequential behavior learning.

Since the embedding-based representation facilitates the grasping of complex behavioral relations, it guarantees temporal consistency. Our framework enhances predictive performance in different downstream applications, including fraud prevention, recommendation systems, and ca, using sequential learning techniques.

B. Autoregressive Modeling of Behavioral Sequences

The introduced methodology is grounded in an autoregressive formulation to sequentially classify user interactions. Its core objective is to maximize the predictive accuracy of an imminent event given its

historical context. Formally, the model is trained to estimate the conditional probability of an event s_t at a specific temporal index:

$$P(s_t | s_{t-n}, \dots, s_{t-1}; \Theta), \quad (4)$$

In this expression, Θ denotes the complete set of learnable model parameters, while the integer n defines the window of preceding events utilized for each prediction.

C. Mitigating Token Proliferation via Concatenation

We introduce a concatenation approach for embeddings to solve the token overload problem of multivariate behavioral sequences. At every time step, the method combines the multiple behavior dimensions into a single token to reduce the embedding dimension. Consequently, it increases efficiency while improving generalization and reducing sparsity.

D. A Reconstruction-Oriented Paradigm for Sequential Forecasting

This work introduces a novel learning paradigm for forecasting subsequent states in behavioral trajectories. Rather than generating the next token as a monolithic unit, the model projects the target event into orthogonal representational subspaces and reconstructs it component-wise.

The associated reconstruction loss is mathematically defined as:

$$L_{\text{rec}} = \frac{1}{T} \sum_{t=1}^T \frac{1}{D} \sum_{d=1}^D L_{\text{CE}}(s_{t,d}, \hat{s}_{t,d}), \quad (5)$$

Here, T signifies the total sequence length, D corresponds to the count of behavioral attributes, and L_{CE} represents the cross-entropy loss function computed independently for each attribute prediction.

E. Parameter Sharing via Weight Tying

To boost model efficiency with a matching input and output embeddings, we incorporate *weight tying* into the framework. This strategy introduces two primary benefits:

1. It would minimize trainable coefficients and accelerate the training process.
2. It works like regularization, stabilizing weight updates and making the encoding and decoding more consistent, being introduced to the weighting factors of various components.

In summary, our sequential behavior modelling pretraining is successful in capturing temporal dependency and dealing with token redundancy as well as next event reconstruction issues. The framework is a strong solution for sequential behavior modelling and prediction due to the use of token concatenation, effects disentangling, and weight tying on an autoregressive architecture.

VI. FINE-TUNING FOR ANOMALY DETECTION USING DIFFERENTIAL AND CONVOLUTIONAL OPERATIONS

The model would be more fit for anomaly detection after pretraining and fine-tuning. As the model goes through the process, it learns to capture the differences in context in sequences while transforming the 2D representation of the sequence into compact 1D, which can successfully detect anomalous behavior.

A. Supervised Fine-Tuning (SFT)

1) *Temporal Stabilization via Differential Operation:* A differential sequence operator is developed to handle non-stationarity in behavior sequences. This method removes long-run trends and seasonal

components from the data to ensure that the model focuses on the major fluctuations that signify anomalies. To increase prediction process efficiency, we consider embeddings' anomalous behaviors using a first-order difference transformation.

2) *Convolutional-Based Anomaly Detection*: After generating the differential sequence, we employ convolution for the construction of an anomaly detection mechanism. The convolutional layers of the CNN can extract local features, have shared weights, are translation invariant etc. The traits encountered by the model enable it to spot unusual behaviors in the sequence.

The hierarchical convolutional structure of the model learns fine-grained anomaly representations important for downstream anomaly detection tasks.

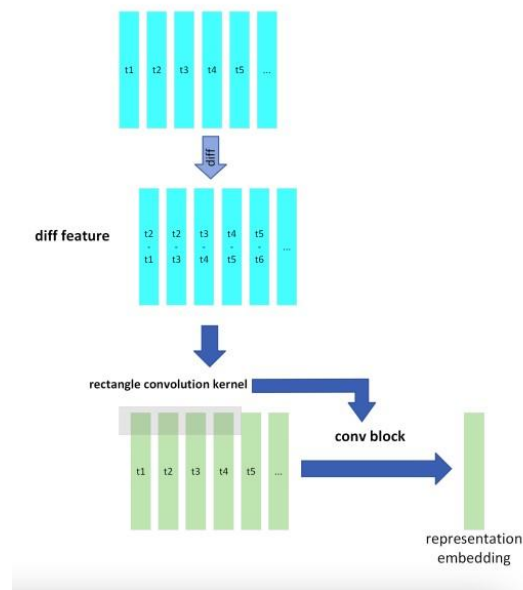


Fig. 2: Comparison of different encoder fine-tuning strategies.

VII. SELF-SUPERVISED FINE-TUNING USING CONTRASTIVE LEARNING

A. Learning Representations via Contrast

Contrastive learning has emerged as a prominent self-supervised technique within natural language processing and computer vision. By training on unannotated data, it enhances a model's transferable capabilities for subsequent specialized tasks.

The core principle involves teaching a model to differentiate between analogous and disparate data instances. This is accomplished by minimizing a contrastive objective function, which reduces the distance between representations of similar pairs while increasing it for dissimilar pairs.

We generate negative samples through a dropout-based augmentation method akin to SimCSE. This technique applies varied dropout masks to the same input sequence, producing multiple perturbed versions that serve as positive pairs. Simultaneously, distinct sequences function as negatives. Learning to distinguish these contrasting examples endows the model with strong representational generality.

B. Contrastive Objective for Sequential Behavior Embeddings

To refine the embeddings of behavioral sequences, we adopt the *InfoNCE loss*, a standard contrastive objective. It is formally expressed as:

$$L_i = -\log \frac{\exp(\text{sim}(v_i, v_i^+))}{\exp(\text{sim}(v_i, v_i^+)) + \sum_{j=1}^{N-1} \exp(\text{sim}(v_i, v_j^-))} \quad (6)$$

Here, $\text{sim}(v_i, v_j)$ represents a similarity metric, usually cosine similarity, between embedding vectors v_i and v_j .

The aggregate contrastive loss is then the mean across all instances:

$$L = \frac{1}{N} \sum_{i=1}^N L_i \quad (7)$$

By optimizing this loss, the model learns to project embeddings of analogous behaviors nearer in the latent space while separating those of dissimilar ones. This process enables the model to internalize the underlying structure of behavioral sequences, thereby improving downstream anomaly identification.

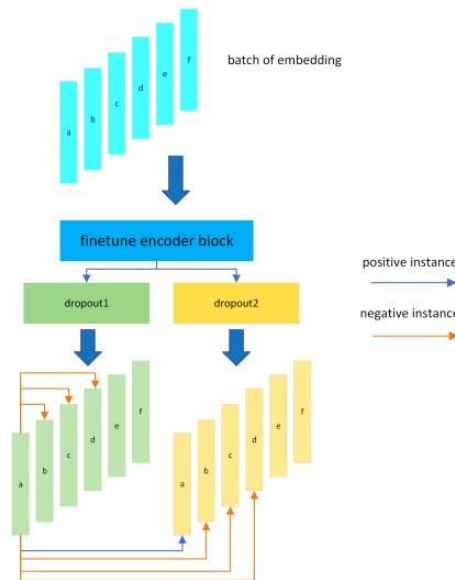


Fig. 3: Architecture of the contrastive learning framework for behavioral sequence analysis.

VIII. EMPLOYING FEW-SHOT LEARNING FOR FRAUD IDENTIFICATION

Within the domain of payment security, we address the challenge of identifying fraudulent activities when labeled examples are scarce using *few-shot learning*. Our framework strengthens anomaly detection systems by utilizing the generalized knowledge embedded within pre-trained models to recognize illicit accounts.

A. Connecting Anomaly Detection with Few-Shot Learning

Transaction monitoring and user complaint analysis are core components of payment risk management aimed at uncovering irregularities. Since confirmed fraudulent accounts represent a minute fraction of all data, this task aligns with a semi-supervised learning paradigm. Consequently, a model's ability to infer unseen fraudulent patterns from limited evidence is critical.

Models initialized with pre-training on extensive datasets typically surpass models trained from initialization in few-shot scenarios. They can extrapolate from minimal demonstrations to recognize novel instances. Therefore, our approach involves an initial phase of broad behavioral pre-training, enabling the model to later deduce new fraud typologies from few examples.

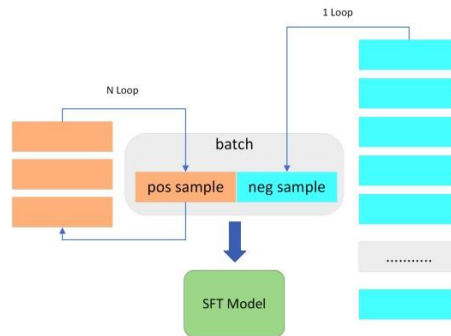


Fig. 4: Depiction of contrastive learning applied to anomaly detection.

B. The Role of Anomaly Detection in Financial Security

A robust payment risk control framework is fundamental to business integrity. Detecting behavioral deviations before they materialize into significant financial loss is essential. The model's proficiency in discerning subtle anomalous patterns is vital for continuous monitoring, requiring it to distill the underlying representation of fraud from a handful of identified cases.

Operationally, the model continuously evaluates transactions against established behavioral profiles, flagging suspicious activities in real-time. This proactive detection helps contain potential fraud before it escalates in scale.

C. Fine-Tuning with Limited Supervised Data under Imbalance

Fraudulent transactions constitute a severe minority within the data, creating a pronounced class imbalance that complicates model training. To address the scarcity of positive (fraudulent) examples, we employ targeted sampling strategies.

We apply oversampling to augment the limited positive samples while utilizing a random subset of the abundant negative (legitimate) samples during training. This differentiated sampling regimen enables the model to learn the discriminative boundary between fraud and normal transactions effectively prior to any re-balancing adjustments.

Our methodology directly tackles the difficulties inherent in fraud detection by integrating few-shot learning principles with pre-trained models on skewed datasets. This enables robust learning from minimal labeled data, resulting in a more adaptable and resilient system for payment risk management.

IX. EXPERIMENTAL EVALUATION

To show the effectiveness of our model, we present ample experiments on various tasks including multi-class classification, binary classification for scoring the anomaly and retrieval on embedding based on contrastive learning. Nine selected behavioral features were experimented upon.

A. Multi-Class Classification Performance

Detecting financial fraud was one of the main evaluation tasks with a multi-class classifier. Due to their inherent low frequency of occurrence, online payment fraud cases are ideal candidates for a few-shot

learning setting. A variety of fraudulent activities encounter a limited number of positive instances, often in the few hundreds. So, they account for a very tiny fraction of the total user base.

The evaluation dataset includes 500,000 negatives and 22,042 positives. The majority of positive cases stem from user complaint records rather than manual review. So, there might be some other false positives. In the future, other datasets with manual verification will be used for this experiment, since multi-class classification needs enough positive examples per class.

Table I provides an overview of the classification performance for different fraud categories.

Fraud Category	Recall (%)	Precision (%)	Positive Sample Ratio (%)
Normal	98.7	99.0	96.1
Free Gift Fraud	13.6	30.6	0.1
Adult Content Fraud	1.3	23.6	0.18
Undelivered Transaction	19.6	27.9	0.58
Gaming Fraud	4.2	25.8	0.35
Other Fraud	28.5	34.6	0.27
Financial Credit Fraud	5.14	27.2	0.077
Part-Time Job Fraud	59.6	49.5	0.34
Dating Fraud	76.2	42.0	1.34

TABLE I: Multi-class classification results for different fraud types.

The results show that while the *Part-Time Job Fraud* with a relatively small number of samples, the high precision and high recall suggest that the model may nevertheless capture more subtle behavioral anomalies in recruitment. This suggests that the model can detect frauds which deviate from normal user behavior.

B. Binary Classification for Abnormal Behavior Scoring

Besides the multi-class classification task, we also tested the model in a binary classification task which generates anomaly scores for user behaviors.

Experimental Design: We found only 800 frauds in a dataset of 80 million high-value transaction users. To evaluate the model's ability to rank, we calculated the score of an anomaly and sorted them. Because of the high-class imbalance, positive instances were only 0.001% of the samples.

The evaluation was checked manually using the complaint logs confirming that it included only genuine fraud cases. False positive chances are eliminated, guaranteeing high reliability of achieved results.

Table II presents the ranking-based anomaly detection performance.

Ranking Threshold	Precision (%)	Recall (%)
Top 1%	0.02	19.16
Top 0.1%	0.05	5.5
Top 0.01%	0.13	1.19

TABLE II: Ranking-based evaluation of abnormal behavior scoring.

The model is proven to highly rank the fraudster accounts, as per the evidence. The dataset has plenty of real users. It is understandable that this is low precision. Nevertheless, the recall values appear to show that many false accounts get predicted in the above average predictions. The model is useful for risk-based prioritization of accounts for further investigation in money laundering scenarios.

X. ANALYSIS OF EXPERIMENTAL OUTCOMES

This section offers a detailed examination of the empirical findings reported in Section 9, with primary emphasis on the multi-class classification performance. Subsequently, the binary classification results for anomaly scoring are evaluated. The analysis highlights the model's capabilities, limitations, and avenues for enhancement.

A. Analysis of Multi-Class Classification Performance

The outcomes displayed in Table I indicate that the model attains a recall of 98.7% and a precision of 99.0% for identifying legitimate transactions. Performance across different fraudulent categories, however, varies substantially.

- **Part-Time Job Fraud** achieves the highest recall (59.6%) among fraudulent classes, suggesting its transactional signatures are relatively distinctive and more readily identifiable.
- **Dating Fraud** also exhibits patterns sufficiently unique for the model to recognize with moderate effectiveness.
- **Free Gift Fraud, Adult Content Fraud, and Financial Credit Fraud** prove more challenging; Free Gift Fraud attains a recall of merely 13.6%. These categories are difficult to distinguish because their transaction profiles closely resemble normal user behavior.
- **Precision Trade-offs** are observed for certain fraud types. Free Gift Fraud and Adult Content Fraud show low precision scores of 30.6% and 23.6%, respectively, indicating a significant number of false positives accompany correct identifications.

These results imply that the model successfully detects fraud categories with well-defined behavioral divergence from normalcy but struggles with those exhibiting substantial overlap with legitimate transactions.

B. Analysis of Binary Anomaly Scoring

The data in Table ?? illustrates the model's proficiency in ranking users by their predicted likelihood of fraudulent activity. Despite fraudulent accounts constituting only 0.001% of the dataset, the ranking-based methodology effectively surfaces them.

- Within the **Top 1%** of ranked users, the model attains a recall of 19.16%, meaning approximately one-fifth of all fraudulent accounts are contained in this highest-risk segment.
- At more stringent thresholds (**Top 0.1% and Top 0.01%**), recall declines to 5.5% and 1.19%, respectively. This underscores the inherent difficulty of pinpointing anomalies within extremely imbalanced data at granular levels.
- Model **precision improves progressively with stricter thresholds**, confirming that the highest-ranked accounts correspond to the most reliably predicted fraud cases.

The scoring mechanism demonstrates efficacy in prioritizing suspicious accounts, though achieving high recall at very narrow percentiles remains a significant challenge.

C. Limitations and Prospective Enhancements

While the model shows promising results, several constraints merit attention:

- **Class Imbalance:** The extreme scarcity of positive samples for certain fraud types (e.g., Financial Credit Fraud) impedes detection accuracy.
- **False Positives:** Low precision for specific fraud categories leads to numerous false alarms, which can adversely affect operational efficiency and user experience.
- **Feature Engineering:** Further refinement of behavioral feature representations could improve discriminatory power across challenging categories.
- **Few-Shot Learning Optimization:** Advancing few-shot learning techniques may bolster detection for fraud types with minimal available training instances.

Future work will explore techniques such as advanced self-supervised learning and adversarial training to increase the model's robustness and generalization capacity for payment risk mitigation.

XI. CONCLUSION

This work presents an innovative framework for behavior pre-training, applicable to anomaly identification and related tasks. The introduced approach provides significant benefits and is suitable for diverse real-world implementations. The principal contributions and findings are detailed below:

- **Scalability and Few-Shot Adaptability:** Our framework possesses a highly scalable design that operates effectively within high-dimensional data spaces. Furthermore, its capacity for few-shot learning is crucial for addressing class imbalance or cold-start scenarios in anomaly identification, particularly where labeled examples are sparse for specific attributes or categories.
- **Versatile Deployment:** The model is suitable for applications such as identifying financial fraud, evaluating risk, and modeling user activity patterns. It can be employed for vector-based similarity search, can be fine-tuned for specialized objectives, or can serve as a feature extractor for traditional supervised algorithms.

To summarize, the proposed framework demonstrates robust capabilities for anomaly detection and serves as an effective pre-training strategy. Future research will focus on augmenting the model's robustness to severe class imbalances, increasing its generalization through advanced self-supervised objectives, and extending its applicability to novel domains.

REFERENCES:

1. Balazs Hidasi, Alexandros Karatzoglou, Linas Baltrunas, and Domonkos' Tikk. Session-based Recommendations with Recurrent Neural Networks.
2. Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N. Gomez, Lukasz Kaiser, and Illia Polosukhin. Attention Is All You Need.
3. Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding.
4. Wang-Cheng Kang and Julian McAuley. Self-Attentive Sequential Recommendation.
5. Fei Sun, Jun Liu, Jian Wu, Changhua Pei, Xiao Lin, Wenwu Ou, and Peng Jiang. BERT4Rec: Sequential Recommendation with Bidirectional Encoder Representations from Transformer.
6. Kun Zhou, Hui Wang, Wayne Xin Zhao, Yutao Zhu, Sirui Wang, Fuzheng Zhang, Zhongyuan Wang, and Ji-Rong Wen. S3-Rec: SelfSupervised Learning for Sequential Recommendation with Mutual Information Maximization.

7. Xu Xie, Fei Sun, Zhaoyang Liu, Shiwen Wu, Jinyang Gao, Bolin Ding, and Bin Cui. Contrastive Learning for Sequential Recommendation.
8. Fajie Yuan, Xiangnan He, Alexandros Karatzoglou, and Liguang Zhang. Parameter-Efficient Transfer from Sequential Behaviors for User Modeling and Recommendation.
9. Chaojun Xiao, Ruobing Xie, Yuan Yao, Zhiyuan Liu, Maosong Sun, Xu Zhang, and Leyu Lin. UPRec: User-Aware Pre-training for Recommender Systems.
10. Yupeng Hou, Shanlei Mu, Wayne Xin Zhao, Yaliang Li, Bolin Ding, and Ji-Rong Wen. Towards Universal Sequence Representation Learning for Recommender Systems.
11. Shijie Geng, Shuchang Liu, Zuohui Fu, Yingqiang Ge, and Yongfeng Zhang. Recommendation as Language Processing (RLP): A Unified Pretrain, Personalized Prompt & Predict Paradigm (P5).
12. Zeyu Cui, Jianxin Ma, Chang Zhou, Jingren Zhou, and Hongxia Yang. M6-Rec: Generative Pretrained Language Models are Open-Ended Recommender Systems.
13. Bowen Tan, Zichao Yang, Maruan Al-Shedivat, Eric P. Xing, and Zhiting Hu. Progressive Generation of Long Text with Pretrained Language Models.
14. Tom B. Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, Sandhini Agarwal, Ariel Herbert-Voss, Gretchen Krueger, Tom Henighan, Rewon Child, Aditya Ramesh, Daniel M. Ziegler, Jeffrey Wu, Clemens Winter, Christopher Hesse, Mark Chen, Eric Sigler, Mateusz Litwin, Scott Gray, Benjamin Chess, Jack Clark, Christopher Berner, Sam McCandlish, Alec Radford, Ilya Sutskever, and Dario Amodei. Language Models are Few-Shot Learners.