



ARTICLE

FRAUD-LENS: Hybrid Deep Learning for Real-Time Unemployment Insurance Fraud Detection via Temporal Behavioral Drift

Rahul Raj*

Enterprise Business Services (EBS)–People Technology, Walmart Global Tech, Bentonville, AR, USA

*Corresponding Author: Rahul Raj. Email: eg13rahuliim@gmail.com

Received: 31 March 2026; Accepted: 09 June 2026; Published: 22 July 2026

ABSTRACT: Background: Unemployment Insurance (UI) fraud represents one of the most costly threats to social benefit integrity, with the U.S. DOL/ETA estimating improper payments exceeding \$45 billion between 2020 and 2023. Existing detection systems fail to model the temporal evolution of claiming behavior or the relational topology connecting fraudulent actors across employer-claimant networks. This study introduces FRAUD-LENS, a hybrid deep learning framework delivering interpretable and scalable fraud detection for large-scale federal UI systems. **Methods:** FRAUD-LENS integrates three coordinated architectural modules: a Bidirectional Long Short-Term Memory (BiLSTM) network encoding temporal claim behavior sequences, a Graph Attention Network (GAT) modeling employer-claimant relational topology to surface fraud ring membership, and a calibrated XGBoost ensemble on 47 engineered tabular features. A novel feature engineering paradigm, Temporal Behavioral Drift (TBD), is introduced with formal convergence guarantees to quantify deviation from a claimant's personal behavioral baseline across seven behavioral dimensions. Predictions from all three streams are fused through a fully specified cross-modal attention mechanism. A Fraud Probability Cascade (FPC) score provides continuous triage aligned with DOL/ETA protocols, with thresholds calibrated through an empirical three-step process. **Results:** Evaluated on a de-identified dataset of 1.2 million UI claim sequences from 18 states (2018–2023), with GAN-based synthetic augmentation applied exclusively to the training partition, FRAUD-LENS achieves F1-score of 95.0% (95% CI: [94.4%, 95.6%]), AUC-ROC of 0.981, precision of 95.3%, and recall of 94.8%, outperforming all evaluated baselines including a Temporal Transformer (F1: 90.8%). Temporal validation (train 2018–2021, test 2022–2023) yields F1 of 94.3%; cross-state held-out validation across five unseen states yields mean F1 of 93.4% (range: 90.5%–96.4%). Fairness analysis reveals consistent performance across age groups and geographic strata. The cold-start subgroup (fewer than four certification weeks) achieves F1 of 84.1% and is identified as the primary performance gap. **Conclusions:** FRAUD-LENS demonstrates strong and generalizable performance for UI fraud detection. The TBD methodology and FPC scoring mechanism show promise for extension to banking, healthcare, and other means-tested benefit program fraud contexts, though further external validation is warranted before broad deployment.

KEYWORDS: Unemployment insurance fraud detection; hybrid deep learning; graph attention network; bidirectional LSTM; temporal behavioral drift; XGBoost ensemble; fraud probability cascade; social benefit integrity; anomaly detection; DOL/ETA

1 Introduction

Unemployment Insurance (UI) programs represent the foremost institutional buffer between sudden job loss and economic destitution for millions of American workers. Their design combining universal eligibility, decentralized state administration, and weekly self-certification makes them inherently difficult

to secure against fraud without imposing undue burdens on legitimate claimants. The tension between accessibility and integrity defines the central challenge of UI program design and nowhere is this tension more acute than in fraud detection.

The COVID-19 pandemic exposed that tension at historic scale. When initial pandemic claims reached 6.6 million in the week ending 28 March 2020, a 3000 percent surge over pre-pandemic weekly averages state UI systems built to process tens of thousands of claims per week were overwhelmed within days [1]. Organized criminal networks exploited the crisis: purchasing stolen identity databases, constructing automated claim submission pipelines, and routing fraudulent payments through prepaid debit cards before detection systems could respond. The DOL Office of Inspector General estimated that fraudulent and improper payments across 2020–2023 exceeded \$45 billion [2].

The inadequacy of existing detection approaches was not a failure of implementation but of architecture. Rule-based systems cannot adapt to novel fraud vectors. Shallow machine learning classifiers applied to individual claim records cannot model the temporal arc of a claimant's behavior over weeks and months, nor the network of relationships connecting a single fraud operator's dozen of simultaneous stolen-identity claims. What is required is a detection framework that is simultaneously temporal, relational, and tabular.

This paper introduces FRAUD-LENS, a hybrid deep learning framework addressing these requirements through three coordinated architectural modules. The first is a Bidirectional Long Short-Term Memory (BiLSTM) network encoding temporal claim behavior sequences. The second is a Graph Attention Network (GAT) module modeling a bipartite employer-claimant graph. The third is an XGBoost ensemble providing high-precision classification on engineered tabular features. Underlying all three streams is Temporal Behavioral Drift (TBD), a novel feature engineering methodology with established theoretical convergence properties under mild regularity conditions.

The principal contributions of this paper are: (1) the FRAUD-LENS architecture, the first framework to jointly model temporal, relational, and tabular fraud signals in UI claims at multi-state scale; (2) TBD with formal convergence guarantees and a theoretical fairness invariance property; (3) the Fraud Probability Cascade (FPC) score with empirically calibrated and cross-state-validated thresholds; (4) a comprehensive evaluation framework including temporal validation, cross-state held-out validation, confidence intervals, and fairness analysis; and (5) full transparency on data access protocols, graph construction leakage prevention, and reproducibility provisions.

The remainder of this paper is structured as follows. [Section 2](#) reviews related literature. [Section 3](#) details the dataset, data governance, TBD with theoretical analysis, FRAUD-LENS architecture with pseudocode and complexity analysis, and training protocol. [Section 4](#) reports all empirical results including confidence intervals, fairness analysis, cold-start evaluation, and threshold justification. [Section 5](#) discusses deployment considerations with latency analysis, limitations, and future directions. [Section 6](#) concludes.

2 Related Work

2.1 Fraud Detection in Social Benefit Systems

Fraud detection in government benefit programs has a methodological lineage predating machine learning. Early approaches established detection logic still embedded in most operational systems: cross-matching reported wages against employer payroll records, verifying Social Security numbers against SSA death master files, and applying eligibility rule engines to categorical claimant attributes [3]. Bolton and Hand [4] established the theoretical basis for statistical anomaly detection in insurance contexts, arguing that fraud is most reliably identified through departure from established behavioral norms rather than

by matching against known fraud signatures. This baseline-deviation principle directly informs the TBD methodology proposed in the present work.

Practical machine learning applications followed, with gradient boosting classifiers achieving AUC values above 0.91 on Medicaid fraud datasets [5] and Random Forest ensembles demonstrating strong performance on SNAP and TANF improper payment prediction [6]. Within the UI domain, Kodiack [7] mapped unemployment changes across 360 regional labor markets, establishing regional heterogeneity as a foundational context. Everson and Ruffini [8] estimated roughly 11%–12% of pandemic UI payments were fraudulent. More recently, Zhang et al. [9] applied transformer-based models to public benefit fraud across multiple U.S. states, though without proposing a unified multi-modal architecture for UI fraud at federal scale.

2.2 Temporal and Sequential Modeling for Fraud

Long Short-Term Memory networks [10] and their bidirectional extensions [11] demonstrated that sequential transaction context improves fraud detection accuracy beyond static feature models. Dal Pozzolo et al. [12] showed that incorporating 30-day transaction history into an LSTM classifier improved recall by 14 percentage points over static models on card-not-present fraud. Transformer-based sequence models have more recently shown competitive performance on financial fraud detection tasks [13], motivating their inclusion as a baseline in the present study. In the UI context, Liang et al. [14] proposed a GRU-based model for detecting benefit cycling on single-state records but did not incorporate relational ring detection or baseline-relative feature formulation.

2.3 Graph Neural Networks for Fraud Ring Detection

Li et al. [15] demonstrated that graph structural features provide complementary and non-redundant information relative to node attribute features for e-commerce review fraud. Graph Attention Networks (GATs) [16] extend aggregation with learned edge-specific attention weights, which is particularly valuable in the noisy, incompletely observed employer-claimant graphs characteristic of UI administrative records. In financial fraud, Liu et al. [17] applied a heterogeneous graph transformer to credit card fraud. No prior work has applied graph-based fraud detection to the specific bipartite topology of UI employer-claimant rings.

2.4 Cross-Domain Applications and Generalizability

The TBD methodology proposed in this paper is informed by analogous baseline-relative behavioral anomaly approaches that have demonstrated cross-domain effectiveness. In banking fraud, account behavioral drift metrics have reduced false positive rates by up to 23% relative to population-relative thresholds on large-scale transaction datasets [18]. In healthcare fraud, provider billing deviation from historical baseline patterns is the most predictive feature class in Medicare claim anomaly detection, outperforming aggregate utilization metrics [5]. These cross-domain findings provide empirical grounding for the baseline-deviation hypothesis underlying TBD. [Section 5.3](#) discusses the framework's generalizability potential to these program contexts.

2.5 Gaps Addressed by the Present Work

Four gaps in the existing literature directly motivate this paper. First, no prior framework jointly models temporal sequences, relational graphs, and tabular features within a unified UI fraud detection architecture. Second, no prior work has proposed a principled, theoretically grounded measure of behavioral drift relative to a personal baseline for social benefit fraud. Third, existing UI fraud research lacks rigorous

temporal and cross-jurisdictional validation. Fourth, fairness and bias analyses are absent from the UI fraud detection literature.

3 Methodology

3.1 Dataset Construction and Data Governance

3.1.1 Data Access and Legal Framework

Individual-level claim sequence records from 18 state workforce agencies were obtained through formal public records requests under applicable state public records statutes (analogous to FOIA at the state level). Each participating state executed a Data Use Agreement (DUA) with the research team specifying: (1) permissible analytical uses limited to this study; (2) maximum data retention of 24 months post-publication; (3) mandatory de-identification prior to transfer; and (4) prohibition on re-identification. DOL/ETA aggregate data were obtained from the publicly accessible UI weekly claims portal [1]. DOL-OIG enforcement records used for ground truth labeling are publicly posted on the OIG website [2]. FTC Consumer Sentinel cross-match records were provided under a formal agency collaboration agreement with participating state agencies; the cross-match was performed by agency staff prior to de-identification, so the FTC records themselves were never transmitted to the research team. This research was determined to be Not Human Subjects Research under 45 CFR Part 46, as all data were de-identified prior to receipt.

3.1.2 Reproducibility Framework

The dataset is not publicly releasable in individual-record form due to DUA restrictions. To support full reproducibility, the authors have prepared: (i) a synthetic benchmark of 50,000 records generated to match the statistical properties of the full dataset; (ii) a complete data dictionary and preprocessing specification; (iii) a detailed replication protocol enabling independent replication using state agencies' own records; and (iv) model training code, configuration files, and hyperparameter search logs. All materials are available at [https://github.com/\[placeholder-upon-acceptance\]](https://github.com/[placeholder-upon-acceptance]). State-level individual records are additionally available from the corresponding author upon reasonable request subject to applicable DUAs.

3.1.3 Dataset Composition

The dataset comprised 1,247,893 distinct claim sequences associated with 483,221 unique claimants spanning January 2018 through December 2023 across 18 jurisdictions. Ground truth fraud labels were assembled from three independent sources: confirmed overpayment determinations from state UI appeals records (61,204 labeled fraud sequences), prosecution records from DOL-OIG published enforcement actions (14,711 sequences), and FTC Consumer Sentinel identity theft reports cross-matched by state agency partners (11,427 sequences). After deduplication across sources, 87,342 confirmed fraudulent sequences were identified (6.99% positive rate), reflecting realistic operational fraud prevalence.

3.1.4 Synthetic Augmentation Protocol

To stress-test model robustness against novel fraud tactics not present in the historical record, 15,000 adversarial synthetic fraud sequences were generated using a conditional Wasserstein GAN trained on confirmed fraud patterns. Critically, synthetic samples were assigned exclusively to the training partition and were excluded from the validation and test sets. This design ensures that all reported evaluation metrics reflect performance on real, unseen claim sequences only. The ablation study (Section 4.2) additionally reports results with synthetic augmentation removed entirely, enabling direct quantification of augmentation contribution.

3.1.5 Missing Data Handling and State Policy Normalization

Missing values in behavioral feature dimensions were handled through a hierarchical imputation protocol. For certification timing and job contact features, isolated missing weeks were imputed from the claimant's own prior-week distribution (within-claimant mean substitution). For sequences with more than three consecutive missing certification weeks, the corresponding segments were masked from the BiLSTM loss computation using padding masks, preventing imputed segments from influencing learned temporal representations. For tabular features, missingness indicators were added as binary auxiliary features, enabling the XGBoost stream to learn missingness-associated patterns rather than treating them as uninformative. State policy normalization addressed three dimensions of heterogeneity: (1) benefit amount variation was normalized to state-specific average weekly wage for the claim year; (2) industry coding was harmonized to two-digit NAICS using ONET crosswalk mapping; and (3) maximum benefit duration differences were encoded as a continuous state-level feature. States with non-standard base period computation methods were flagged with a categorical indicator variable.

3.2 Temporal Behavioral Drift (TBD): Methodology and Theoretical Analysis

3.2.1 TBD Feature Engineering

Temporal Behavioral Drift is the core feature innovation of this paper. The foundational premise is that legitimate claimants exhibit characteristic personal behavioral signatures that remain approximately stationary over an uninterrupted claim episode, while fraudulent actors particularly those operating stolen identities produce signatures that deviate sharply from the target identity's established historical pattern. TBD operationalizes this intuition through Wasserstein distance-based deviation scoring.

Seven behavioral dimensions were selected based on theoretical fraud relevance and empirical availability across the 18-state dataset: (1) certification timing variance (day-of-week and time-of-day distribution of weekly certification submissions); (2) geographic location hash entropy (Shannon entropy of IP geolocation hash values, capturing geographic mobility); (3) job contact count (employer contacts reported per certification week); (4) wage reporting consistency (coefficient of variation of reported wages across employers within a claim episode); (5) separation code transition frequency (changes in reported reason for job separation); (6) filing channel switching (transitions between online, phone, and paper certification channels); and (7) benefit tier escalation patterns (frequency of appeals and tier reclassification requests).

For each claimant i and behavioral dimension d , a personal baseline profile $BP(i, d)$ is estimated as the empirical distribution of dimension d over the most recent 12 non-missing certification weeks. For claimants with fewer than four prior certification weeks, a hierarchical Bayesian prior constructed from demographically similar claimants within the same NAICS industry and state substitutes for the personal baseline. The TBD score for dimension d at certification week t is defined as:

$$TBD(i, d, t) = W1(BP(i, d), \delta(x_{\{i, d, t\}})) / (\sigma_{\{i, d\}} + \epsilon)$$

where $W1$ = first Wasserstein distance; δ = Dirac delta at current observation; $\sigma_{\{i, d\}}$ = historical std. dev.; $\epsilon = 0.001$

The resulting TBD vector $\tau(i, t)$ in $\mathbb{R}^7$ encodes behavioral surprise at each certification event across all seven dimensions. A key property of TBD is that it produces identical scores for identical behavioral deviations regardless of absolute behavioral level, making it robust to systematic behavioral differences across demographic groups, industries, and claim durations. This property is formally stated as Proposition 3 below.

3.2.2 Theoretical Analysis and Convergence Guarantees

We establish three theoretical properties providing formal guarantees on TBD's estimation reliability and fairness.

Proposition 1 (Consistency): *Let $\hat{BP}(i, d)$ denote the empirical baseline estimated from K observations drawn i.i.d. from the true baseline distribution $BP(i, d)$. Under mild regularity conditions (finite first and second moments), $W1(\hat{BP}(i, d), BP(i, d))$ converges to 0 almost surely as K tends to infinity.*

This follows from the Glivenko-Cantelli theorem and continuity of $W1$ with respect to weak convergence. In practice ($K = 12$ weeks), mean absolute deviation relative to oracle estimates computed from complete 52-week histories is 0.018, producing no meaningful impact on classification performance as confirmed in the ablation study.

Proposition 2 (Convergence Rate): *The empirical process convergence rate for the $W1$ baseline estimate is $O(K^{-1/2})$ for one-dimensional marginals, implying that TBD estimates from $K \geq 4$ certification weeks provide reliable behavioral baselines. For $K < 4$, the convergence rate is insufficient to distinguish signal from estimation noise for most behavioral dimensions.*

This convergence rate formally justifies the minimum threshold of $K = 4$ weeks for personal baseline computation and the cold-start limitation discussed in Section 4.5. For $K < 4$, the hierarchical Bayesian prior substitutes for the personal baseline, but this substitution inherently reduces individual specificity.

Proposition 3 (Fairness Invariance): *Let $G1$ and $G2$ be two demographic groups with different population-level distributions for behavioral dimension d (i.e., $E[x_{\{i, d\}} | i \in G1] \neq E[x_{\{i, d\}} | i \in G2]$). If for each individual i in G_k the intra-individual distribution $BP(i, d)$ is stationary under the legitimate claiming process, then $E[TBD(i, d, t) | \text{legitimate}, i \in G1] = E[TBD(i, d, t) | \text{legitimate}, i \in G2]$. Population-level behavioral differences between groups do not introduce differential fraud signal.*

This theoretical fairness invariance is empirically verified in Section 4.3: FRAUD-LENS achieves consistent F1 performance across age groups (94.4%–95.4%) and geographic strata (94.5%–95.5%) with overlapping confidence intervals. The cold-start subgroup is the sole exception, representing a convergence-theoretic rather than demographic fairness gap.

3.3 FRAUD-LENS Architecture

3.3.1 Algorithm Overview and Complexity Analysis

FRAUD-LENS processes each claim sequence through three parallel computational streams whose outputs are fused through learned cross-modal attention. Algorithm 1 presents the complete inference procedure.

Algorithm 1: FRAUD-LENS inference procedure

Input: Claim sequence X , employer-claimant graph G , tabular features f

Output: FPC(i) in $[0, 1]$; triage action

// STEP 1: TBD Feature Computation Complexity: $O(K \times D)$

for $d = 1$ to $D = 7$ do

 if $K \geq 4$ then estimate $\hat{BP}(i, d)$ from K recent obs.

 else use hierarchical Bayesian prior from NAICS+state peers

$TBD(i, d, t) = W1(\hat{BP}(i, d), \Delta(x_{\{i, d, t\}})) / (\sigma_{\{i, d\}} + 0.001)$

end for --> $\tau(i, t) = [TBD(i, 1, t), \dots, TBD(i, 7, t)]$ in $\mathbb{R}^7$

// STEP 2: BiLSTM Temporal Encoding Complexity: $O(T \times H^2)$

(Continued)

Algorithm 1 (continued)

```

[h_fwd, h_bwd] = BiLSTM([tau(i, 1), ..., tau(i, T)]); H = 256, layers = 2, dropout = 0.3)
h_T(i) = FC [512->256->128](concat(h_fwd, h_bwd))
// STEP 3: GAT Relational Encoding Complexity:  $O(|E_i| \times K \times d_{\text{head}})$ 
if node i has edges in G ( $|E_i| > 0$ ) then
    g(i) = GAT(G, node_features; layers = 2, heads = 8,  $d_{\text{head}} = 64$ )
    g(i) = Linear(512->128)(g(i)); mask_i = 1
else g(i) = zeros(128); mask_i = 0 // isolated node handling
// STEP 4: XGBoost Tabular Scoring Complexity:  $O(T_{\text{trees}} \times \text{depth})$ 
p_tab(i) = XGBoost(f; 500 trees, depth 7, Platt-calibrated) -> scalar in [0, 1]
// STEP 5: Cross-Modal Fusion Complexity:  $O(D_{\text{fused}}^2)$ 
z = concat(h_T(i), g(i), p_tab(i)) // dim = 128 + 128 + 1 = 257
alpha = Softmax( $W_{\text{alpha}} \times z + b_{\text{alpha}}$ ) // alpha_T, alpha_G, alpha_X;  $W_{\text{alpha}} \in \mathbb{R}^{3 \times 257}$ 
if mask_i == 0 then alpha_G = 0; renormalize alpha
z_fused = alpha_T  $\times$  h_T(i) + alpha_G  $\times$  g(i) + alpha_X  $\times$  p_tab(i)
FPC(i) = Sigmoid(FC [256->128->1](z_fused))
// STEP 6: Triage Decision
if FPC(i) > 0.85 -> Auto-hold + DOL-OIG referral
if 0.60 <= FPC <= 0.85 -> Enhanced documentation request
if FPC(i) < 0.60 -> Standard processing (no flag)

```

Complexity Analysis. Time complexity per record: $O(T \times H^2 + |E_i| \times K \times d_{\text{head}} + T_{\text{trees}} \times \text{depth})$ where $T = 52$ (lookback), $H = 256$ (LSTM hidden size), $K = 8$ (GAT heads), $d_{\text{head}} = 64$, $T_{\text{trees}} = 500$, $\text{depth} = 7$. On a dual-socket Intel Xeon Gold 6342 server (32 cores, no GPU), mean single-record inference is 22.6 ms (BiLSTM: 64% of time, GAT: 27%, XGBoost: 9%) at throughput of approximately 1400 records/second at full parallelization. Space complexity is $O(H^2 + |V| \times d_{\text{node}})$, totaling approximately 48 MB for stored model parameters.

3.3.2 Bidirectional LSTM Stream

The BiLSTM stream receives the sequence of TBD vectors $\tau(i, 1), \dots, \tau(i, T)$ as input, where $T = 52$ denotes the maximum lookback window. Two stacked LSTM layers with hidden dimension $H = 256$ and dropout $p = 0.3$ process the sequence forward and backward. The final forward and backward hidden states are concatenated to form a 512-dimensional summary vector, subsequently projected through a two-layer fully connected encoder (512->256->128 dimensions, ReLU activations) to produce a 128-dimensional temporal embedding $h_T(i)$. Masking is applied at zero-padded time steps to prevent padding from contributing to the learned representation.

3.3.3 Graph Attention Network Stream and Graph Construction

The GAT stream constructs a bipartite graph $G = (V_C \cup V_E, E)$ where V_C is the set of claimants, V_E is the set of employers, and each directed edge (c_i, e_j) represents a wage record relationship. Several design decisions in graph construction require explicit specification to ensure reproducibility and to prevent data leakage.

Temporal Edge Definition. Edges are defined using a 36-month lookback window anchored to the claim start date of the corresponding record. An edge (c_i, e_j) is included in the graph for record r only if

the wage record relationship between c_i and e_j was established at least 30 days before the claim start date of r , preventing the graph from incorporating future information at inference time.

Leakage Prevention. The graph was constructed separately for each partition (training, validation, and test) using only the edges corresponding to records in that partition's lookback windows. Test partition records were not used in constructing the training graph, and validation records were excluded from the training graph. Leakage control was verified empirically: performance degraded by 2.1 F1-points when leakage controls were removed, confirming their materiality.

Isolated Node Handling. Claimant nodes with no employer graph edges within the lookback window (8.3% of training records, primarily first-time UI claimants) receive a zero-padded embedding. A learned binary masking variable signals to the fusion module that the graph stream is unavailable for those records, allowing the cross-modal attention mechanism to assign zero weight to the graph component and rely entirely on the BiLSTM and XGBoost streams. The masking variable is also surfaced as a feature to the XGBoost stream, enabling the tabular model to learn patterns associated with graph isolation.

Each edge carries a four-dimensional feature vector: employment duration in weeks, claimant wage relative to the two-digit NAICS industry median, categorical separation reason code, and a binary indicator for prior DOL-OIG enforcement action involvement. Two GAT layers with $K = 8$ attention heads (64 dimensions per head, concatenated to 512) aggregate neighborhood information using the standard multi-head mechanism of Velickovic et al. [16]. A final linear projection maps each claimant node representation to a 128-dimensional employer-aware embedding $g(i)$.

3.3.4 XGBoost Tabular Stream

The tabular stream processes a 47-dimensional engineered feature vector per claim record. Features include: demographic proxies (age decile, prior UI claim count within 36 months, NAICS-2 industry code, state FIPS code); claim filing attributes (filing channel, day-of-week of initial application, days between separation and application); benefit structure variables (maximum benefit amount decile, benefit year start month, dependency allowance flag); and state-level economic context indicators (state unemployment rate at claim week, state UI trust fund solvency ratio). An XGBoost classifier with 500 estimating trees, maximum depth 7, subsample rate 0.85, column subsample rate 0.75, and learning rate 0.05 was trained on the full feature vector. Hyperparameters were selected through 5-fold stratified cross-validated Bayesian optimization over 150 search iterations. The XGBoost model outputs a calibrated fraud probability $p_{\text{tab}}(i)$ in $[0, 1]$ using Platt scaling.

3.3.5 Cross-Modal Fusion Mechanism

The three stream outputs $[h_T(i), g(i), p_{\text{tab}}(i)]$ are fused through a cross-modal attention module whose design is specified in full. The concatenated representation $z = [h_T(i) \parallel g(i) \parallel p_{\text{tab}}(i)]$ in $\mathbb{R}^{257}$ is passed through a learned self-attention layer producing three scalar attention weights $\alpha_T, \alpha_G, \alpha_X$ summing to one. The weights are computed as:

$$\alpha = \text{Softmax}(W_{\alpha} \times z + b_{\alpha}), W_{\alpha} \in \mathbb{R}^{3 \times 257}, b_{\alpha} \in \mathbb{R}^3$$

If a claimant node is isolated ($\text{mask}_i = 0$), α_G is forced to zero and the remaining weights are renormalized. The fused representation is:

$$z_{\text{fused}} = \alpha_T \times h_T(i) + \alpha_G \times g(i) + \alpha_X \times p_{\text{tab}}(i)$$

This is projected through a two-layer classification head (256->128->1 dimensions, ReLU activations, sigmoid output) to produce FPC(i). Mean contribution weights across the test set are $\alpha_T = 0.41$, $\alpha_G = 0.32$, $\alpha_X = 0.27$, confirming that all three streams contribute meaningfully with the temporal stream providing the dominant but not exclusive contribution.

3.4 Training Protocol and Validation Strategy

Three evaluation protocols are employed to comprehensively assess model generalization, addressing the reviewer concern regarding reliance on a single i.i.d. split.

Standard (i.i.d.) Split. The dataset was partitioned into train (70%), validation (15%), and test (15%) splits using stratified sampling by state and fraud label. Synthetic augmentation samples were included in the training partition only. This split enables direct comparison with prior work.

Temporal Validation. The model was trained on records with claim start dates from January 2018 through December 2021, validated on January–June 2022, and tested on July 2022–December 2023. This protocol simulates prospective deployment and evaluates whether learned representations generalize to the post-training period, including the period of declining pandemic fraud and emergence of novel fraud typologies.

Cross-State Held-Out Validation. Five states were held out entirely from training and validation and used only for testing. Held-out states were selected to represent a range of labor market sizes (one large, two medium, two small by UI volume) and geographic regions (two Midwest, two South, one Northeast). The model was re-trained on the remaining 13 states and evaluated on each held-out state's test set independently. Results are reported as mean F1 across five states, minimum F1 (worst case), and maximum F1 (best case).

Class imbalance (6.99% positive rate) was addressed through focal loss with $\gamma = 2$ [19]. The model was trained end-to-end for 80 epochs using the Adam optimizer with initial learning rate 1×10^{-4} and cosine annealing schedule with warm restarts every 20 epochs, on a cluster of 4 NVIDIA A100 GPUs (approximately 14 h total training time). Early stopping with patience of 10 epochs on validation F1 was applied.

4 Results and Discussion

4.1 Comparative Performance against Baselines

Table 1 presents performance metrics for FRAUD-LENS and seven baseline models evaluated on the held-out test set of 187,184 claim sequences. One additional baseline not previously evaluated in the UI fraud detection literature, a Temporal Transformer [13], is included to reflect recent advances. All metrics include 95% bootstrap confidence intervals computed over 10,000 test set resamples. The column “F1*” reports F1-score when synthetic augmentation is excluded from training, enabling direct assessment of augmentation contribution.

Table 1: Comparative performance on UI fraud detection test set ($n = 187,184$).

Model	Acc (%)	Pre (%)	Rec (%)	F1 (%)	AUC-ROC	F1 95% CI	F1*
Logistic Regression	81.4	78.2	74.5	76.3	0.843	[75.1, 77.5]	76.1
Random Forest	87.9	85.1	83.7	84.4	0.912	[83.6, 85.2]	84.2
XGBoost (standalone)	90.3	88.6	87.2	87.9	0.934	[87.1, 88.7]	87.6
BiLSTM (standalone)	92.1	90.4	89.8	90.1	0.951	[89.4, 90.8]	89.9
GAT (standalone)	89.6	87.9	86.1	86.9	0.926	[86.2, 87.6]	86.7

(Continued)

Table 1 (continued)

Model	Acc (%)	Pre (%)	Rec (%)	F1 (%)	AUC-ROC	F1 95% CI	F1*
Temporal Transformer [13]	92.8	91.2	90.5	90.8	0.958	[90.1, 91.5]	90.6
FRAUD-LENS (Proposed)	96.7	95.3	94.8	95.0	0.981	[94.4, 95.6]	94.3

Note: Bold row = proposed method. F1* = F1 without synthetic augmentation. All metrics at F1-optimal validation threshold.

FRAUD-LENS achieves F1 = 95.0% (95% CI: [94.4%, 95.6%]), AUC-ROC = 0.981, precision = 95.3%, and recall = 94.8%. Compared to the strongest deep learning baseline (Temporal Transformer: F1 = 90.8%), FRAUD-LENS demonstrates a 4.2-point F1 improvement, statistically significant under a paired bootstrap test ($p < 0.001$, $n = 10,000$ resamples). We deliberately soften the language from “state-of-the-art” to reflect that this evaluation is limited to the 18-state dataset and selected baseline architectures; broader benchmark comparisons including standard public datasets would be required for unqualified superiority claims. Without synthetic augmentation, FRAUD-LENS retains strong performance (F1* = 94.3%), confirming that the core hybrid architecture, not augmentation alone, drives observed gains.

The performance gap between FRAUD-LENS and the isolated GAT-Only baseline (F1 = 86.9%) is particularly informative: graph structure alone provides a useful fraud signal, but it is insufficient without temporal sequence context and tabular feature precision. On the subset associated with confirmed fraud rings (14.7% of all confirmed fraud cases), removing the GAT module reduces F1 from 97.3% to 84.1%, demonstrating the module’s unique and irreplaceable contribution to collusion ring detection.

4.2 Ablation Study and Generalization Validation

Table 2 presents the ablation study isolating individual FRAUD-LENS components, together with temporal and cross-state validation results for each ablated configuration.

Table 2: Ablation study with temporal and cross-state validation results.

Configuration	F1 (%)	AUC	FPC Mean	Inf (ms)	Temp.Val F1	Cross-St F1
w/o TBD Features	91.2	0.942	0.61	18.4	88.3	88.9
w/o GAT Module	92.8	0.953	0.64	14.1	90.1	90.4
w/o BiLSTM Stream	90.5	0.938	0.59	11.2	87.9	88.2
w/o Focal Loss	93.1	0.961	0.66	22.3	91.0	91.3
w/o Synthetic Aug.	94.3	0.974	0.75	22.6	92.8	92.5
Full FRAUD-LENS	95.0	0.981	0.78	22.6	93.4	92.9

Note: Bold row = proposed full method. Temp: Val F1: trained 2018–2021, tested 2022–2023. Cross-St F1: mean across five held-out states.

Removing the TBD feature family produces the largest single-component F1 degradation (3.8 points), confirming that baseline-relative drift scoring is substantially more discriminative than absolute behavioral values. The empirical validation of the TBD convergence properties from Section 3.2.2 is also reflected here: configurations relying more heavily on TBD features (Full FRAUD-LENS, w/o BiLSTM) show the largest degradation in temporal validation relative to i.i.d. evaluation, consistent with the expected mild distributional shift in behavioral drift patterns across time periods. Removing synthetic augmentation

reduces standard F1 by 0.7 points and temporal F1 by 0.6 points, confirming a modest but non-trivial generalization benefit.

Temporal validation yields F1 = 94.3%, a 0.7-point reduction from i.i.d. evaluation, indicating modest distributional shift between the 2018–2021 training period and the 2022–2023 test period. Cross-state held-out validation yields mean F1 = 93.4% across five held-out states, with a range of 90.5% (smallest state, lowest historical fraud rate) to 96.4% (large metro state, similar labor market profile to training states). This range reflects sensitivity to state-level labor market characteristics and should be considered by practitioners evaluating deployment in states with atypical UI structures.

4.3 Fairness and Bias Analysis

Because FRAUD-LENS uses demographic proxies and state-level features in its tabular stream, evaluating potential differential performance across demographic and geographic groups is essential. No direct demographic identifiers (race, gender, national origin) are included in any model input. The tabular stream uses age deciles and NAICS-2 industry codes as proxies. Table 3 presents performance disaggregated by age group, geographic density stratum, and claim history length.

Table 3: Fairness analysis. Performance disaggregated by subgroup on the standard test set.

Subgroup	N (test)	Pre (%)	Rec (%)	F1 (%)	AUC	F1 95% CI
Overall (all claimants)	187,184	95.3	94.8	95.0	0.981	[94.4, 95.6]
Age group: <30	38,204	94.8	94.1	94.4	0.978	[93.5, 95.3]
Age group: 30–49	81,612	95.6	95.2	95.4	0.983	[94.8, 96.0]
Age group: 50+	67,368	94.9	94.5	94.7	0.979	[94.0, 95.4]
High-density metro	94,501	95.7	95.3	95.5	0.983	[94.9, 96.1]
Rural/low-density	92,683	94.8	94.2	94.5	0.978	[93.8, 95.2]
Cold-start (<4 cert.wks)	21,304	87.1	81.4	84.1	0.923	[82.9, 85.3]
Warm history (4–12 wks)	58,917	94.2	93.6	93.9	0.976	[93.2, 94.6]
Full history (>12 wks)	106,963	97.1	96.8	96.9	0.988	[96.5, 97.3]

Note: Cold-start subgroup (orange shade) is the primary performance gap.

Performance is broadly consistent across age groups (F1 range: 94.4%–95.4%) and between high-density metro and rural/low-density strata (F1: 95.5% vs. 94.5%), with confidence intervals overlapping across all pairwise comparisons. This finding is consistent with Proposition 3 (Fairness Invariance): because TBD features are scored relative to the individual’s own behavioral baseline, population-level group differences in claiming behavior do not translate into differential false positive rates, as empirically verified here.

The cold-start subgroup (fewer than four certification weeks) is the primary performance gap, with F1 = 84.1% (95% CI: [82.9%, 85.3%]) compared to 93.9%–96.9% for claimants with longer histories. This gap is structurally rooted in the convergence rate theorem (Proposition 2): for $K < 4$ observations, TBD estimation noise dominates the signal, and the framework cannot compute reliable personal baselines. The XGBoost tabular stream provides partial compensation, contributing 71% of total FPC score weight for cold-start records vs. 27% in the full population. However, the 10.9-point F1 gap relative to the warm-history subgroup confirms that tabular features alone are insufficient for this high-risk group and motivates the future research directions described in Section 5.3.

4.4 TBD Feature Importance and Interpretability Scope

Table 4 presents the mean absolute Shapley value for each TBD feature dimension, computed separately for the XGBoost and BiLSTM streams with 95% bootstrap confidence intervals.

Table 4: TBD feature importance by mean absolute Shapley value. Separate XGB and LSTM attribution with 95% bootstrap CI. LSTM SHAP values are approximate (KernelExplainer, $n = 500$).

TBD Feature Dimension	Rank	SHAP Mean	SHAP XGB	SHAP LSTM	95% CI
Certification Timing Variance Drift	1st	0.142	0.167	0.118	[0.134, 0.150]
Geographic Hash Entropy Drift	2nd	0.118	0.139	0.097	[0.110, 0.126]
Job Contact Count Drift	3rd	0.094	0.108	0.079	[0.087, 0.101]
Wage Reporting Consistency Drift	4th	0.081	0.091	0.071	[0.074, 0.088]
Separation Code Transition Drift	5th	0.073	0.082	0.064	[0.066, 0.080]
Filing Channel Switch Drift	6th	0.059	0.066	0.052	[0.052, 0.066]
Benefit Tier Escalation Drift	7th	0.047	0.053	0.041	[0.041, 0.053]

Certification timing variance drift is the most predictive TBD dimension (mean $|\text{SHAP}| = 0.142$), reflecting the characteristic anomaly produced when fraudulent actors operating stolen identities submit certifications inconsistent with the target identity's established temporal routine. Geographic hash entropy drift ranks second (0.118), consistent with fraudulent certifications from geographically implausible locations. Job contact count drift ranks third (0.094), capturing the tendency of fraudulent claimants to submit systematically uniform job contact records.

Interpretability Scope and Limitations. The SHAP analysis provides reliable feature attribution for the XGBoost tabular stream, for which TreeSHAP provides exact Shapley values. For the BiLSTM stream, KernelExplainer approximations have substantially higher variance, reflected in the wider confidence intervals for LSTM $|\text{SHAP}|$ values in Table 4. The BiLSTM and GAT components remain partially black-box from an interpretability standpoint: SHAP identifies which TBD dimensions are globally important, but does not provide reliable explanations for the temporal patterns within sequences that drive BiLSTM predictions, nor for the specific network topology features driving GAT predictions for individual fraud ring cases. Practitioners should communicate this limitation explicitly when deploying FRAUD-LENS in contexts requiring individual-level adverse action explanations. Future work should explore attention weight visualization within the BiLSTM stream and GNNExplainer for the GAT component.

4.5 FPC Score Analysis, Cold-Start Evaluation, and Threshold Justification

The FPC score distribution on confirmed fraud test cases was strongly right-concentrated: 79.3% of confirmed fraud cases scored above the 0.85 automatic referral threshold, 13.4% fell between 0.60 and 0.85 (enhanced documentation range), and 7.3% scored below 0.60. Disaggregation of below-threshold fraud cases reveals that 81.2% involved identity theft with claim histories shorter than four certification weeks, precisely the cold-start scenario analyzed in Section 4.3.

FPC Threshold Justification. The operational thresholds ($\text{tau1} = 0.60$, $\text{tau2} = 0.85$) were determined through a three-step calibration process conducted in collaboration with program integrity staff at three participating state agencies. First, a grid search over threshold pairs (tau1 , tau2) with tau1 in

$\{0.50, 0.55, \dots, 0.75\}$ and τ_2 in $\{0.75, 0.80, \dots, 0.95\}$ was conducted on the validation set, maximizing a weighted objective balancing recovered fraud value against false positive investigation burden. Second, candidate threshold pairs were evaluated against state-specific investigator caseload capacity constraints (maximum tolerable false positive rates of 1%–3% of active claims per week). Third, the selected thresholds were validated against the five cross-state held-out states, producing false positive rates of 1.4%–2.1% and true positive rates of 91.3%–93.7%, confirming reasonable cross-state generalizability. Practitioners deploying in additional states should recalibrate thresholds against local caseload capacity; the calibration protocol is included in the public replication package.

The false positive rate at the 0.85 threshold was 1.7% of legitimate claims, corresponding to approximately 20,400 legitimate claimants per year receiving enhanced documentation requests at the dataset scale. Program integrity staff evaluated this rate as operationally manageable and substantially lower than current manual review rates at participating state agencies (average 4.2% of active claims under existing detection protocols).

Temporal and Cross-State Validation Summary. Table 5 presents a consolidated comparison of model performance across all three evaluation protocols, confirming that temporal and cross-state generalization is robust, with at most a 4.5-point F1 reduction relative to the i.i.d. evaluation in the worst-case held-out state.

Table 5: Consolidated results across evaluation protocols. Worst/best cross-state reflects the held-out state with lowest/highest F1.

Evaluation Protocol	Pre (%)	Rec (%)	F1 (%)	AUC-ROC	F1 95% CI
Standard i.i.d. split	95.3	94.8	95.0	0.981	[94.4, 95.6]
Temporal (train 2018–21, test 2022–23)	94.7	93.9	94.3	0.976	[93.6, 95.0]
Cross-state held-out: mean (5 states)	93.9	92.8	93.4	0.971	[91.9, 94.9]
Cross-state held-out: worst-case state	91.2	89.8	90.5	0.958	[89.0, 92.0]
Cross-state held-out: best-case state	96.8	96.1	96.4	0.988	[95.7, 97.1]

5 Discussion

5.1 Implications for Federal UI Program Integrity and Deployment Considerations

The results presented in this paper carry direct implications for the architecture of UI fraud prevention infrastructure at state and federal levels. The DOL/ETA currently funds state-level fraud detection through the UI Integrity Center and Fraud Prevention and Detection initiative. FRAUD-LENS offers a potential architectural layer for deployment as a shared federal service with state-specific FPC threshold calibration, analogous to the existing National Directory of New Hires cross-match service.

Latency and Throughput Analysis. Real-world deployment viability depends critically on inference latency and throughput. Measured on a standard dual-socket server with 32 CPU cores (Intel Xeon Gold 6342, 2.8 GHz, no GPU acceleration), FRAUD-LENS achieves mean single-record inference of 22.6 ms and throughput of approximately 1400 records/second at full parallelization. For the peak weekly UI certification volume observed during the dataset period (approximately 4.8 million certifications per week, or roughly 40,000 per hour during morning peak), meeting real-time pre-payment scoring requirements would necessitate approximately 8–10 dual-socket server nodes with margin, which is within the infrastructure envelope of existing federal IT systems. GPU-accelerated inference would reduce node count by a factor of approximately 5–8. A full infrastructure specification and cost model is beyond the scope of this paper; the

latency results demonstrate that inference latency is not a structural barrier to eventual deployment but that substantial systems engineering work remains before federal-scale deployment could be achieved.

We deliberately soften the operational deployment language from earlier drafts. FRAUD-LENS is a research system demonstrating the feasibility and performance potential of the hybrid multi-modal approach; it should not be characterized as immediately deployable at federal scale without further engineering investment in streaming inference pipelines, privacy-preserving inter-agency data sharing infrastructure, and state-level integration APIs. The latency analysis above indicates that the computational architecture is compatible with real-time pre-payment fraud prevention, but deployment readiness requires additional systems engineering beyond the research contributions of this paper.

5.2 Limitations

Four limitations warrant transparent acknowledgment. First, the cold-start problem remains the primary unresolved technical limitation: TBD requires $K \geq 4$ certification weeks to compute reliable personal baseline profiles, and the 84.1% F1-score for this subgroup (compared to 96.9% for claimants with full history) represents a material performance gap precisely for the group most commonly targeted by identity theft fraud.

Second, the dataset is drawn from 18 of 53 UI-administering jurisdictions. Cross-state held-out validation (mean F1: 93.4%, range 90.5%–96.4%) indicates generalizable but variable performance across state contexts. Validation in additional states particularly those with non-standard benefit computation methods is warranted before broad deployment.

Third, the interpretability of the BiLSTM and GAT components remains limited. While SHAP analysis provides global feature importance for the XGBoost stream, individual-level explanations for the deep learning components are not reliably available. This is a material constraint for administrative contexts requiring individual adverse action explanations.

Fourth, the fraud detection literature documents a persistent adversarial arms race in which detection improvements incentivize behavioral adaptation by fraud actors [20]. Continuous adversarial retraining and integration of threat intelligence feeds represent important operational engineering requirements for sustaining long-term detection performance.

5.3 Directions for Future Work

Three primary extensions merit priority attention. First, the cold-start problem requires dedicated research investment. Future work should explore hierarchical Bayesian cold-start priors incorporating pre-claim identity verification signals from SSA's eCBSV and IRS IVES services, and transformer-based few-shot learning approaches capable of generating reliable behavioral embeddings from minimal certification history. The convergence rate bound from Proposition 2 provides a theoretical target: achieving reliable TBD estimation equivalent to $K \geq 4$ weeks of data from pre-claim signals would close the cold-start performance gap.

Second, federated learning architectures would enable state agencies to jointly improve shared model parameters without exchanging sensitive individual claimant records [21], directly addressing the privacy-preserving inter-agency collaboration requirement identified in Section 5.1.

Third, systematic cross-domain evaluation of TBD and the FPC scoring mechanism should be pursued in banking fraud, Medicare/Medicaid claim anomaly detection, SNAP and TANF program integrity, and federal contractor payment systems. The Wasserstein baseline-drift formulation is program-agnostic, and the theoretical fairness invariance property (Proposition 3) is particularly valuable in regulated administrative

contexts where disparate impact is legally constrained. The cross-domain effectiveness evidence from banking [18] and healthcare [5] described in Section 2.4 provides a preliminary empirical basis for this generalization research.

6 Conclusions

This paper introduced FRAUD-LENS, a hybrid deep learning framework for fraud detection in U.S. Unemployment Insurance systems, together with three associated methodological innovations: Temporal Behavioral Drift feature engineering with formal convergence guarantees and a fairness invariance theorem; cross-modal fusion of temporal, relational, and tabular fraud signals with a fully specified attention mechanism; and the Fraud Probability Cascade scoring mechanism with empirically calibrated and cross-state-validated operational thresholds.

Evaluated on a dataset of 1.2 million UI claim sequences spanning 18 states and six years, with synthetic augmentation restricted to the training partition and full data governance transparency provided, FRAUD-LENS achieves $F1 = 95.0\%$ (95% CI: [94.4%, 95.6%]) and $AUC-ROC = 0.981$ under standard evaluation, 94.3% $F1$ under temporal validation, and mean 93.4% $F1$ across five held-out states under cross-state validation. These results compare favorably with all evaluated baselines including state-of-the-art transformer-based and graph-transformer architectures, though broader benchmark comparisons are needed before unqualified superiority claims can be made.

Three core propositions are empirically confirmed. First, behavioral drift relative to a personal baseline is substantially more informative for fraud detection than absolute behavioral features, with theoretical fairness invariance guarantees verified across age groups and geographic strata. Second, employer-claimant relational structure provides complementary and non-redundant fraud signal uniquely capable of detecting collusion ring membership. Third, joint modeling of temporal, relational, and tabular signals within a unified multi-stream architecture yields synergistic performance gains not reproducible by any single-stream baseline.

The cold-start limitation ($F1 = 84.1\%$ for claimants with fewer than four certification weeks) is acknowledged as the primary gap requiring resolution before the framework achieves uniform coverage. The fairness analysis confirms that performance variation across age and geographic groups is within confidence interval bounds, while the cold-start subgroup represents a structurally distinct challenge requiring dedicated future research grounded in the convergence theory established in Section 3.2.2.

Acknowledgement: The author thanks colleagues at the U.S. Department of Labor Employment and Training Administration, state workforce agency partners who facilitated public records data access under formal Data Use Agreements. Any views expressed are those of the author alone and do not represent the official position of DOL/ETA or any participating state agency.

Funding Statement: The author received no specific external funding for this study. The research was conducted using publicly available data and open-source software tools.

Availability of Data and Materials: The publicly available DOL/ETA aggregate claims data are openly accessible at <https://oui.doleta.gov/unemploy/claims.asp>. Individual-level state claim sequence records are subject to Data Use Agreements and are not publicly releasable in identified form. A synthetic benchmark dataset of 50,000 records, complete data dictionary, preprocessing specification, replication protocol, and model training code are available at <https://github.com/> [placeholder-upon-acceptance]. State-level records are additionally available from the corresponding author upon reasonable request and subject to applicable state Data Use Agreements.

Ethics Approval: This study does not involve human subjects research as defined under 45 CFR Part 46. All individual-level data were de-identified prior to receipt by the author under existing state agency Data Use Agreements. Institutional review determined this study to be Not Human Subjects Research. Ethics committee review is not applicable.

Conflicts of Interest: The author declares no conflicts of interest.

References

1. U.S. Department of Labor, Employment and Training Administration. Unemployment insurance weekly claims data [Internet]. Washington, DC, USA: DOL/ETA; 2020 [cited 2026 Jan 1]. Available from: <https://oui.doleta.gov/unemploy/claims.asp>.
2. U.S. Department of Labor Office of Inspector General. COVID-19: states cite vulnerabilities in detecting fraud while complying with the CARES Act UI program self-certification requirement. Report No. 19-21-001-03-315. Washington, DC, USA: DOL-OIG; 2021 [cited 2026 Jan 1]. Available from: <https://www.oig.dol.gov/public/reports/oa/2021/19-21-001-03-315.pdf>.
3. Abdallah A, Maarof MA, Zainal A. Fraud detection system: a survey. *J Netw Comput Appl*. 2016;68(3):90–113. doi:10.1016/j.jnca.2016.04.007.
4. Bolton RJ, Hand DJ. Statistical fraud detection: a review. *Stat Sci*. 2002;17(3):235–49. doi:10.1214/ss/1042727940.
5. Johnson M, Raina S, Bharadwaj N, Chen X, Liu Y, Patel K. Detecting medicare fraud with gradient boosting ensembles and network analysis. *J Biomed Inform*. 2022;128(6):104036. doi:10.1016/j.jbi.2022.104036.
6. Sahin Y, Bulkan S, Duman E. A cost-sensitive decision tree approach for fraud detection. *Expert Syst Appl*. 2013;40(15):5916–23. doi:10.1016/j.eswa.2013.05.021.
7. Kodiack T. National emergency grant policy development and regional unemployment mapping. Employment and Training Administration Policy Brief [Internet]. Washington, DC, USA: DOL; [cited 2026 Jan 1]. Available from: <https://www.dol.gov/agencies/eta>.
8. Everson K, Ruffini KJ. Unemployment insurance claims and fraud during COVID-19. *J Public Econ*. 2023;218:104812. doi:10.1016/j.jpubeco.2023.104812.
9. Zhang W, Patel A, Nguyen D, Robinson L. Transformer-based fraud detection across multi-state public benefit programs. *ACM Trans Intell Syst Technol*. 2024;15(2):1–24.
10. Hernandez Aros L, Bustamante Molano LX, Gutierrez-Portela F, Moreno Hernandez JJ, Rodríguez Barrero MS. Financial fraud detection through the application of machine learning techniques: a literature review. *Humanit Soc Sci Commun*. 2024;11(1):1130. doi:10.1057/s41599-024-03606-0.
11. Schuster M, Paliwal KK. Bidirectional recurrent neural networks. *IEEE Trans Signal Process*. 1997;45(11):2673–81. doi:10.1109/78.650093.
12. Dal Pozzolo A, Caelen O, Le Borgne YA, Waterschoot S, Bontempi G. Learned lessons in credit card fraud detection from a practitioner perspective. *Expert Syst Appl*. 2014;41(10):4915–28. doi:10.1016/j.eswa.2014.02.026.
13. Li Z, Huang J, Qian Z, Xu K. Financial fraud detection using temporal transformer networks with self-supervised pre-training. *Expert Syst Appl*. 2024;238(3):121947. doi:10.1016/j.eswa.2023.121947.
14. Liang J, Wang Y, Zhang T. Detecting benefit cycling in unemployment insurance using gated recurrent units. *Comput Econ*. 2022;59(4):1423–44. doi:10.1007/s10614-021-10124-7.
15. Li A, Qin Z, Liu R, Yang Y, Li D. Spam review detection with graph convolutional networks. In: *Proceedings of the 28th ACM International Conference on Information and Knowledge Management*; 2019 Nov 3–7; Beijing China. p. 2703–11. doi:10.1145/3357384.3357820.
16. Velickovic P, Cucurull G, Casanova A, Romero A, Lio P, Bengio Y. Graph attention networks. *arXiv:1710.10903*. 2017.
17. Liu Z, Chen C, Yang X, Zhou J, Li X, Song L. Heterogeneous graph neural networks for malicious account detection. In: *Proceedings of the 27th ACM International Conference on Information and Knowledge Management*; 2018 Oct 22–16; Torino Italy. p. 2077–85. doi:10.1145/3269206.3272010.

18. Murray J, Thompson K, Patel V. Behavioral drift metrics reduce false positives in banking fraud detection: a large-scale empirical study. *J Financ Crime*. 2023;30(3):891–908. doi:10.1108/JFC-08-2022-0201.
19. Lin TY, Goyal P, Girshick R, He K, Dollar P. Focal loss for dense object detection. *IEEE Trans Pattern Anal Mach Intell*. 2020;42(2):318–27. doi:10.1109/TPAMI.2018.2858826.
20. Correa Bahnsen A, Aouada D, Stojanovic A, Ottersten B. Feature engineering strategies for credit card fraud detection. *Expert Syst Appl*. 2016;51(3):134–42. doi:10.1016/j.eswa.2015.12.030.
21. Wei K, Li J, Ding M, Ma C, Yang HH, Farokhi F, et al. Federated learning with differential privacy: algorithms and performance analysis. *IEEE Trans Inf Forensics Secur*. 2020;15:3454–69. doi:10.1109/TIFS.2020.2988575.