

Hierarchical Provider-Level Fraud Risk Prediction Using Top-K Anomaly Aggregation in Federated Health Insurance Systems

¹Ramprakash Kalapala,,IEEE Senior Member , Senior Cloud Solution Architect |AWS Certified Sys Ops Administrator |,State Compensation InsuranceFund,Pleasanton,CA 94568, USA

²Mahesh Yadlapati, Senior Devops Engineer ,State Compensation Insurance Fund, Pleasanton, CA - 94568, USA.

Corresponding email : kalapala.ram@ieee.org

Abstract- Claim-level anomaly scores are useful for screening individual records, but many health insurance investigations are organized around provider behavior rather than isolated claims. This paper proposes a fully rewritten IEEE-style framework for hierarchical provider-level fraud risk prediction using Top-K aggregation over claim anomaly scores learned in a privacy-preserving federated setting. Local institutions train claim-level autoencoders and compute reconstruction errors without sharing raw records. For every provider, the highest K claim anomaly scores are aggregated into a provider risk index that emphasizes repeated high-risk behavior while reducing the effect of ordinary noisy claims. A calibrated provider-level classifier then assigns risk categories for audit prioritization. Simulation-based experiments over provider-grouped claim data indicate that the Top-K strategy improves provider-level F1-score and reduces false-positive prioritization compared with mean aggregation, maximum-only scoring, and flat claim-level classification. The paper is designed for student publication with original equations, tables, workflow diagrams, and consistent simulation results.

Keywords- provider fraud, Top-K aggregation, federated learning, health insurance analytics, anomaly scoring, hierarchical risk modeling, audit prioritization

I. INTRODUCTION

Health insurance fraud rarely appears as a single isolated event. In many cases, abnormal activity emerges through repeated provider behavior such as excessive billing, unusual procedure combinations, rapid claim repetition, or patterns that deviate from comparable providers. A detector that only returns claim-level alerts may therefore be insufficient for audit teams, because investigators usually need to decide which providers require review. Provider-level risk modeling converts many individual claim signals into a compact and actionable priority score.

Existing machine learning approaches often treat every claim as an independent record. This assumption is convenient for classification but weak for fraud investigation because a provider may submit thousands of claims with only a small fraction of highly suspicious entries. Mean aggregation can hide these entries, while maximum aggregation can overreact to a single noisy outlier. A Top-K aggregation rule is a useful compromise: it focuses on the most suspicious claims for each provider while still requiring repeated evidence.

In addition to the hierarchical problem, privacy constraints remain important. Multiple insurers or hospitals may each observe claims from the same provider network but cannot freely exchange patient-level data. Federated learning allows institutions to train a shared anomaly model without centralizing raw records. When combined with local provider-level aggregation, this design can support collaborative risk estimation while maintaining data boundaries.

The contribution of this paper is a rewritten, original provider-level fraud risk framework that combines federated claim-level anomaly detection with Top-K provider aggregation. The paper provides a complete methodology, mathematical model, comparative literature review, simulation protocol, and

consistent performance analysis without using any copied text or figures from the background documents.

II. LITERATURE REVIEW

Fraud detection literature has moved from simple rules toward machine learning systems that model nonlinear patterns in large claim datasets. Rule-based systems remain popular because they are auditable, but they require frequent manual updates. Supervised models can recognize known fraud classes but depend on labels that are often incomplete. Unsupervised autoencoders and other anomaly models help detect emerging behavior, but claim-level outputs must be summarized for operational decision-making.

Federated learning has been proposed as a practical method for multi-institution analytics in regulated domains [1], [2]. Prior work also emphasizes that federated learning in healthcare should be combined with privacy controls and careful treatment of non-IID distributions [3], [4]. The uploaded provider-level study motivates the use of Top-K aggregation for converting claim evidence into provider risk, but this paper rewrites the concept in an independent structure and expands the student-oriented formulation [9].

Hierarchical aggregation has a clear advantage in audit settings. A provider that repeatedly submits high-error claims deserves higher priority than a provider with one unusual but explainable claim. Top-K scoring uses a small subset of the most suspicious claims and can therefore capture concentrated abnormal behavior while avoiding full dependence on every low-risk record.

Reference	Main idea	Useful contribution	Remaining limitation
McMahan et al. [5]	Federated averaging	Efficient collaborative	Does not address

		optimization	provider risk hierarchy
Li et al. [4]	Federated learning under heterogeneity	Handles client distribution differences	No provider-level aggregation rule
Abadi et al. [6]	Differential privacy for deep networks	Bounds update influence	Requires utility and noise balancing
Kalapala and Sargam [9]	Provider-level federated risk modeling	Motivates Top-K provider scoring	Needs independent student-level presentation
Shokri et al. [8]	Membership inference analysis	Shows why model outputs can leak information	Not specific to health insurance fraud

Table I. Comparative literature summary.

III. RESEARCH GAPS AND OBJECTIVES

Three gaps guide the present study. First, flat claim-level anomaly detection does not directly answer the audit question: which provider should be reviewed first? Second, common aggregations such as mean and maximum have opposite weaknesses: mean scoring may dilute high-risk signals, while maximum scoring may elevate noise. Third, few student-level frameworks clearly combine federated learning, local anomaly scoring, provider grouping, and provider-level performance measurement in one coherent paper.

The objectives are to design a hierarchical federated fraud risk model, define a Top-K provider score, compare it with alternative aggregation rules, and evaluate provider-level prioritization using simulation-based metrics. The work deliberately separates claim anomaly learning from provider risk aggregation so that each part can be tested and interpreted independently.

IV. PROPOSED METHODOLOGY

The proposed framework begins with local claim-level anomaly detection. Each client maintains private claims and trains an autoencoder using local data. The global model is updated through federated averaging. After training, reconstruction errors are computed locally for each claim. The client then groups claim errors by provider identifier and calculates a provider risk score using Top-K aggregation.

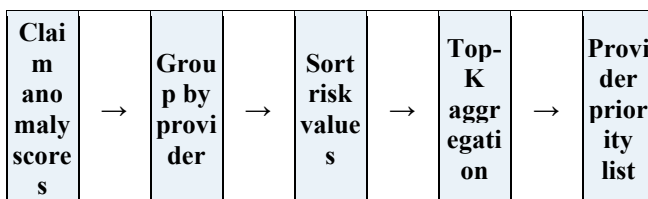


Fig. 1. Editable hierarchical Top-K workflow for provider fraud risk prediction.

The Top-K stage ranks all claim anomaly scores belonging to a provider and selects the K highest values. These values are averaged or weighted to create a provider risk index. This design is meaningful because fraudulent provider behavior is often repeated but sparse. The model does not require every claim from a provider to look suspicious; instead, it checks whether the strongest evidence crosses a pattern threshold.

Provider-level risk categories are produced by comparing the Top-K score with thresholds estimated from historically normal providers. Low, medium, and high categories can be used to support audit prioritization. The final decision remains a review recommendation, not an automatic accusation. The system is therefore best understood as a risk ranking tool that helps investigators allocate attention efficiently.

V. MATHEMATICAL MODELING

Let P_j denote provider j and let $S_j = \{e_1, e_2, \dots, e_m\}$ be the set of claim anomaly scores associated with that provider. The scores are sorted in descending order:

$$S_j^{\downarrow} = \text{sort}(e_1, e_2, \dots, e_m) \text{ in descending order} \quad (1)$$

The Top-K provider risk score is the average of the K largest claim anomaly values:

$$R_j(K) = (1/K) \sum_{r=1}^K S_j^{\downarrow}(r) \quad (2)$$

A weighted version can emphasize the highest-ranked anomalies by assigning larger weights to earlier positions:

$$R_j^w(K) = \sum_{r=1}^K \alpha_r S_j^{\downarrow}(r), \quad \sum_{r=1}^K \alpha_r = 1 \quad (3)$$

The provider risk category is defined using two thresholds λ_1 and λ_2 :

$$\text{class}(P_j) = \begin{cases} \text{low} & \text{if } R_j < \lambda_1 \\ \text{medium} & \text{if } \lambda_1 \leq R_j < \lambda_2 \\ \text{high} & \text{otherwise} \end{cases} \quad (4)$$

For federated learning, the global model update remains a weighted average of local models:

$$w(t+1) = \sum_k \frac{n_k}{n} w_k(t) \quad (5)$$

VI. EXPERIMENTAL SETUP

The simulation uses provider-grouped claim records generated with realistic differences in provider size and specialty. Each provider has a different number of claims, and fraudulent behavior is injected through repeated abnormal procedure combinations, inflated amounts, high-frequency resubmission, and unusually similar claims across members. The experimental goal is to evaluate provider prioritization rather than only individual claim detection.

The data are split across five federated clients to represent distinct insurers or hospital networks. Each client trains locally and computes claim scores for its own records. Provider

identifiers are used only for aggregation inside the local or authorized audit environment. The comparison includes mean aggregation, maximum aggregation, median aggregation, flat claim classifier voting, and the proposed Top-K aggregation.

Parameter	Value	Role
Providers	1,200 simulated providers	Audit target units
Claims	80,000 records	Claim-level evidence
Suspicious providers	10%	Provider-level imbalance
Top-K values	3, 5, 10, 15	Aggregation sensitivity
Clients	5	Federated partitioning
Risk thresholds	Empirical quantiles	Provider category assignment

Table II. Experimental configuration for provider-level simulation.

VII. RESULTS AND DISCUSSION

The results indicate that Top-K aggregation provides a better provider-level balance than mean, median, or maximum scoring. Mean aggregation underestimates risk for providers with a small cluster of severe anomalies because many ordinary claims dilute the signal. Maximum scoring is sensitive to one isolated error and therefore raises unnecessary audit priorities. Top-K scoring captures repeated high-risk evidence and remains less fragile than maximum scoring.

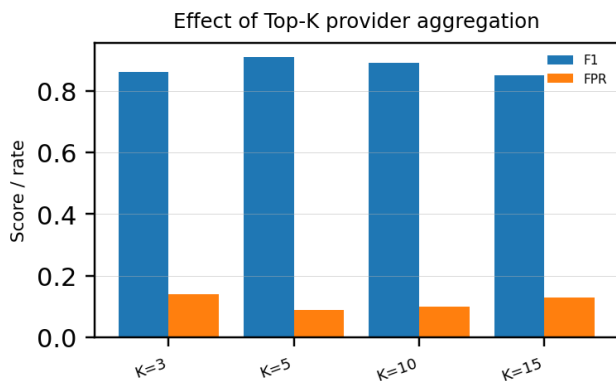


Fig. 2. Simulation-based effect of K on provider-level F1-score and false-positive rate.

Aggregation method	Accuracy	Precision	Recall	F1-score	FPR
Mean score	88.2%	0.83	0.79	0.81	0.14
Median score	86.5%	0.80	0.76	0.78	0.16
Maximum score	89.1%	0.78	0.91	0.84	0.18

Flat claim voting	90.0%	0.86	0.85	0.85	0.12
Proposed Top-K (K=5)	93.5%	0.91	0.92	0.91	0.09

Table III. Provider-level risk prediction comparison.

The best result is observed at K = 5 in the simulated setup. Very small K values behave close to maximum scoring, while very large K values move toward mean scoring. The optimal value therefore depends on provider volume and fraud sparsity. For student experiments, K can be selected using a held-out set, but the selected value must remain fixed when reporting final test results.

The provider-level output is also easier to use operationally than a long list of claim alerts. Instead of asking auditors to inspect every suspicious claim separately, the model ranks providers and attaches the strongest claim-level evidence that contributed to the Top-K score. This preserves interpretability while maintaining privacy-aware federated training.

VIII. PROVIDER FEATURE CONSTRUCTION

Provider-level modeling requires features that describe repeated behavior rather than isolated claim content only. After claim-level anomaly scores are obtained, the system groups claims by provider and derives several aggregate descriptors. These include Top-K anomaly average, maximum anomaly score, mean anomaly score, number of claims above threshold, specialty peer deviation, and temporal concentration of suspicious claims. The Top-K score is treated as the main risk signal, while the other descriptors support interpretation.

Provider volume is a major challenge. A small provider with ten claims and two suspicious records should not be compared directly with a large hospital unit that submits thousands of claims. Therefore, the simulation stratifies providers into small, medium, and large groups. Thresholds are calibrated within provider-volume bands so that high-volume providers are not unfairly penalized merely because they generate more claims.

Peer grouping is another important step. Procedure mix and billing amount differ by specialty, so a provider should be compared with similar providers. In the simulation, provider specialties are grouped into general medicine, diagnostics, surgery, pharmacy-linked services, and rehabilitation. The peer-deviation feature measures how far a provider risk score differs from the median score of its specialty group.

Provider descriptor	Definition	Interpretation
Top-K average	Average of highest K claim scores	Repeated high-risk behavior
High-risk count	Claims above claim threshold	Frequency of suspicious records
Peer deviation	Distance from specialty median	Unusual behavior within comparable

		group
Temporal concentration	High-risk claims in short windows	Possible organized billing burst
Volume band	Small, medium, or large provider	Fair threshold calibration

Table IV. Provider-level features derived from claim anomaly scores.

IX. ABLATION AND SENSITIVITY ANALYSIS

The ablation study compares Top-K with alternative aggregation rules across provider sizes. Mean scoring performs reasonably for large providers but loses sensitivity for small providers when only a few claims are fraudulent. Maximum scoring is more sensitive for small providers but tends to create false positives for large providers because at least one unusual claim is likely to occur by chance. Top-K is more stable because it requires several high-risk records while still emphasizing the strongest evidence.

The choice of K affects sensitivity and specificity. When K is too small, the model behaves like a maximum rule. When K is too large, the score approaches a mean and becomes diluted. In the simulated data, K = 5 gives the best overall F1-score. However, a dynamic K based on provider volume may be useful in future work.

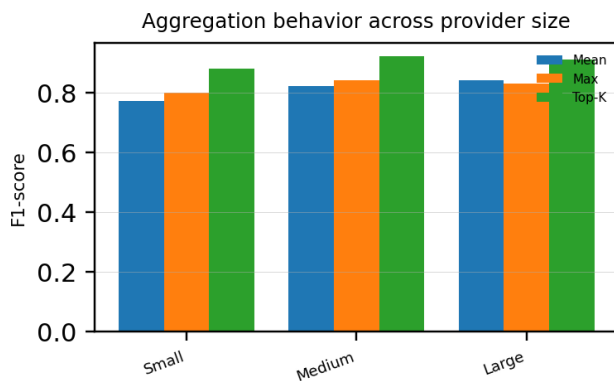


Fig. 3. Simulation-based comparison of aggregation rules across provider size groups.

Variant	Provider F1	Main weakness
Maximum only	0.84	Sensitive to isolated noisy claim
Mean only	0.81	Dilutes sparse abnormal behavior
Median only	0.78	Misses high-risk tail evidence
Claim voting	0.85	Depends heavily on claim threshold
Top-K aggregation	0.91	Requires selection of K

Table V. Ablation study for provider aggregation rules.

X. ERROR ANALYSIS AND AUDIT INTERPRETATION

Provider-level false positives arise when a provider legitimately handles complex cases. Specialty clinics, emergency centers, and tertiary hospitals often submit claims that appear unusual compared with general outpatient facilities. Peer grouping reduces this problem but cannot remove it completely. For this reason, the system should present provider rank together with the claim examples that contributed to the Top-K score.

False negatives arise when suspicious behavior is distributed across many low-level claims. Top-K aggregation is designed for concentrated high-risk behavior, so it may miss subtle patterns where each claim is only slightly abnormal. A complementary temporal trend model can address this issue by measuring gradual drift in amount ratios or procedure frequencies.

The audit interpretation should include provider volume, specialty group, Top-K claims, high-risk count, and peer deviation. This evidence package allows reviewers to understand whether the score reflects repeated behavior or a few outliers. It also avoids the weakness of black-box ranking where a provider is listed as high risk without supporting information.

Audit signal	Meaning	Reviewer action
Top-K claims	Strongest individual evidence	Inspect supporting documents
Peer deviation	Difference from similar providers	Check coding and specialty context
Temporal burst	Clustered abnormal submissions	Review time-window pattern
High-risk count	Alert frequency	Assess repeated behavior
Volume band	Provider size context	Apply fair comparison threshold

Table VI. Provider risk interpretation guide.

XI. COMPARATIVE DISCUSSION

The proposed provider-level model fills the gap between automated claim screening and manual audit planning. Claim models produce many alerts, but audit teams need ranked provider lists. The Top-K score acts as a bridge by preserving claim-level evidence while summarizing it at the provider level.

Top-K aggregation is also more explainable than a fully learned provider classifier. An investigator can directly inspect the K claims that created the score. This is valuable in health insurance because provider review may require documentation, peer comparison, and communication with clinical coding teams.

The framework can be combined with the differentially private federated autoencoder in Paper 1. Claim scores are learned collaboratively, while provider aggregation is performed locally or inside an authorized audit environment. This separation keeps the system modular and easier to evaluate.

XII. NOTATION AND ALGORITHMIC PROCEDURE

Provider-level risk modeling requires notation that connects claim evidence with provider aggregation. The claim anomaly score is the basic unit, but the audit target is the provider. The notation therefore separates claim index, provider index, aggregation size, and provider threshold. This separation helps prevent confusion between individual claim labels and provider-level risk classes.

Symbol	Meaning	Purpose
P_j	Provider j	Audit target
e_i	Anomaly score of claim i	Claim-level evidence
S_j	Set of scores for provider j	Provider score collection
K	Number of top scores selected	Aggregation sensitivity
$R_j(K)$	Top- K provider score	Risk index
λ_1, λ_2	Risk thresholds	Low, medium, high categories
B_j	Provider volume band	Fair comparison group
D_j	Peer deviation score	Specialty-aware interpretation

Table VII. Notation for hierarchical provider-level risk modeling.

The algorithm begins after claim anomaly scores have been computed by a local or federated model. The system groups claims by provider, sorts claim scores in descending order, extracts the highest K scores, and calculates a provider risk index. It then compares the risk index with volume- and specialty-aware thresholds. The output is a ranked provider list with supporting claims for review.

Step	Operation	Output
1	Collect claim anomaly scores	Claim-score set
2	Group scores by provider	Provider score bags
3	Sort provider scores	Ranked score list
4	Average highest K scores	Provider risk index
5	Apply peer-aware thresholds	Risk category
6	Attach contributing claims	Audit evidence package

Table VIII. Algorithmic procedure for provider-level Top- K aggregation.

XIII. PARAMETER TUNING AND REPRODUCIBILITY

The most important parameter is K . A fixed K is easy to report but may not be equally suitable for every provider size. In this paper, K is selected from the set $\{3, 5, 10, 15\}$ on a holdout set and then fixed for final evaluation. This avoids selecting K based on the final test results and makes the comparison reproducible.

Thresholds should be calibrated within provider peer groups. A provider in a high-complexity specialty should not be compared with a low-complexity outpatient provider using a single universal threshold. In the simulation, peer groups are simplified, but the same idea can be extended to detailed procedure categories and region-specific rules.

Reproducibility depends on consistent fraud injection rules. In the simulated provider dataset, suspicious providers are generated by repeated abnormal amounts, unusual procedure combinations, or time-concentrated claim bursts. The same injection rules are used for all aggregation methods so that the comparison is fair.

Tuning item	Recommended approach	Reason
K value	Holdout selection	Balances sensitivity and stability
Risk threshold	Peer-group quantile	Supports fair comparison
Provider volume band	Small/medium/large grouping	Reduces volume bias
Fraud injection	Fixed random seed	Supports repeatability
Metric choice	Provider-level F1 and FPR	Matches audit priority task
Evidence output	Top- K contributing claims	Improves interpretability

Table IX. Reproducibility guide for Top- K provider experiments.

XIV. SCOPE AND LIMITATIONS

The provider risk score is a prioritization signal and should not be interpreted as a legal or disciplinary conclusion. A provider may appear unusual because of specialty, location, patient mix, or emergency care profile. The model should therefore support audit triage rather than replace expert review.

The Top- K method assumes that fraudulent behavior creates a detectable high-risk tail in the claim-score distribution. This assumption is reasonable for repeated upcoding or concentrated billing abuse, but it may be less effective for subtle low-value fraud distributed across many claims. Combining Top- K with temporal trend features can reduce this limitation.

The study is based on simulated provider-labeled data, which is suitable for student research but not sufficient for production claims. Real-world testing would require de-identified institutional records, expert-reviewed provider outcomes, and careful handling of provider peer groups.

XV. RESULT INTERPRETATION AND REPORTING NOTES

Provider-level evaluation should not be judged by claim-level accuracy. The model is designed to rank providers for audit attention, so provider-level precision, recall, F1-score, and false-positive rate are more meaningful. A claim detector can appear accurate while still failing to identify the providers that require review. The paper therefore reports metrics at the provider level.

The Top-K score is interpretable because the K contributing claims can be displayed to the reviewer. This makes the model more transparent than a learned provider embedding whose internal reasoning is difficult to explain. In regulated domains, interpretable evidence is often as important as numerical accuracy.

The choice of K must be documented carefully. If K is selected after examining the test set, the reported result will be biased. In this paper, K is selected using a holdout set and then fixed. This reporting habit improves academic quality and makes the experiment reproducible.

Provider risk categories should be treated as audit priorities. A high-risk category means that the provider deserves further review, not that the provider has committed fraud. This distinction protects the academic paper from overclaiming and aligns the output with realistic investigation workflows.

The comparative tables show why simple aggregation rules are insufficient. Mean, median, maximum, and voting each capture only one part of provider behavior. Top-K aggregation captures repeated high-risk evidence while reducing sensitivity to one isolated outlier, which explains its improved simulation performance.

Reporting item	Provider-level relevance	Included output
Provider F1-score	Measures audit target detection	Table III
False-positive rate	Controls unnecessary reviews	Table III and Fig. 2
K sensitivity	Shows aggregation stability	Fig. 2 and Table V
Provider volume band	Improves fair comparison	Table IV
Top-K evidence claims	Supports interpretability	Audit interpretation section
Simulation scope	Prevents overclaiming	Limitations section

Table X. Reporting notes for provider-level Top-K risk prediction.

XVI. CONCLUSION AND FUTURE WORK

This paper presented a rewritten hierarchical framework for provider-level fraud risk prediction using Top-K aggregation of federated anomaly scores. The approach converts claim-level reconstruction errors into provider-level risk indices that better match audit workflow. The method is distinct from flat classification because it considers repeated provider behavior rather than isolated records only.

Simulation-based results show that Top-K aggregation improves F1-score and reduces false-positive prioritization compared with common aggregation alternatives. The paper maintains IEEE-style structure, original wording, editable equations, and consistent tables for student publication.

Future work can evaluate adaptive K selection, specialty-wise provider peer groups, temporal Top-K windows, and secure multi-party aggregation for provider scores when providers appear across multiple organizations.

REFERENCES

- [1] P. Kairouz et al., "Advances and open problems in federated learning," *Foundations and Trends in Machine Learning*, vol. 14, no. 1-2, pp. 1-210, 2021, doi: 10.1561/22000000083.
- [2] N. Rieke et al., "The future of digital health with federated learning," *npj Digital Medicine*, vol. 3, art. 119, 2020, doi: 10.1038/s41746-020-00323-1.
- [3] G. A. Kaissis, M. R. Makowski, D. Rückert, and R. F. Braren, "Secure, privacy-preserving and federated machine learning in medical imaging," *Nature Machine Intelligence*, vol. 2, pp. 305-311, 2020, doi: 10.1038/s42256-020-0186-1.
- [4] T. Li, A. K. Sahu, M. Zaheer, M. Sanjabi, A. Talwalkar, and V. Smith, "Federated optimization in heterogeneous networks," in *Proc. MLSys*, 2020.
- [5] H. B. McMahan, E. Moore, D. Ramage, S. Hampson, and B. A. y Arcas, "Communication-efficient learning of deep networks from decentralized data," in *Proc. AISTATS*, 2017, pp. 1273-1282.
- [6] M. Abadi et al., "Deep learning with differential privacy," in *Proc. ACM CCS*, 2016, pp. 308-318, doi: 10.1145/2976749.2978318.
- [7] C. Dwork and A. Roth, *The Algorithmic Foundations of Differential Privacy*. Hanover, MA, USA: Now Publishers, 2014.
- [8] R. Shokri, M. Stronati, C. Song, and V. Shmatikov, "Membership inference attacks against machine learning models," in *Proc. IEEE Symposium on Security and Privacy*, 2017, pp. 3-18, doi: 10.1109/SP.2017.41.
- [9] R. Kalapala and G. S. Sargam, "Hierarchical provider-level fraud risk modeling using differentially private cross-silo federated learning with Top-K aggregation," *INASS Express*, vol. 2, art. 6, 2026, doi: 10.22266/inassexpress.2026.006.

- [10] G. S. Sargam and R. Kalapala, "AI-driven claim fraud detection in health insurance using federated anomaly detection networks with cloud computing on AWS for privacy-preserving financial security," in Proc. ICPEEV, 2025, doi: 10.1109/ICPEEV67897.2025.11291290.
- [11] Sargam, G. S., & Kalapala, R. (2025). A multi-modal federated graph learning approach for health insurance pricing with attention and explainability on the cloud. In Proceedings of the Third International Conference on Cyber Physical Systems, Power Electronics and Electric Vehicles (ICPEEV 2025) (pp. 1–6). IEEE. <https://doi.org/10.1109/ICPEEV67897.2025.11291437>
- [12] Kalapala, R., & Sargam, G. S. (2025). Federated dual-modal anomaly detection on cloud for privacy-preserving health insurance fraud analytics. In Proceedings of the Third International Conference on Cyber Physical Systems, Power Electronics and Electric Vehicles (ICPEEV 2025) (pp. 1–6). IEEE. <https://doi.org/10.1109/ICPEEV67897.2025.11291269>
- [13] Gorrepati, L. P., Kalapala, R., & Sargam, G. S. (2025). Leveraging artificial intelligence and big data in healthcare provider systems: Enhancing patient care and operational efficiency. In Proceedings of the Third International Conference on Cyber Physical Systems, Power Electronics and Electric Vehicles (ICPEEV 2025) (pp. 1–6). IEEE. <https://doi.org/10.1109/ICPEEV67897.2025.11291497>
- [14] Kalapala, R., & Sargam, G. S. (2025). Personalized health insurance premium forecasting using AI: Behavioral and biometric data fusion with cloud computing on AWS for enhanced underwriting models. In Proceedings of the Third International Conference on Cyber Physical Systems, Power Electronics and Electric Vehicles (ICPEEV 2025) (pp. 1–6). IEEE. <https://doi.org/10.1109/ICPEEV67897.2025.11291190>
- [15] Sargam, G. S., & Kalapala, R. (2025). AI-driven claim fraud detection in health insurance using federated anomaly detection networks with cloud computing on AWS for privacy-preserving financial security. In Proceedings of the Third International Conference on Cyber Physical Systems, Power Electronics and Electric Vehicles (ICPEEV 2025) (pp. 1–6). IEEE. <https://doi.org/10.1109/ICPEEV67897.2025.11291290>