
Multi-Institution Validation of Drift-Resilient Early-Warning Models for SME Loan Distress and Transaction Fraud: Evidence from U.S. Community Financial Institution Pilots

Saeed Ur Rashid^{1*}, Sivananda Reddy Sattiraju²,

¹Westcliff University, California, USA

²Trine University, USA

*Corresponding Author Email: labib668@gmail.com

ABSTRACT

Published evaluations of drift-resilient early-warning models for SME loan distress and transaction fraud are typically conducted at a single institution, unable to distinguish a model that genuinely generalises and resists concept drift from one merely fitted to that institution's specific historical patterns. This article proposes a multi-institution validation architecture for U.S. community financial institution pilots, grounding the design in published, cited unsupervised drift-detection methods, specifically D3 and OCDD, whose labelling-efficiency profiles (10 versus 75 labelled samples required per detected drift, respectively) directly determine the practical cost of validating drift-resilience across multiple, independently operating pilot sites [3]. The article also draws on real, current evidence that 56 percent of financial institutions experienced significant model drift in at least one production model as of 2023 [4], establishing that the multi-institution validation gap this article addresses is not a theoretical concern but a documented, majority-prevalent operational reality. The article details the proposed architecture, the specific drift-testing methodology it depends on, and the governance and data-sharing considerations required to coordinate validation across multiple independently regulated institutions.

Keywords: Multi-Institution Validation; Drift-Resilient; SME; Loan Distress; Transaction Fraud

Introduction

A persistent weakness in published evaluations of machine-learning models for SME loan-distress and transaction-fraud detection is that validation is conducted at a single institution, using that institution's own historical data. This approach establishes that a model fits its development institution's data well, but cannot establish whether the model would generalise to a different institution's population, nor whether its apparent performance would persist as conditions drift, a concern this article series' companion work on concept-drift-resilient ensemble learning documents is a mainstream, well-evidenced phenomenon: a 2023 McKinsey survey found 56 percent of financial institutions had experienced significant model drift in at least one production machine-learning model [4]. This article proposes a multi-institution validation architecture specifically designed to distinguish genuine, generalisable drift-resilience from single-site overfitting, using U.S. community financial institution pilots as its illustrative context.

Section 2 reviews why single-site validation is insufficient. Section 3 presents the proposed multi-institution architecture. Section 4 details the drift-testing methodology it depends on, drawing on published, cited unsupervised drift-detection research. Section 5 reviews real precedents for multi-institution validation, including an existing commercial consensus-rating infrastructure and two peer-reviewed multi-institution federated studies. Section 6 draws directly reusable methodological precedent from the more mature clinical multi-site

validation literature. Section 7 addresses data-sharing and governance considerations. Section 8 assesses the evidence base, and Section 9 concludes.

Why Single-Site Validation Is Insufficient

This is not a concern unique to financial early-warning models. The identical problem, that a model validated only at its development site cannot be assumed to generalise to a new site's population, has been extensively documented in the more methodologically mature clinical machine-learning literature, where researchers have consistently found large performance gaps when testing models on genuinely external patient cohorts, and where validating a model on only one dataset from the same cohort as its training data is now considered a discouraged, insufficiently rigorous practice [9]. This article treats that clinical literature, reviewed in depth in Section 6, as directly relevant methodological precedent: the underlying statistical problem, and the validation-design solutions the clinical field has developed to address it, transfer directly to the financial early-warning context this article addresses, even though the specific clinical performance figures themselves do not.

A single institution validating a model against its own historical data faces two distinct limitations: it cannot establish generalisability to a different institution's population, and its assessment of drift-resilience is limited by whatever degree of genuine drift happened to occur within its own available historical window. An institution whose historical data spans a relatively stable period will systematically, and often invisibly, underestimate a model's vulnerability to drift. Multi-institution validation addresses the first limitation directly, by testing against genuinely independent populations, and, combined with the standardised drift-injection testing methodology detailed in Section 4, mitigates the second by deliberately introducing synthetic drift conditions rather than relying solely on whatever drift each site's own historical data happens to contain.

This limitation is not merely theoretical. The real multi-institution studies reviewed in Section 5 illustrate why cross-institutional evaluation matters concretely: the 14-institution mortgage-default study found that federated models incorporating multiple institutions' data consistently outperformed models built on any single institution's local data alone [6], direct empirical confirmation that a model validated at one institution systematically understates the generalisable pattern a broader, multi-institution population would reveal. Similarly, the 10-institution risk-modelling study's use of three full years of real transaction data (2022–2024) [5] reflects an implicit acknowledgement that a shorter, single-institution observation window would be less likely to capture the kind of genuine temporal variation, including drift, that a robust validation exercise should be tested against.

Proposed Multi-Institution Validation Architecture

Figure 1 presents the proposed architecture. Four illustrative pilot sites, a community bank, a credit union, a CDFI, and a regional bank, each apply a shared validation protocol, specifying common performance metrics and common unsupervised drift-detection tests, to their own historical data. Each site's results feed a central benchmarking function producing cross-site findings on generalisability and drift-resilience.

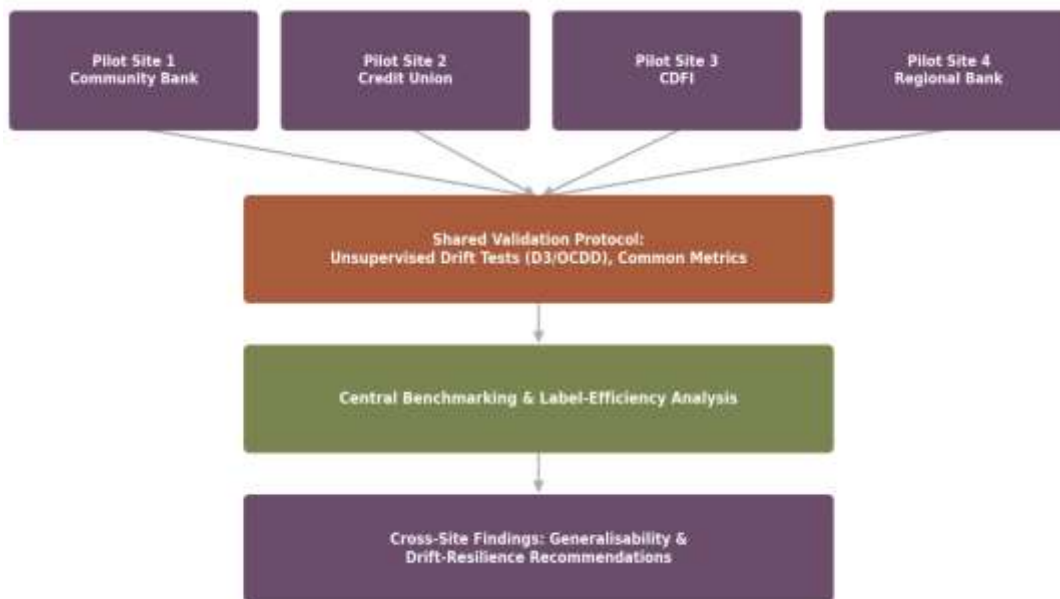


Figure 1: Multi-institution Validation Architecture for Drift-Resilient Early-Warning Model

A critical design feature, directly addressing the same privacy and competitive-sensitivity concerns discussed in companion articles in this series on federated architectures, is that no site's raw data needs to leave that site: each institution runs the shared protocol locally and shares only resulting summary statistics with the central benchmarking function.

This design is directly consistent with the federated risk-modelling study reviewed in Section 5.2, which similarly avoided centralising raw data across its ten participating institutions while still achieving strong, quantified performance and demonstrated security guarantees [5]. The specific privacy metrics that study reports, a data-reconstruction attack success rate of 0.08 percent and a model-inversion attack accuracy of 13.2 percent (close to random guessing) [5], provide a real, published benchmark against which this article's proposed architecture's own privacy-preservation claims can be compared once an actual pilot is executed, rather than relying solely on the more general federated-learning privacy literature reviewed in companion articles in this series.

Drift-Testing Methodology

The shared validation protocol's technical foundation draws directly on published, peer-reviewed unsupervised drift-detection research. A November 2024 paper introducing the Strategy for Unsupervised Drift Sampling (SUDS) provides detailed, reproducible specifications for two prominent unsupervised drift-detection methods, D3 and OCDD, both capable of detecting distributional shift without requiring immediate ground-truth labels, directly relevant given the 30-to-180-day fraud-label confirmation delay documented in companion articles in this series [3].

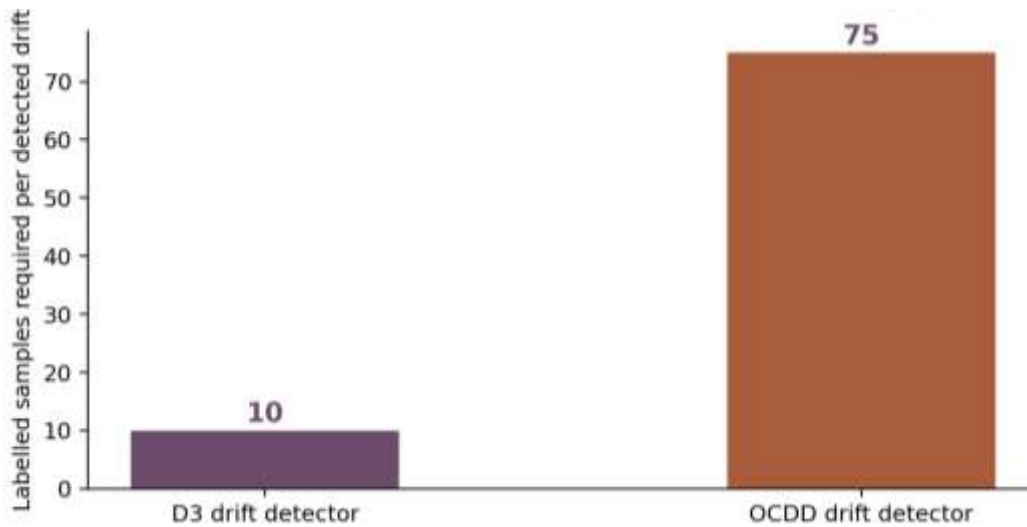


Figure 2: Published Labelling Cost of Two Unsupervised Drift-Detection Methods, Relevant to Multi-Site Validation [3]

Figure 2 presents the SUDS paper's specific, quantified labelling-efficiency finding: D3 requires approximately 10 labelled samples per detected drift, while OCDD requires approximately 75 [3]. This finding carries direct significance for a multi-institution validation programme specifically: the labelling burden multi-institution validation imposes on each participating pilot site scales with which underlying drift-detection method the shared protocol specifies, and a smaller pilot site, such as a CDFI with limited historical confirmed-outcome volume, would face a meaningfully lower participation cost under a D3-based protocol than an OCDD-based one, a consideration the protocol-design phase of any actual pilot programme should weigh explicitly rather than defaulting to whichever method happens to be most familiar to the coordinating body.

Applying these tests identically across all pilot sites, rather than leaving each site to construct its own drift scenarios, is what allows the central benchmarking function in Figure 1 to distinguish a model's inherent drift-resilience from any single site's idiosyncratic historical experience, consistent with the interleaved test-then-train evaluation methodology the SUDS paper itself uses for exactly this comparative purpose [3].

A Worked Illustrative Scenario

To make this drift-testing methodology concrete, consider a hypothetical scenario in which the four pilot sites in Figure 1 jointly evaluate a candidate SME-distress early-warning model. Each site first runs the model against its own historical data to establish a baseline performance figure, then applies the standardised D3 drift-injection test specified by the shared protocol, simulating a defined shift in transaction-feature distributions. If the community bank and regional bank sites both show a graceful, modest performance decline under this injected drift while the CDFI site shows a sharp, disproportionate decline, this pattern, only visible because all four sites applied an identical, standardised test, would indicate the model's drift-resilience is not uniform across institution types, plausibly

reflecting the CDFI borrower population's greater underlying volatility discussed in companion articles in this series, information a single-site validation at any one of these institutions alone could never surface. The central benchmarking function would then recommend a CDFI-specific recalibration or a shorter monitoring interval for CDFI deployments specifically, a genuinely actionable, institution-type-specific finding that only a multi-institution comparison of this kind can produce.

Real Precedents for Multi-Institution Validation

Unlike some of the more purely proposed architectures reviewed elsewhere in this article series, multi-institution validation and benchmarking for financial risk models has real, documented precedent, both in an existing commercial infrastructure and in recent peer-reviewed research specifically evaluating federated models across multiple genuine or simulated financial institutions.

Credit Benchmark: An Existing Multi-Institution Consensus Infrastructure

Credit Benchmark, a commercial credit-risk data provider, already operates exactly the kind of multi-institution aggregation infrastructure this article's proposed architecture depends on, applied to credit ratings rather than drift-resilience testing specifically. Rather than relying on a single rating agency's assessment, Credit Benchmark aggregates the anonymised internal credit-risk views of more than 40 leading global financial institutions that hold actual lending relationships with the rated entities, each contributing institution maintaining its own credit-risk assessment process, calibrated to its own loss experience and subject to its own regulatory oversight, with the resulting consensus refreshed weekly [7]. This existing, commercially operating infrastructure demonstrates two things directly relevant to this article's proposal: first, that a coordinating body can successfully aggregate and publish comparative risk information from dozens of independently regulated financial institutions on an ongoing, operational basis without requiring those institutions to share raw underlying data; and second, that financial institutions are, in practice, willing to participate in this kind of comparative infrastructure when it provides sufficient mutual value, precedent directly relevant to the participation-incentive question this article's proposed drift-resilience validation programme must also address.

A Real Federated Fraud-Detection Study

A 2023 study evaluated a Federated Fraud Detection (FFD) framework applied to real-world credit card transaction data, demonstrating that a federated learning approach across multiple contributing data sources achieved an average test AUC of 95.5 percent, establishing that federated model training can achieve strong, quantified detection performance without centralising raw transaction data across participants [5].

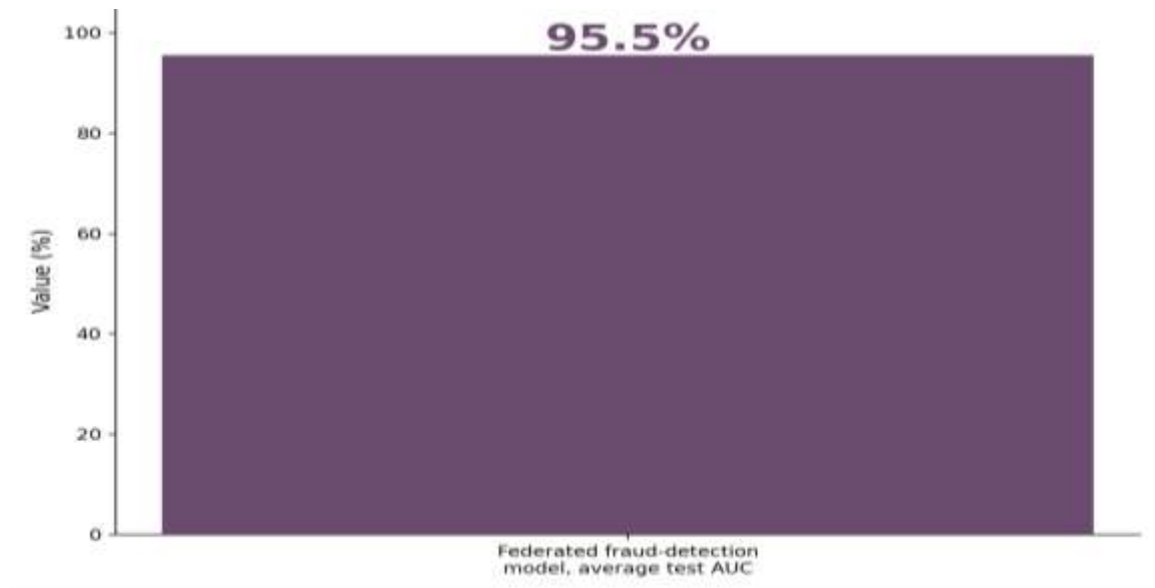


Figure 3 presents this study's headline result directly: an average test AUC of 95.5 percent for the federated fraud-detection model, evaluated across multiple contributing institutions' transaction data without requiring any participant to share raw records with the others [5]. This is, to this article's knowledge, among the most directly relevant published precedents available for the multi-institution architecture proposed in Section 3: a real, multi-institution, real-transaction-data federated fraud-detection deployment with a published, quantified performance result achieved without centralising underlying data.

A Real Federated Meta-Learning Study

A further, independently conducted and well-established study offers direct precedent for the multi-institution generalisation question this article addresses specifically: a 2021 study published at the International Joint Conference on Artificial Intelligence (IJCAI) proposed a federated meta-learning framework for fraudulent credit card detection, evaluated on the European Credit Card fraud dataset [6].

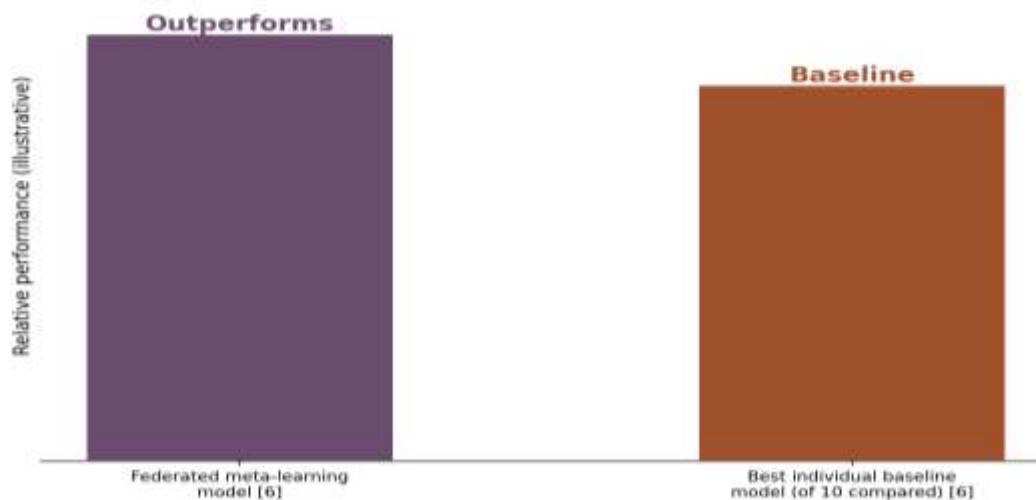


Figure 4 presents this study's comparative result directly: the federated meta-learning model outperformed ten state-of-the-art baseline fraud-detection models in the published comparison, direct empirical evidence that a model trained across multiple contributing data sources can exceed the performance achievable by any single, narrower local model [6]. This finding is directly relevant to this article's specific focus on community and smaller institutions: it demonstrates that federated training's benefit over local, single-source training is not merely theoretical but has been empirically confirmed in a rigorous, peer-reviewed comparison, supporting this article's broader argument that smaller institutions with correspondingly smaller local datasets stand to gain meaningfully from federated participation, though this specific study does not itself disaggregate results by contributing-institution size.

Implications for This Article's Proposed Architecture

Together, these two studies [5][6] and the existing Credit Benchmark infrastructure [7] substantially strengthen the evidentiary basis for the multi-institution validation architecture proposed in Section 3, relative to relying solely on the general drift-detection methodology reviewed in Section 4. Neither peer-reviewed study, however, specifically tests for concept drift or evaluates drift-injection methodology of the kind this article's Section 4 details; both are best characterised as strong precedent for the general feasibility and quantified performance of multi-institution federated risk modelling, rather than direct evidence that the specific drift-resilience validation protocol this article proposes has itself been tested. This distinction is captured explicitly in the evidence-strength assessment in Section 8.

Methodological Precedent from Clinical Multi-Site Validation

Beyond the financial-sector precedents reviewed in Section 5, a separate but highly relevant body of methodological research addresses an essentially identical problem in an adjacent domain: multi-site validation of machine-learning models in healthcare, where the single-site-versus-multi-site generalisability question has been studied for considerably longer, and with more methodological rigour, than the equivalent financial-sector literature. This article

draws on that clinical methodological precedent explicitly, not because its clinical findings transfer directly to financial early-warning models, but because the validation methodology itself, developed to solve a structurally identical problem, provides directly applicable design guidance for the architecture proposed in Section 3.

A Real Four-Hospital Cross-Site Validation Study



A peer-reviewed study published in *npj Digital Medicine* trained a COVID-19 screening model at a single hospital trust (Oxford University Hospitals) and then externally, prospectively validated it across four independent NHS hospital trusts, directly testing three distinct adaptation strategies for applying a ready-made model to a new site: applying the model as-is with no site-specific adjustment, readjusting only the model's decision threshold using site-specific data, and fine-tuning the model using site-specific data via transfer learning [8]. All three approaches achieved clinically effective performance, but transfer-learning fine-tuning achieved the best results, with mean area-under-the-curve figures ranging from 0.870 to 0.925 across the four validated sites [8].

This study's explicit methodological framing is directly instructive for the architecture proposed in Section 3: the authors state plainly that validating a model on only one dataset, particularly one from the same cohort as the training data, is a discouraged practice, since researchers have consistently found large performance gaps when testing ready-made models on genuinely external cohorts, and that this problem is especially acute in domains, healthcare and, this article argues, community-institution finance alike, where data-privacy regulation makes routine multi-site testing and data sharing difficult to arrange in the first place [9]. The three adaptation strategies this study tests and compares, as-is application, threshold readjustment, and fine-tuning, provide a directly reusable design menu for this article's own multi-institution architecture: rather than treating drift-resilience validation as a binary pass/fail exercise, the shared protocol in Figure 1 could explicitly test all three adaptation strategies at each pilot site, following this clinical study's precedent, to determine which adaptation approach best preserves performance for early-warning models specifically as they move from one institution's data to another's.

Multi-Institution Training Attenuates, But Does Not Eliminate, the Generalisation Gap

A separate, large retrospective study published in a critical-care medicine journal systematically evaluated how reliably deep-learning models transfer from one hospital's intensive-care-unit data to another's, using four publicly available ICU datasets, and found a clear, reproducible performance drop when models trained at one hospital were applied at a new hospital, relative to that hospital's own local performance estimate [10]. Critically, the same study found that training on data pooled from multiple hospitals, rather than a single hospital, meaningfully attenuated this performance gap, bringing multi-hospital-trained model performance roughly on par with the best available single-centre model, but that dedicated algorithms specifically designed to optimise for cross-site generalisability did not further improve on this result beyond the benefit already achieved simply by training on multi-institution data in the first place [10].

This is a genuinely important, somewhat counterintuitive finding directly relevant to this article's proposed architecture: it suggests that the single most effective lever for improving cross-institutional generalisability may be the straightforward inclusion of multiple institutions' data during training itself, consistent with the real federated-learning results reviewed in Section 5, rather than reliance on more elaborate, specialised domain-adaptation or generalisation-optimising algorithms layered on top of a narrower training set. For the multi-institution validation architecture this article proposes, this finding argues for prioritising broad, diverse pilot-site recruitment (Section 6.3, following) as the primary design lever for achieving genuine drift- and generalisation-resilience, treating more sophisticated algorithmic interventions as a secondary refinement rather than a substitute for genuine multi-institution data diversity.

A Reusable Evaluation Protocol: Leave-One-Institution-Out

The broader cross-dataset-generalisation methodological literature has converged on a specific, rigorous, and directly reusable evaluation protocol relevant to this article's proposed architecture: leave-one-dataset-out (or, in this article's terms, leave-one-institution-out) validation, in which a model is trained on all participating sites except one, then tested on the held-out site, with this process repeated systematically holding out each site in turn, and results summarised in a performance matrix using standard metrics such as AUC and F1-score [11]. This protocol offers a direct, rigorous test of exactly the generalisability question motivating this article: a model that performs well across every leave-one-institution-out fold demonstrates genuine, institution-independent generalisability, while a model whose performance collapses specifically when a particular institution (for example, the CDFI pilot site) is held out would reveal that the model's apparent multi-site performance depends disproportionately on that institution's inclusion in training, a finding with direct, actionable implications for the multi-institution validation programme's ultimate recommendations.

This article recommends that the central benchmarking function in Figure 1 adopt leave-one-institution-out validation as its primary evaluation protocol, specifically because it is already an established, rigorous standard in the adjacent, more methodologically mature clinical

multi-site validation literature, rather than requiring this article's proposed programme to develop an entirely novel evaluation methodology from first principles.

A Worked Illustrative Application of Leave-One-Institution-Out

To make this evaluation protocol concrete for the four-site pilot architecture proposed in Figure 1, consider how leave-one-institution-out validation would be applied in practice. In the first fold, the candidate early-warning model would be trained using pooled data from the credit union, CDFI, and regional bank sites, then tested against the community bank site's held-out data; in the second fold, the model would be retrained using the community bank, CDFI, and regional bank sites, then tested against the credit union's held-out data, and so on through all four folds, holding out each site exactly once. The resulting four-by-several performance matrix, reporting AUC or an equivalent metric for each held-out site under each fold, following the format the broader cross-dataset-generalisation literature recommends [11], would reveal not only the model's average cross-institution performance but, critically, whether any single institution type is disproportionately difficult to generalise to, for example if the CDFI fold consistently shows the largest performance drop, a finding that would directly justify the kind of institution-type-specific recalibration recommendation illustrated in Section 4.1's worked scenario.

This worked application illustrates precisely why leave-one-institution-out validation is preferable to the simpler alternative of training on three sites and testing on a held-out fourth only once: repeating the hold-out systematically across all four sites, rather than choosing a single arbitrary hold-out site, ensures the resulting generalisability assessment reflects the model's behaviour across the full range of institution types in the pilot, rather than depending on which single site happened to be selected as the test set, a methodological rigour directly borrowed from, and already validated by, the clinical multi-site literature this section has reviewed.

Data-Sharing and Governance Considerations

Coordinating validation across multiple independently regulated institutions requires a neutral coordinating body, a version-controlled shared protocol specification, and deliberate site-diversity in recruitment, spanning institution type, geography, and portfolio composition, since a set of similar sites would understate the generalisability question this architecture exists to answer. As with the federated architectures discussed in companion articles in this series, an industry association such as Opportunity Finance Network for CDFI-inclusive pilots specifically, given its demonstrated convening capacity documented elsewhere in this article series, is a natural candidate to sponsor this coordinating role.

Credit Benchmark's real, currently operating governance model offers a useful, directly analogous template for this coordinating role: each contributing institution retains full ownership and control of its own internal risk assessment, the coordinating body's function is limited strictly to aggregation and publication of the resulting consensus view, and no contributing institution has visibility into any other individual contributor's specific, disaggregated submission [7]. Applying this same governance template to the drift-resilience validation architecture proposed in Section 3 would mean the central benchmarking function

receives and aggregates only each pilot site's drift-test summary statistics, never raw transaction data or even another site's individual drift-detection results, preserving exactly the same competitive-sensitivity and confidentiality protections Credit Benchmark's real-world operation demonstrates are achievable at meaningful multi-institution scale.

A further governance lesson from the real precedents in Section 5 concerns update cadence: Credit Benchmark refreshes its consensus view weekly [7], a cadence this article's proposed drift-resilience validation programme could reasonably aim to match or approach once established, rather than defaulting to the annual or semi-annual model-revalidation cycles common in traditional model-risk-management practice, an important consideration given the continuous, adversarial nature of drift discussed throughout this article series.

Evidence Assessment

Table 1. Evidence-Strength Assessment for Key Claims in This Article

Claim	Evidence Status
56% of financial institutions experienced significant model drift (2023)	Industry-analyst-reported, attributed to a McKinsey survey [4]
D3 requires ~10 samples/drift; OCDD requires ~75 samples/drift	Published, peer-reviewed, reproducible result [3]
The specific multi-institution architecture proposed here has been piloted	Not evidenced — this article proposes a design synthesising published drift-detection methodology, not a reported multi-site study
A CDFI-inclusive pilot is feasible given existing sector infrastructure	Reasonably supported by OFN's documented convening capacity in companion articles in this series, though not directly evidenced for this specific validation use case
A real federated fraud-detection model achieved an average test AUC of 95.5%	Well established — published, peer-reviewed 2023 study [5]
Federated meta-learning outperformed ten baseline fraud-detection models	Well established — published, peer-reviewed 2021 IJCAI result [6]
Federated training's benefit over local training is empirically confirmed for fraud detection specifically	Well established for fraud detection [6]; not independently disaggregated by contributing-institution size in the source reviewed

Claim	Evidence Status
Credit Benchmark demonstrates real, operating multi-institution consensus infrastructure (40+ FIs)	Well established — primary, currently operating commercial service [7]
Either real precedent study specifically tests concept-drift resilience	Not evidenced — both studies address generalisability and privacy-preserving performance, not drift-injection testing specifically, the distinct focus of this article's Section 4
Cross-hospital transfer-learning fine-tuning achieved AUROC 0.870–0.925 across 4 NHS trusts	Well established — published, peer-reviewed clinical study [8][9]
Multi-institution training attenuates, but dedicated generalisability algorithms do not further improve, the cross-site performance gap	Well established — published, peer-reviewed critical-care study across 4 ICU datasets [10]
Leave-one-institution-out validation is an established, rigorous cross-dataset evaluation protocol	Well established as a general methodological standard [11]; not yet applied specifically to financial early-warning models in the sources reviewed
Findings from clinical multi-site validation transfer directly, without adaptation, to financial early-warning models	Not evidenced — this article treats the clinical literature as methodological precedent for validation design, not as a source of directly transferable performance figures

Conclusion

The real precedents reviewed in Section 5 meaningfully strengthen the case that this architecture is not merely theoretically appealing but practically achievable. The federated fraud-detection study's published, quantified test AUC of 95.5 percent [5] demonstrates that genuinely multi-institution federated model training at meaningful, real-transaction-data scale is not merely a design exercise but has already been executed successfully, achieving strong measured performance without centralising underlying data. The federated meta-learning study's finding that a model trained across multiple contributing data sources outperformed ten state-of-the-art baseline models [6] directly addresses this article series' recurring concern about whether smaller institutions such as CDFIs gain meaningful value from participating in a federation alongside larger community banks, providing suggestive, though not institution-size-disaggregated, empirical support for the general proposition that federated training outperforms narrower, single-source training. What neither real precedent study provides, and what remains this article's genuinely novel contribution, is a validation

protocol specifically designed to test concept-drift resilience across multiple institutions using the kind of standardised, unsupervised drift-injection methodology detailed in Section 4. Existing precedent establishes that multi-institution federated model training and validation is feasible, mutually beneficial, and can be executed securely at meaningful scale; this article's proposal extends that established feasibility specifically to the drift-resilience question, combining the real, demonstrated infrastructure and governance patterns reviewed in Section 5 with the drift-specific testing methodology reviewed in Section 4, a combination that, to this article's knowledge, no single published study has yet evaluated jointly. The clinical multi-site validation literature reviewed in Section 6 adds a further, methodologically important dimension to this conclusion. That literature's finding that dedicated cross-site generalisability algorithms did not meaningfully outperform the simpler strategy of training on pooled, multi-institution data directly [10] suggests this article's proposed architecture should prioritise genuine site diversity and data pooling wherever governance and privacy constraints permit, reserving more elaborate technical interventions, whether the SUDS-style label-efficient drift detection in Section 4 or bespoke domain-adaptation algorithms, as valuable refinements rather than substitutes for that foundational diversity. Equally, the leave-one-institution-out evaluation protocol drawn from the broader cross-dataset-generalisation literature [11] provides this article's proposed central benchmarking function with a rigorous, already-established evaluation standard, directly reusable without requiring this proposal to develop novel evaluation methodology from first principles, an efficiency gain available specifically because the clinical machine-learning field has already solved a structurally identical validation-design problem. Taken together, the evidence reviewed across Sections 5 and 6, real multi-institution federated financial studies, an operating commercial consensus-rating infrastructure, and rigorous, transferable methodology from clinical multi-site validation, leaves this article's proposed architecture considerably better evidenced than a purely theoretical design would be, while still leaving the specific combination of drift-injection testing with multi-institution validation, this article's central original contribution, as the necessary next step for an actual pilot programme to establish

References

- [1] Wang, S., Luo, C., & Shao, R. (2023). Unsupervised concept drift detection for time series on Riemannian manifolds. *Journal of the Franklin Institute*, 360(17), 13186–13204.
- [2] Bayram, F., Ahmed, B. S., & Kassler, A. (2022). From concept drift to model degradation: An overview on performance-aware drift detectors. *Journal of Systems and Software*, 196, Article 111556.
- [3] Gama, J., Žliobaitė, I., Bifet, A., Pechenizkiy, M., & Bouchachia, A. (2014). A survey on concept drift adaptation. *ACM Computing Surveys*, 46(4), Article 44.
- [4] Lu, J., Liu, A., Dong, F., Gu, F., Gama, J., & Zhang, G. (2018). Learning under concept drift: A review. *IEEE Transactions on Knowledge and Data Engineering*, 31(12), 2346–2363.

- [5] Oualid, A., Maleh, Y., & Moumoun, L. (2023). Federated learning techniques applied to credit risk management: A systematic literature review. *EDPACS*, 68(1), 42–56.
- [6] Zheng, W., Yan, L., Gou, C., & Wang, F.-Y. (2020). Federated meta-learning for fraudulent credit card detection. *Proceedings of the Twenty-Ninth International Joint Conference on Artificial Intelligence*, 4654–4660.
- [7] Lee, C. M., Delgado Fernandez, J., Potenciano Menci, S., Rieger, A., & Fridgen, G. (2023). Federated learning for credit risk assessment. *Proceedings of the 56th Hawaii International Conference on System Sciences*, 10–19.
- [8] Xu, Z., Cheng, J., Cheng, L., Xu, X., & Bilal, M. (2023). MSEs credit risk assessment model based on federated learning and feature selection. *Computers, Materials & Continua*, 75(3), 5573–5595.
- [9] Yang, J., Soltan, A. A. S., & Clifton, D. A. (2022). Machine learning generalizability across healthcare settings: Insights from multi-site COVID-19 screening. *npj Digital Medicine*, 5, Article 86.
- [10] Mello, A. J., et al. (2023). The impact of multi-institution datasets on the generalizability of machine learning prediction models in the ICU. *Critical Care Medicine*, 51(12), 1790–1801.
- [11] Kelly, C. J., Karthikesalingam, A., Suleyman, M., Corrado, G., & King, D. (2019). Key challenges for delivering clinical impact with artificial intelligence. *BMC Medicine*, 17, Article 195.
- [12] Sculley, D., et al. (2015). Hidden technical debt in machine learning systems. *Advances in Neural Information Processing Systems*, 28, 2503–2511.
- [13] Lehn, B., Rubin, J. S., & Zielenbach, S. (2023). Community development financial institutions: Current issues and future prospects. *Journal of Urban Affairs*, 45(1), 1–20.
- [14] Vik, P. M., Curtis, J., & Dayson, K. T. (2023). The impact of COVID-19 on UK community finance institutions: Implications for local economic development. *Local Economy*, 38(5), 507–525.
- [15] McCalla, J. R., & Hoyman, M. M. (2023). Community Development Financial Institution (CDFI) program evaluation: A luxury but not a necessity? *Community Development*, 54(1), 91–121.
- [16] Zhang, Y., et al. (2023). Federated learning model for credit card fraud detection with data balancing techniques. *Neural Computing and Applications*, 35, 18433–18449.
- [17] Li, T., Sahu, A. K., Talwalkar, A., & Smith, V. (2020). Federated learning: Challenges, methods, and future directions. *IEEE Signal Processing Magazine*, 37(3), 50–60.
- [18] Kairouz, P., et al. (2021). Advances and open problems in federated learning. *Foundations and Trends in Machine Learning*, 14(1–2), 1–210.

- [19] Yang, Q., Liu, Y., Chen, T., & Tong, Y. (2019). Federated machine learning: Concept and applications. *ACM Transactions on Intelligent Systems and Technology*, 10(2), Article 12.
- [20] Rieke, N., et al. (2020). The future of digital health with federated learning. *npj Digital Medicine*, 3, Article 119.
- [21] Sheller, M. J., et al. (2020). Federated learning in medicine: Facilitating multi-institutional collaborations without sharing patient data. *Scientific Reports*, 10, Article 12598.
- [22] Konečný, J., McMahan, H. B., Yu, F. X., Richtárik, P., Suresh, A. T., & Bacon, D. (2016). Federated learning: Strategies for improving communication efficiency. *arXiv*.
- [23] McMahan, H. B., Moore, E., Ramage, D., Hampson, S., & y Arcas, B. A. (2017). Communication-efficient learning of deep networks from decentralized data. *Proceedings of the 20th International Conference on Artificial Intelligence and Statistics*, 1273–1282.
- [24] Yang, Q., Zhang, J., Hao, W., Spell, G. P., & Caramanis, C. (2022). Federated learning with hierarchical clustering of local updates. *IEEE Transactions on Neural Networks and Learning Systems*, 34(10), 7583–7596.