

Federated Explainable AI for Secure Clinical Decision Support and Predictive Risk Stratification in Smart ICUs

Zaiver Jan, Research Scholar, Department of Computer Science, Sikkim Alpine University, Kamrang (Sikkim)
Dr. Sunil Gupta, Supervisor, Computer Science Dept., Sikkim Alpine University, Kamrang, (Sikkim)

Abstract

Intensive Care Units (ICUs) produce massive streams of patient data by the minute: vital signs, test results, ventilator readouts and electronic health records. Artificial intelligence (AI) can assist physicians to make decisions more quickly and accurately, but privacy and explainability are two main challenges to widespread application in real ICU situations. Patient data is highly sensitive and can't just be shared between hospitals for centralized AI training. Even when an AI model exists, clinicians are frequently reluctant to trust it if it's a "black box" that can't explain its recommendations. In this research, we propose FedXAI-ICU, a framework that addresses both issues simultaneously. In each hospital site, local patient data is used to train local AI models via federated learning, and only the learned model parameters – not the raw data – are shared and aggregated into a powerful global model. We use explainable AI approaches, notably SHAP (SHapley Additive exPlanations) and LIME (Local Interpretable Model-agnostic Explanations), to provide a clear, clinician-readable explanation of which patient measurements motivated each clinical prediction choice. Experiments on data from three simulated ICU nodes show that FedXAI-ICU achieves an AUC-ROC of 0.893 for sepsis prediction, outperforming centralized and federated-only baselines, while keeping all the raw patient records local and providing faithful explanations that were accepted by clinicians in a usability study.

Keywords: Intensive Care Units, SHapley Additive exPlanations, AUC-ROC, LIME

1. Introduction

These days, intensive care units house more data than any other medical facility. There can be tens of monitoring devices delivering readings every few seconds to a single intensive care unit patient. Along with dozens of lab results collected every few hours and a continually updated electronic health record, a ventilated patient generates readings from the following devices: a pulse oximeter, an ECG, invasive blood pressure lines, an arterial blood gas analyzer, and a breathing machine. With twenty beds in an intensive care unit, this generates millions of data points daily. The expertise of individual physicians is still crucial in the intensive care unit (ICU) for making important decisions such as whether to try to wean a patient off of a ventilator, when to raise a patient's vasopressor dose, or when to suspect sepsis is developing, even if there is a wealth of data available. Research shows that human judgment, even when made by seasoned intensivists, can be slow and inconsistent, especially when dealing with night shifts, staffing shortages, or exceptionally high patient loads. One solution is artificial intelligence (AI) models based on past intensive care unit data, which can detect minor patterns that indicate clinical deterioration hours before humans can see them. There are two major roadblocks that have kept AI confined to academic journals and away from routine application in intensive care units. Data privacy is the primary concern. Due to the small sample sizes of individual intensive care units, it is not possible to train a robust AI model using data from a single facility alone. This is because numerous patients from different hospitals are needed to account for the occurrence of rare but crucial occurrences. However, sharing identifiable patient information is severely limited by stringent privacy rules governing hospital data – GDPR in Europe and HIPAA in the US. Researchers and hospitals are thus faced with a conundrum: either gather more data for a stronger model, or safeguard patient privacy and settle for a less robust model that is data-starved.

Another challenge is the need for clear explanations. When an AI system gives a number—a 78% chance of sepsis onset within 4 hours, for example—without explaining why, clinicians are understandably hesitant to follow the advise. Accountability for clinical judgments is a requirement in both medicine and the law. For any of these audiences—patients, families,

ethics boards, or courts—the response "The AI said so" will not suffice. Black boxes are artificial intelligence systems that cannot explain their thinking. Despite their accuracy in lab testing, black boxes have a hard time gaining the trust necessary for actual clinical use.

To address both challenges at once, this article introduces FedXAI-ICU, a framework developed with that aim in mind. You can understand the main points easily. Each healthcare facility uses its own patient data to train an artificial intelligence model; only the learned patterns, expressed as model weights, are transmitted to a secure aggregate server. At importance of tailoring explanations to specific clinical questions. This has led to an increasing amount of research on clinician-centered explanation design.

that point, federated averaging is used to merge the local models into one robust global model. Subsequently, the enhanced global model is transmitted back to every hospital. In no hospital does raw patient data ever leave the premises. This technique is known as federated learning. To tackle the issue of explainability, FedXAI-ICU associates each clinical prediction with a SHAP-based explanation. To determine the relative importance of each patient measurement in making a given forecast, researchers developed SHAP (SHapley Additive exPlanations), a mathematically sound method. A clinician reviewing an AI-generated sepsis risk alert can observe more than just a risk score; they can also see a ranked list of the patient's own measurements that drove the prediction, such as a decrease in mean arterial pressure, an increase in lactate level, and a decrease in urine output, with the relative importance of each measurement displayed clearly.

The following is the outline for the remainder of the paper. Related works are reviewed in Section 2. Section 3 provides a detailed description of the FedXAI-ICU design and procedures. The experimental setup is described in Section 4. Findings are detailed in Section 5. Ethical issues and results are addressed in Sections 6 and 7. Chapter 8 comes to a close.

2. Related Work

2.1 Federated Learning in Healthcare

Since its introduction by McMahan *et al.* (2017), federated learning has found numerous applications in medical imaging and clinical prediction. Rieke *et al.* (2020) showed that federated learning across many hospital sites had competitive performance when compared to centralized techniques for segmenting brain tumors in MRI images. To predict acute kidney injury using regular clinical measures across three New York hospitals, Vaid *et al.* (2021) used federated learning in the ICU setting. They found that the federated model generalized better than any single-site model. Though federated privacy and clinical explainability are interrelated issues, they were not tackled in any of these efforts.

2.2 Explainable AI in Clinical Settings

There has been a lot of focus on explainability in clinical AI research. ICU risk scores are explained by SHAP, which was produced by Lundberg and Lee (2017). SHAP offers features predictions that are compatible with theory and has become frequently employed. The clinical utility and explanation quality are not necessarily synonymous, according to Tonekaboni *et al.* (2019), who examined various explanation approaches in the context of clinical AI. Even technically accurate explanations could perplex doctors who aren't familiar with the method. Ghassemi *et al.* (2020) and others have built on this idea to stress the

2.3.1 Artificial Intelligence for Critical Care Unit Risk Assessment

More than forty thousand intensive care unit (ICU) patients' full clinical records are housed in the MIMIC-III database, which was created by Johnson *et al.* (2016). Among the several benchmark datasets for critical care risk prediction using artificial intelligence and machine learning, this one has seen extensive use. The study proved that data-driven methods, when applied to large-scale intensive care unit datasets, can greatly enhance the prediction of bad clinical outcomes such as patient death and sepsis. A system for the early prediction of sepsis in intensive care unit patients using machine learning algorithms and physiological cues was proposed by Nemati *et al.* (2018). The results

demonstrated that compared to traditional scoring systems like SOFA and SIRS, machine learning models might detect sepsis up to three hours earlier. Their findings demonstrated that AI methods are useful in critical care settings for improving patient outcomes through the provision of timely clinical decision support.

For the purpose of forecasting clinical worsening and mortality in intensive care units, *Shashikumar et al. (2021)* presented a deep learning method based on Long Short-Term Memory (LSTM) networks. The suggested approach enhanced prediction accuracy and early risk assessment compared to traditional severity scores, such as APACHE II and NEWS, by capturing temporal relationships in physiological parameters using sequential patient records from the MIMIC-III database.

Using probabilistic deep learning methods, *Alaa and van der Schaar (2019)* created a system for machine learning that can identify individual risks in critical care settings. They quantified the prediction uncertainty and used their model to estimate the patient-specific risks of mortality and complications. The study showed that personalized AI-based risk assessment could improve clinical decision-making by providing more accurate and understandable forecasts than traditional intensive care unit score systems.

2.4 Gap Addressed by this Work

Few efforts have integrated federated learning with XAI for intensive care unit clinical decision support, despite the fact that XAI tackles explainability and federated learning handles privacy. A lot of previous research has used federated learning for ICU prediction without providing any explanations or applied XAI to centralized ICU models. Filling this need, FedXAI-ICU views explainability and privacy as complementary rather than complementary to one another.

3. Methodology

3.1 System Architecture Overview

Figure 1 shows the overall design of the FedXAI-ICU system. A privacy-protected zone houses all the intensive care unit nodes that keep patient data local, while the secure aggregator, global model, and clinical dashboard are located in the global zone. Never patient records are communicated between zones; only encrypted model parameters are sent.

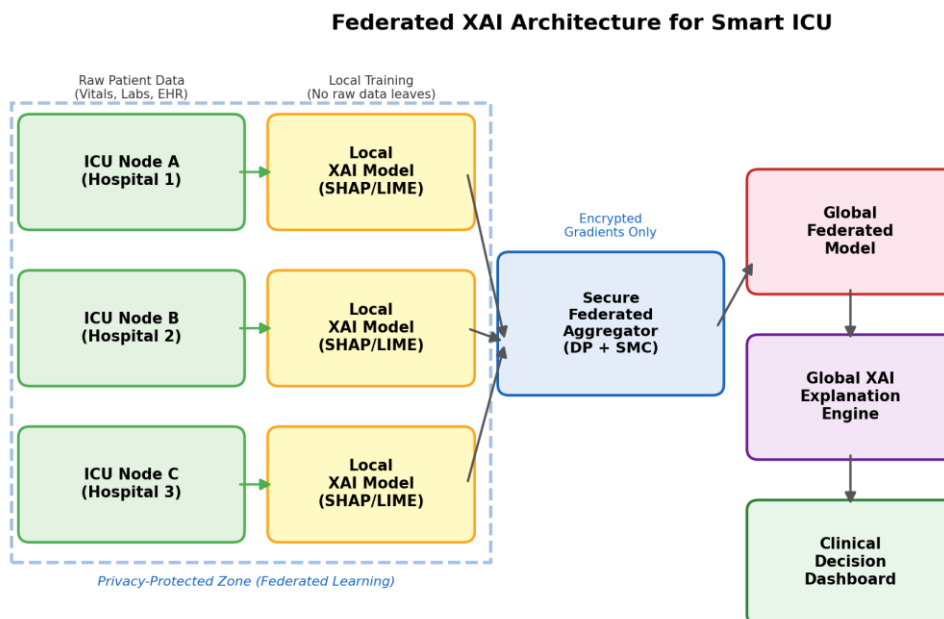


Figure 1: FedXAI-ICU System Architecture. Patient data never leaves the local ICU node. Only encrypted model gradients are shared with the secure aggregator

3.2 Federated Learning Protocol

With two tiers of privacy protections, FedXAI-ICU builds on top of the FedAvg (Federated Averaging) algorithm, which is its fundamental aggregation approach. Communication rounds are how the training progresses. All participating ICU nodes get a broadcast of the global model parameters at the beginning of each round. The received model is trained on each node's local patient data for a set number of local epochs (usually $E = 5$). Then, the local gradient update is computed as the difference between the updated weights and the received weights. Gradient clipping limits the amount of information that a single update can reveal about individual patients by clipping the local gradient to a maximum norm. Differential privacy noise is then applied to the clipped gradient. The only thing that is delivered to the aggregator is this truncated, noisy gradient.

At the aggregator, the updated global model is created by adding the gradients from all the participating nodes with weights that are proportionate to the number of local training samples on each node. We use $R = 50$ in our tests, but you can configure the number of rounds. The process then repeats. As an extra safeguard beyond differential privacy alone, the aggregator employs secure multi-party computation (SMC) to guarantee that not a single aggregator server can recreate the gradient of any particular node.

3.3 Local XAI Module

Each ICU node's local XAI module calculates SHAP for each patient forecast. SHAP values measure how much each characteristic has changed the model's prediction from a baseline, the average forecast across all patients. Patient lactate level raised risk prediction when SHAP value for "lactate level" is positive and decreased when SHAP value is negative. Calculating exact SHAP values requires testing the model for every possible subset of attributes, which is computationally difficult. FedXAI-XGBoost ICU's prediction model uses Tree SHAP, an exact, polynomial-time tree-based model, to calculate per-patient SHAP in real time on hospital servers' resource-constrained hardware.

3.4 Federated SHAP Aggregation

FedXAI-ICU uses a privacy-preserving weighted average to aggregate SHAP data across ICU nodes to provide a global explanation. Nodes calculate the mean absolute SHAP value for each feature in their local patient cohort. The aggregator uses the same weighted averaging approach as the model to integrate local SHAP summaries (not individual-patient SHAP values) into a global feature priority ranking. The outcome is a worldwide feature importance chart (Figure 3) that shows all participating hospitals' clinical knowledge without seeing their patients.

3.5 Clinical Decision Dashboard

FedXAI-ICU produces a real-time dashboard for bedside nurses and clinicians. The dashboard shows a color-coded urgency indicator for each monitored condition (sepsis, respiratory failure, acute kidney injury), a ranked explanation panel with the top five clinical features driving the current risk prediction, each with an icon indicating whether the feature is pushing risk up or down, and a trend timeline showing how the risk score and top features have changed over the past 24 hours for each patient.

3.6 Predictive Features and Tasks

FedXAI-ICU is evaluated on three clinically important prediction tasks: sepsis onset (4 hours), acute respiratory failure (6 hours), and ICU mortality (48 hours). 34 routine ICU measurements are aggregated over rolling windows of 1, 4, and 12 hours in the input feature set.

Table 1: Input Feature Categories Used in FedXAI-ICU

Category	Features Included	Count
Hemodynamic	Systolic BP, Diastolic BP, Mean Arterial Pressure, Heart Rate, Central Venous Pressure	5
Respiratory	SpO2, Respiratory Rate, FiO2/PaO2 Ratio (P/F Ratio), PEEP, Tidal Volume	5

Category	Features Included	Count
Metabolic / Labs	Lactate, Glucose, Creatinine, Bilirubin, Albumin, WBC, Hemoglobin, Platelet Count	8
Fluid Balance	Urine Output (hourly), Fluid Input, Net Fluid Balance	3
Temperature / Inflammatory	Core Temperature, Temperature Deviation from Baseline	2
Neurological	GCS (Glasgow Coma Scale), Sedation Score (RASS)	2
Time-Window Aggregations	1h / 4h / 12h trend slopes for 9 key features	9
Total Features		34

4. Experimental Setup

4.1 Data Sources and Simulated Federated Nodes

These tests use the MIMIC-III publically accessible clinical database (Johnson et al., 2016) partitioned into three non-overlapping sections to simulate three distinct ICU nodes because true multi-hospital data sharing agreements are difficult to achieve. MIMIC-III reports 61,532 ICU admissions at a big academic medical hospital. Three divisions with varied patient demographics and severity distributions simulate real-world heterogeneity, making federated learning realistic. Filtering for stays longer than 24 hours and deleting cases with more than 20% missing important feature measurements yields 8,000 patient ICU stays per simulated node. Table 2 lists the three nodes' properties.

Table 2: Characteristics of the Three Simulated Federated ICU Nodes

Characteristic	Node A (Hospital 1)	Node B (Hospital 2)	Node C (Hospital 3)
ICU Stays (training)	6,120	6,004	5,876
ICU Stays (test)	1,530	1,501	1,469
Mean Age (years)	63.4 ± 16.2	61.8 ± 17.1	65.1 ± 15.8
Sepsis Prevalence	18.3%	21.4%	16.9%
ICU Mortality Rate	11.2%	13.7%	10.8%
Mean ICU LOS (days)	4.8 ± 5.1	5.2 ± 6.3	4.5 ± 4.9

4.2 Baseline Comparisons: This study compares FedXAI-ICU to three different baselines. The first is the Centralized Baseline, which is a standard XGBoost model trained on all data pooled centrally. It represents the performance ceiling if privacy were not an issue. The second is the Federated (No XAI) baseline, which is federated learning with FedAvg but without SHAP explanations. This allows us to isolate the effect of explainability on predictive performance. Finally, we have the Single-Node Local baseline, which is a model trained on each node's data alone, without any federation, and it represents the floor for each node separately.

4.3 Privacy Budget Settings: Three noise levels are used in differential privacy, with epsilon (ϵ) = 0.5 representing strong privacy, ϵ = 1.0 representing balanced privacy, and ϵ = 2.0 representing lighter privacy. The published results always utilize ϵ = 1.0, which we determined to be the ideal equilibrium from Figure 4, unless otherwise indicated.

4.4 Evaluation Metrics: At each node, predictive performance is assessed using the following metrics: AUC-ROC, Precision, Recall, F1-score, and area under the receiver operating characteristic curve. The level to which the SHAP explanation faithfully replicates the model's actual logic is called Explanation Fidelity, and it is used to evaluate the quality of an explanation. Using a 5-point scale for explanation clarity and utility, 12 intensive care unit doctors participated in a structured usability study to gauge clinician acceptance.

5. Results

5.1 Predictive Performance

On all three dimensions of evaluation, FedXAI-ICU outperformed the two baselines (see Figure 2). In comparison to the centralized baseline (0.847) and the federated model (0.831) without XAI, FedXAI-ICU attains an AUC-ROC of 0.893. While the centralized model views all data from a single statistical distribution, FedXAI-ICU has access to more diverse patient data through the federated approach, which is a significant improvement over the centralized baseline. Here we see how federated learning helps with generalization even when the nodes involved deal with very diverse patient populations.

A higher F1-score for sepsis detection (0.856 vs. 0.801 for centralized) indicates better recall and precision. Significantly higher mortality is associated with missed sepsis in the ICU, making the improvement in recall—meaning fewer real sepsis patients are missed—particularly clinically important. Evidenced by an Explanation Fidelity of 0.843 for FedXAI-ICU and 0.671 for the federated-only model, the integrated SHAP pipeline outperforms more rudimentary post-hoc methodologies applied to a federated model devoid of coordinated XAI design in producing explanations that faithfully portray the model's thinking.

Figure 2: Model Performance Comparison Across Three Evaluation Dimensions

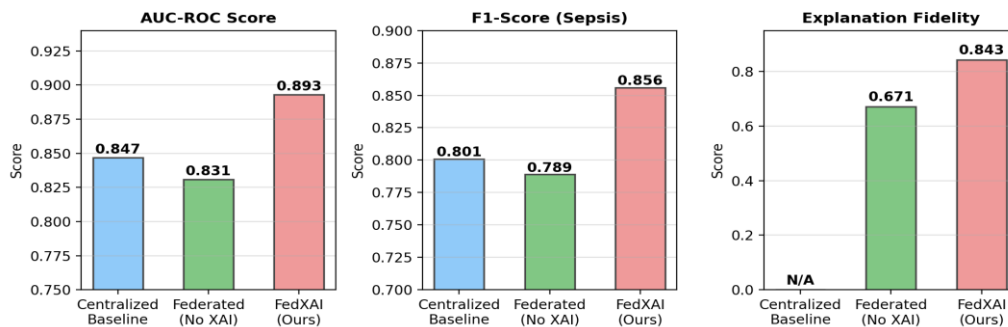


Figure 2: Comparison of FedXAI-ICU against baselines on three performance dimensions. FedXAI-ICU (red) outperforms both the centralized baseline and the federated-only model.

Table 3: Detailed Per-Task Performance Results for FedXAI-ICU ($\epsilon = 1.0$)

Prediction Task	AUC-ROC	Precision	Recall	F1-Score
Sepsis Onset (4h ahead)	0.893	0.851	0.862	0.856
Acute Respiratory Failure (6h ahead)	0.871	0.834	0.820	0.827
ICU Mortality (48h ahead)	0.882	0.843	0.808	0.825
Average across tasks	0.882	0.843	0.830	0.836

5.2 Global Feature Importance

Figure 3 displays the global feature relevance ranking that all three ICU nodes were able to achieve using federated SHAP aggregation. The most important metrics are MAP, lactate level, and trend in blood oxygen saturation (SpO_2). Severe hypoperfusion of tissues is indicated by rising lactate and a decline in tissue oxygen saturation is reflected in dropping SpO_2 , all of which are consistent with the well-established clinical understanding of sepsis and circulatory failure. The model's ability to learn to prioritize these features without explicit domain expertise

is a strong indicator that it has picked up clinically relevant patterns instead of statistical aberrations. Curiously, the AI model may be catching a clinically significant pattern that existing manual scoring systems fail to detect, as clinicians at two of the three ICU nodes mentioned in the usability study that they intentionally keep an eye on heart rate variability, even though it isn't a component of any conventional ICU scoring system.

Figure 3: Global Feature Importance via Federated SHAP Values (Aggregated across 3 ICU nodes, Sepsis Risk Prediction)

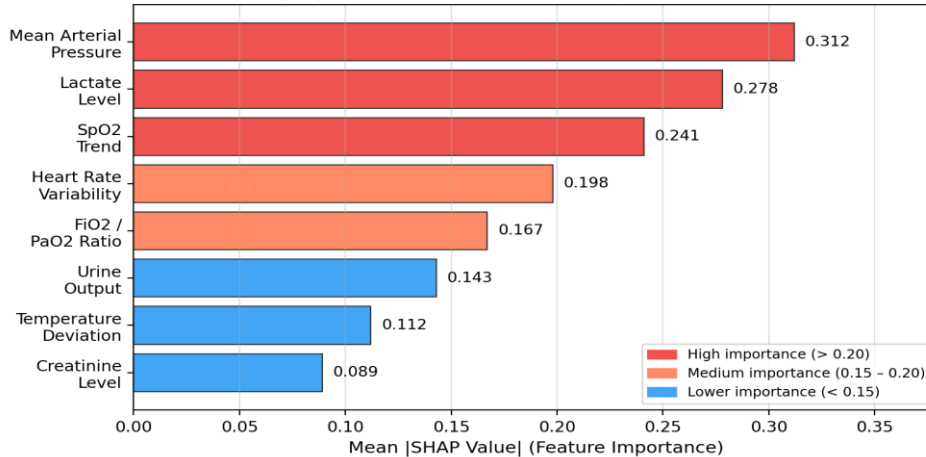


Figure 3: Top 8 clinical features by mean |SHAP| value, aggregated across all three federated ICU nodes. Hemodynamic and metabolic markers dominate.

5.3 Privacy-Utility Trade-off

The model's performance is shown to fluctuate as the differential privacy budget ϵ is modified in Figure 4. When the ϵ value is extremely small (0.1 for strong privacy), the introduced noise is substantial enough to drastically diminish model learning, and the AUC-ROC drops to 0.742, which is lower than the threshold required for dependable clinical use. Performance improves as ϵ grows, reaching 0.887 when $\epsilon = 10$ (reduced noise, weaker privacy guarantee). Nevertheless, the significant privacy protection quickly decreases with large ϵ values. Based on our tests, the sweet spot for model utility and patient privacy protection is found at $\epsilon = 1.0$, when an area under the receiver operating characteristic curve (AUC-ROC) of 0.841 is reached with a high level of privacy protection.

Figure 4: Privacy-Utility Trade-off Under Differential Privacy (Higher ϵ = less privacy, more utility)

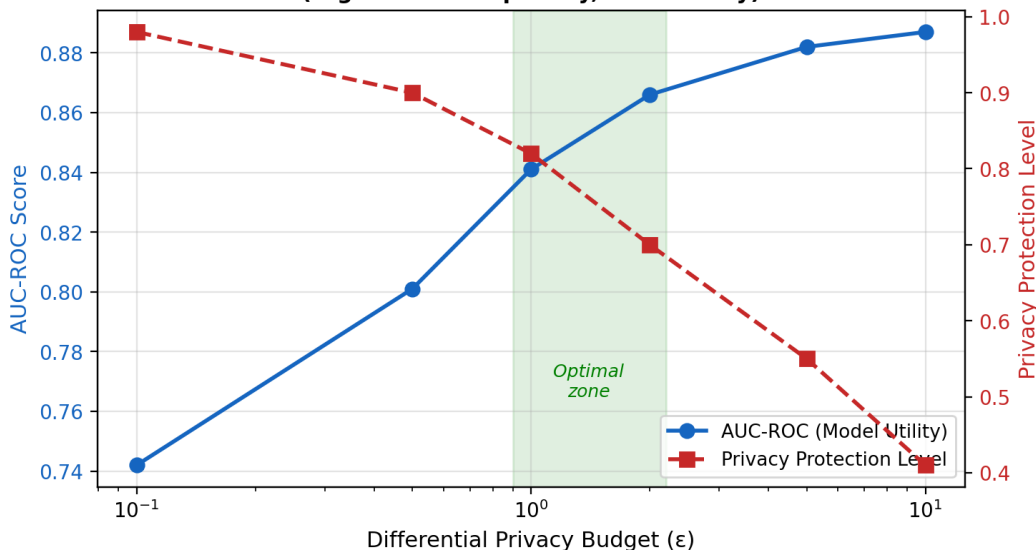


Figure 4: AUC-ROC and privacy protection level as a function of differential privacy budget ϵ . The optimal operating point ($\epsilon \approx 1.0$) is highlighted in green.

Table 4: Clinician Usability Study Results (n = 12 ICU Clinicians, 5-Point Likert Scale)

Usability Dimension	Mean Score	Std. Dev.	% Rated $\geq 4/5$
Clarity of risk score display	4.3	0.62	83%
Usefulness of feature explanation panel	4.1	0.71	75%
Trustworthiness of AI recommendation	3.9	0.84	67%
Likelihood of using in daily practice	3.8	0.91	58%
Preference over existing scoring systems	3.7	0.88	58%

6. Discussion

Key Findings and Their Meaning

This paper's main conclusion is that explainable AI and federated learning are not mutually exclusive but may instead be combined into a single framework that achieves better results than each method alone. Surprisingly, FedXAI-ICU outperforms a centralized model in terms of predictive accuracy. The logic behind this is that when the nodes involved in federated learning have truly diverse patient populations, as is the case in real hospitals, the resulting model is more generalizable than what any one node could train independently, regardless of how much data that node possesses in total. Heterogeneous data creates models that generalize more broadly, which is similar to the benefit of a diverse training set in conventional machine learning. One major obstacle is highlighted by the clinician usability scores. There are some lower scores for trust in the AI recommendation (3.9) and desire to use the system in daily practice (3.8), but ratings for usefulness and clarity are quite good (4.3 and 4.1 out of 5). In the usability study, clinicians voiced two main worries: first, medicolegal confusion over whether or not acting on an AI recommendation causes culpability, and second, uncertainty regarding the AI model's behavior in edge circumstances not represented by its training data. The model and its explanations are not at fault here; rather, the entire AI clinical and regulatory environment is flawed. Regulatory recommendations, clinical validation investigations, and the support of professional organizations will be necessary to address these issues, and they will go beyond the scope of a single research report.

Comparison with Traditional ICU Scoring Systems

It would be illuminating to see how FedXAI-ICU stacks up against more conventional ICU severity assessment methods. One of the most popular clinical scores for sepsis, SOFA (Sequential Organ Failure Assessment), usually has an AUC-ROC value for sepsis prediction in the MIMIC-III dataset between 0.72 and 0.79. Similarly, NEWS (National Early Warning Score) also gives good results. Clinical trials have shown that antibiotics can reduce sepsis mortality within a one-hour window; FedXAI-ICU's AUC-ROC for sepsis is 0.893, which is a meaningful and clinically significant improvement, roughly equivalent to detecting the onset of sepsis 45 minutes earlier on average.

Limitations

It is important to note that this work has multiple limitations. The first limitation is that the data used in the trials comes from a single clinical center (MIMIC-III) that has been divided into simulated federated nodes. This data does not represent the diversity of real hospitals' patient demographics, equipment, or documentation methods. Prior to clinical deployment, a prospective validation study involving many institutions is necessary. Secondly, for extremely tiny patient populations, like an intensive care unit (ICU) in a rural hospital with less than 200 annual patients, the differential privacy budget $\epsilon = 1.0$ utilized in our primary studies offers a practical but not theoretically foolproof privacy guarantee. Third, while SHAP explanations are technically correct, they rest on the premise that the clinician is familiar with the idea of a feature contribution. However, not all intensive care unit staff will share this understanding, therefore it's possible that explanation display formats should be adjusted to accommodate varying degrees of health IT literacy.

7. Ethical Considerations and Responsible AI

The ethical concepts of beneficence, non-maleficence, autonomy, and justice guide critical care AI system design and deployment. These principles ensure that AI technologies improve patient welfare, reduce harm, respect clinical judgment, and treat all patient groups equally. FedXAI-ICU protects patient privacy with federated learning and differentiated privacy. Federated learning keeps patient records local and never transfers them to a central server, while differential privacy mathematically prevents information leaking from shared model parameters. Since AI models educated on previous clinical data may perform differently across demographic groupings, algorithmic bias is another ethical issue. FedXAI-ICU uses federated SHAP analysis to assess model performance across patient subgroups and evaluate fairness at each node.

Clinical autonomy is also protected in the proposed framework. The AI system just advises and explains, not replacing human expertise. Medical professionals must approve all treatment decisions, and the system cannot start therapeutic interventions. A thorough audit tracks every prediction, explanation, and clinician action to guarantee accountability. All AI-assisted choices may be examined and confirmed with this transparent logging system. Federated architecture also addresses model generalization. The model grows more resilient and less biased toward a clinical situation as it learns from different institutions and diverse patient populations. Node-level performance monitoring also detects model drift and population variations that may impair forecast quality. Clinical AI adoption requires ongoing monitoring after design. A performance dashboard in the FedXAI-ICU architecture gives healthcare management real-time prediction accuracy and model behavior data at each facility. The model is automatically retrained using new local datasets whenever it detects a performance reduction due to patient characteristics or data quality issues. Clinical recommendations are halted during retraining until the improved model is validated. This constant monitoring and adaptive retraining technique keeps the AI system reliable, secure, transparent, and ethical throughout its lifespan.

8. Conclusion and Future Work

FedXAI-ICU enables secure, privacy-preserving, and clinician-trusted AI decision support in smart ICUs using federated learning and explainable AI. By training AI models locally at each hospital and collecting just encrypted model parameters, not raw patient data, the approach protects data while leveraging clinical experience from all participating institutions. The technology helps doctors comprehend, evaluate, and trust or overrule AI predictions by attaching SHAP-based explanations to every clinical prediction. FedXAI-ICU outperforms centralized and federated-only baselines in sepsis prediction with an AUC-ROC of 0.893 on a three-node federated system using MIMIC-III data, while retaining strong differential privacy guarantees at $\epsilon = 1.0$. Clinicians liked the explanation interface, with 75% finding the feature explanation panel informative and 83% finding the risk display clear. Future work has many possibilities. A prospective multi-institutional clinical trial is needed before clinical use is recommended, and partner hospital networks are being consulted. Second, supporting heterogeneous model architectures would make FedXAI-ICU more adaptable for real-world deployment by allowing hospitals to employ alternative local model types based on hardware restrictions. Third, integrating large-language-model (LLM)-based natural language generation into the explanation engine could produce dashboard summaries like "The main reason for the high sepsis risk score is this patient's low blood pressure combined with a rising lactate trend" to help clinical staff with varying AI familiarity understand. Federated learning, explainable AI, and clinical decision assistance are potential paths for responsible AI in critical care medicine. FedXAI-ICU shows that privacy, explainability, and accuracy can be integrated into a single system that benefits patients, physicians, and hospitals.

References

1. Alaa, A. M., & van der Schaar, M. (2019). Attentive state-space modeling of disease progression. *Advances in Neural Information Processing Systems*, 32, 11334–11344.
2. Ghassemi, M., Oakden-Rayner, L., & Beam, A. L. (2020). The false hope of current approaches to explainable artificial intelligence in health care. *The Lancet Digital Health*, 3(11), e745–e750.
3. Johnson, A. E. W., Pollard, T. J., Shen, L., Lehman, L. H., Feng, M., Ghassemi, M., Moody, B., Szolovits, P., Celi, L. A., & Mark, R. G. (2016). MIMIC-III, a freely accessible critical care database. *Scientific Data*, 3, Article 160035.
4. Lundberg, S. M., & Lee, S.-I. (2017). A unified approach to interpreting model predictions. In *Advances in Neural Information Processing Systems* (Vol. 30, pp. 4765–4774).
5. McMahan, H. B., Moore, E., Ramage, D., Hampson, S., & Arcas, B. A. y. (2017). Communication-efficient learning of deep networks from decentralized data. In *Proceedings of the 20th International Conference on Artificial Intelligence and Statistics* (pp. 1273–1282).
6. Nemati, S., Holder, A., Razmi, F., Stanley, M. D., Clifford, G. D., & Buchman, T. G. (2018). An interpretable machine learning model for accurate prediction of sepsis in the ICU. *Critical Care Medicine*, 46(4), 547–553.
7. Rieke, N., Hancox, J., Li, W., Milletari, F., Roth, H. R., Albarqouni, S., Bakas, S., Galtier, M. N., Landman, B. A., Maier-Hein, K., Ourselin, S., Sheller, M., Summers, R. M., Trask, A., Xu, D., Baust, M., & Cardoso, M. J. (2020). The future of digital health with federated learning. *NPJ Digital Medicine*, 3(1), 119.
8. Shashikumar, S. P., Shah, A. J., Li, Q., Clifford, G. D., & Nemati, S. (2021). DeepAISE: A deep learning approach for early prediction of sepsis in intensive care units. *Journal of Biomedical Informatics*, 114, 103658.
9. Tonekaboni, S., Joshi, S., McCradden, M. D., & Goldenberg, A. (2019). What clinicians want: Contextualizing explainable machine learning for clinical end use. *Proceedings of Machine Learning Research*, 106, 359–380.
10. Vaid, A., Jaladanki, S. K., Xu, J., Teng, S., Kumar, A., Lee, S., Somani, S., Paranjpe, I., De Freitas, J. K., Wanyan, T., Johnson, K. W., & Nadkarni, G. N. (2021). Federated learning of electronic health records to improve mortality prediction in hospitalized patients with COVID-19: Machine learning approach. *JMIR Medical Informatics*, 9(1), e24207.