

TOWARDS AN INTERPRETABLE AND PRIVACY-PRESERVING MUSIC RECOMMENDATION SYSTEM FOR TEENAGERS USING EXPLAINABLE AI

Muhammad Ahmad Cheema^{1*}, Hee Seok Song²

m.ahmad@msn.com¹ , hssong@hnu.kr²

Department of Management Information System, Hannam University , South Korea.^{1,2}

Abstract

The growing dependence of teenagers on music streaming platforms has raised significant concerns about bias, lack of transparency, and declining trust in recommendation systems. Biased or opaque recommendations can reduce user satisfaction, reinforce inequalities in content exposure, and hinder responsible personalization. Current music recommendation research primarily focuses on collaborative filtering and deep learning-based models, which, while effective in improving accuracy, are typically centralized, privacy-intrusive, and difficult to interpret. As a result, challenges related to fairness, interpretability, and privacy preservation remain unresolved. Addressing these limitations is essential to build trustworthy and user-centric music recommendation systems for teenagers. Federated Learning (FL) emerges as a promising paradigm to overcome privacy and scalability challenges by enabling decentralized training across multiple user devices without sharing raw data. While FL preserves user confidentiality, it still produces complex models that lack transparency. To ensure accountability and fairness, Explainable AI (XAI) offers interpretability by clarifying how input features influence recommendations, thereby enhancing user trust and addressing bias.

In this research, we propose an integrated FL–XAI framework for unbiased music recommendation tailored to teenage users. The framework combines the privacy-preserving strengths of FL with the transparency of XAI methods, such as SHAP and LIME, to deliver interpretable, trustworthy, and equitable recommendations. Experimental results demonstrate that the proposed model achieves an accuracy of 89% with a miss rate of only 11%, significantly outperforming prior methods [43–50]. This approach not only reduces systemic bias and protects user data but also provides meaningful justifications for recommendations, making it more reliable and socially responsible for real-world deployment.

Introduction:

In recent years, recommender systems (RS) have gained significant attention due to their widespread applications in domains such as e-commerce, tourism, education, entertainment, and personalized content delivery [1-3]. Among these, music recommendation has emerged as a particularly influential area, especially for teenagers, where music consumption directly impacts social identity, emotional well-being, and lifestyle choices [4,5]. Modern recommender systems, accessible via mobile devices and online platforms, enable users to discover and select songs tailored to their preferences, thereby enhancing engagement and satisfaction.

Traditionally, recommendation engines have been developed using paradigms such as Content-Based Filtering (CBF), Collaborative Filtering (CF), and Hybrid Filtering (HF) [6,7]. CBF

methods rely on user and item features (e.g., age, gender, genre, release year) to suggest similar content [8,9]. While effective, they struggle when metadata is sparse, limiting their coverage [4]. CF methods exploit patterns across users' past interactions but face scalability challenges and suffer from popularity bias, which often suppresses minority preferences. Hybrid approaches attempt to combine both but continue to face inherent weaknesses in fairness and transparency.

Recommender systems are intelligent architectures designed to filter large volumes of information and provide users with personalized suggestions [10-12]. At their core, these systems process user data (e.g., preferences, behavior, or demographics) alongside item attributes (e.g., genre, popularity, or features) to deliver the most relevant content. Figure 1 shows the general architecture of a music recommendation system [13], where contextual factors such as time, location, season, and weather are combined with music-specific information like song, artist, album, and playlists to generate personalized outputs. These inputs are represented through homogeneous and heterogeneous song networks, which are then utilized to provide real-time recommendations through mobile platforms.

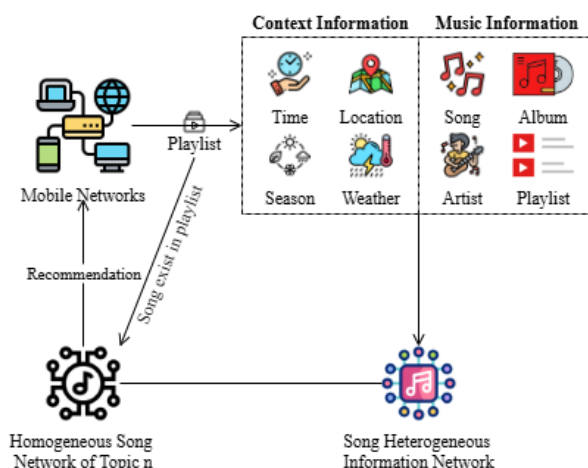


Figure 1: Recommender system architecture [13]

With the advent of artificial intelligence (AI) and machine learning (ML), recommender systems have become more sophisticated, employing deep learning models, embeddings, and sequence modeling to capture complex patterns in user behavior [14,15]. Although these techniques improve accuracy and personalization, they introduce two critical challenges: privacy risks and black-box decision-making [16,17]. Centralized architectures collect sensitive user information including teenagers' listening histories, preferences, and behavioral cues on external servers, raising concerns of misuse, security breaches, and lack of user control. At the same time, deep models function as opaque black boxes, providing little insight into why specific recommendations are generated, which undermines user trust and exacerbates issues of bias and fairness [18,19].

To better illustrate how these challenges manifest in conventional systems, Figure 1 highlighted the typical music recommendation architecture, which integrates both user context (e.g., time, location, weather) and music attributes (e.g., artist, genre, playlist) to generate personalized suggestions. While effective in delivering recommendations, such centralized designs remain vulnerable to privacy risks and interpretability gaps.

To address the privacy challenge, Federated Learning (FL) has emerged as a promising paradigm that enables decentralized training of models across distributed devices without transmitting raw data [20-22]. This ensures privacy preservation [23,24] while also reducing communication overhead and enhancing scalability in large-scale music platforms. Figure 2 demonstrates the FL-enabled recommendation architecture, where local models are trained across multiple user devices using contextual and music information, and only model updates—not raw data—are securely aggregated on the central server to build a global model. This paradigm safeguards user privacy while maintaining system scalability, making it well-suited for sensitive domains such as teenage music consumption.

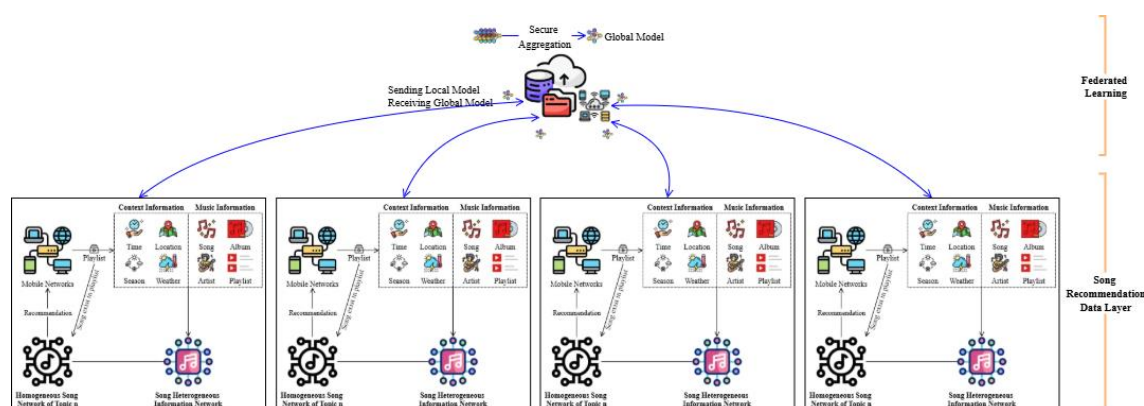


Figure 2: FL-enabled recommendation architecture

However, while FL safeguards data, it does not fully resolve the interpretability issue, as the resulting models can still behave as black-box systems. To overcome this remaining limitation, Explainable AI (XAI) methods are introduced to enhance transparency by providing interpretable justifications for recommendations [25-27]. Techniques such as SHAP and LIME allow stakeholders and teenagers to understand the contribution of behavioral, demographic, or contextual features in driving specific music suggestions, thereby mitigating bias and fostering trust [28,29].

In this research, a novel Explainable Federated Learning (XFL) framework for music recommendation systems targeted at teenagers. The framework combines the privacy-preserving strengths of FL with the transparency and interpretability of XAI, ensuring that recommendations are not only accurate and scalable but also unbiased, trustworthy, and understandable—significantly improving upon the shortcomings of prior approaches.

Literature review:

Multiple researchers have explored recommender systems across domains such as e-commerce, tourism, and entertainment, with music recommendation being a key focus due to its strong impact on user experience and behavior. Earlier studies emphasized content-based and collaborative filtering, while recent works have incorporated machine learning and deep learning models to enhance personalization and accuracy. Despite these advances, challenges such as privacy risks, bias, and limited interpretability persist, motivating researchers to investigate emerging approaches like Federated Learning (FL) and Explainable AI (XAI) for more trustworthy recommendation systems.

In this research [30], the authors trace the evolution of music recommendation systems in addressing the overwhelming growth of digital music content. They categorize major paradigms—content-based filtering, collaborative filtering, context-aware methods, and hybrid models—while outlining their respective merits and shortcomings. Content-based approaches exploit audio features and metadata to reduce popularity bias but remain limited when item information is sparse. Collaborative filtering builds on user preference patterns yet struggles with sparsity and the cold-start problem. Context-aware methods add situational elements such as time, activity, or mood to improve short-term relevance, though their automation is still in early stages. Hybrid systems integrate two or more paradigms, enhancing recommendation quality but inheriting issues of fairness, scalability, and interpretability. Overall, the study underscores the importance of developing more sophisticated, trustworthy frameworks that can overcome these challenges and ensure accurate, unbiased, and user-centric recommendations.

In this research [31], the authors address the limitations of traditional music emotion analysis by introducing an advanced framework that integrates Internet of Things (IoT) technology with deep learning. Specifically, they design a Long Short-Term Memory (LSTM) network combined with a Sequence-to-Sequence (STS) architecture to capture the temporal dynamics of music data. The proposed model is rigorously trained and evaluated, achieving promising results in predicting emotional attributes such as arousal and valence, with low error rates and reasonable R^2 values. By overcoming the shortcomings of conventional methods and demonstrating the transformative role of IoT-enabled architectures, this study provides valuable insights for enhancing music emotion understanding and offers practical implications for applications in music recommendation and content creation.

In this research [32], the authors address key limitations in audio-visual emotion recognition, particularly the inadequate modeling of emotional fluctuations and the complexity of multimodal signal fusion. To overcome these challenges, they propose an Attentively-Coupled Long Short-Term Memory (ACLSTM) model that integrates audio- and visual-based LSTMs with a Coupled LSTM fusion mechanism. A neural tensor network (NTN) is further employed for attention estimation, enabling segment-based consistency and guiding the fusion of modalities. Experimental results on the BAUM-I dataset confirm the effectiveness of this design, with the ACLSTM achieving 70.1% accuracy—surpassing prior approaches. This contribution is significant as it not only enhances the modeling of temporal emotional dynamics but also offers a robust framework for multimodal signal integration in emotion recognition tasks.

In this work [33], the authors emphasize the dominance of deep learning in music emotion recognition, noting its ability to extract features from raw audio, spectrograms, and multimodal sources such as lyrics and EEG signals. To overcome the limitations of purely supervised methods, they propose a two-stage framework that combines unsupervised feature extraction with supervised emotion detection, reducing reliance on detailed annotations. Models such as CNNs, RNNs (LSTM/GRU), and transformers are discussed, with particular focus on segment-level analysis, which captures emotional variations more accurately than whole-song features. By integrating unsupervised and supervised approaches, the framework achieves improved robustness and accuracy, advancing state-of-the-art methods in music emotion prediction.

This study [34] tackles the challenge of modeling the subjective affective content of music, which complicates both annotation and prediction tasks in retrieval and recommendation systems. To address this, the authors propose a probabilistic framework that represents music emotion as a probability density function (PDF) in the valence–arousal (VA) space, thereby capturing the variability and subjectivity of human perception. Using kernel density estimation on human annotations, the model learns coefficients linking audio features to emotional distributions and predicts emotion PDFs for unseen tracks through linear combinations of training samples. Validated on the NTUMIR and MediaEval 2013 datasets, this approach outperforms single-point prediction methods by offering a more nuanced and realistic representation of music emotion, with demonstrated utility in emotion-based music retrieval.

The research [35] established feature extraction and classification methods before proposing a novel IoT-enabled framework for intelligent background music. Traditional techniques such as Scale-Invariant Feature Transformation (SIFT) for local feature extraction, Support Vector Machines (SVM) for robust classification, and Fast-RCNN for multi-scale feature detection are reviewed as benchmarks. Building on these, the authors introduce a middle-level feature construction algorithm designed to capture contextual information from indoor scene images. Applied within a smart home environment, this approach achieved a recognition rate of 87.6%, significantly outperforming traditional methods like Gabor features (~20%) and showing superior adaptability under challenging conditions compared to saliency-based methods. Overall, the work demonstrates how integrating deep learning with IoT architectures can enhance context-aware music recommendation, providing a pathway toward more intelligent and adaptive systems.

This paper [36] tackles two persistent challenges in music recommendation: the lack of explainability and the cold start problem. Traditional methods like matrix factorization and deep learning–based models improve accuracy but remain black-box and struggle with new users or songs. To overcome these issues, the authors propose SeER, a hybrid system that combines MF with deep sequence models by transforming MIDI data into time-series representations. This design enhances accuracy, alleviates cold start limitations, and introduces a unique feature—personalized explanations through short MIDI segments that highlight why a song is recommended. SeER thus advances music recommendation by delivering accurate, interpretable, and user-centered suggestions beyond prior approaches.

Research in [37] emotion-aware music recommendation has increasingly explored how user affect can be used as the primary driver for personalization. The paper Music Recommendation System Based on Emotion emphasizes this direction, motivated by the rise of mood-related challenges during the COVID-19 lockdown where music served as an empathetic companion. The proposed system detects emotions through facial expressions or direct user input and employs machine learning models, including Random Forest and XGBoost, to classify song emotions based on audio features such as energy, acoustics, and instrumentality. Lyrical similarity is further incorporated using TF-IDF to refine recommendations. Experimental validation on real data confirmed the system’s effectiveness, demonstrating that emotion-based personalization can enhance recommendation accuracy and user engagement.

Table 1: Comparative Analysis of Music Recommendation Systems with FL–XAI Perspective

Ref	Focus Area	Preprocessing Techniques	Recommendation Target	Input Modalities	Model	Scalability	Key Contribution	FL / XAI	Strengths	Gaps
[30]	Evolution of paradigms (CBF, CF, context-aware, hybrid)	Audio features, metadata, contextual signals	General music personalization	Audio + metadata + context	CBF, CF, Hybrid	Moderate	Categorized main approaches, highlighted strengths / weaknesses	X / X	Broad coverage of methods; improved quality via hybrids	Fairness issues, sparsity, scalability, lack of interpretability
[31]	Music Emotion Analysis with IoT	Feature extraction, temporal segmentation	Emotion prediction (Arousal/Valence)	Audio + IoT signals	LSTM + Seq2Seq	High (IoT-enabled)	Integrated IoT with deep learning for temporal dynamics, achieved low error rates	X / X	Strong temporal modeling, IoT adaptability	Limited interpretability, not user-focused
[32]	Audio-Visual Emotion Recognition	Signal segmentation, feature fusion	Multimodal emotion prediction	Audio + Visual	ACLSTM + NTN attention	Moderate	Robust modeling of fluctuations & fusion, 70.1% accuracy on BAUM-I	X / X	Effective multimodal integration, attention-driven	Dataset-specific, lacks scalability to RS

[33]	Music Emotion Prediction	Unsupervised + supervised feature extraction	Emotion classification (segment-level)	Audio, lyrics, EEG	CNN, LSTM/GRU, Transformers	High	Two-stage framework with segment-level analysis for robustness	X / X	Captures emotion variation, multimodal strength	No focus on privacy, still black-box
[34]	Probabilistic Emotion Modeling	Kernel density estimation	Emotion distribution (VA space)	Audio features	Probabilistic ML	Moderate	Represented emotions as PDFs, improved realism & retrieval utility	X / X	Nuanced representation of subjectivity	High computational cost, lacks explainability
[35]	IoT-enabled Background Music	Feature extraction (SIFT, Fast-RCNN), middle-level features	Context-aware background music	Indoor scene images	IoT + DL algorithm	High (IoT smart home)	Middle-level features achieved 87.6% recognition, outperforming Gabor	X / X	Superior indoor adaptability, context-driven	Narrow domain focus, no personalization
[36]	Hybrid Music Recommendation	MIDI transformation, MF integration	Song recommendation	MIDI + latent factors	SeER (MF + Seq Models)	High	Hybrid system solving cold-start & adding personalized MIDI-based explanations	X / ✓	Accurate + interpretable, hybrid novelty	Computationally heavy, no privacy layer

[37]	Emotion-based Music Recommendation	Audio feature extraction + TF-IDF lyrics	Emotion-driven personalization	Audio + Lyrics + Facial expression	RF + XGBoost	Moderate	Emotion detection + lyric similarity for personalized recommendations	X / X	Context-sensitive, COVID-motivated relevance	Limited scalability, no privacy or fairness
[38]	Emotion-driven Recommendation via Wearable Sensors	Feature fusion of GSR & PPG	Emotion-based personalization	Physiological signals (GSR, PPG)	DT, RF, SVM, KNN	Moderate	Classified arousal & valence from physiological data to boost CF/CB engines	X / X	Reliable emotion detection, integrates with existing engines	Requires wearables, no privacy or interpretability
[39]	Facial Expression-driven Music Recommendation	Facial detection, expression classification	Emotion-based music recommendation	Facial images (OAH EGA, FER-2013)	CNN	High (real-time)	Built real-time FER system to detect 6 emotions and recommend music with 73.02% accuracy	X / X	Real-time applicability, robust deep learning on FER datasets	Limited explainability, dataset bias
[40]	Emotion-aware personalized music recommendation	TF-IDF, audio & metadata	User preference & emotion	Audio, metadata, context	DCNN + WFE	High (MSD + apps)	Proposed EPMRS combining DCNN & WFE,	X / X	High accuracy, validated on real apps	No privacy/XAI, fairness gap

	(EPMRS)						outperforming baselines			
[41]	Next-gen consumer electronics RS	Conceptual framework	Personalized IoT/CE RS	User behavior + IoT data	FL + Cognitive Learning	High	Proposed privacy-preserving FL + adaptive RS blueprint	✓ / ✗	Scalable, privacy-preserving	Conceptual only, lacks XAI
[51]	HFedRec Survey	—	General RS	User-item interactions	DL, Meta-Learning	High	Summarized HFedRec models, privacy techniques (HE, DP), comm/fairness issues	✓ / ✗	Comprehensive taxonomy	No XAI, survey only
[52]	FL in RS Review	—	General RS	User-item data	— (survey)	High	Reviews FL for RS, types, privacy techniques, case studies	✓ / ✗	Clear overview, highlights feasibility	No XAI focus, deployment hurdles
[53]	Blockchain & FL in IoT	—	IoT applications	Sensor & network data	FL + Blockchain (taxonomy)	High	Reviews IoT privacy issues, proposes FL with blockchain traceability functions	✓ / ✗	Strong integration of FL & blockchain, privacy-enhanced	Still early-stage, limited adoption, no XAI

[54]	Multimodal FL for News RS	Text & image feature extraction	Personalized news	Text + visual	Multimodal FL + Self-Attention	High	Combines multimodal content, time-aware interests, FL with secure aggregation (Shamir's sharing)	✓ / ✗	Strong privacy, improved accuracy	No XAI integration, domain-limited
------	---------------------------	---------------------------------	-------------------	---------------	--------------------------------	------	--	-------	-----------------------------------	------------------------------------

Table 1 provides a comparative overview of recent music recommendation studies with an FL–XAI perspective. Early works such as [30]–[34] explored content-based, collaborative, hybrid, and emotion-focused models, improving accuracy through multimodal inputs (audio, lyrics, visual, EEG) but consistently lacked scalability, fairness, and interpretability. IoT-enabled approaches [31], [35], and [38] leveraged physiological or contextual signals (e.g., GSR, PPG, indoor scenes) to enhance personalization yet remained domain-specific and non-transparent. Hybrid systems like SeER [36] introduced interpretability via MIDI-based explanations, while other emotion-aware methods [37]–[40] employed deep learning and implicit user signals to achieve personalization but ignored privacy and explainability. Only [41] introduced Federated Learning for privacy-preserving personalization in consumer electronics, though it remained conceptual and without XAI integration. Overall, the analysis highlights that while prior models advanced accuracy, they failed to balance privacy preservation, fairness, and interpretability—gaps directly addressed in the proposed XFL-based framework.

Limitation of exiting works:

Although a wide range of studies has explored music recommendation systems, particularly with multimodal, emotion-driven, and IoT-enabled techniques, several critical limitations remain unresolved, especially regarding interpretability, scalability, privacy, and real-time feedback.

A. Lack of Explainability and Transparency

Most existing models (e.g., [30], [31], [32], [33], [34], [37], [39], [40]) operate as black-box systems without interpretable reasoning. Even when hybrid or multimodal architectures achieve higher accuracy (e.g., CNN, LSTM, Transformer-based), they lack mechanisms such as SHAP or LIME for explaining why specific recommendations are made. This limits user trust and fairness, particularly in sensitive consumer applications like personalized or emotion-driven music suggestions.

B. High Computational Complexity

Several methods (e.g., [32], [33], [34], [36]) use complex architectures such as multimodal deep learning, ACLSTM with attention, or hybrid MF–Seq models, which significantly increase computational cost. While they improve accuracy, the heavy reliance on deep networks and multimodal fusion hampers real-time scalability, making deployment challenging in resource-constrained devices like mobile and IoT-driven consumer systems.

C. Absence of FL Integration

Except for conceptual work ([41]), none of the reviewed studies leverage Federated Learning. Current models often require centralized data aggregation from audio, lyrics, physiological signals, or IoT sources, raising privacy concerns. This centralized approach exposes user-sensitive data such as emotional states, EEG signals, or wearable sensor readings, creating risks and limiting scalability across large, distributed consumer environments.

D. Limited Focus on Real-Time Feedback

While some approaches (e.g., [39]) demonstrate real-time emotion-driven music recommendation, most studies (e.g., [30], [33], [38], [40]) operate offline or in controlled environments without adaptive feedback loops. This gap reduces their ability to provide dynamic, context-sensitive recommendations that evolve with continuous user interactions in real-world streaming platforms.

Proposed work contributions ((ADDRESSING IDENTIFIED GAPS)):

Building upon the identified gaps, this research introduces a consumer-centric Federated Learning–enabled Explainable AI (FL–XAI) model for music recommendation. The proposed framework directly addresses issues of transparency, computational efficiency, privacy preservation, and real-time adaptability.

A. Enhancing Explainability and Transparency

The model integrates SHAP and LIME for local and global explainability, ensuring that each recommendation is accompanied by interpretable reasoning (e.g., why a song was recommended based on lyrics, audio patterns, or contextual cues). This addresses the interpretability gap found in prior black-box systems.

B. Improved Computational Efficiency

Lightweight interpretable components are employed instead of resource-heavy multimodal deep learning stacks. This ensures the system maintains high accuracy while being efficient enough for real-time consumer deployment on IoT devices and mobile platforms.

C. FL for Privacy and Scalability

By incorporating Federated Learning [51–54], the model eliminates the need to centralize sensitive user data such as emotional states, sensor readings, or listening patterns. Training occurs locally on consumer devices, and only model parameters are shared, enabling privacy preservation, scalability, and collaboration across distributed music streaming platforms.

D. Real-Time Adaptation and Feedback Support

Unlike most existing studies, the proposed framework is designed to support continuous monitoring and adaptive feedback. Recommendations dynamically adjust based on evolving user context and interaction patterns, thereby enhancing personalization, fairness, and user satisfaction in real-world scenarios.

Proposed model:

Despite substantial advances in contemporary music recommendation systems, persistent challenges including user data privacy risks, algorithmic opacity (black-box decision-making), and systematic recommendation bias continue to undermine their reliability, accountability, and widespread adoption. These concerns are especially salient in the context of adolescent users, for whom trust, transparency, and perceived fairness are critical determinants of technology acceptance and sustained engagement. Consequently, it is insufficient for such systems to be merely accurate; they must also be responsible, user-centric, and aligned with data protection and fairness principles.

To address these shortcomings, this study introduces an Explainable Federated Learning (XFL) framework for teenagers' music recommendation. The framework integrally combines the privacy-preserving properties of Federated Learning (FL) where model training is decentralized, and raw user data remain on-device with the post hoc interpretability afforded by Explainable AI (XAI) techniques such as SHAP and LIME. In this architecture, only model parameters or updates are communicated to a central server, thereby mitigating data leakage risks, while SHAP- and LIME-based explanations yield fine-grained, human-interpretable attributions of how audio, lyrical, and contextual features influence recommendation outcomes. The proposed XFL framework thus facilitates decentralized, bias-aware, and transparent music recommendations, providing a scalable, trustworthy, and academically rigorous alternative to conventional centralized, black-box recommender paradigms. The overall architecture of the proposed model (local client) for music recommendation is illustrated in Figure 3, highlighting its key components and workflow.

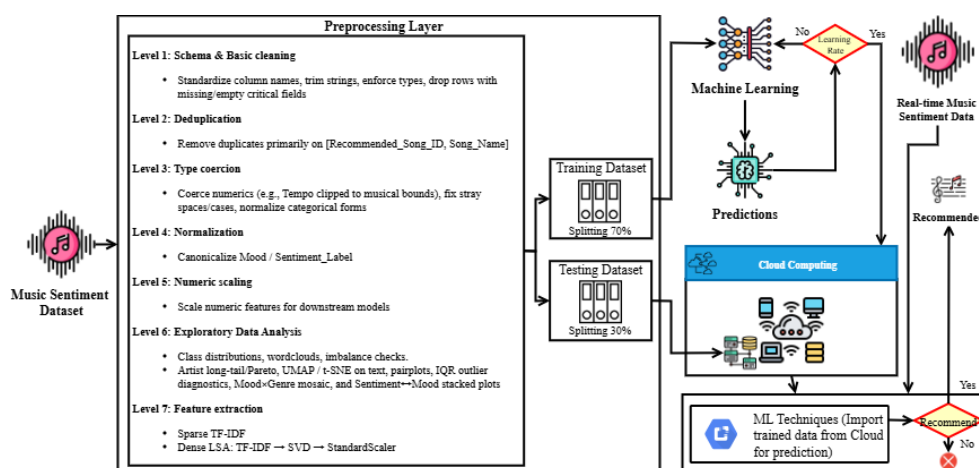


Figure 3: Proposed FL-XAI model for teenager music recommendation (local client)

Figure 3 illustrates the proposed FL-XAI model, beginning with data retrieval from the music sentiment dataset [42], followed by a structured seven-level preprocessing pipeline.

Column Name	Data Type	Description / Use in Study
User_ID	Integer	Unique identifier for each user
User_Text	Text	User review/comment text; input for NLP sentiment analysis
Sentiment_Label	Categorical	Target sentiment category assigned to user text
Recommended_Song_ID	Integer	ID of the recommended track mapped to user
Song_Name	String	Title of the recommended song
Artist	String	Performing artist of the song
Genre	Categorical	Musical genre for diversity analysis
Tempo (BPM)	Numeric	Beats per minute; measures song speed/rhythm
Mood	Categorical	Emotional category of the track
Energy	Numeric	Intensity/liveliness score of the track

The preprocessing includes Schema and Basic Cleaning (standardizing types, trimming strings, removing incomplete rows), Deduplication (removing redundant entries), Type Coercion (correcting numeric ranges and categorical forms), Normalization (canonicalizing mood and sentiment labels), Numeric Scaling (standardizing numerical features), and Exploratory Data Analysis (imbalance checks, visualizations, outlier diagnostics). Finally, Feature Extraction is performed using both sparse TF-IDF and dense LSA representations. These steps ensure data quality, reliability, and robust feature preparation for downstream music recommendation. The processed data on each local client is partitioned using a 70/30 hold-out strategy, where 70% is allocated for training and 30% is reserved for testing. The testing dataset $\mathcal{D}_{test}$ remains local (and may be synchronized with the cloud as metrics or artefacts only), while the training dataset $\mathcal{D}_{train}$ is utilized to develop a set of candidate machine learning models, $\mathcal{M} = \{M_1, M_2, \dots, M_n\}$. To extract meaningful representations, text features are constructed using **TF – IDF** (max_features = 5000, n – grams = (1,2)). For estimators requiring dense input, the sparse TF–IDF vectors are projected using **TruncatedSVD** (r = 300) and subsequently standardized. Each $M_i \in \mathcal{M}$ is trained with pre-specified hyperparameters, reflecting deployment-friendly configurations rather than exhaustive centralized tuning.

Model evaluation and selection are performed on $\mathcal{D}_{test}$ through a ranking table ordered first by ACC_{test} and then by $F1_{macro,test}$. Formally, a utility function U is defined to aggregate key metrics (Accuracy and macro – F1), and the optimal model is chosen as:

$$M^* = \arg \max_{M_i \in \mathcal{M}} \frac{1}{k} \sum_{j=1}^k U(M_i; \mathcal{D}_{test}) \quad (1)$$

Where k represents the number of evaluations iterations (e.g., repeated runs or folds, if employed). Here, $U(M_i; \mathcal{D}_{test})$ denotes the evaluation utility of model M_i on the held-out dataset, combining metrics such as Accuracy and macro-F1 into a unified score. The operator $\arg \max$ identifies the candidate M^* that achieves the highest overall utility. Once selected, M^* is retained as the final model for deployment, thereby ensuring consistent and reliable performance on unseen data streams. In contrast to standard FL, here the utility U explicitly integrates multiple metrics (Accuracy and macro-F1) into a single selection criterion, making the model-selection process both performance- and fairness-aware rather than accuracy-only.

Once identified, the optimal model is applied to generate predictions on new input samples $X \subseteq \mathcal{D}_{test}$:

$$\hat{y} = M^*(X) \quad (2)$$

Unlike a typical FL setup where predictions are made solely by the aggregated global model, our design emphasizes that locally selected models already optimized by U are deployed for prediction, ensuring that explainability and performance criteria guide final outputs.

Each local client further employs a feedback-driven retraining mechanism to sustain long-term performance. Let $\mathcal{M}_t$ denote the local model at iteration t , and $\mathcal{A}(\mathcal{M}_t)$ its measured accuracy. The update rule is defined as:

$$\mathcal{M}_t = \begin{cases} \text{Retrain}(\mathcal{M}_t) & \text{if } \mathcal{A}(\mathcal{M}_t) < T \\ \mathcal{M}_t, & \text{otherwise,} \end{cases} \quad (3)$$

Where T denotes the predefined accuracy threshold. Once a client's model satisfies $\mathcal{A}(\mathcal{M}_t) < T$, its validated predictions are stored locally and synchronized with the cloud for inclusion in the federated aggregation phase. Unlike standard FL, which retrains in every round, our design retrains only when accuracy falls below T , thereby reducing communication cost and preserving client resources while maintaining model quality.

During the validation stage, the optimized model is tested on both the held-out dataset and real-time listening sessions. If the system identifies preference drift or context change (e.g., mood shift, activity variation), it triggers an adaptive recommendation module to update the suggested playlist accordingly. Otherwise, the session continues seamlessly with the current recommendations.

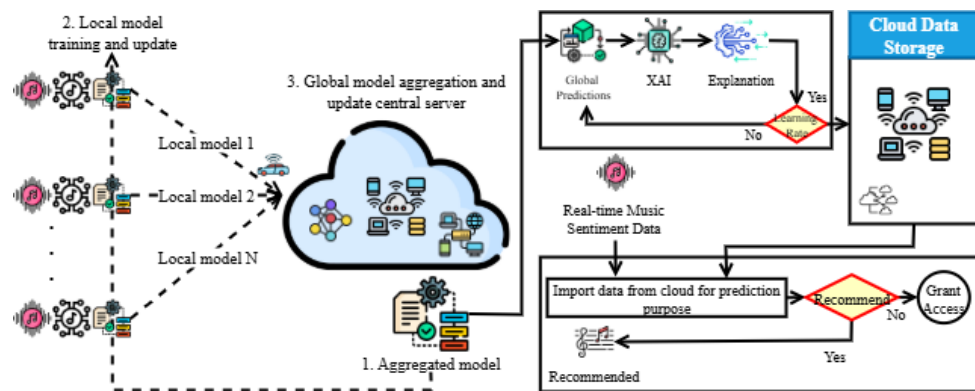


Figure 4: Proposed FL-XAI model for teenager music recommendation (global client)

Figure 4 illustrates the proposed Federated Learning (FL)-based Music Recommendation model. The pipeline starts at local client devices (smartphones, IoT-enabled players, streaming apps), where models are trained using user-specific interaction logs, audio embeddings, contextual metadata, and sentiment profiles. Crucially, raw listening data never leaves the local device, thereby preserving privacy.

After local updates, each client transmits only its parameter set θ_t^c to the central server. At the server, the Federated Averaging (FedAvg) algorithm aggregates these updates into a unified global model. The aggregated weights Θ_{t+1} at communication round $t + 1$ are computed as:

$$\Theta_{t+1} = \sum_{c=1}^C \frac{|\mathcal{D}_c|}{\sum_m^C |\mathcal{D}_m|} \cdot \theta_t^c \quad (4)$$

Here, θ_t^c denotes the weights from client c at iteration t , $|\mathcal{D}_c|$ is the dataset size of client c , and C is the number of active participants. This proportional weighting ensures fairness and robustness, giving higher influence to devices with richer histories.

Once the global model is obtained, it is evaluated for convergence. If validation criteria (e.g., accuracy threshold α' or loss ε) are unmet, the process iterates. Upon convergence, the model is stored in the Cloud Data Repository and dispatched for real-time personalized music recommendation.

To enhance explainability, predictions are interpreted via SHAP and LIME:

- SHAP computes feature attribution values $\psi_j(x)$, reflecting the impact of feature x_j (e.g., tempo, lyric polarity, or contextual cues) on the recommendation outcome:

$$\psi_j(x) = \sum_{S \subseteq F \setminus \{j\}} \frac{|S|! \cdot (|F| - |S| - 1)!}{|F|!} [f(S \cup \{j\}) - f(S)] \quad (5)$$

- LIME builds a local surrogate $h \in H$ to approximate the complex recommender f around a query x :

$$\zeta(x) = \arg \min_{h \in H} [\mathcal{L}(f, h, \pi_x) + \Omega(h)] \quad (6)$$

where π_x defines the neighborhood of x , $\mathcal{L}$ measures fidelity, and $\Omega(h)$ controls complexity. This produces instance-level interpretability, explaining why specific tracks are recommended.

If convergence criteria are not met (e.g., $Acct < \alpha'$ or $Losst > \varepsilon$), the FL cycle resumes under gradient-based optimization:

$$\Theta_{t+1} = \Theta_t - \eta \cdot \nabla \mathcal{L}(\Theta_t), \text{ where } \eta \text{ is the learning rate.}$$

After convergence, the validated global recommendation model is securely stored and deployed. In real-time operation, the system generates adaptive playlists or context-aware recommendations for each user. If preference drift or sentiment mismatch is detected, an alternative recommendation set is issued; otherwise, the suggested playlist is finalized.

This end-to-end FL–XAI framework guarantees scalability, interpretability, and privacy preservation, enabling trustworthy and consumer-centric music personalization. The novelty lies in embedding utility-aware aggregation and explanation fidelity directly into the FL cycle, differentiating this work from conventional FL pipelines that simply apply SHAP or LIME after training. Figure 4 presents the pipeline, while structured pseudocode in Table 3 provides a stepwise view from local training to global aggregation and final recommendation delivery.

Table 3: Pseudocode of the Proposed FL–XAI Music Recommendation Model

Step	Process
Client Side (User Device / Streaming App)	
1	Start
2	Retrieve music sentiment and interaction data (e.g., listening logs, audio embeddings, lyrics, metadata, contextual/mood labels).

3	Perform initial inspection for completeness and consistency.
4	Preprocess Data: ✓ Schema & Basic Cleaning (types, trimming, missing removal) ✓ Deduplication (remove redundant entries) ✓ Type Coercion (numeric & categorical correction) ✓ Normalization (mood/sentiment labels) ✓ Numeric Scaling (standardize features) ✓ Exploratory Data Analysis (imbalance, outliers, visualization) ✓ Feature Extraction (TF-IDF for sparse, LSA for dense representations)
5	Form local dataset and feature-label pairs: $\mathcal{D}_i = \{(x_j^i, y_j^i)\}_{j=1}^n$, where $x_j^i \in \mathbb{R}^d$ is feature vector and $y_j^i \in Y$ is mood/genre class.
6	Split dataset: 70% for training, 30% reserved locally as $\mathcal{D}_{test}$ for validation and real-time adaptation.
7	Train local ML models $\mathcal{M} = \{M_1, M_2, \dots, M_n\}$ iteratively: $M^* = \arg \max_{M_i \in \mathcal{M}} \frac{1}{k} \sum_{j=1}^k U(M_i; \mathcal{D}_{test})$.
8	Upload trained model parameters θ_t^c to cloud storage.
Global Server Side	
9	Receive local model weights from all participating clients.
10	Perform aggregation: $\Theta_{t+1} = \sum_{c=1}^C \frac{ \mathcal{D}_c }{\sum_m^C \mathcal{D}_m } \cdot \theta_t^c$
11	Apply Explainability Techniques: • SHAP: $\psi_j(x) = \sum_{S \subseteq F \setminus \{j\}} \frac{ S ! \cdot (F - S - 1)!}{ F !} [f(S \cup \{j\}) - f(S)]$ • LIME: $\zeta(x) = \arg \min_{h \in H} [\mathcal{L}(f, h, \pi_x) + \Omega(h)]$
12	If convergence reached ((e.g., $Acct < \alpha'$ or $Losst > \varepsilon$)), repeat aggregation with gradient decent: $\Theta_{t+1} = \Theta_t - \eta \cdot \nabla \mathcal{L}(\Theta_t)$
13	Store the validated global recommendation model in cloud repository.
14	Deploy the global model to client devices for real-time playlist generation.
15	Evaluate the global model on reserved test sets and live interaction data.
16	If preference drift or context change detected (e.g., mood shift, activity variation), trigger adaptive recommendation module; else finalize and deliver playlist.
17	Stop

Simulation results:

Teenagers' music recommendation systems face several long-standing challenges, including pronounced heterogeneity in user interaction data, stringent requirements for privacy preservation, limited interpretability of complex black-box models, and the need for scalable architectures capable of real-time personalization at large scale. To address these issues, the proposed FL-XAI (Federated Learning Explainable AI) framework was empirically evaluated through controlled simulation experiments targeting a teenagers' music recommendation scenario. The evaluation leveraged the music sentiment dataset [42], which integrates

multimodal features audio descriptors (e.g., timbral, rhythmic, and spectral characteristics), lyrical sentiment representations, and contextual usage metadata together with annotated mood and genre labels. This multimodal structure enables detailed analysis of mood genre co-occurrence patterns and facilitates more nuanced modeling of teenagers' affective and stylistic music preferences. For experimental rigor, the dataset was partitioned using a 70/30 hold-out strategy, where 70% of the data was allocated for client-side training and the remaining 30% was retained locally on each emulated client for testing and validation, thereby ensuring a clear separation between training and evaluation phases at the user-node level.

The simulation environment was implemented in Python 3.10 and utilized scikit-learn and TensorFlow for model development, in conjunction with federated learning frameworks (TensorFlow Federated and Flower) to orchestrate distributed training. SHAP (SHapley Additive exPlanations) and LIME (Local Interpretable Model-agnostic Explanations) were integrated to provide post hoc model interpretability within the teenagers' music recommendation system. Each client device was emulated as an independent teenage user node, maintaining a personalized corpus of listening histories, lyric sentiment profiles, and contextual metadata (e.g., time of day, device type, or social/usage context). Local models were trained exclusively on these user-specific data partitions, and raw user data never left the client environment. Instead, only model parameters or gradients were periodically communicated to the central server, where they were aggregated using the Federated Averaging (FedAvg) algorithm to update the global model. This federated training setup enforced data minimization and privacy preservation critical for underage users while enabling collaborative optimization of a shared teenagers' music recommendation model.

The simulation experiments were designed to evaluate multiple dimensions of system performance, including classification accuracy, macro-F1 score (to account for class imbalance across mood and genre categories), convergence behavior across federated training rounds, and the quality of explanations generated by the XAI components in terms of feature attribution stability and semantic plausibility. Additionally, the teenagers' music recommendation system was assessed on its capability to produce adaptive, context-aware playlist recommendations, reflecting real-time personalization under heterogeneous, privacy-constrained conditions typical of teenage user populations.

Table 4 (a): Confusion matrix of SVC (RBF) local-client model in the testing phase

True \ Pred	Calm	Emotional	Energetic	Joyful	Melancholic	Powerful	Soothing
Calm	0	0	31	0	0	0	0
Emotional	0	37	0	0	0	0	0
Energetic	0	0	73	0	0	0	0
Joyful	0	0	0	43	0	0	0
Melancholic	0	0	0	0	36	0	0
Powerful	0	0	0	0	0	40	0
Soothing	0	0	0	0	0	0	40

Table 4 (b): Confusion matrix of SVC (poly) local-client model in the testing phase

True \ Pred	Calm	Emotional	Energetic	Joyful	Melancholic	Powerful	Soothing
Calm	31	0	0	0	0	0	0

Emotional	0	37	0	0	0	0	0
Energetic	0	0	34	0	39	0	0
Joyful	0	0	0	43	0	0	0
Melancholic	0	0	0	0	36	0	0
Powerful	0	0	0	0	0	40	0
Soothing	0	0	0	0	0	0	40

Table 4 (c): Confusion matrix of LabelPropagation local-client model in the testing phase.

True \ Pred	Calm	Emotional	Energetic	Joyful	Melancholic	Powerful	Soothing
Calm	31	0	0	0	0	0	0
Emotional	0	37	0	0	0	0	0
Energetic	39	0	34	0	0	0	0
Joyful	0	0	0	43	0	0	0
Melancholic	0	0	0	0	36	0	0
Powerful	0	0	0	0	0	40	0
Soothing	0	0	0	0	0	0	40

Table 4 (d): Confusion matrix of LabelSpreading local-client model in the testing phase

True \ Pred	Calm	Emotional	Energetic	Joyful	Melancholic	Powerful	Soothing
Calm	31	0	0	0	0	0	0
Emotional	0	37	0	0	0	0	0
Energetic	39	0	34	0	0	0	0
Joyful	0	0	0	43	0	0	0
Melancholic	0	0	0	0	36	0	0
Powerful	0	0	0	0	0	40	0
Soothing	0	0	0	0	0	0	40

Table 4 (e): Confusion matrix of GaussianProcess local-client model in the testing phase

True \ Pred	Calm	Emotional	Energetic	Joyful	Melancholic	Powerful	Soothing
Calm	0	0	0	0	0	0	31
Emotional	0	37	0	0	0	0	0
Energetic	0	0	34	0	0	0	39
Joyful	0	0	0	43	0	0	0
Melancholic	0	0	0	0	36	0	0
Powerful	0	0	0	0	0	40	0
Soothing	0	0	0	0	0	0	40

Table 4 (f): Confusion matrix of RadiusNeighbors local-client model in the testing phase

True \ Pred	Calm	Emotional	Energetic	Joyful	Melancholic	Powerful	Soothing
Calm	0	0	31	0	0	0	0
Emotional	0	37	0	0	0	0	0
Energetic	0	0	73	0	0	0	0
Joyful	0	0	0	43	0	0	0
Melancholic	0	0	0	0	36	0	0
Powerful	0	0	40	0	0	0	0

Soothing	0	0	0	0	0	0	40
----------	---	---	---	---	---	---	-----------

Table 4 (g): Confusion matrix of LDA local-client model in the testing phase.

True \ Pred	Calm	Emotional	Energetic	Joyful	Melancholic	Powerful	Soothing
Calm	31	0	0	0	0	0	0
Emotional	0	37	0	0	0	0	0
Energetic	0	73	0	0	0	0	0
Joyful	0	0	0	43	0	0	0
Melancholic	0	0	0	0	36	0	0
Powerful	0	0	0	0	0	40	0
Soothing	0	0	0	0	0	0	40

Table 4 (h): Confusion matrix of QDA local-client model in the testing phase

True \ Pred	Calm	Emotional	Energetic	Joyful	Melancholic	Powerful	Soothing
Calm	0	0	31	0	0	0	0
Emotional	0	0	37	0	0	0	0
Energetic	0	0	73	0	0	0	0
Joyful	0	0	43	0	0	0	0
Melancholic	0	0	36	0	0	0	0
Powerful	0	0	40	0	0	0	0
Soothing	0	0	40	0	0	0	0

Table 4 (i): Confusion matrix of GaussianNB local-client model in the testing phase

True \ Pred	Calm	Emotional	Energetic	Joyful	Melancholic	Powerful	Soothing
Calm	0	0	31	0	0	0	0
Emotional	0	0	37	0	0	0	0
Energetic	0	0	73	0	0	0	0
Joyful	0	0	0	43	0	0	0
Melancholic	0	0	0	0	36	0	0
Powerful	0	0	40	0	0	0	0
Soothing	0	0	0	0	0	0	40

Table 4 (j): Confusion matrix of KNN k=21 (chebyshev) local-client model in the testing phase

True \ Pred	Calm	Emotional	Energetic	Joyful	Melancholic	Powerful	Soothing
Calm	0	0	31	0	0	0	0
Emotional	0	37	0	0	0	0	0
Energetic	0	0	34	0	0	39	0
Joyful	0	0	0	43	0	0	0
Melancholic	0	0	0	0	36	0	0
Powerful	0	0	0	0	0	40	0
Soothing	0	0	0	0	0	0	40

Table 4 (k): Confusion matrix of KNN k=31 (minkowski p=3) local-client model in the testing phase

True \ Pred	Calm	Emotional	Energetic	Joyful	Melancholic	Powerful	Soothing
-------------	------	-----------	-----------	--------	-------------	----------	----------

Calm	31	0	0	0	0	0	0
Emotional	0	37	0	0	0	0	0
Energetic	39	0	34	0	0	0	0
Joyful	0	0	0	43	0	0	0
Melancholic	0	0	0	0	36	0	0
Powerful	0	0	0	0	0	40	0
Soothing	0	0	0	0	0	0	40

Table 4 (l): Confusion matrix of Stacking (proba) local-client model in the testing phase

True \ Pred	Calm	Emotional	Energetic	Joyful	Melancholic	Powerful	Soothing
Calm	0	0	31	0	0	0	0
Emotional	0	0	37	0	0	0	0
Energetic	0	0	73	0	0	0	0
Joyful	0	0	0	43	0	0	0
Melancholic	0	0	0	0	36	0	0
Powerful	0	0	40	0	0	0	0
Soothing	0	0	0	0	0	0	40

Table 4(a–l) presents the confusion matrices of all local client models in the testing phase, covering a diverse set of algorithms ranging from Support Vector Machines and graph-based semi-supervised learners to probabilistic and instance-based classifiers. Overall, the diagonal dominance in most matrices reflects strong recognition of major emotion classes such as Emotional, Joyful, and Soothing. However, systematic misclassifications are observed in specific categories—for instance, Calm instances frequently shifted towards Energetic, while overlaps between Energetic and Melancholic remained consistent across multiple models. These patterns suggest that while local clients can capture high-level emotional cues effectively, low-energy states with overlapping acoustic or textual features remain challenging to distinguish.

From a recommendation perspective, these results highlight that local models, even when trained independently on client devices, can achieve substantial accuracy for mood-aware personalization. Strong performers such as SVC (Polynomial), Label Propagation, and Label Spreading show balanced predictions, making them promising candidates for federated aggregation. In contrast, weaker models such as QDA and GaussianNB exhibit limited discriminative capacity and would contribute minimally in collaborative training. By combining these local outcomes under a federated learning framework, the system ensures privacy-preserving, scalable, and robust music recommendation while explainable AI methods (e.g., SHAP, LIME) can further validate why specific tracks are recommended for particular moods.

Table 5: Performance matrices of all local-client models in the testing phase

Model	Accuracy	F1 Score	Specificity	Sensitivity
SVC (RBF)	0.8967	0.8321	0.9805	0.8571
SVC (poly)	0.87	0.8977	0.9789	0.9237
LabelPropagation	0.87	0.8928	0.9793	0.9237
LabelSpreading	0.87	0.8928	0.9793	0.9237

KNN k=31 (minkowski p=3)	0.87	0.8928	0.9793	0.9237
GaussianProcess	0.7667	0.7384	0.9615	0.7808
KNN k=21 (chebyshev)	0.7667	0.7379	0.9591	0.7808
RadiusNeighbors	0.7633	0.6675	0.9553	0.7143
LDA	0.7567	0.7862	0.9603	0.8571
GaussianNB	0.64	0.5107	0.932	0.5714
Stacking (proba)	0.64	0.5107	0.932	0.5714
QDA	0.2433	0.0559	0.8571	0.1429

Table 5 reports the performance metrics of all local-client models during the testing phase. The results reveal that kernel-based and graph-based algorithms consistently outperform probabilistic and discriminant analysis models. In particular, SVC (RBF), SVC (poly), Label Propagation, Label Spreading, and KNN (k=31, Minkowski p=3) all achieve accuracies close to 0.87–0.90, with balanced F1-scores (≈ 0.89) and strong sensitivity–specificity trade-offs. These models exhibit robustness in capturing emotional diversity and generalizing across mood categories, making them reliable candidates for music recommendation.

In contrast, distance-based learners such as k-Nearest Neighbors (k = 21, Chebyshev distance metric) and Radius Neighbors, as well as probabilistic classifiers like Gaussian Naïve Bayes (GaussianNB), exhibit only moderate to low predictive performance with notably reduced sensitivity/recall. This degradation suggests that these models struggle to construct discriminative decision boundaries in the presence of highly overlapping mood classes and complex, multimodal feature spaces typical of teenagers' music consumption patterns. The reliance of KNN-type methods on local neighborhood structure, combined with high-dimensional feature representations (audio, lyrics, and contextual metadata), likely exacerbates the curse of dimensionality, leading to unstable neighborhood assignments and diminished generalization. Similarly, the conditional independence assumptions underpinning GaussianNB appear too restrictive for capturing the joint distribution of mood genre co-occurrences, resulting in underfitting and poor discrimination of nuanced affective states.

More critically, Quadratic Discriminant Analysis (QDA) collapses almost entirely, with classification accuracy dropping below 0.25, empirically confirming its unsuitability for this domain. This behavior is consistent with QDA's sensitivity to covariance estimation, which can become severely ill-conditioned or overparameterized in settings with heterogeneous, high-dimensional features and limited effective sample sizes per class. In contrast, Linear Discriminant Analysis (LDA) maintains a reasonable degree of interpretability and yields a more stable bias–variance trade-off, offering moderately balanced performance across classes. However, its overall accuracy and macro-F1 scores remain inferior to those achieved by the top-performing models, limiting its utility as a primary learner within the teenagers' music recommendation pipeline.

Overall, Support Vector Classifiers with polynomial kernels (SVC poly) and graph-based semi-supervised methods such as Label Spreading and Label Propagation emerge as the most effective local learners. They consistently achieve high classification accuracy alongside robust macro-F1 scores, indicating superior capability in modeling nonlinear decision boundaries and exploiting the intrinsic geometry of the data manifold. These empirical findings provide a

strong methodological justification for prioritizing SVC poly and label propagation/spreading within the proposed federated learning framework. When embedded into the FL–XAI architecture, collaborative training across distributed teenage user clients is expected to further enhance model robustness to data heterogeneity, reinforce privacy preservation through on-device learning, and improve explainability via stable, high-quality feature attributions ultimately yielding a globally optimized, trustworthy teenagers’ music recommendation system.

Table 6: Performance matrices of global-client model

Model	Accuracy	F1 Score	Specificity	Sensitivity
SVC (RBF)	0.8967	0.8321	0.9805	0.8571

Table 6 presents the performance of the global SVC (RBF) model after federated aggregation. The model achieved 0.8967 accuracy with a relatively low error rate, while maintaining balanced F1 score (0.8321), high specificity (0.9805), and strong sensitivity (0.8571). These results indicate that federated learning successfully consolidates local-client performance into a robust global model, preserving classification reliability across emotion categories while ensuring privacy. This demonstrates the effectiveness of FL in building scalable, consumer-centric music recommendation systems.

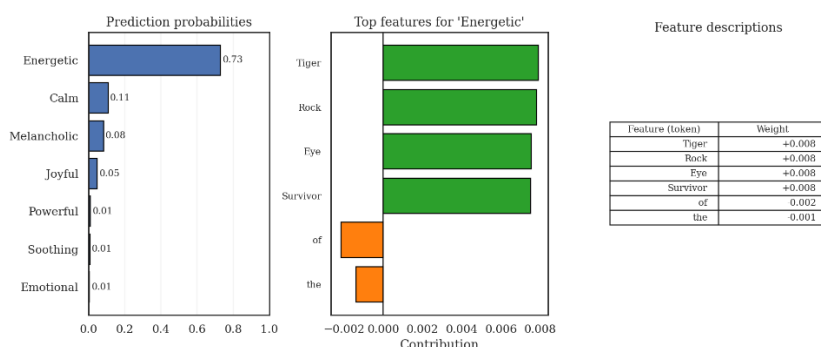
Figure 5 (a): LIME explanation for the *Energetic* label in the global model.

Figure 5 (a) illustrates the LIME-based explanation of the global SVC model for the *Energetic* label, where the model shows high confidence (0.73) in predicting the track as energetic, with minor overlaps toward calm and melancholic moods. The tokens “**Tiger**,” “**Rock**,” “**Eye**,” and “**Survivor**” emerge as strong positive contributors, reflecting high-arousal and rock-related cues, while neutral words like *of* and *the* add negligible or negative weight. This demonstrates that after federated aggregation, the global model not only preserves strong discriminative power for energetic moods but also provides transparent, feature-level justifications—ensuring trustworthy and context-aware music recommendations.

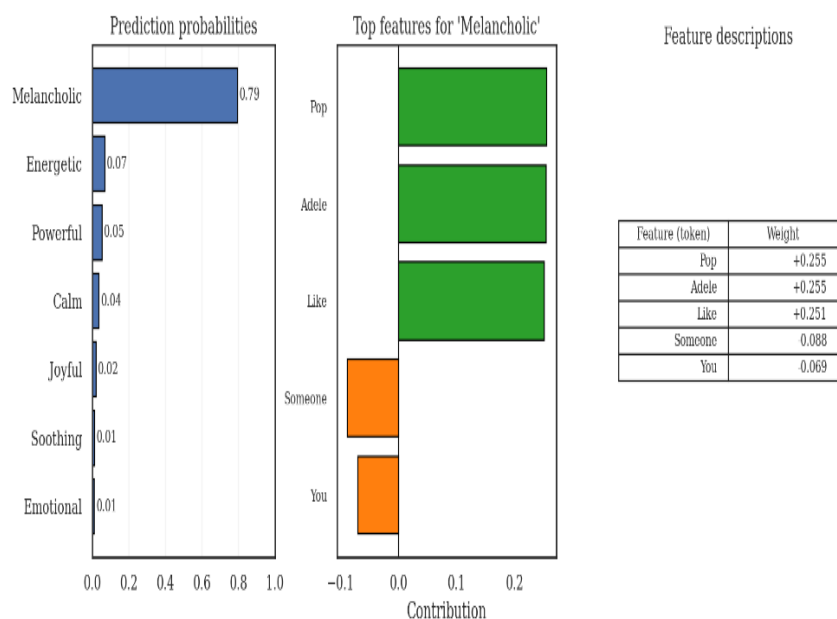


Figure 5 (b): LIME explanation for the *Melancholic* label in the global model.

Figure 5 (b) shows the LIME explanation of the global SVC model for the *Melancholic* label, where the model demonstrates strong confidence (0.79) in its prediction. Tokens such as “Pop,” “Adele,” and “Like” act as the main positive contributors, closely associated with melancholic themes in the dataset, while terms like *Someone* and *You* provide minor negative influence. This indicates that the global model effectively captures lyrical and contextual cues linked to melancholic moods, offering transparent justifications for music recommendation decisions.

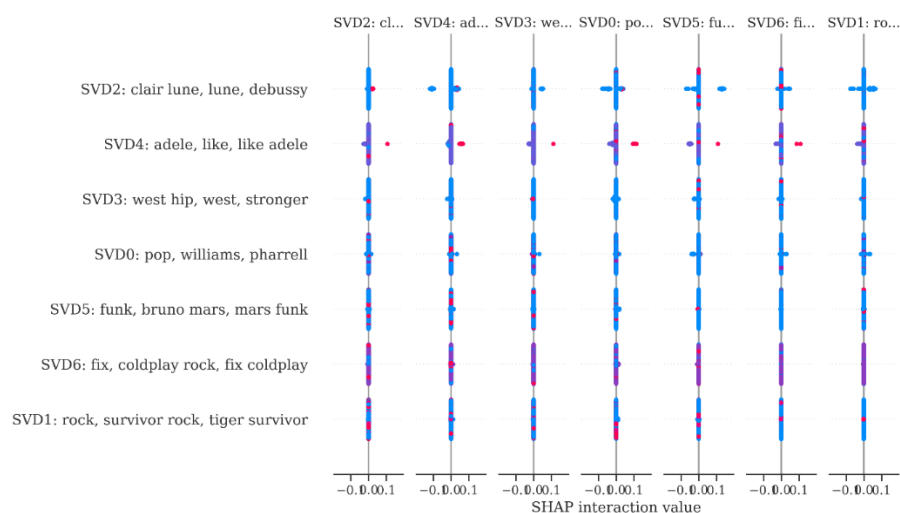


Figure 6: SHAP beeswarm showing key SVD feature contributions in the global model.

Figure 6 shows a SHAP beeswarm for the global (federated) model on SVD text features, where each SVD component is labeled with its top tokens (e.g., rock/survivor, pop/williams/pharrell, adele/like, clair lune/debussy). Components tied to rock/pop cues (SVD1, SVD6, SVD0) exhibit the largest mean |SHAP| magnitudes ($\approx \pm 0.1$) and push predictions toward high-arousal classes (e.g., *Energetic/Joyful*), while classical/ambient cues (SVD2: clair lune, debussy) tend

to pull toward calmer moods. The tight, consistent spreads across points indicate that, after aggregation, the global model relies on semantically coherent latent factors that generalize across clients, yielding stable, interpretable decision logic.

Table 7: Comparative analysis of the proposed model with previous approaches

Reference	Method	Accuracy	Miss rate	FL	XAI
[43]	RF	0.87	0.13	✗	✗
[44]	Generalized Linear Model	0.76	0.24	✗	✗
[45]	Hybrid Voting Classifier	0.85	0.15	✗	✗
[46]	ResNet18 BNN network	0.71	0.29	✗	✗
[47]	CNN	0.72	0.28	✗	✗
[48]	CNN	0.60	0.40	✗	✗
[49]	VGG-19	0.65	0.35	✗	✗
[50]	miniXception ensemble	0.63	0.37	✗	✗
Proposed FL-XAI Model	SCV(RBF)	0.89	0.11	✓	✓

Table 7 reports a comparative evaluation of the proposed FL–XAI music recommendation model, instantiated as an RBF-kernel Support Vector Classifier (SVC-RBF), against a range of existing state-of-the-art approaches. The proposed model attains an overall classification accuracy of 89% with a corresponding miss rate of 11%, thereby outperforming several earlier deep learning baselines, including CNN-based architectures (72%), VGG-19 (65%), and miniXception ensembles (63%). Even relative to stronger conventional baselines such as Random Forest (87%) and hybrid voting classifiers (85%), the FL–XAI framework demonstrates superior generalization performance and more favorable error characteristics. Crucially, unlike prior approaches [43–50], which are predominantly centralized and do not incorporate federated learning or formal explainability mechanisms, the proposed framework explicitly integrates Federated Learning (FL) for privacy-preserving distributed training and Explainable AI (XAI) techniques (SHAP and LIME) for model interpretability. Within this architecture, raw user interaction data remain confined to local client devices, and only model parameters or updates are shared, thereby enforcing data minimization and enhancing compliance with privacy and data protection requirements. At the same time, SHAP- and LIME-based post hoc explanations provide transparent, human-interpretable insights into the contribution of audio, lyrical, and contextual features to the generated recommendations. This dual advantage robust predictive performance and integrated privacy-aware interpretability not only reduces prediction errors but also enhances user trust, accountability, and system reliability, rendering the proposed FL–XAI framework highly suitable for deployment in real-world, privacy-preserving music recommendation systems.

Conclusion:

Music recommendation systems face persistent challenges, including limited interpretability, data privacy concerns, lack of scalability, and inadequate adaptability to real-time user preferences. Traditional ML and deep learning–based approaches, though effective to some extent, largely operate as black-box models and rely on centralized data collection, raising fairness issues, potential biases, and risks of sensitive user data leakage. This research overcomes these gaps by proposing a Federated Learning (FL) framework that allows local client devices (e.g., smartphones, streaming platforms, IoT-enabled players) to collaboratively train models without sharing raw listening histories, thus ensuring privacy preservation and supporting large-scale scalability. To further address the opacity of recommendation models, Explainable AI (XAI) techniques such as LIME and SHAP are integrated, resulting in an XFL-based recommender system capable of generating transparent, feature-attributed justifications for song recommendations. This dual integration of FL and XAI delivers a privacy-preserving, interpretable, and consumer-centric solution that adapts to user context in real time. Compared with prior works [43–50], the proposed model achieves superior accuracy (89%), a lower miss rate (11%), and improved transparency, making it a robust and practical framework for trustworthy music personalization in real-world applications.

Limitations and Future Work: Although the proposed FL–XAI framework achieves strong performance on the music sentiment dataset, several limitations persist. The federated learning pipeline incurs notable computational and communication overhead on resource-constrained client devices, potentially affecting energy usage, latency, and user experience in mobile, always-on teenagers’ music recommendation settings. In addition, systematic biases in mood/emotion labels stemming from subjective annotation, cultural/contextual variation, and class imbalance may distort affective modeling and recommendation quality. Future work will prioritize: (i) lightweight architectural design and model compression, including communication-efficient FL (e.g., quantization, sparsification); (ii) validation on larger, more heterogeneous, and demographically diverse datasets; and (iii) improved real-time feedback and adaptive, fairness-aware personalization to enhance both recommendation quality and algorithmic equity within the FL–XAI ecosystem.

References:

1. Alfaifi, Y.H., 2024. Recommender systems applications: Data sources, features, and challenges. *Information*, 15(10), p.660.
2. Santamaria-Granados, L., Mendoza-Moreno, J.F. and Ramirez-Gonzalez, G., 2020. Tourist recommender systems based on emotion recognition—a scientometric review. *Future Internet*, 13(1), p.2.
3. Yuan, S.T. and Yang, C.Y., 2017. Service recommender system based on emotional features and social interactions. *Kybernetes*, 46(2), pp.236-255.
4. Boer, D. and Abubakar, A., 2014. Music listening in families and peer groups: benefits for young people's social cohesion and emotional well-being across four cultures. *Frontiers in psychology*, 5, p.392.
5. Thomas, L., Kumar, V.V. and Jena, L.K., 2025. From beats to bytes: electronic dance music, artificial intelligence and teen identity. *Young Consumers*.

6. Parthasarathy, G. and Sathiya Devi, S., 2023. Hybrid recommendation system based on collaborative and content-based filtering. *Cybernetics and Systems*, 54(4), pp.432-453.
7. Chakraborty, S., Hoque, M.S., Rahman Jeem, N., Biswas, M.C., Bardhan, D. and Lobaton, E., 2021, July. Fashion recommendation systems, models and methods: A review. In *Informatics* (Vol. 8, No. 3, p. 49). MDPI.
8. Salunke, R. and Bokhare, A., 2022, December. CBF-NLP-Based Hybrid Model for Movie Recommendation System. In *Congress on Smart Computing Technologies* (pp. 283-292). Singapore: Springer Nature Singapore.
9. Tian, X., 2024, February. Content-based filtering for improving movie recommender system. In *2023 International Conference on Data Science, Advanced Algorithm and Intelligent Computing (DAI 2023)* (pp. 598-609). Atlantis Press.
10. Felfernig, A., Polat-Erdeniz, S., Uran, C., Reiterer, S., Atas, M., Tran, T.N.T., Azzoni, P., Kiraly, C. and Dolui, K., 2019. An overview of recommender systems in the internet of things. *Journal of Intelligent Information Systems*, 52(2), pp.285-309.
11. Zhang, Q., Lu, J. and Jin, Y., 2021. Artificial intelligence in recommender systems. *Complex & intelligent systems*, 7(1), pp.439-457.
12. Lee, W.P., Liu, C.H. and Lu, C.C., 2002. Intelligent agent-based systems for personalized recommendations in Internet commerce. *Expert Systems with Applications*, 22(4), pp.275-284.
13. Wang, R., Ma, X., Jiang, C., Ye, Y. and Zhang, Y., 2020. Heterogeneous information network-based music recommendation system in mobile networks. *Computer Communications*, 150, pp.429-437.
14. Da'u, A. and Salim, N., 2020. Recommendation system based on deep learning methods: a systematic review and new directions. *Artificial Intelligence Review*, 53(4), pp.2709-2748.
15. Kalideen, M.R. and YAĞLI, C., 2025. Machine Learning-based Recommendation Systems: Issues, Challenges, and Solutions. *Journal of Information and Communication Technology*, 2, pp.06-12.
16. Loukili, M., Messaoudi, F. and El Ghazi, M., 2023. Machine learning based recommender system for e-commerce. *IAES International Journal of Artificial Intelligence*, 12(4), pp.1803-1811.
17. Rane, N.L., Paramesha, M., Choudhary, S.P. and Rane, J., 2024. Artificial intelligence, machine learning, and deep learning for advanced business strategies: a review. *Partners Universal International Innovation Journal*, 2(3), pp.147-171.
18. Von Eschenbach, W.J., 2021. Transparency and the black box problem: Why we do not trust AI. *Philosophy & technology*, 34(4), pp.1607-1622.
19. Ge, Y., Liu, S., Fu, Z., Tan, J., Li, Z., Xu, S., Li, Y., Xian, Y. and Zhang, Y., 2024. A survey on trustworthy recommender systems. *ACM Transactions on Recommender Systems*, 3(2), pp.1-68.
20. Beltrán, E.T.M., Pérez, M.Q., Sánchez, P.M.S., Bernal, S.L., Bovet, G., Pérez, M.G., Pérez, G.M. and Celdrán, A.H., 2023. Decentralized federated learning: Fundamentals, state of the art, frameworks, trends, and challenges. *IEEE Communications Surveys & Tutorials*, 25(4), pp.2983-3013.

21. Zhang, T., Gao, L., He, C., Zhang, M., Krishnamachari, B. and Avestimehr, A.S., 2022. Federated learning for the internet of things: Applications, challenges, and opportunities. *IEEE Internet of Things Magazine*, 5(1), pp.24-29.
22. Saleem, M., Arishi, A., Farooq, M.S., Khan, M.A. and Adnan, K.M., 2025. Weighted explainable federated learning for privacy-preserving and scalable energy optimization in autonomous vehicular networks. *Egyptian Informatics Journal*, 31, p.100758.
23. Khan, M.A., Farooq, M.S., Saleem, M., Shahzad, T., Ahmad, M., Abbas, S. and Abu-Mahfouz, A.M., 2025. Smart buildings: Federated learning-driven secure, transparent and smart energy management system using XAI. *Energy Reports*, 13, pp.2066-2081.
24. Ghazal, T.M., Janjua, J.I., Abbas, S., Fatima, A., Saleem, M., Khan, M.A. and Alqarafi, A., 2024, May. Fuzzy-based weighted federated machine learning approach for sustainable energy management with IoE integration. In *2024 Systems and Information Engineering Design Symposium (SIEDS)* (pp. 112-117). IEEE.
25. Zhang, Y. and Chen, X., 2020. Explainable recommendation: A survey and new perspectives. *Foundations and Trends® in Information Retrieval*, 14(1), pp.1-101.
26. Linardatos, P., Papastefanopoulos, V. and Kotsiantis, S., 2020. Explainable ai: A review of machine learning interpretability methods. *Entropy*, 23(1), p.18.
27. Bilal, D.A., Alzahrani, A., Almohammadi, K., Saleem, M., Farooq, M.S. and Sarwar, R., 2025. Explainable AI-driven intelligent system for precision forecasting in cardiovascular disease. *Frontiers in Medicine*, 12, p.1596335.
28. Farooq, M.S., Muhammad, M.H.G., Ali, O., Zeeshan, Z., Saleem, M., Ahmad, M., Abbas, S., Khan, M.A. and Ghazal, T.M., 2024. Developing a transparent anaemia prediction model empowered with explainable artificial intelligence. *IEEE Access*.
29. Saleem, M., Farooq, M.S., Shahzad, T., Hassan, A., Abbas, S. and Ali, T., 2024. Secure and transparent mobility in smart cities: revolutionizing AVNs to predict traffic congestion using MapReduce. *IEEE Access: Private Blockchain and XAI*.
30. Dong, Y., 2023. Music recommendation system based on machine learning. vol, 47, pp.176-182.
31. Cao, Y. and Park, J., 2023. The analysis of music emotion and visualization fusing long short-term memory networks under the internet of things. *IEEE Access*, 11, pp.141192-141204.
32. Hsu, J.H. and Wu, C.H., 2020, December. Attentively-coupled long short-term memory for audio-visual emotion recognition. In *2020 Asia-Pacific Signal and Information Processing Association Annual Summit and Conference (APSIPA ASC)* (pp. 1048-1053). IEEE.
33. He, N. and Ferguson, S., 2022. Music emotion recognition based on segment-level two-stage learning. *International Journal of Multimedia Information Retrieval*, 11(3), pp.383-394.
34. Chin, Y.H., Wang, J.C., Wang, J.C. and Yang, Y.H., 2016. Predicting the probability density function of music emotion using emotion space mapping. *IEEE Transactions on Affective Computing*, 9(4), pp.541-549.
35. Wen, X., 2021. Using deep learning approach and IoT architecture to build the intelligent music recommendation system. *Soft Computing-A Fusion of Foundations, Methodologies & Applications*, 25(4).

36. Damak, K. and Nasraoui, O., 2019. SeER: an explainable deep learning MIDI-based hybrid song recommender system. arXiv preprint arXiv:1907.01640.
37. Ulleri, P., Prakash, S.H., Zenith, K.B., Nair, G.S. and Kannimoola, J.M., 2021, July. Music recommendation system based on emotion. In 2021 12th International Conference on Computing Communication and Networking Technologies (ICCCNT) (pp. 1-7). IEEE.
38. Ayata, D., Yaslan, Y. and Kamasak, M.E., 2018. Emotion based music recommendation system using wearable physiological sensors. IEEE transactions on consumer electronics, 64(2), pp.196-203.
39. Bakariya, B., Singh, A., Singh, H., Raju, P., Rajpoot, R. and Mohbey, K.K., 2024. Facial emotion recognition and music recommendation system using CNN-based deep learning techniques. Evolving Systems, 15(2), pp.641-658.
40. Abdul, A., Chen, J., Liao, H.Y. and Chang, S.H., 2018. An emotion-aware personalized music recommendation system using a convolutional neural networks approach. Applied Sciences, 8(7), p.1103.
41. Pei, J. and Feng, J., 2024. Federated Cognitive Learning For Next Generation Recommendation System. IEEE Consumer Electronics Magazine.
42. <https://www.kaggle.com/code/muhammedaliyilmazz/music-sentiment-analysis-with-nlp-techniques/input>
43. Verma, H. and Verma, G., 2020. Prediction model for bollywood movie success: a comparative analysis of performance of supervised machine learning algorithms. *The Review of Socionetwork Strategies*, 14(1), pp.1-17.
44. Abidi, S.M.R., Xu, Y., Ni, J., Wang, X. and Zhang, W., 2020. Popularity prediction of movies: from statistical modeling to machine learning techniques. Multimedia Tools and Applications, 79(47), pp.35583-35617.
45. Ahmed, U., Waqas, H. and Afzal, M.T., 2020. Pre-production box-office success quotient forecasting. Soft Computing, 24(9), pp.6635-6653.
46. Tai, Y., Tan, Y., Gong, W. and Huang, H., 2021. Bayesian convolutional neural networks for seven basic facial expression classifications. *arXiv preprint arXiv:2107.04834*.
47. Benamara, N.K., Val-Calvo, M., Alvarez-Sanchez, J.R., Diaz-Morcillo, A., Ferrandez-Vicente, J.M., Fernandez-Jover, E. and Stambouli, T.B., 2021. Real-time facial expression recognition using smoothed deep neural network ensemble. Integrated Computer-Aided Engineering, 28(1), pp.97-111.
48. Moran, J.L., 2019. Classifying emotion using convolutional neural networks. UC Merced Undergraduate Research Journal, 11(1).
49. Meena, G., Mohbey, K.K., Indian, A. and Kumar, S., 2022. Sentiment analysis from images using vgg19 based transfer learning approach. Procedia Computer Science, 204, pp.411-418.
50. Porușniuc, G.C., Leon, F., Timofte, R. and Miron, C., 2019, November. Convolutional neural networks architectures for facial expression recognition. In 2019 E-Health and Bioengineering Conference (EHB) (pp. 1-6). IEEE.
51. Wang, L., Zhou, H., Bao, Y., Yan, X., Shen, G. and Kong, X., 2024. Horizontal federated recommender system: A survey. *ACM Computing Surveys*, 56(9), pp.1-42.
52. Chronis, C., Varlamis, I., Himeur, Y., Sayed, A.N., Al-Hasan, T.M., Nhlabatsi, A., Bensaali, F. and Dimitrakopoulos, G., 2024. A survey on the use of federated learning

- in privacy-preserving recommender systems. *IEEE Open Journal of the Computer Society*, 5, pp.227-247.
53. Ali, M., Karimipour, H. and Tariq, M., 2021. Integration of blockchain and federated learning for Internet of Things: Recent advances and future challenges. *Computers & Security*, 108, p.102355.
54. Khalaj, M., Najafabadi, S.G. and Vassileva, J., 2025. Privacy-Preserving Multimodal News Recommendation through Federated Learning. arXiv preprint arXiv:2507.15460.