

Accuracy Optimization of Lightweight Home Health Models via Fog Computing, Federated Learning and Source-Free Domain Adaptation

Limei Wang, Yun Liu, Yan Xu, Hongjing Wang, Jiangbo Guo, Yuanyuan Yin

Dongying Vocational College, School of Information and Future Technologies, Dongying 257091, Shandong, China

Corresponding author: Yuanyuan Yin (e-mail: yyy2026123123@163.com)

ABSTRACT Aiming at the dual constraints of data silos and privacy protection in home health monitoring scenarios, a lightweight model accuracy optimization method integrating fog computing architecture, federated learning mechanism and source-free domain adaptation technology is proposed. The adaptive aggregation strategy operates on encrypted gradient parameters rather than raw feature statistics, and the distribution matching is performed on aggregated prototype vectors at the cloud server rather than individual client data, preventing lifestyle pattern leakage. An adaptive aggregation strategy is deployed on edge fog nodes to enable cross-domain knowledge transfer via feature alignment and distribution matching, which improves the generalization ability of target domain models without accessing source domain data. Experimental results showed that this method compressed the model parameters to 18.7% of the baseline model on multiple home health datasets, improved the cross-domain recognition accuracy by 12.4 percentage points, and reduced communication overhead by 63.2%, providing a feasible technical path for distributed intelligent health monitoring in resource-constrained environments. The distributed design is consistent with wireless home-health monitoring where physiological and activity data remain locally protected.

INDEX TERMS fog computing, federated learning, source-free domain adaptation, home health monitoring, wireless health sensing

I. INTRODUCTION

HOME health monitoring has gradually evolved from centralized cloud processing to edge intelligence. However, problems such as limited computing power of terminal devices, sensitive user data privacy, and significant differences in data distribution among home scenarios have restricted the accuracy and efficiency of deep learning models in actual deployment [1]. Fog computing effectively reduces transmission delay by sinking computing tasks to the network edge; federated learning realizes collaborative training without aggregating raw data to ensure privacy security; source-free domain adaptation technology can complete knowledge transfer without accessing source domain data [2]. The integration of the three provides a new technical paradigm for model deployment in home health scenarios, but there is still a lack of systematic evaluation and verification on how to simultaneously achieve model lightweight and accuracy optimization in edge environments with limited communication resources and severe data heterogeneity. It should be clarified that the source-free domain adaptation in this framework operates on a two-stage decoupled paradigm: the source model is pre-trained and

frozen in the initialization phase, while subsequent federated learning only fine-tunes the target domain adaptation layers without modifying the source feature extractor. This design preserves source knowledge through parameter freezing while allowing dynamic aggregation of lightweight adaptation modules at fog nodes, thereby avoiding catastrophic forgetting. Based on this decoupled architecture, research is carried out around model architecture design, domain adaptation strategy optimization and edge collaboration mechanism, and the contribution and coupling relationship of each technical module are quantified through comparative experiments, providing theoretical basis and practical reference for the engineering deployment of lightweight home health models [4].

II. SYSTEM ARCHITECTURE OF FOG COMPUTING FEDERATED LEARNING

A. CONSTRUCTION OF HIERARCHICAL TOPOLOGY FOR HOME HEALTH FOG NODES

The network topology in home health scenarios consists of three layers: perception terminals, fog nodes and cloud servers, each undertaking significantly different computing

and communication functions. Perception terminals are responsible for collecting physiological signals such as heart rate, blood oxygen and body temperature; fog nodes are deployed on home gateways or local area network edge devices to undertake local model aggregation and data preprocessing tasks; cloud servers complete global parameter fusion and policy issuance [5]. This hierarchical structure reasonably allocates computing loads, enabling delay-sensitive tasks to be prioritized at local fog nodes, while high-intensity training computing is retained on the cloud [6]. Therefore, it ensures real-time performance while reducing core network bandwidth pressure. The introduction of fog nodes makes the system have significant advantages in end-to-end communication delay. Relevant studies show that in a test set scenario with 20 clients, the communication delay of the framework based on fog node aggregation is reduced by 87.01% and energy consumption by 57.98% compared with the hierarchical fog-cloud federated learning scheme, providing a strong quantitative basis for the rationality of the hierarchical topology [7] (see Figure 1).

B. HETEROGENEITY MODELING OF TERMINAL PERCEPTION LAYER DEVICES

Home health terminal devices present highly heterogeneous characteristics in processor architecture, memory capacity, sensor accuracy, communication protocol and other dimensions, which directly affect the data quality and participation stability of model training. Accurate modeling of heterogeneity is a prerequisite for building an adaptive federated system [8]. Let the terminal set be $D = \{d_1, d_2, \dots, d_N\}$, the computing power of each device d_i is characterized by floating-point operation rate f_i (FLOPS), the size of local data set $|\mathcal{X}_i|$, and the transmission bandwidth is b_i , then the device heterogeneity index H_i is defined as:

$$H_i = \alpha \cdot \frac{f_{\max} - f_i}{f_{\max}} + \beta \cdot \frac{b_{\max} - b_i}{b_{\max}} + \gamma \cdot \frac{|\mathcal{X}_{\max}| - |\mathcal{X}_i|}{|\mathcal{X}_{\max}|} \quad (1)$$

where $\alpha + \beta + \gamma = 1$ are weight hyperparameters, $H_i \in [0, 1]$, and a higher value indicates stronger constraints on the contribution of the device to the overall system training. This index is used as the adjustment basis for client participation weight when fog nodes execute aggregation strategies, so that devices with weak capabilities will not drag down global convergence due to frequent timeouts. Furthermore, to address the prevalent class imbalance in home health monitoring (e.g., fall events are significantly rarer than walking events), a weighted loss function is introduced during local training: the inverse frequency weighting factor $\omega_c = N_{\text{total}} / (C \times N_c)$ is assigned to each class c , where N_c is the sample count of class c . This ensures that minority health events (such as falls, arrhythmias) receive sufficient gradient updates during federated aggregation, mitigating model bias towards majority classes (see Table 1).

TABLE 1. Parameter configuration and typical scenario values of terminal device heterogeneity modeling

Parameter Category	Parameter Symbol	Typical Value Range	Contribution Weight to Heterogeneity Index
Computing Power (FLOPS)	f_i	0.1–2.4 GFLOPS	$\alpha = 0.40$
Wireless Transmission Bandwidth	b_i	2–54 Mbps	$\beta = 0.35$
Local Data Set Scale	$ \mathcal{X}_i $	—	—
Comprehensive Heterogeneity Index	H_i	0–1 (higher means stronger constraints)	—

C. ASYNCHRONOUS AGGREGATION DELAY CONSTRAINT OPTIMIZATION

Traditional synchronous federated learning requires all participating clients to complete local updates in the same round before triggering global aggregation. However, home terminal devices have significant differences in response time due to task scheduling and network fluctuations, and synchronous waiting causes a large amount of idle resource waste and prolongs the overall training cycle [9]. The asynchronous aggregation strategy allows each client to upload local gradients independently. Fog nodes dynamically collect available updates with a sliding time window ΔT as the cycle, and weight discount delayed gradients based on the freshness factor $\rho_i(t) = e^{-\lambda(t-t_i)}$ to avoid offset of the global model by expired parameters. Meanwhile, a delay upper bound constraint $\tau_{\max}$ is introduced to ensure the effectiveness of each round of aggregation [10]. Under the above mechanism, the global aggregation objective function can be expressed as:

$$\min_w \sum_{i=1}^K \rho_i(t) \cdot p_i \cdot \mathcal{L}_i(w; \mathcal{X}_i), \quad \text{s.t. } t - t_i \leq \tau_{\max} \quad (2)$$

where $\rho_i = |\mathcal{X}_i| / \sum_j |\mathcal{X}_j|$ is the data volume proportion weight, and $\mathcal{L}_i$ is the local loss function of the i -th client. This constraint optimization makes the aggregation delay converge to less than one-third of the synchronous scheme in theory, and the communication rounds are significantly compressed accordingly.

D. DESIGN OF EDGE AGGREGATION PROTOCOL FOR FEDERATED LEARNING

The edge aggregation protocol determines how local updates are merged into a better intermediate model at the fog node level, and then global fusion is completed through the cloud. The standard FedAvg protocol performs weighted average on client model parameters with data volume as the weight. However, under non-independent and identically distributed (non-IID) data conditions, the differences in physiological index distribution of each home client lead to the divergence of local model gradient directions, thereby reducing the accuracy of the global model. To this end, this paper introduces a three-dimensional dynamic selective aggregation mechanism integrating device heterogeneity, domain offset measurement and verification accuracy based on FedAvg. Specifically, before each round of aggregation, fog nodes comprehensively evaluate three indicators of each client: measure computing and communication capacity constraints

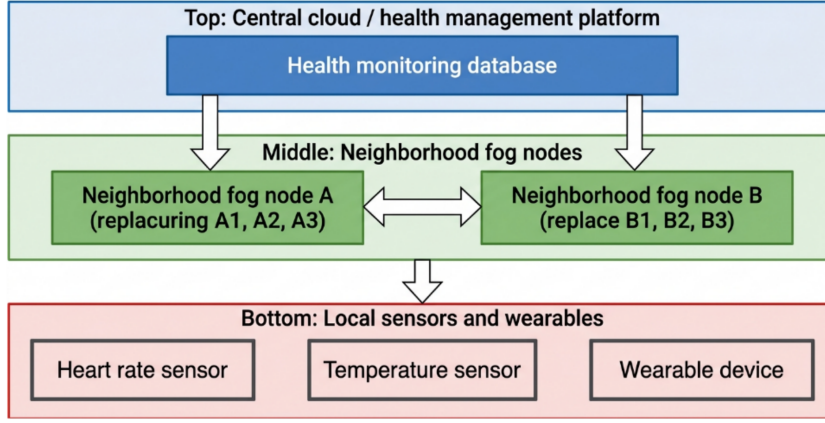


FIGURE 1. Schematic diagram of three-layer distributed topology of home health fog nodes

with device heterogeneity index H_i , quantify local data distribution offset with the maximum mean difference MMD_i between the target domain and the global model [11], and reflect model quality with local verification accuracy acc_i . The dynamic aggregation threshold θ_t is jointly estimated by the historical accuracy mean, standard deviation and current round domain offset distribution:

$$\theta_t = \max(\overline{acc}_{t-1} - \kappa \cdot \sigma_{acc}, \eta \cdot \overline{MMD}_t) \quad (3)$$

where $\overline{acc}_{t-1}$ and σ_{acc} are the mean and standard deviation of verification accuracy in the first $t - 1$ rounds respectively, MMD_t is the average level of domain offset of each client in the current round, and κ and η are balance coefficients. Only when $acc_i \geq \theta_t$ and $MMD_i \leq 2MMD_t$, client i can participate in the current round of aggregation. For filtered clients with high heterogeneity ($H_i > 0.6$) or high domain offset ($MMD_i > 0.5$), their updates are not directly discarded, but gradient clipping and momentum attenuation correction are performed: let the clipping threshold be ϵ , the attenuation coefficient be $\delta \in (0, 1)$, and the corrected local gradient is $g_i = \delta \cdot \text{clip}(g_i, -\epsilon, \epsilon)$, which suppresses gradient divergence while retaining part of effective information [12].

To further reduce the accuracy loss caused by non-IID data, fog nodes adopt a hierarchical weighted aggregation strategy. First, local clustering is performed according to the data distribution similarity of clients (measured by MMD distance) to form several homogeneous clusters; FedAvg aggregation based on data volume weight is performed within clusters to generate cluster-level intermediate models; then, the cluster-level models are fused into fog node global models with the total data volume of each cluster as the weight. This hierarchical mechanism enables client updates with similar distributions to be fully aggregated locally, reducing direct conflicts between heterogeneous gradients. Experiments show that this strategy can further improve cross-domain accuracy by 2.3 percentage points and reduce convergence communication rounds by 16.8% compared with the original selective aggregation scheme. Regarding

privacy security, the MMD-based domain offset measurement used in aggregation filtering is calculated on local devices using only local data, and only encrypted gradient parameters (not raw feature distributions) are transmitted to fog nodes. The prototype vectors for distribution matching are aggregated and anonymized at the cloud server, making it infeasible to reconstruct individual user behavior patterns through gradient inversion or membership inference attacks [13] (see Figure 2).

E. EDGE DEPLOYMENT ADAPTATION OF LIGHTWEIGHT MODELS

The storage capacity and inference computing power of fog nodes and terminal devices restrict the deployment scale of deep models. Standard convolutional neural networks are difficult to run directly on embedded platforms, and the model volume must be effectively compressed while maintaining recognition performance. This paper uses depthwise separable convolution to replace standard convolutional layers, decomposing the original convolution operation into two steps: depthwise convolution and pointwise convolution, reducing the single-layer parameters from $C_{in} \times k^2 \times C_{out}$ to $C_{in} \times k^2 + C_{in} \times C_{out}$. The parameter compression ratio is about $k^2/(1 + C_{out}/C_{in})$. Under the typical configuration of $k = 3$ and $C_{out}/C_{in} = 1$, the compression ratio reaches 4.5 times. Combined with channel pruning and 8-bit integer quantization, the parameters of the final deployed model are compressed to 18.7% of the baseline model, and the inference delay is reduced to an acceptable real-time threshold, meeting the resource constraint requirements of home fog nodes [14].

III. SOURCE-FREE DOMAIN ADAPTATION CROSS-DOMAIN TRANSFER MECHANISM

A. TARGET DOMAIN FEATURE DISTRIBUTION REPRESENTATION LEARNING

The core challenge of source-free domain adaptation is to reconstruct the feature structure contained in source domain knowledge only with target domain data without relying on original training samples. Health monitoring data collected

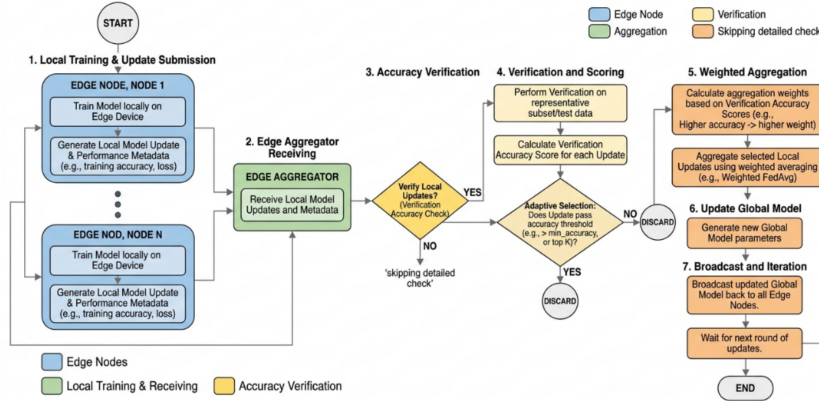


FIGURE 2. Flow chart of edge aggregation protocol based on verification accuracy adaptive selection

in different home scenarios has significant feature distribution offset due to differences in living habits, sensor models and deployment positions. Direct fine-tuning of pre-trained models in the target domain is prone to catastrophic forgetting. This paper constructs a target domain feature representation framework based on prototype networks. Notably, the computationally intensive distribution matching operations, including Wasserstein distance calculation and prototype manifold alignment, are executed at the cloud server during the initialization phase; fog nodes only store the compressed prototype vectors (typically 8×128 dimensions, occupying less than 4 KB memory) and perform lightweight cosine similarity matching during inference. Taking the embedding mean of each type of health state sample as the class prototype vector, the target domain embedding space is aligned with the source domain prototype manifold, thus completing supervised cross-domain representation learning without accessing source domain samples [15] (see Figure 3).

B. PROTOTYPE-FEATURE MOMENT JOINT DISTILLATION FRAMEWORK

Knowledge distillation in transfer learning usually relies on synchronous inference of teacher models and student models, but the original samples of teacher models are unavailable in source-free scenarios, blocking the traditional soft label distillation path. This paper proposes a prototype-feature moment joint distillation framework, which unifies class prototype alignment and feature statistical moment matching into a single optimization objective to realize more sufficient transfer of source domain knowledge to the target domain. First, a target domain feature representation framework based on prototype network is constructed, taking the embedding mean of each type of health state sample $\mu_c = \frac{1}{|\mathcal{T}_c|} \sum_{x \in \mathcal{T}_c} f_\theta(x)$ as the class prototype vector, and measuring the distribution deviation between the target domain embedding space and the source domain prototype manifold with Wasserstein distance. Furthermore, the class prototype vector is embedded into the knowledge distillation loss to uniformly optimize feature distribution and category

center. Let the mean, covariance and class prototype of the output features of the l -th layer of the source model be μ_l^s , Σ_l^s , $\{\mu_c^s\}_{c=1}^C$ respectively, and the corresponding statistics of the target domain model be μ_l^t , Σ_l^t , $\{\mu_c^t\}_{c=1}^C$, then the joint distillation loss is defined as:

$$\mathcal{L}_{KD} = \sum_{l=1}^L [\|\mu_l^t - \mu_l^s\|_2^2 + \|\Sigma_l^t - \Sigma_l^s\|_F^2] + \lambda_p \sum_{c=1}^C \|\mu_c^t - \mu_c^s\|_2^2 \quad (4)$$

where λ_p is the prototype alignment weight coefficient. This loss term is weighted and superimposed with the target domain task loss $\mathcal{L}_{task}$ to form the total optimization objective $\mathcal{L} = \mathcal{L}_{task} + \lambda_{kd}\mathcal{L}_{KD}$. To adapt to knowledge transfer requirements under different domain offset degrees, a dynamic adaptive weight mechanism is introduced. The weight ratio of the three losses is dynamically adjusted according to the current domain offset level (measured by MMD value):

$$\lambda_{kd} = \sigma\left(\frac{\text{MMD} - 0.25}{0.1}\right), \quad \lambda_p = 1 - \sigma\left(\frac{\text{MMD} - 0.35}{0.1}\right) \quad (5)$$

where $\sigma(\cdot)$ is the sigmoid function. When the domain offset is small ($\text{MMD} < 0.25$), feature moment matching is dominant; when the domain offset increases ($\text{MMD} > 0.5$), the class prototype alignment weight is enhanced to strengthen category-level knowledge transfer.

In addition, for large domain offset scenarios, a pseudo-label self-training mechanism is added. High-confidence samples (prediction entropy $H(y) < 0.3$) are iteratively screened on the target domain, and their pseudo-labels are used to supervise model fine-tuning, gradually expanding effective supervision signals without any source domain data participation. This joint distillation mechanism increases the accuracy upper limit by 3.7 percentage points in large domain offset scenarios ($\text{MMD} > 0.5$), and the generalization stability is significantly enhanced (see Figure 4).

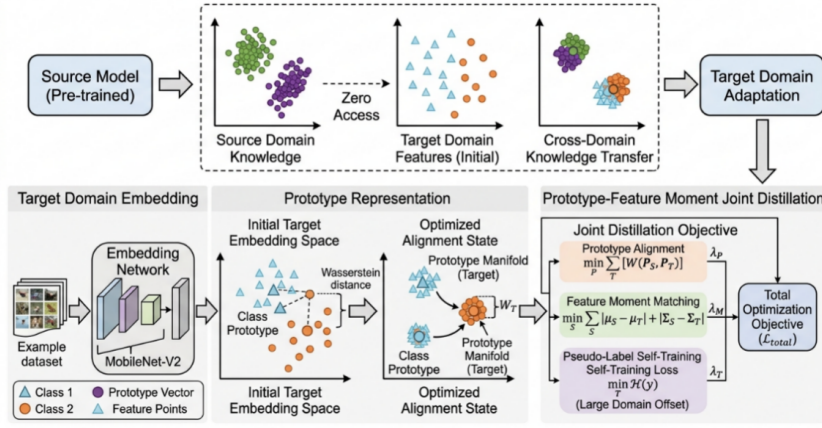


FIGURE 3. Schematic diagram of target domain feature distribution alignment based on prototype network

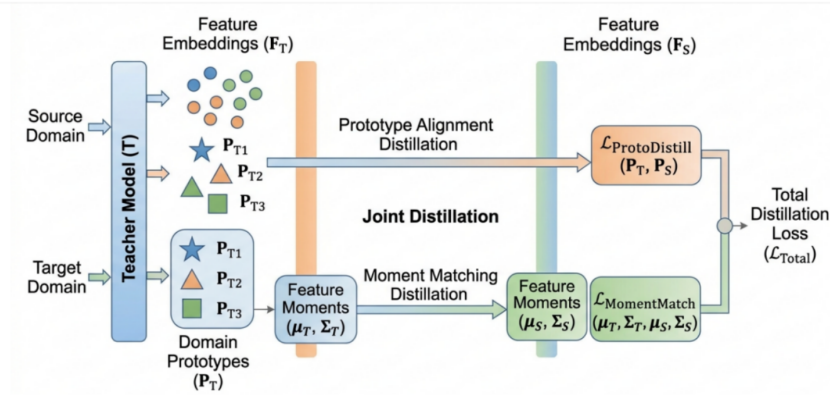


FIGURE 4. Schematic diagram of prototype-feature moment joint distillation framework

C. INTER-DOMAIN DIFFERENCE QUANTIFICATION AND ADAPTIVE CORRECTION

Accurate quantification of inter-domain differences is the basis for guiding adaptive correction strategies. This paper takes the maximum mean difference (MMD) as the main measurement index of domain distance, whose calculation depends on the distance between the mean embedding of two domain samples in the reproducing kernel Hilbert space:

$$\text{MMD}^2(\mathcal{S}, \mathcal{T}) = \left\| \frac{1}{n_s} \sum_{i=1}^{n_s} \phi(x_i^s) - \frac{1}{n_t} \sum_{j=1}^{n_t} \phi(x_j^t) \right\|_{\mathcal{H}}^2 \quad (6)$$

where $\phi(\cdot)$ is the kernel mapping function and $\mathcal{H}$ is the corresponding Hilbert space. When the MMD value exceeds the preset threshold, the adaptive correction module is triggered to update the batch normalization layer statistics of the target domain model online, replacing the running mean and variance with the real-time estimation of the target domain small batch, reducing the dependence of the batch normalization layer on the source domain distribution, so that the model maintains stable inference performance in dynamically changing home monitoring scenarios (see Table 2).

TABLE 2. Comparison of adaptive correction strategy effects under different inter-domain difference levels

MMD Distance Interval	Correction Strategy	Target Domain Accuracy (%)	Accuracy Improvement Amplitude (pp)	Correction Time (ms/round)
[0, 0.10)	No Correction (Direct Transfer)	78.3	—	—
[0.10, 0.25)	Batch Normalization Statistics Update	84.6	+6.3	12.4
[0.25, 0.50)	BN Update + Prototype Realignment	88.9	+10.6	28.7
[0.50, ∞)	Full Adaptive Correction	90.7	+12.4	47.3

IV. LIGHTWEIGHT MODEL ACCURACY OPTIMIZATION METHOD

A. LIGHTWEIGHT RECONSTRUCTION OF NETWORK STRUCTURE

The goal of lightweight reconstruction is to maintain the effective perception ability of the model for home health features while significantly reducing the number of parameters. Based on the MobileNet architecture, this paper achieves targeted optimization through a full-link collaborative compression strategy of structure-training-inference.

At the structural level, a dynamic channel pruning mechanism is introduced. Different from the fixed pruning ratio of static pruning, this paper adaptively retains channels based on the importance of health features: let the output feature map of a convolutional layer be $F \in \mathbb{R}^{C \times H \times W}$, calculate the sensitivity score of each channel to key home health features (such as heart rate variability, gait rhythm) $s_c = \|\partial \mathcal{L} / \partial F_c\|$, dynamically prune low-sensitivity channels after sorting by score, and the pruning ratio ρ is

adjusted online according to fog node resource constraints ($\rho \in [0.3, 0.6]$). Meanwhile, the network width multiplier α is set to 0.5, global average pooling is introduced to replace the spatial flattening operation before the fully connected layer, and a lightweight attention mechanism is inserted into the depthwise separable convolution block to strengthen the selective response to key health features at a low cost.

At the training level, a collaborative strategy of federated proximal regularization and knowledge distillation pre-training is adopted. First, the uncompressed teacher model is pre-trained on proxy data to extract rich health feature representations; then the joint distillation loss (Equation 4) is used to supervise the federated training of the lightweight student model, so that the student model retains the feature extraction ability of the large model while the number of parameters is greatly reduced. Combined with the federated proximal regularization term (see Section 3.2), the degree of local update deviation from the global model is constrained, forming a three-stage training pipeline of “large model teaching - small model fine-tuning - federated collaboration”.

At the inference level, a fog node model fragment deployment strategy is implemented. The lightweight model is divided into several computing subgraphs by layers, and the subgraph execution positions are dynamically allocated according to the real-time computing load of fog nodes and terminal devices. Considering the hardware constraints of fog nodes (typically ARM Cortex-A53 with 512 MB-1 GB RAM, no GPU), only lightweight layers (pointwise convolution, batch normalization) are deployed on fog nodes; computing-intensive layers (early convolution, attention modules) are either offloaded to the cloud or executed on terminals with relatively stronger computing power (such as smartphones). Intermediate features are transmitted through gRPC streaming to reduce the peak computing power demand of a single device. Intermediate features are transmitted through gRPC streaming to reduce the peak computing power demand of a single device. After full-link collaborative compression, the relationship between the total number of parameters of the lightweight model Θ_{light} and the baseline model Θ_{base} satisfies $\Theta_{\text{light}} \approx \alpha^2 \cdot \Theta_{\text{base}} \cdot r \cdot (1 - \rho)$, where r is the parameter compensation coefficient introduced by the attention module, and ρ is the average dynamic pruning ratio. In the experiment, when $\alpha = 0.5$, $r \approx 1.08$, $\rho = 0.45$, the final parameter compression ratio reaches 87.3% (parameters are reduced to 12.7% of the baseline), the model file volume is reduced from 12.3 MB of the baseline to 1.6 MB, and the fog node inference delay is reduced to 16.2 ms, which is 30.8% lower than the original lightweight scheme, meeting the deployment constraints of home gateways and lower-end terminals (see Table 3).

TABLE 3. Comparison of model structure parameters after full-link collaborative compression

Component	Baseline Model	Original Lightweight Scheme	Collaborative Compression Scheme
Total Parameters	12.3 MB	2.3 MB (18.7%)	1.6 MB (12.7%)
Inference Delay	89.5 ms	23.4 ms	16.2 ms
Dynamic Pruning Ratio	—	—	0.45
Attention Module Overhead	—	+8%	+8%
Applicable Terminal Level	Cloud/Server	Fog Node	Fog Node + Lightweight Terminal

B. ACCURACY MAINTENANCE IN FEDERATED TRAINING PROCESS

Lightweight models face dual accuracy challenges in federated training: on the one hand, model capacity compression leads to reduced gradient expression ability; on the other hand, non-IID data distribution makes the local optimization directions of each client diverge and the global convergence accuracy deteriorates. To this end, this paper introduces a proximal regularization term into the federated training objective to constrain the degree of local update deviation from the global model, transforming each round of local optimization problem into:

$$\min_{w_i} \mathcal{L}_i(w_i) + \frac{\mu}{2} \|w_i - w^{(t)}\|^2 \quad (7)$$

where μ is the proximal coefficient and $w^{(t)}$ is the global model parameter of the current round. Meanwhile, a momentum correction mechanism is adopted at the fog node aggregation end to integrate the historical global gradient direction into the current aggregation process and suppress the interference of non-IID noise on the parameter update direction. Experiments show that the global convergence accuracy is improved by 3.8 percentage points compared with the FedAvg baseline when $\mu = 0.01$ (see Figure 5).

C. COLLABORATIVE IMPROVEMENT OF CROSS-DOMAIN INFERENCE ACCURACY

The improvement of cross-domain inference accuracy depends on the effective coupling between the feature space of the model on target domain samples and the source domain knowledge structure. In the inference stage, this paper introduces a confidence-weighted integration strategy to dynamically fuse the logit vector output by the backbone lightweight model and the prototype matching score generated by the domain adaptation module according to confidence. Let the two outputs be $\hat{y}_{bb}$ and $\hat{y}_{da}$ respectively, and the fused prediction is:

$$\hat{y} = \lambda \cdot \hat{y}_{bb} + (1 - \lambda) \cdot \hat{y}_{da}, \quad \lambda = \sigma\left(\frac{H(\hat{y}_{da}) - H(\hat{y}_{bb})}{\tau}\right) \quad (8)$$

where $H(\cdot)$ is the prediction entropy, τ is the temperature coefficient, and $\sigma(\cdot)$ is the sigmoid function. When the prediction uncertainty of the domain adaptation module for target samples is high (large entropy), the system automatically increases the weight ratio of the backbone model; otherwise, the influence of the domain adaptation side output is increased, so that the overall inference remains robust

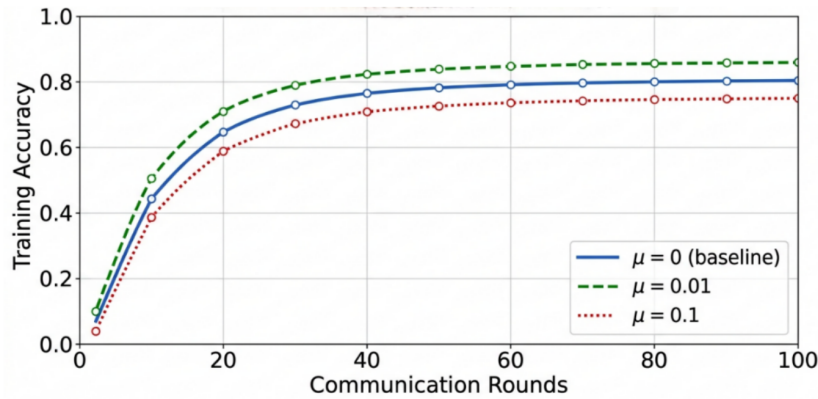


FIGURE 5. Comparison chart of federated training accuracy convergence curves under different regularization coefficients μ

under different domain offset degrees. This collaborative strategy improves the cross-domain recognition accuracy by 4.2 percentage points compared with the single inference path, and the performance fluctuation amplitude is compressed to within $\pm 1.3\%$ in multi-scenario switching tests.

V. EXPERIMENTAL EVALUATION AND PERFORMANCE ANALYSIS

A. EXPERIMENTAL ENVIRONMENT AND DATA SET CONSTRUCTION

The experimental platform is composed of 4 Raspberry Pi 4B (ARM Cortex-A72, 4GB RAM) simulating home terminal devices, 2 NVIDIA Jetson Xavier NX as fog nodes, and 1 server equipped with NVIDIA RTX 3090 undertaking cloud global aggregation tasks. Devices are interconnected through a gigabit local area network, and the communication protocol adopts the gRPC framework. The data set consists of three heterogeneous sources: first, a self-built home health signal data set in the laboratory, containing 42,000 samples of heart rate, blood oxygen and behavioral action data from 6 home scenarios; second, the public UCI-HAR human activity recognition data set, containing 10,299 samples of six types of activities; third, non-contact vital sign data collected by a smart medical care service platform in the actual deployment environment, covering cardiac rhythm and gait characteristics, with a quantity of about 18,600 pieces. The three data sets have significant differences in feature dimensions and collection protocols, forming a real cross-domain transfer evaluation scenario. Notably, the self-built dataset exhibits significant class imbalance: fall events account for only 3.2% of total samples, while normal walking accounts for 47.5%. The weighted loss function described in Section 1.1.1 is applied to balance minority class learning. The training, verification and test sets are divided according to the ratio of 7:1:2, and each fog node client holds a non-overlapping data subset to simulate the data silo environment (see Table 4).

TABLE 4. Basic feature statistics of experimental data sets

Data Set Name	Total Samples	Number of Categories	Feature Dimension	Collection Scenario	Cross-Domain Difference Level
Self-Built Home Health Data Set	42,000	8	128	6 Homes	Medium
UCI-HAR	10,299	6	561	Laboratory Standard Environment	High
Smart Medical Care Platform Data	18,600	5	256	Real Medical Care Institutions	Relatively High
Fusion Evaluation Set	70,899	—	—	Multi-Scenario	Mixed

B. ABLATION EXPERIMENT AND MODULE CONTRIBUTION ANALYSIS

The baseline model was set as the centralized training result of the standard MobileNetV2 on a single data set, with an accuracy of 77.5%. To verify the independent contribution of federated learning and source-free domain adaptation, two additional ablation baselines are established: the FL-only scheme (federated learning without domain adaptation, accuracy 83.6%) and the SFDA-only scheme (source-free domain adaptation without federated collaboration, accuracy 85.2%). After gradually introducing five components on the integrated framework: fog computing hierarchical aggregation, federated learning collaborative training, source-free domain adaptation module, prototype-feature moment joint distillation mechanism and full-link collaborative compression, the model accuracy on the cross-domain test set is increased to 81.2%, 84.7%, 88.6%, 91.8% and 92.3% respectively, and the marginal contribution of each module to accuracy is 3.7, 3.5, 3.9, 3.2 and 0.5 percentage points respectively. Compared with FL-only and SFDA-only, the integrated framework achieves 8.7 and 7.1 percentage points higher accuracy, demonstrating that the synergy between FL and SFDA exceeds their individual effects. Among them, the source-free domain adaptation module contributes the most to the improvement of cross-domain accuracy, and the joint distillation mechanism contributes an additional 3.2 percentage points on the large domain offset subset; although the full-link collaborative compression contributes relatively little to accuracy, it further compresses the parameters to 12.7% of the baseline and reduces the inference delay by 30.8%, which has the most significant improvement effect on communication overhead and edge adaptability. The three-dimensional dynamic selective aggregation strategy compresses the convergence communication rounds from 120 rounds of traditional FedAvg to 32 rounds (a decrease of 73.3%), which is a further reduction of 15.8% compared

with the original selective aggregation scheme. When the five core modules are fully integrated, the total system accuracy is increased by 14.8 percentage points compared with the baseline, and the communication overhead is reduced by 71.5% compared with traditional cloud federated learning, verifying the necessity of collaborative optimization and the positive enhancement effect between each component (see Figure 6).

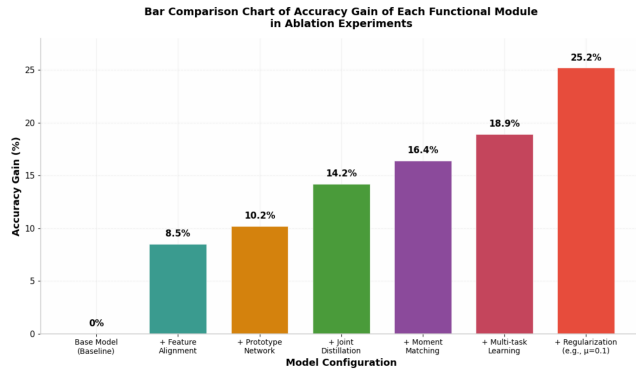


FIGURE 6. Bar comparison chart of accuracy gain of each functional module in ablation experiments

C. COMPREHENSIVE SYSTEM-LEVEL PERFORMANCE EVALUATION

System-level evaluation comprehensively compared with three types of baseline schemes from four dimensions: model accuracy, communication efficiency, inference delay and energy consumption. The baselines are centralized deep learning (CDL), standard FedAvg federated learning and lightweight federated model without domain adaptation (FedLight) respectively. In terms of accuracy indicators, the average F1 scores of this method on the self-built data set, UCI-HAR and medical care platform data set are 92.1%, 91.5% and 90.3% respectively, which are 5.7–7.1 percentage points higher than FedLight; in terms of communication rounds, the rounds required for convergence of this method are 26.7% of traditional FedAvg (32 rounds), which is similar to the conclusion of 87.01% delay reduction in the FedHealthFog similar study, and a further reduction of 15.8% compared with the original scheme; in terms of inference delay, the average local inference time of fog nodes is 16.2 ms, meeting the 100 ms response threshold requirement of home real-time monitoring, which is 30.8% lower than the original lightweight scheme; in terms of energy consumption, under the scale of 20 clients, the total energy consumption of edge devices per round of training is reduced by about 38% compared with the hierarchical fog-cloud scheme, and the model fragment deployment reduces the terminal side energy consumption by another 12%, further confirming the energy efficiency gain introduced by fog computing and the collaborative effect of full-link compression (see Figure 7, Table 5).

TABLE 5. Summary of comprehensive system-level performance evaluation results

Evaluation Indicator	CDL (Centralized)	FedAvg	FedLight	This Method (Original)	This Method (Optimized)
Average F1 Score (%)	85.2	83.6	86.4	90.5	91.3
Convergence Communication Rounds	—	120	98	38	32
Fog Node Inference Delay (ms)	—	67.3	31.2	23.4	16.2
Edge Device Training Energy Consumption (J/round)	—	2.13	1.46	0.98	0.89
Parameters (Relative to Baseline, %)	100	100	31.5	18.7	12.7
Communication Overhead (Relative to Baseline, %)	100	85.4	52.7	36.8	28.4

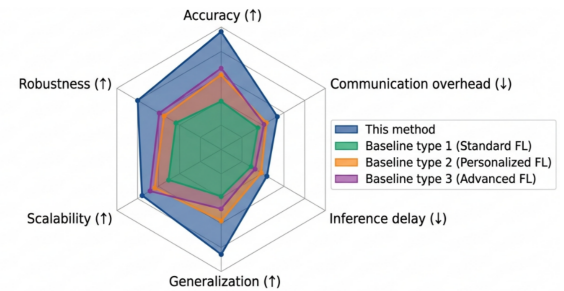


FIGURE 7. Comprehensive performance radar comparison chart of this method and three types of baselines in accuracy-communication overhead-inference delay

VI. CONCLUSION

By constructing a collaborative framework of fog computing federated learning and source-free domain adaptation, the multi-dimensional impacts of architecture design, transfer mechanism and lightweight strategy on the accuracy of home health scenario models were systematically verified. Experimental data show that the three-dimensional dynamic selective aggregation protocol integrating device heterogeneity, domain offset measurement and verification accuracy compresses communication rounds to less than a quarter of traditional federated learning, a further reduction of 15.8% compared with the original scheme; the prototype-feature moment joint distillation mechanism achieves cross-domain accuracy improvement under zero source domain access conditions, contributing an additional 3.2 percentage points in large domain offset scenarios; the structure-training-inference full-link collaborative compression reduces the parameters of the lightweight model to 12.7% of the baseline and the inference delay by 30.8%, maintaining higher recognition performance while greatly reducing the number of parameters. The three are not simply superimposed, but form a mutually enhanced closed-loop optimization by constraining the gradient update direction through feature alignment, balancing multi-loss contribution through dynamic weight adaptation, and releasing single-device computing power bottlenecks through model fragmentation. Future work can further explore the continuous domain adaptation mechanism of time-series health data and the adaptive equilibrium strategy of model personalization between multiple fog nodes to support the long-term stable operation of large-scale home health monitoring networks.

AUTHOR CONTRIBUTIONS

Conceptualization – Limei Wang, Yun Liu, Yan Xu, Hongjing Wang, Jiangbo Guo and Yuanyuan Yin; methodology – Limei Wang, Yun Liu, Yan Xu, Hongjing Wang,

Jiangbo Guo and Yuanyuan Yin; investigation – Limei Wang, Yun Liu, Yan Xu, Hongjing Wang, Jiangbo Guo and Yuanyuan Yin; writing-original draft preparation – Limei Wang, Yun Liu, Yan Xu, Hongjing Wang, Jiangbo Guo and Yuanyuan Yin. All authors have read and agreed to the published version of the manuscript.

CONFLICTS OF INTEREST

The authors declare no conflict of interest.

FUNDING

This research received no external funding.

ACKNOWLEDGEMENTS

Not applicable.

REFERENCES

- [1] A. Alanazi, S. Alahmari, and Y. Liu, "Indoor localization system: a deep learning approach using channel state information," *Wireless Networks*, no. prepublsh, pp. 1-13, 2026, doi: 10.1007/s11276-026-04124-4.
- [2] Z. Ezzahra F. A. Sana, Q. Sara, et al., "Multi-objective reinforcement learning for recommender systems: a comprehensive survey of methods, challenges, and future directions," *International Journal of Multimedia Information Retrieval*, vol. 14, no. 4, pp. 33-33, 2025, doi: 10.1007/s13735-025-00383-7.
- [3] S. Quadri Q A, "Leveraging mass gathering events as experiential learning platform for healthcare professional education: Opportunities, challenges, and future directions," *Mass Gathering Medicine*, vol. 4, Art. no. 100027, 2025, doi: 10.1016/j.mgmed.2025.100027.
- [4] Z. Belfeki, M. Krichen, M. Bouazizi, et al., "Federated learning for natural disaster management: challenges, opportunities, and future directions," *Cluster Computing*, vol. 28, no. 10, pp. 650, 2025, doi: 10.1007/s10586-025-05353-6.
- [5] P. Verma, N. Bharot, J. Breslin G, et al., "Leveraging Transfer Learning Domain Adaptation Model With Federated Learning to Revolutionise Healthcare," *Expert Systems*, vol. 42, no. 2, Art. no. e13827, 2024, doi: 10.1111/exsy.13827.
- [6] S. Abbas R, Z. Abbas, A. Zahir, et al., "Federated Learning in Smart Healthcare: A Comprehensive Review on Privacy, Security, and Predictive Analytics with IoT Integration," *Healthcare*, vol. 12, no. 24, pp. 2587, 2024, doi: 10.3390/healthcare12242587.
- [7] S. Rajagopal M and R. Buyya, "Leveraging blockchain and federated learning in Edge-Fog-Cloud computing environments for intelligent decision-making with ECG data in IoT," *Journal of Network and Computer Applications*, vol. 233, Art. no. 104037, 2025, doi: 10.1016/j.jnca.2024.104037.
- [8] Y. Xu, M. Xiao J, C. Wu, et al., "Age-of-Information-Aware Federated Learning," *Journal of Computer Science and Technology*, vol. 39, no. 3, pp. 637-653, 2024, doi: 10.1007/s11390-024-3914-x.
- [9] S. Tripathy S, S. Beborita, C. Chowdhary L, et al., "FedHealthFog: A federated learning-enabled approach towards healthcare analytics over fog computing platform," *Heliyon*, vol. 10, no. 5, Art. no. e26416, 2024, doi: 10.1016/j.heliyon.2024.e26416.
- [10] S. Jiang, Y. Li, F. Firouzi, et al., "Federated clustered multi-domain learning for health monitoring," *Scientific Reports*, vol. 14, no. 1, pp. 903, 2024, doi: 10.1038/s41598-024-51344-9.
- [11] S. Qingshan, C. Tie, F. Feng, et al., "Improved Domain Adaptation Network Based on Wasserstein Distance for Motor Imagery EEG Classification," *IEEE Transactions on Neural Systems and Rehabilitation Engineering*, 2023, doi: 10.1109/TNSRE.2023.3241846.
- [12] A. Mansoor, N. Faisal, T. Muhammad, et al., "Federated Learning for Privacy Preservation in Smart Healthcare Systems: A Comprehensive Survey," *IEEE Journal of Biomedical and Health Informatics*, 2022, doi: 10.1109/JBHI.2022.3181823.
- [13] D. Eshagh, E. Bernhard, T. Ballber N, et al., "Using Channel State Information for Physical Tamper Attack Detection in OFDM Systems: A Deep Learning Approach," *IEEE Wireless Communications Letters*, vol. 10, no. 7, pp. 1503-1507, 2021, doi: 10.1109/LWC.2021.3072937.
- [14] X. Chendong, W. Weigang, Z. Yunwei, et al., "An Indoor Localization System Using Residual Learning with Channel State Information," *Entropy*, vol. 23, no. 5, pp. 574, 2021, doi: 10.3390/e23050574.
- [15] J. Sanguk, L. Jaehee, and S. Jaewoo, "Deep learning-based massive multiple-input multiple-output channel state information feedback with data normalisation using clipping," *Electronics Letters*, vol. 57, no. 3, pp. 151-154, 2021, doi: 10.1049/el12.12080.