

Fair Federated Learning for Trustworthy Service Provisioning in Disaggregated Optical Networks

Saroj Kumar Panda*, Tania Panayiotou[†], Sadananda Behera*, and Georgios Ellinas[†]

* Dept. of Electronics and Communication Eng., National Institute of Technology Rourkela, Odisha, India

[†] KIOS CoE and Dept. of Electrical and Computer Eng., University of Cyprus, Nicosia, Cyprus

Abstract—Federated learning (FL), that keeps data isolated at local clients, has emerged as an effective and promising solution for mitigating privacy concerns in disaggregated networks where multiple operators independently manage different segments. Ensuring fairness in traffic prediction and resource allocation under heterogeneous traffic datasets remains, however, challenging. Fair FL, through optimization of q -fair FL (q -FFL) function, seeks to balance accuracy across clients by tuning a single inequality aversion parameter, q . As this requires exploration of a computationally prohibitive parameter space, a metaheuristic-guided q -FFL framework is introduced, which effectively reduces the search space while also providing a flexible and cost-efficient framework for collaborative optical networks to best meet their targeted quality-of-service (QoS) criteria. It is demonstrated that metaheuristic-guided q -FFL effectively balances global traffic accuracy and cross-operator fairness, while reducing the parameter search space by up to 70%. An elastic optical network is considered to demonstrate the impact of fairer allocations on collaborative service provisioning demonstrating up to 54% improvement in QoS-based fairness across operators.

I. INTRODUCTION

With the rapid proliferation of advanced technologies such as 6G networks, the Internet of Things, and applications such as virtual reality, smart healthcare, and smart cities, the demand for seamless and high-quality network connectivity has risen exponentially. This surge in traffic places enormous pressure on network service providers to enhance their infrastructure while maintaining cost efficiency. As a result, the use of traffic-driven AI-powered frameworks has become crucial for network optimization, enabling more effective planning and reducing energy consumption and operational costs [1]. In parallel, optical network infrastructures have been evolving from proprietary, closed systems to disaggregated architectures [2], enabling programmability and automation, and ultimately forming service-oriented networks.

In these networks, AI trustworthiness for collaborative service provisioning goes beyond raw predictive accuracy [1]. Specifically, operators need decisions that are reliable, risk-aware, and equitable under data drift, heterogeneous segments, and strict QoS constraints. In disaggregated networks, this entails, among others, uncertainty calibration so that traffic-driven optimizers can quantify the risk of under-/over-provisioning and set margins accordingly [3], as well as privacy and fairness, so that learning does not systematically advantage or disadvantage particular operators or network segments [4], [5]. In this work, the focus is on fairness in collaborative learning across operators through FL objectives that effectively balance global accuracy and cross-client equity trade-off, while reducing the parameter search space. FL is

attractive in disaggregated settings because it mitigates data-privacy concerns by keeping raw data local, while enabling collaborative training across operators [6]. However, ensuring cross-client fairness remains a challenge under heterogeneous and imbalanced datasets.

Federated learning in disaggregated networks has previously been studied in [4], [5], [7]–[11]. The vast majority of those works mainly focus on demonstrating FL's ability to achieve performance accuracies comparable to the centralized learning paradigm. Our previous work, focusing on fair FL for collaborative resource allocation [4], [5], adopts q -fair FL (q -FFL) [6], which introduces an inequality-aversion parameter, q , to trade global accuracy for equity. A drawback is that identifying an acceptable operating point typically requires sweeping multiple q values, each necessitating a full FL training run, making the accuracy–fairness tuning process computationally expensive, especially as the number of participants grows. This work introduces a simulated annealing (SA)–guided fair FL scheme that adaptively explores the q -space, significantly reducing the tuning cost, while effectively balancing the global accuracy–fairness trade-off. The impact of fairer traffic predictions across operators on service provisioning is examined considering an elastic optical network (EON).

II. SA-GUIDED FAIR FEDERATED LEARNING

In q -FFL the global objective function to be minimized is given by $f_q = \sum_{k=1}^m \frac{p_k}{q+1} [F_k(w)]^{q+1}$, where m is the number of clients, F_k and n_k are the loss function and the number of samples available at the k th client, respectively, $n = \sum_k n_k$ is the total number of samples, $p_k = \frac{n_k}{n}$, and q is a scalar that controls the trade-off between model accuracy and fairness. When $q=0$, q -FFL reduces to the traditional FL objective, maximizing efficiency in model accuracy (i.e., results in the least-fair scheme). As q increases, more uniformity is imposed on the training loss, improving the fairness across clients. Setting q to a large enough value results in min-max fairness [12], which is considered to be the fairest scheme as it minimizes the maximum loss encountered at a client.

In a nutshell, during training, the loss function gives more weight to clients with higher losses, aiming to reduce performance disparities across clients. Specifically, each client (i.e., operator) trains a local model on its local traffic traces and computes its own loss function F_k , raised to the power of q , so that operators with higher losses have a proportionally greater impact. Then, a server (i.e., orchestrator) aggregates the weighted updates from all clients. As the updates from clients with higher losses are weighted more heavily, the global model

is updated in a way that improves fairness across the operators. The process of local training, aggregation, and global model updates, is repeated until convergence is achieved. The main limitation of this scheme is that to appropriately balancing the global accuracy-fairness trade-off requires iterations over several q values, significantly increasing the computational overhead associated with training all the q -FFL-models.

To reduce the q -search space this work introduces simulated annealing (SA) for q -selection. The adopted SA is a deterministic metaheuristic that explores the search space under a cooling schedule, focusing evaluations on progressively better regions. Specifically, q -tuning is cast as an optimization problem over the composite objective:

$$J(q) = w_1 \cdot \frac{CV(q) - \gamma_{min}}{\gamma_{max} - \gamma_{min}} + w_2 \cdot \frac{G(q) - \delta_{min}}{\delta_{max} - \delta_{min}}, \quad (1)$$

that balances global accuracy $G(q) = \sum_{k=1}^m \frac{1}{m} F_k$ and cross-operator fairness calculated via the coefficient of variation

$CV(q) = \sqrt{\frac{m^2 \sum_k (F_k - \frac{1}{m} \sum_k F_k)^2}{m-1 (\sum_k F_k)^2}}$. In essence, Eq. (1) is a weighted sum of normalized global loss (i.e., mean squared error (MSE)) and CV-based fairness [13] obtained after q -FFL training, where γ, δ are the normalization parameters, and w_1, w_2 (with $w_1 + w_2 = 1$) are the weights used to denote the importance of fairness/accuracy according to the operators' targets (i.e., with $w_1 = 1/w_2 = 0$ the fairest model is opted, with $w_2 = 1/w_1 = 0$ the goal is to maximize global accuracy, while when $w_1 = 0.5/w_2 = 0.5$ the aim is to strike a balance between global model accuracy and fairness).

Starting from $q_0 \in [q_{min}, q_{max}]$ and temperature T_0 , SA iteratively explores the state space, concentrating evaluations on promising regions of the q -space, significantly reducing the number of runs compared to exhaustive sweeps. It should be noted that in this setting, each evaluation corresponds to a single FL training run at the candidate q . At each iteration i , a candidate solution is generated as $q' = \Pi_{[q_{min}, q_{max}]}(q + \beta_i)$, where Π denotes the projection onto the feasible search space and $\beta_i \in \{\beta_{min}, \beta_{max}\}$ is randomly sampled. A move is accepted if the improvement is $\Delta J = J(q') - J(q) \leq 0$; otherwise, it is rejected. The temperature decreases according to a geometric cooling schedule, $T_{i+1} = \alpha * T_i$ and the search terminates once either N iterations are completed or the temperature falls below the minimum threshold T_{min} .

III. FEDERATED TRAFFIC DATASETS

For each operator k , a different federated dataset $D_k = \{\mathbf{x}_t^{(k)'}\}_{t=1}^{n_k}$ is created, where $\mathbf{x}_t^{(k)'}$ is a traffic sequence that includes the past and present traffic observations, $\mathbf{x}_t^{(k)'} = [x_{t-z}^{(k)}, \dots, x_{t-1}^{(k)}, x_t^{(k)}]$, and $\mathbf{x}_t^{(k)}$ is the next traffic observation that the federated model is trained to predict, $\mathbf{x}_t^{(k)} = x_{t+1}^{(k)}$, where z is the number of past observations, t is the current time instant, and past and present traffic data are collected at regular intervals of h time units. For federated dataset generation, the real-world traffic dataset of the 12-node, 30-link Abilene backbone network (<http://sndlib.zib.de/home.action>) is utilized, providing bit-rate information every 5-minutes for every node pair, for a 6-month period. For simplicity, it is assumed that each operator manages a single network node (e.g., node k). Hence, each D_k is created to capture the

aggregated bit-rate at node k every 5 minutes. Traffic patterns for each sequence in D_k are then created with $z=70$, following the sliding window approach. In total, 5 federated datasets are created, with each D_k varying in size to evaluate the performance of q -FFL on imbalanced datasets. Specifically, $D_1=3000, D_2=2000, D_3=8000, D_4=5000, D_5=7500$. To introduce heterogeneity (form non-iid datasets), noise is sampled from various distributions (Gaussian, Weibull, Lognormal, Gamma) and infused into the traffic traces, while datasets are normalized according to a common, large enough value, to keep the CV over the scaled and unscaled values consistent.

IV. SA-GUIDED MODEL TRAINING AND EVALUATION

For q -FFL, both global and local models are long-short-term memory (LSTM) neural networks, with 2 hidden layers, capable of capturing the non-linear long-term dependencies in network traffic sequences [1]. To optimize q -FFL, the fair FedAvg algorithm described in [4] is applied with a batch size=256, learning rate=0.001, and MSE local loss functions. Each dataset D_k is partitioned so that the last 100 sequences are retained for testing, while from the remaining sequences 80% are used for training and 20% for validation.

Parameter q is optimized according to the SA metaheuristic following the $J(q)$ objective (Eq. (1)). For the SA-guided model training phase, the state space is set according to $q_{min} = 0, q_{max} = 3, T_0 = 1.0, T_{min} = 0.1, \alpha = 0.99, \beta_{min} = -0.375, \beta_{max} = 0.375$ with a step size of 0.0625, and $N=100$. At each SA iteration, the q -FFL model is optimized once, according to 10 training iterations, for each candidate q generated. After SA optimization, the q^* -FFL model for which $J(q^*)$ is minimum is selected. Note that $\gamma_{max}, \gamma_{min}, \delta_{max}$, and δ_{min} were set to 0.78, 0.19, 0.02 and 0.01 (obtained empirically).

The choice of SA parameters (i.e., α and β range) significantly impacts the effectiveness of the proposed method. These are effectively tuned by considering a small state space (i.e., q values ranging from 0 to 3, with the step size set to 0.1875). Note that $T_0=1, T_{min}=0.1$, are set as default values throughout the parameter tuning process, while α is tuned over the 0.8-0.99 range. Further, the q -FFL models trained through the parameter tuning process are stored and reused throughout, thus gaining advantage during SA optimization for q -selection.

In this work, $J(q)$ is optimized with the objective of (i) balancing the global accuracy-fairness trade-off (i.e., $w_1=0.5$), (ii) prioritizing fairness over global accuracy (i.e., $w_1=0.75$) and (iii) approximating min-max fairness ($w_1=1$). Table I summarizes the results obtained as SA evolves. Note that the q values explored are listed in ascending order, even though in practice the order they are explored depends on the SA evolution. For all cases, SA returns the best q^* values resulting in the minimum $J(q^*)$; that is, $q^* = 0.3125, 0.6875, 0.9375$ for $w_1 = 0.5, 0.75, 1$, respectively. As expected, for higher w_1 values, higher q^* values are returned (i.e., as q increases fairer solutions are expected, equivalently resulting in lower $CV(q^*)$ [4]). In all cases, SA explored a range of 16 to 22 different q values, resulting in up to 67% exploration reduction, as opposed to the case where all 49 q values are explored. Subsequently, this also implies a significant reduction in training time, computational resources, and energy consumption. Importantly, SA-guided q -FFL is capable

TABLE I
SA EVOLUTION: q VALUES EXPLORED WITH RESULTING $G(q)$, $CV(q)$, $J(q)$ FOR $w_1 = 0.5, 0.75, 1$

$w_1 = 0.5$				$w_1 = 0.75$				$w_1 = 1$			
q	$G(q)$	$CV(q)$	$J(q)$	q	$G(q)$	$CV(q)$	$J(q)$	q	$G(q)$	$CV(q)$	$J(q)$
0	0.0121	1.005	0.8006	0.3125	0.0116	0.3020	0.1826	0.5625	0.0141	0.2242	0.058
0.0625	0.0107	0.7860	0.5410	0.375	0.0122	0.2737	0.1622	0.4375	0.0128	0.2537	0.1526
0.125	0.0102	0.5346	0.3061	0.4375	0.0128	0.2537	0.1526	0.6875	0.0150	0.2016	0.0197
0.1875	0.0105	0.4092	0.2128	0.5	0.0135	0.2386	0.1494	0.75	0.0154	0.1963	0.0108
0.25	0.0110	0.3427	0.1802	0.5625	0.0141	0.2242	0.1410	0.8125	0.0157	0.1930	0.0051
0.3125	0.0116	0.3020	0.1752	0.6875	0.0150	0.2016	0.1406	0.875	0.016	0.1912	0.0020
0.375	0.0122	0.2737	0.1825	0.75	0.0154	0.1963	0.1433	0.9375	0.0162	0.1907	0.0012
0.4375	0.0128	0.2537	0.1972	0.8125	0.0157	0.1930	0.1471	1	0.0164	0.1915	0.0025
0.5	0.0135	0.2386	0.2164	0.875	0.0160	0.1912	0.1517	1.0625	0.0166	0.1933	0.0056
0.5625	0.0141	0.2242	0.2348	0.9375	0.0162	0.1907	0.1507	1.125	0.0167	0.1956	0.0096
0.625	0.0146	0.2104	0.2474	1	0.0164	0.1915	0.1629	1.1875	0.0169	0.1985	0.0144
0.6875	0.0150	0.2016	0.2614	1.0625	0.0166	0.1933	0.1696	1.25	0.0170	0.2018	0.0200
0.75	0.0154	0.1963	0.2759	1.125	0.0167	0.1956	0.1766	1.375	0.0173	0.2090	0.0332
0.8126	0.0157	0.1930	0.2891	1.375	0.0173	0.2090	0.2068	1.5	0.0175	0.2162	0.0445
0.8125	0.0157	0.1930	0.2891	1.5	0.01756	0.2162	0.2225	1.875	0.0185	0.2337	0.0740
0.875	0.0160	0.1912	0.3013	1.875	0.0185	0.2337	0.2681	-	-	-	-
0.9375	0.0162	0.1907	0.3127	-	-	-	-	-	-	-	-
1	0.0164	0.1915	0.3127	-	-	-	-	-	-	-	-
1.0625	0.0166	0.1933	0.3336	-	-	-	-	-	-	-	-
1.125	0.0167	0.1956	0.3436	-	-	-	-	-	-	-	-
1.375	0.0173	0.2090	0.3814	-	-	-	-	-	-	-	-
1.5	0.0175	0.2162	0.40004	-	-	-	-	-	-	-	-
1.875	0.0185	0.2337	0.4622	-	-	-	-	-	-	-	-

TABLE II
IMPACT OF SA-GUIDED q^* -FFL SOLUTIONS ON PROACTIVE RSA

w_1/q^*	O (Gbps / SSs)	U (Gbps / SSs)	CV_{QoS}
0/0	30.3 / 1	12.15 / 0.37	0.34
0.5/0.3125	41.22 / 1.32	5.5 / 0.17	0.24
0.75/0.6875	52.51 / 1.7	2.2 / 0.07	0.18
1/0.9375	55.88 / 1.8	1.68 / 0.05	0.15

of exploring optimization functions that differently prioritize accuracy-fairness terms in $J(q)$, using pre-trained q -FFL models; indicatively, $J(q)$ with $w_1 = 0.5$ is optimized by exploring an additional seven q values compared to $J(q)$ with $w_1 = 0.75$, thus providing a flexible and cost-efficient framework for an operator to best meet (tune) targeted QoS.

V. TRAFFIC-DRIVEN RESOURCE ALLOCATION

Table II reports the effect of the selected q^* -FFL models on proactive routing and spectrum allocation (RSA), measured via over-/under-provisioning (O/U), in Gbps and spectrum slots (SSs), utilizing the bit-rate predictions in test datasets.

Additionally, $CV_{QoS} = \sqrt{\frac{m^2}{m-1} \frac{\sum_k (u_k + o_k - \frac{1}{m} \sum_k (u_k + o_k))^2}{(\sum_k (u_k + o_k))^2}}$ is computed to measure the fairness of the spectrum allocations (i.e., the QoS-based fairness) as q^* increases, where u_k and o_k are defined as the absolute difference between the true traffic demand (x^k) and the predicted demand ($\hat{x}^k$) in Gbps, where u_k applies when $x^k < \hat{x}^k$ and o_k applies when $x^k > \hat{x}^k$.

For each scenario, the conventional RSA algorithm (i.e., shortest path routing and first-fit spectrum allocation) runs on an EON [Abilene topology, BPSK/QPSK/8-QAM/16-QAM modulation formats, 320 frequency slots/link (12.5 GHz spacing, 10.5 Gbaud)], with each client mapped to a s - d pair. As the RSA results in zero blocking for all scenarios examined, O/U metrics are calculated as the average positive/negative prediction errors across clients and test instances. In this problem setting, relative to the least-fair ($q=0$) baseline, fairer models increase overprovisioning by up to 45% and significantly reduce underprovisioning by up to 86%, leading to significant improvement in QoS violations. Importantly, fairness in the achievable QoS across clients (CV_{QoS}) improves by up to 54%, mitigating the negative side-effects on overprovisioning.

VI. CONCLUSIONS

An SA-guided q -FFL framework is introduced that optimizes a composite accuracy-fairness objective, providing a flexible yet computationally efficient way to address heterogeneity in federated traffic datasets, while enabling tunable QoS for traffic-driven collaborative service provisioning in optical networks. Future work will focus on client grouping techniques to address scalability and on contention diversity across groups of clients.

ACKNOWLEDGMENTS

This work was supported by the Republic of Cyprus through the Deputy Ministry of Research, Innovation, and Digital Policy.

REFERENCES

- [1] T. Panayiotou *et al.*, “Survey on machine learning for traffic-driven service provisioning in optical networks,” *IEEE Communications Surveys & Tutorials*, vol. 25, no. 2, pp. 1412–1443, 2023.
- [2] S. Chen *et al.*, “The evolution of open and disaggregated optical networks: From open line system to open box system,” in *Proc. OFC*, 2024, p. Th1G.6.
- [3] T. Panayiotou and G. Ellinas, “Addressing traffic prediction uncertainty in multi-period planning optical networks,” in *Proc. OFC*, 2022.
- [4] S. K. Panda *et al.*, “Federated traffic-driven resource allocation in disaggregated networks,” in *Proc. ICTON*, 2025.
- [5] —, “A fair federated learning framework for collaborative network traffic prediction and resource allocation,” in *Proc. IEEE ICC*, 2025.
- [6] T. Li *et al.*, “Fair resource allocation in federated learning,” *arXiv:1905.10497 [cs.LG]*, 2019.
- [7] S. Behera *et al.*, “Federated learning for network traffic prediction,” in *Proc. IFIP Networking Conference*, 2024.
- [8] G. Drainakis *et al.*, “From centralized to federated learning: Exploring performance and end-to-end resource consumption,” *Computer Networks*, vol. 225, p. 109657, 2023.
- [9] K. Abdelli *et al.*, “Federated learning approach for lifetime prediction of semiconductor lasers,” in *Proc. OFC*, 2022.
- [10] M. Asad *et al.*, “Federated learning versus classical machine learning: A convergence comparison,” *arXiv:2107.10976 [cs.LG]*, 2021.
- [11] N. Hashemi *et al.*, “Vertical federated learning for privacy-preserving ML model development in partially disaggregated networks,” in *Proc. ECOC*, 2021.
- [12] M. Mohri *et al.*, “Agnostic federated learning,” in *Proc. ICML*, 2019.
- [13] T. Panayiotou and G. Ellinas, “Balancing efficiency and fairness in resource allocation for optical networks,” *IEEE Trans. Netw. Serv. Manag.*, vol. 21, no. 1, pp. 389–401, 2024.