



Private, Fair and Secure Collaborative Learning Framework for Human Activity Recognition

Debaditya Roy*
Royal Institute of Technology (KTH)
Qamcom Research and Technology
AB
Stockholm, Sweden
droy@kth.se

Ahmed Lekassays*
Qatar Computing Research Institute
Doha, Qatar
alekssays@hbku.edu.qa

Šarūnas Girdzijauskas
Royal Institute of Technology (KTH)
Stockholm, Sweden
sarunasg@kth.se

Elena Ferrari
University of Insubria
Varese, Italy
elena.ferrari@uninsubria.it

Barbara Carminati
University of Insubria
Varese, Italy
barbara.carminati@uninsubria.it

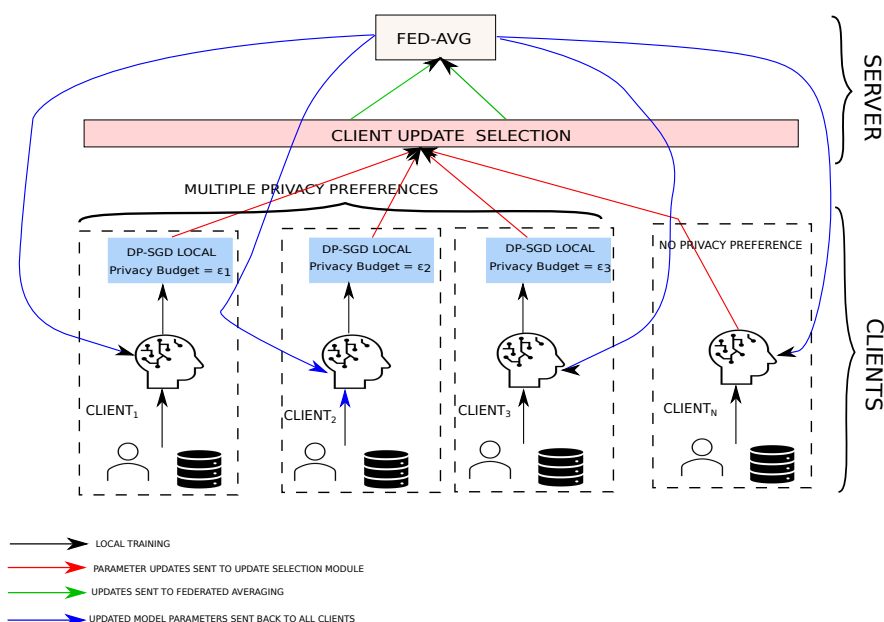


Figure 1: Differentially Private, Fair and Secure Collaborative Framework for Human Activity Recognition

ABSTRACT

Federated learning (FL), a decentralized machine learning technique, enhances privacy by enabling multiple devices to collaboratively train a model without transferring data to a central server. FL is used in Human Activity Recognition (HAR) problems, where multiple users generating private wearable data share models with a

*Equal Contribution



This work is licensed under a Creative Commons Attribution International 4.0 License.

UbiComp/ISWC '23 Adjunct, October 08–12, 2023, Cancun, Quintana Roo, Mexico
© 2023 Copyright held by the owner/author(s).
ACM ISBN 979-8-4007-0200-6/23/10.
<https://doi.org/10.1145/3594739.3610675>

server to learn a useful global model. However, FL may compromise data privacy through model information sharing during training. Moreover, it adheres to a one-size-fits-all approach toward data privacy, potentially neglecting varied user preferences in collaborative scenarios such as HAR. In response to these challenges, this paper presents a collaborative learning framework integrating differential privacy (DP) and FL, thus providing a tailored approach to privacy protection. While some existing works integrate DP and FL, they do not allow clients to have different privacy preferences. In this work, we introduce a framework that allows different clients to have different privacy preferences and hence more flexibility in terms of privacy. In our framework, DP adds individualized noise to individual clients' gradient updates for privacy. However, such noised updates can also be interpreted as an attack on the FL system. Defending against these attacks might result in excluding honest

private clients altogether from training, posing a *fairness* concern. On the other hand, not having any defensive measures might allow malicious users to attack the system, posing a *security* issue. Thus, to address security and fairness, our framework incorporates a *client selection* strategy that protects the global model from malicious clients and provides fair model access to honest private clients. We have demonstrated the effectiveness of our system on a HAR dataset and provided insights into our framework's privacy, utility, and fairness.

CCS CONCEPTS

• **Human-centered computing** → **Ubiquitous computing**; • **Computing methodologies** → **Supervised learning by classification**; • **Security and privacy** → **Privacy protections**.

KEYWORDS

human activity recognition; differential privacy; federated learning neural networks; security

ACM Reference Format:

Debaditya Roy, Ahmed Lekassays, Šarūnas Girdzijauskas, Elena Ferrari, and Barbara Carminati. 2023. Private, Fair and Secure Collaborative Learning Framework for Human Activity Recognition. In *Adjunct Proceedings of the 2023 ACM International Joint Conference on Pervasive and Ubiquitous Computing & the 2023 ACM International Symposium on Wearable Computing (UbiComp/ISWC '23 Adjunct)*, October 08–12, 2023, Cancun, Quintana Roo, Mexico. ACM, New York, NY, USA, 7 pages. <https://doi.org/10.1145/3594739.3610675>

1 INTRODUCTION

Federated learning (FL), a decentralized machine learning method, allows numerous devices to collaboratively train a model without transferring data to a central server, thus promoting privacy [10]. FL has been successfully applied in the ubiquitous Human Activity Recognition (HAR) scenario, where different users generate data from their wearable sensors [4, 8, 14]. By using FL in a wearable HAR use case, a certain degree of privacy can be offered to the users. However, FL still requires sharing information (through models) among devices during training, potentially compromising sensitive data privacy. Besides the potential privacy risks inherent in sharing model information, another limitation of FL is its one-size-fits-all approach to data privacy. Specifically, it only offers all users the same type of privacy by withholding data within the user's infrastructure. However, when multiple users train collaboratively, different users may have varying privacy preferences. E.g., some users' wearable device data may be highly sensitive due to medical conditions, while others might be less so. Therefore, it is crucial to accommodate these varied privacy needs, which native FL cannot do. Incorporating differential privacy (DP) in FL can address this issue by allowing individuals to dictate their data's privacy protection level. While there have been some efforts to incorporate differential privacy (DP) into FL through *local differential privacy* (LDP) in [17], the work did not consider different privacy preferences per user. In this work, we propose a framework that allows individual clients to have different privacy preferences through LDP and train them using FL to learn a global HAR model. Thus, users with highly sensitive data can assign privacy preferences, while users with lesser

sensitive data might opt for lower or no privacy at all. The individual clients use LDP to perturb the model weights and send them to the federated server for aggregation. In this setup, since clients add noise to model updates to achieve privacy, sometimes they can be interpreted as malicious, attacking the federated system. This is particularly true for clients with strict privacy requirements who add heavy noise for higher LDP. Implementing defense in FL might result in the exclusion of such honest but private clients from the training and pose a *fairness* concern. On the other hand, not implementing defense would allow malicious clients to attack the FL system, raising a *security* concern. Hence, our novel collaborative framework uses a *client update selection* approach to manage the participation of all clients in federated training and provide fair and secure access to the global model. In particular, the strategy *client update selection* uses cosine similarity to compare the incoming client weights to the global model weight. If it is below a predefined threshold, that client is declined participation in the federated averaging only for that round. The assumption behind this idea is that malicious clients will diverge away from honest clients in the weight space. Hence, imposing the above approach does not allow malicious clients to participate in the aggregation and hamper the FL. Although honest private clients can be declined participation in a round based on the above method, they receive the global model and a fair chance to participate across all rounds.

Our proposed framework is shown in Figure 1, where multiple users with different or no privacy preferences train together collaboratively. For achieving LDP, we have used differentially private stochastic gradient descent (DPSGD) [1], where we add noise to the gradient updates. In a DP setup, the privacy requirement is quantified through the *privacy budget* (ϵ), where high privacy indicates a low privacy budget and vice-versa. In our framework for each LDP client, ϵ_i (where i is a client number) represents the privacy budget that each user wants for their local training. The *client update selection* module selects the relevant clients for *federated averaging* during the aggregation. Finally, the global model post-federated averaging is sent back to all the users for training in subsequent rounds. The main contributions of this paper are as follows.

- (1) We present a novel collaborative framework that can accommodate multiple privacy preferences originating from different user needs and learn an effective global model to solve the downstream HAR task. Furthermore, we also propose a *client update selection* module that manages client participation in model training in a secure and fair manner by leveraging threshold-based cosine similarity;
- (2) We demonstrate the effectiveness of the collaborative learning framework on a popular HAR dataset and discuss the privacy-utility trade-off keeping fairness and security into account; To the best of our knowledge, this is the first collaborative framework that uses *differential privacy* and *federated learning* on HAR to incorporate multiple privacy preferences. Furthermore, in this work, we begin the very first discussion on the notion of fairness in terms of client participation in the federated training for such a framework.

In the next section, we discuss the related works, followed by the threat model and methods. Next, we present the evaluation section, concluding with the future directions of this work.

2 RELATED WORKS

2.1 HAR with federated learning

Human Activity Recognition (HAR) leverages machine learning to process accelerometer, gyroscope, and magnetometer data for applications in healthcare, sports, and security [9, 14]. With the proliferation of mobile and wearable devices, privacy becomes a critical concern, as the data may contain user-sensitive information [7]. Federated learning (FL), a privacy-preserving distributed machine learning approach, has been increasingly adopted for HAR [6, 10]. Yang et al. [20] and Sozinov et al. [15] demonstrated FL's effectiveness in preserving accuracy and privacy in HAR tasks. However, current research must still address security and privacy issues, like enhancing model robustness against malicious clients and adapting FL to diverse user privacy needs. In our work, we propose methods to address privacy preferences from multiple as well as protection from malicious clients while being fair to honest clients.

2.2 Federated Learning with DP

Differential Privacy is a framework that protects user privacy when their data is used for any downstream task. Usually, it is achieved through controlled noise addition in different parts of the pipeline, such as data processing, model training, etc. When used together, *DP* and *FL* can complement each other to provide stronger privacy guarantees. In the last couple of years, researchers have utilized a combination of both to provide stronger privacy guarantees [3, 5, 17, 19, 21]. In [19], Wei et al. propose a method for the noises of the local model parameters before federated aggregation to achieve differential privacy. Truex et al. [17] demonstrate local differential privacy in a federated setting with gradient perturbation. However, they provide α -condensed-local-differential-privacy guarantees. Hu et al. [5] propose a federated personalized learning approach with differential privacy guarantees on mobile IoT data. In their approach, they rely on the Gaussian mechanism of differential privacy to add noise (drawn from Gaussian) to local updates. Choudhury et al. [3] use differential privacy and federated learning on sensitive health data. In particular, they approximate differential privacy through objective function perturbation. While the above works have shown successful integration of *DP* into *FL* scenarios, they did not take into account multiple privacy budgets (or preferences) for different clients.

In our work, we want to introduce a generic privacy-preserving collaborative learning framework. The framework accommodates different types of privacy preferences originating from different users. For collaborative learning, we utilize *FL*, and for privacy preservation, we utilize *LDP*. While previous works have used *DP* in an *FL* setup, they did not account for multiple privacy preferences through multiple privacy budgets in *DP*. Furthermore, our use of gradient-based noise addition with differentially private stochastic gradient descent [1] allowed us to use deep learning methods that have been proven effective in wearable sensor-based HAR tasks [14]. Our framework integrates all of these mechanisms to ensure safe and private collaborative learning. Also, our cosine similarity-based client selection mechanism ensures that we subdue clients with suboptimal updates (due to high noise) and learn an optimal global representation with a good privacy-utility and fairness tradeoff.

3 THREAT MODEL

Federated learning was initially suggested as an initial measure to protect users' data privacy by conducting local training and sharing only model weights among users. However, this approach has been found to have weaknesses, such as vulnerability to inference attacks [16]. Moreover, training models in federated environments are susceptible to various security threats, including poisoning attacks, wherein attackers attempt to manipulate the models to make incorrect predictions for specific classes or reduce overall performance across all classes [2]. In this paper, we considered untargeted attacks, and we leave targeted attacks for future work. We discuss the possible ways to incorporate state-of-the-art defenses while preserving privacy in our solution in Section 6.

Assumptions. In a federated learning setting, a central server aggregates model weights shared by the different clients. We assume that the server is honest but curious. In other words, it follows the aggregation protocol, but it is curious to know about the data of the clients. In addition, we assume that the majority of clients are honest in the sense that the majority of clients follow the protocol and do not try to decrease the performance of the global model. However, a minority of clients try to attack the global model by submitting random weights. Furthermore, we assume that the model weights are shared with the central server through a secure tunnel (i.e., using HTTPS). In the context of this paper, our solution advocates flexible privacy preferences. Thus, clients are free to choose the privacy budget and technique they want.

Attacker's Goal and Capability The server aims to infer the data of each participant. In addition, the server has access to the shared model weights of each participant. Regarding malicious clients, they have access to the global model, and they submit random weights to decrease the performance of the global model.

4 METHODS

The goal of the paper is to propose a privacy-preserving collaborative framework that allows multiple users to participate and learn a global model for *human activity recognition*. The core proposition of the privacy-preserving framework lies in its ability to quantify multiple privacy preferences. In the framework, individual users having different privacy requirements in terms of strictness try to learn a global model together. For collaboration among users, we have used *federated learning*, and for privacy preservation, we have used *local differential privacy*. Furthermore, we also ensure that clients with high privacy requirements do not affect the overall model and still get a fair chance to participate in the federated training.

Collaborative training using federated learning

Individual users hold private data originating from wearable sensors. In this setting, users do not want to share the data with each other and still want to learn a global model. To achieve this goal, we have used popular federated learning. There is a central orchestrator that creates a global model and shares it with individual users. In a single federated round, the individual users use their local data to train the received model. At the end of each round, the users share their local model weights with the central server. The federated server aggregates the received weights using averaging

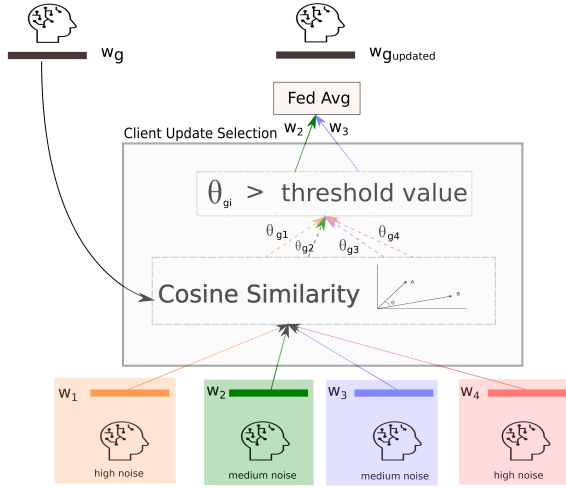


Figure 2: Client Update Selection and Federated Averaging

according to

$$w_{avg} = \frac{1}{N} \sum_{i=1}^N w_i, \quad (1)$$

where w_i represents the local weight updates received from N clients. w_{avg} is used to update the global model, which is resent to the individual clients for retraining in the subsequent rounds till convergence.

Privacy preservation through local differential privacy

The use of federated learning for collaborative training is widespread, and it offers some level of data protection by keeping the data private to individual users. However, it offers the same level of protection to all clients. Furthermore, it cannot quantify the privacy gain we get from using such a system. In a practical setting, different users have different privacy requirements, and this must be quantified to the users. Using differential privacy allows us to quantify privacy for individual clients and accommodate different privacy preferences for different users. Furthermore, when it is used to perturb client updates in a federated learning system, it achieves *LDP* for each client. In our framework, we use the differentially private stochastic gradient descent (DPSGD) method proposed by Abadi et al. [1] to ensure local differential privacy for individual clients.

This method works on a gradient-based learning approach. In this method, the gradient updates are perturbed by noise addition, which helps to obscure the information. Thus, to use *DPSGD*, our framework adheres to gradient descent-based learning in the local training process. While training a local model (i.e., the copy received from the federated server) with the available local data, as per *DPSGD*, each client adds noise to its gradient updates. This noise allows the model to have local differential privacy (LDP).

DPSGD quantifies a privacy loss each time it trains the model with a single example. During the full training process, the loss of privacy in multiple examples is composed according to *Moment's accountant technique* introduced by Abadi et al. [1]. The above technique allows us to numerically quantify the loss of privacy during the training process and calls it *privacy budget* (represented as ϵ). In *DPSGD*, we can ensure that each client adheres to a particular privacy budget ϵ_i throughout the local training, where i represents the client number. In Figure 1, the first three clients train with a privacy budget of $\epsilon_1, \epsilon_2, \epsilon_3$, respectively. The privacy budget directs the tuning of the amount of noise added to the parameter updates before they are sent to the server for aggregation. However, the N^{th} client in Figure 1 does not have any privacy requirements, and hence it sends the parameter updates without any perturbation. The federated averaging mechanism receives the *LDP* models and averages them before sending them back for the next round of training.

Fair and secure selection through similarity measures

Training models without the validation of incoming weights would lead to manipulating their predictions. Model weights might be malicious as they could be crafted to initiate (individually or collaboratively) *untargeted attacks*. We defend against untargeted attacks using similarity measures with the intuition that the malicious model weights will diverge from the direction of the honest model updates. It is worth noting that such defenses can discard model weights coming from honest clients with strict privacy budgets. The weights, perceived as random due to substantial noise injection, could seem like an untargeted attack. However, disallowing such clients to participate in the training process poses a *fairness concern* as it would be unfair toward honest clients with strict privacy requirements. Hence, we try to devise a mechanism that would allow strictly private clients to participate in the training process without hampering the utility of the downstream task severely. To be fair, secure, and have utility, we must satisfy three goals.

- *Goal 1:* All the clients must get access to the global model.
- *Goal 2:* They must at least have the opportunity to participate in the training.
- *Goal 3:* The server must be able to discard useless weights during aggregation, which can be an untargeted attack or a heavily noised and privatized weight update.

The first two contribute towards fairness, while the last one ensures utility and security. In order to satisfy the goals, we devise a *client update selection* module in the federated server that inspects the incoming weights and either selects or rejects them for federated averaging. Specifically, we observe the cosine similarity between the weights of the global model and the weights received from each client according to,

$$\theta_{gi} = \frac{w_g \cdot w_i}{\|w_g\| \|w_i\|}, \quad (2)$$

where w_g represents the weights of the global model and w_i represents the weights of the i^{th} client. If the cosine similarity value, θ_{gi} , is above a certain threshold, then it participates in federated averaging. The federated averaging yields updated global model

weight $w_{updated}$ given by,

$$w_{updated} = \frac{1}{N_{sel}} \sum_{i=1}^{N_{sel}} w_i, \quad (3)$$

where $N_{sel} < N$ represents the number of selected clients for whom $\theta_{gi} > threshold$. The phenomenon is also depicted in Figure 2, with four clients with different noise parameters (based on privacy preferences) having weights w_1, w_2, w_3, w_4 . In the diagram, the cosine similarity-based thresholding selects w_2, w_3 for the federated averaging that creates $w_{updated}$, the new global model weight. Usually, the clients whose weights have been perturbed too much by the noise exhibit a low cosine similarity and are excluded from the federated averaging. Nevertheless, the global model is also circulated back to them so that they can participate in the training process again. The above strategy satisfies *Goal 1* and *Goal 2* introduced earlier, thus providing fairness. Malicious users try the attack the system, and the *client update selection* rejects their updates protecting the global model, satisfying *Goal 3*. Furthermore, since the above method also rejects random weights, the aggregated global model is the best possible representation of the downstream task, thus preserving utility. Although the global model is sent back to malicious users, since it is composed of parameters perturbed by *local DP*, it cannot infer user data from the global model. Thus our work provides a comprehensive privacy-preserving and fair framework for HAR tasks.

Interesting to observe that the θ parameter is quite significant in determining the selection of clients. In the present version of the work, it is a manually chosen and empirically validated hyperparameter with a value of 0.48. However, in the future version of this work, we plan to learn adaptively.

5 EVALUATION

Through our evaluation, we want to probe a practical system that incorporates users with realistic privacy preferences and observe the utility of the HAR task as well as draw some conclusions on privacy preservation and fairness. An acceptable range of privacy budget is defined between 0.5 to 2, where 0.5 represents the strictest privacy budget [13]. In a practical system, different users will have different privacy preferences, and some users will be non-private. For our experiments, we also define other combinations of privacy budgets for different clients (strictest to most relaxed) to understand different tradeoffs. The strictest privacy system is called *All-DPSGD Strict*, where all clients have a strict privacy budget of 0.5. *All-DPSGD Mixed* has a random mix of acceptable privacy budgets for all clients. *50% DPSGD Strict* represents a scenario with 50% of clients being private with a budget of 0.2. *50% DPSGD Mixed* represents the system where 50% of clients are private with mixed privacy budgets (within an acceptable range).

Since we want to propose such a system for the HAR task, we demonstrate our system on a popular HAR dataset called PAMAP2 [12]. It is a multivariate sensor dataset for HAR tasks. It consists of twelve activities recorded from six subjects. For our framework, we consider each subject as a client in the collaborative learning system. We set aside 30% of the data from each client to construct a test set. The final accuracy score is reported on this test set.

Through our experiment, we want to observe the effect of the *client update selection* module on the HAR task and the general *privacy utility tradeoff*.

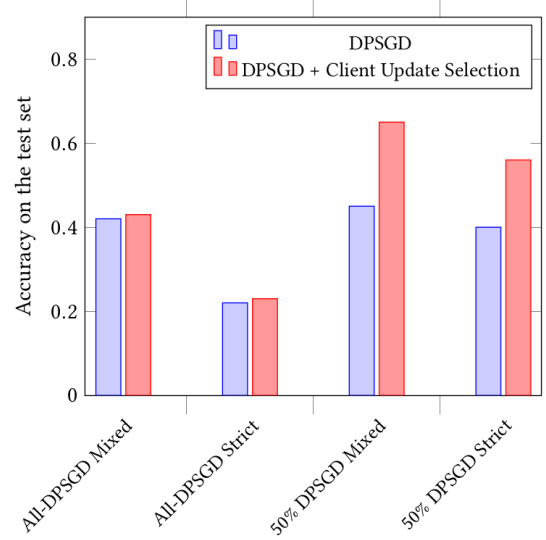


Figure 3: Effect of Client Update Selection: Accuracy on the test set for models. The framework is tested with two groups of privacy budget and two groups of DP settings. *50%-DPSGD* represents when 50% of clients are training with differential privacy. *Mixed* represents a mixed setting of privacy budgets by clients. *Strict* represents the strictest privacy budget set by the client. E.g., *All-DPSGD Strict* represents a setting where all clients are trained with differential privacy with the strictest privacy budget.

Fairness through client-update selection: In this experiment, we want to observe the test-set accuracy of the different types of collaborative systems proposed earlier with and without the *client-update selection*. For this version of our, we empirically validate the *threshold* value for client selection based on cosine similarity to 0.48. As observed from Figure 3, when we have a mix of private and non-private clients, the *client-update selection* can reject noisy updates much more effectively, improving the overall accuracy. This is because the non-private clients can contribute to the optimal learning of the global model in the initial federated rounds. Furthermore, since we incorporate fairness by sending back the global model even to the rejected clients, they can participate in the training at later stages, improving the overall accuracy of the test set. Although without *client-update selection*, we can guarantee fairness through participation in all rounds, this is detrimental to the HAR performance, as seen in Figure 3. Moreover, without *client-update selection*, we cannot defend against malicious weight perturbation attacks. Therefore, the *client-update selection* offers us a well-balanced solution for fairness, utility, and security in a collaborative learning system with private and non-private users.

However, when all the clients are private, we do not observe any improvement through the client selection module. Primarily

because of the fact that when most updates are noisy, the learning signal is distorted significantly, and hence it does not really matter whether we select some clients for the federated averaging. Note that the privacy budget range we have selected for our systems is rather strict, and hence they add substantial noise to the updates. The experiment also let us infer in order to have a usable collaborative system, we must include a mix of clients having different privacy budgets. Too much strictness can cause sub-optimal HAR performance.

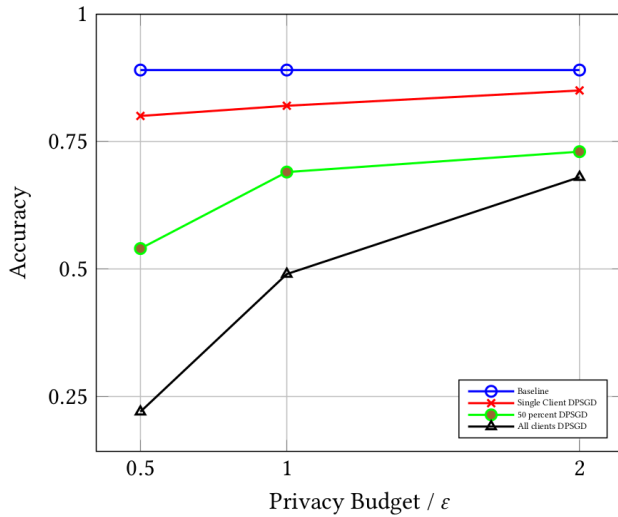


Figure 4: Privacy-utility tradeoff for collaborative learning system with different privacy budgets for PAMAP2 Dataset. The blue line represents the non-private federated baseline.

Privacy-utility tradeoff: In this experiment, we want to observe the general trends of how privacy budget (ϵ) and test set accuracy correlates with each other in our collaborative system for HAR. To do so, we simplify our system a bit and have the same privacy budget for multiple clients that want privacy. Figure 4 shows ϵ in the x-axis, and the y-axis represents accuracy on the test set. When we only have one private client (red line), we have the closest accuracy to the non-private baseline (blue line). Also, we observe the increasing accuracy with increasing privacy budgets for all the instances. This graph gives us some intuition into how we can set privacy budgets for clients to have the desired performance. As an example, if all the clients have a privacy budget of 1, then we have 0.5 as test accuracy. For some HAR applications, this classification performance might not be acceptable, and hence in those cases, the privacy budget might have to be relaxed for some clients to achieve a better utility. Our initial experiments with this framework allow some insights into how to set up the framework.

6 FUTURE WORKS AND CONCLUSION

In this work, we propose the basic framework for human activity recognition (HAR), where multiple users with different privacy preferences can collaborate to learn a global model in a fair and secure

manner. We demonstrated the effectiveness of the framework on a popular HAR benchmark dataset. With this work, we contribute to the broad research avenue that combines differential privacy and federated learning for ubiquitous computing. This allows us to formulate concrete future directions to strengthen the proposed collaborative framework. We plan to expand to other HAR datasets to demonstrate the privacy-utility tradeoff of our framework. In particular, we want to add human activity recognition datasets with a large number of clients. This would allow us to test the limits of our collaborative framework while being very similar to industrial use cases with many different users. We have proposed a basic yet effective approach based on cosine similarity to select client updates such that it incorporates security and fairness in our framework. However, other advanced defenses based on Trusted Execution Environment [11] by crafting the models to reduce their sizes and homomorphic encryption [18] will be incorporated into the framework. This would allow it to be truly flexible in terms of privacy preferences from users as well as different security solutions that can be plugged into the framework. Furthermore, in this work, we begin the discussion around the notion of fairness in terms of client participation in the proposed framework. However, it is inherently hard to define a metric of such fairness, a challenge that we wish to take on in the upcoming iterations of this work.

ACKNOWLEDGMENTS

This project has received funding from the European Union's Horizon 2020 research and innovation programme under the Marie Skłodowska-Curie grant agreement No 813162. This paper's content reflects only the views of its author(s). The European Commission/ Research Executive Agency is not responsible for any use that may be made of the information it contains. We would also like to thank Qamcom Research and Technology for funding the project.

REFERENCES

- [1] Martin Abadi, Andy Chu, Ian Goodfellow, H. Brendan McMahan, Ilya Mironov, Kunal Talwar, and Li Zhang. 2016. Deep Learning with Differential Privacy. In *Proceedings of the 2016 ACM SIGSAC Conference on Computer and Communications Security (Vienna, Austria) (CCS '16)*. Association for Computing Machinery, New York, NY, USA, 308–318. <https://doi.org/10.1145/2976749.2978318>
- [2] Sana Awan, Bo Luo, and Fengjun Li. 2021. CONTRA: Defending Against Poisoning Attacks in Federated Learning. In *Computer Security – ESORICS 2021: 26th European Symposium on Research in Computer Security, Darmstadt, Germany, October 4–8, 2021, Proceedings, Part I* (Darmstadt, Germany). Springer-Verlag, Berlin, Heidelberg, 455–475. https://doi.org/10.1007/978-3-030-88418-5_22
- [3] Olivia Choudhury, Aris Gkoulalas-Divanis, Theodoros Salonidis, Issa Sylla, Yoonyoung Park, Grace Hsu, and Amar Das. 2020. Differential Privacy-enabled Federated Learning for Sensitive Health Data. arXiv:1910.02578 [cs.LG]
- [4] Yu Guan and Thomas Plötz. 2017. Ensembles of deep lstm learners for activity recognition using wearables. *Proceedings of the ACM on interactive, mobile, wearable and ubiquitous technologies* 1, 2 (2017), 1–28.
- [5] Rui Hu, Yuanxiong Guo, Hongning Li, Qingqi Pei, and Yanmin Gong. 2020. Personalized federated learning with differential privacy. *IEEE Internet of Things Journal* 7, 10 (2020), 9530–9539.
- [6] Jakub Konečný, H. Brendan McMahan, Daniel Ramage, and Peter Richtárik. 2016. Federated Optimization: Distributed Machine Learning for On-Device Intelligence. arXiv:1610.02527 [cs.LG]
- [7] Nicholas D. Lane and Petko Georgiev. 2015. Can Deep Learning Revolutionize Mobile Sensing?. In *Proceedings of the 16th International Workshop on Mobile Computing Systems and Applications (Santa Fe, New Mexico, USA) (HotMobile '15)*. Association for Computing Machinery, New York, NY, USA, 117–122. <https://doi.org/10.1145/2699343.2699349>
- [8] Oscar D Lara and Miguel A Labrador. 2012. A survey on human activity recognition using wearable sensors. *IEEE communications surveys & tutorials* 15, 3 (2012), 1192–1209.

- [9] Oscar D Lara and Miguel A Labrador. 2012. A survey on human activity recognition using wearable sensors. *IEEE communications surveys & tutorials* 15, 3 (2012), 1192–1209.
- [10] Brendan McMahan, Eider Moore, Daniel Ramage, Seth Hampson, and Blaise Aguera y Arcas. 2017. Communication-Efficient Learning of Deep Networks from Decentralized Data. In *Proceedings of the 20th International Conference on Artificial Intelligence and Statistics (Proceedings of Machine Learning Research, Vol. 54)*, Aarti Singh and Jerry Zhu (Eds.). PMLR, Florida, USA, 1273–1282. <https://proceedings.mlr.press/v54/mcmahan17a.html>
- [11] Fan Mo, Hamed Haddadi, Kleomenis Katevas, Eduard Marin, Diego Perino, and Nicolas Kourtellis. 2021. PPFL: Privacy-Preserving Federated Learning with Trusted Execution Environments. In *Proceedings of the 19th Annual International Conference on Mobile Systems, Applications, and Services (Virtual Event, Wisconsin) (MobiSys '21)*. Association for Computing Machinery, New York, NY, USA, 94–108. <https://doi.org/10.1145/3458864.3466628>
- [12] Attila Reiss and Didier Stricker. 2012. Introducing a New Benchmarked Dataset for Activity Monitoring. In *Proceedings of the 2012 16th Annual International Symposium on Wearable Computers (ISWC) (ISWC '12)*. IEEE Computer Society, USA, 108–109. <https://doi.org/10.1109/ISWC.2012.13>
- [13] Lucas Rosenblatt, Joshua Allen, and Julia Stoyanovich. 2022. Spending Privacy Budget Fairly and Wisely.
- [14] Debaditya Roy, Sarunas Girdzijauskas, and Serghei Socolovschi. 2021. Confidence-calibrated human activity recognition. *Sensors* 21, 19 (2021), 6566.
- [15] Konstantin Sozinov, Vladimir Vlassov, and Sarunas Girdzijauskas. 2018. Human activity recognition using federated learning. In *2018 IEEE Intl Conf on Parallel & Distributed Processing with Applications, Ubiquitous Computing & Communications, Big Data & Cloud Computing, Social Computing & Networking, Sustainable Computing & Communications (ISPA/IUCC/BDCloud/SocialCom/SustainCom)*. IEEE, Melbourne, Australia, 1103–1111.
- [16] Stacey Truex, Nathalie Baracaldo, Ali Anwar, Thomas Steinke, Heiko Ludwig, Rui Zhang, and Yi Zhou. 2019. A Hybrid Approach to Privacy-Preserving Federated Learning. In *Proceedings of the 12th ACM Workshop on Artificial Intelligence and Security (London, United Kingdom) (AISec'19)*. Association for Computing Machinery, New York, NY, USA, 1–11. <https://doi.org/10.1145/3338501.3357370>
- [17] Stacey Truex, Ling Liu, Ka-Ho Chow, Mehmet Emre Gursoy, and Wenqi Wei. 2020. LDP-Fed: Federated Learning with Local Differential Privacy. In *Proceedings of the Third ACM International Workshop on Edge Systems, Analytics and Networking (Heraklion, Greece) (EdgeSys '20)*. Association for Computing Machinery, New York, NY, USA, 61–66. <https://doi.org/10.1145/3378679.3394533>
- [18] Wenhao Wang, Yichen Jiang, Qintao Shen, Weihao Huang, Hao Chen, Shuang Wang, Xiaofeng Wang, Haixu Tang, Kai Chen, Kristin Lauter, and Dongdai Lin. 2019. Toward Scalable Fully Homomorphic Encryption Through Light Trusted Computing Assistance. *arXiv:1905.07766 [cs.CR]*
- [19] Kang Wei, Jun Li, Ming Ding, Chuan Ma, Howard H Yang, Farhad Farokhi, Shi Jin, Tony QS Quek, and H Vincent Poor. 2020. Federated learning with differential privacy: Algorithms and performance analysis. *IEEE Transactions on Information Forensics and Security* 15 (2020), 3454–3469.
- [20] Qiang Yang, Yang Liu, Tianjian Chen, and Yongxin Tong. 2019. Federated machine learning: Concept and applications. *ACM Transactions on Intelligent Systems and Technology (TIST)* 10, 2 (2019), 1–19.
- [21] Yang Zhao, Jun Zhao, Mengmeng Yang, Teng Wang, Ning Wang, Lingjuan Lyu, Dusit Niyato, and Kwok-Yan Lam. 2020. Local differential privacy-based federated learning for internet of things. *IEEE Internet of Things Journal* 8, 11 (2020), 8836–8853.