

Split-U-Net: Preventing Data Leakage in Split Learning for Collaborative Multi-Modal Brain Tumor Segmentation

Holger R. Roth, Ali Hatamizadeh, Ziyue Xu,
Can Zhao, Wenqi Li, Andriy Myronenko, Daguang Xu

NVIDIA, Bethesda, USA

Abstract. Split learning (SL) has been proposed to train deep learning models in a decentralized manner. For decentralized healthcare applications with vertical data partitioning, SL can be beneficial as it allows institutes with complementary features or images for a shared set of patients to jointly develop more robust and generalizable models. In this work, we propose “Split-U-Net” and successfully apply SL for collaborative biomedical image segmentation. Nonetheless, SL requires the exchanging of intermediate activation maps and gradients to allow training models across different feature spaces, which might leak data and raise privacy concerns. Therefore, we also quantify the amount of data leakage in common SL scenarios for biomedical image segmentation and provide ways to counteract such leakage by applying appropriate defense strategies.

Keywords: Split learning · Vertical federated learning · Multi-modal brain tumor segmentation · Data inversion.

1 Introduction

Collaborative and decentralized techniques to train artificial intelligence (AI) models have been gaining popularity, especially in healthcare applications where data sharing to build centralized datasets is particularly challenging due to patient privacy and regulatory concerns [20]. Federated learning (FL) [16] and split learning (SL) [7] are two approaches that can be useful depending on the nature of the data partitioning [32]. In the healthcare and biomedical imaging sector, data is often “horizontally” partitioned such that each participating site, i.e., a hospital, possesses some data/features and optionally corresponding labels for their set of patients. Horizontal FL (HFL) algorithms like federated averaging [16] typically train models initialized from a current “global” model independently on each participant and frequently update the global model with the model gradients sent by each site. In contrast, so-called “vertical” data partitioning allows sites with complementary features but from an overlapping set of patients to collaborate [32]. This vertical FL (VFL) scenario could be useful where different sites possess features that need to be securely combined in order to train a joined AI model, e.g., one hospital has imaging while the other one has lab results or the

diagnoses for the same set of patients. Here, SL can be used to train models when using deep learning (DL) methods for VFL. During training, SL splits the forward pass of a DL model into two or more parts and exchanges intermediate features or activation maps and gradients between participating sites to complete a training step. Therefore features from different sites can be combined in later parts of the network and the model can be trained across institutional borders [28].

In biomedical image segmentation, VFL could be useful to combine different image modalities of the same patient in order to train joined segmentation models collaboratively. This scenario is what we explore in this work by studying SL as a collaborative technique to learn a tumor segmentation model for multi-modal MRI images. For this purpose, we propose “Split-U-Net”, a modification to the popular U-Net [21] architecture, to allow its use in a VFL setup. Figure 1 illustrates the situation where four sites would like to jointly train a multi-modal segmentation model given their corresponding images. Only one site possesses the label mask and computes the loss to be optimized. Previous works on SL in healthcare appli-

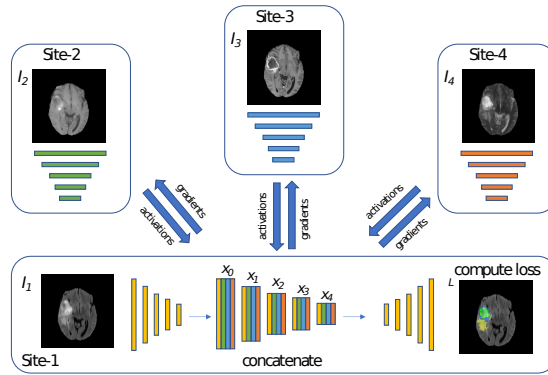


Fig. 1: Split learning set up with Split-U-Net.

cations have focused on classification and regression tasks [28, 19, 8]. One example of splitting U-Net for single modality semantic segmentation was described in [18], but to the best of our knowledge, our work is the first to apply SL in a multi-modal vertical data partitioning scenario for biomedical image segmentation across multiple parties.

We show that SL can be used successfully for this task and also investigate potential security implications that arise from sharing intermediate features between collaborating sites by implementing an effective inversion attack. Prior works on inversion attacks in SL were mainly focused on images of small sizes (MNIST or CIFAR-10) [17, 5, 11, 12] and are theoretical in nature. However, it is important for the medical imaging community to understand the potential benefits and security considerations for applying SL in healthcare applications. Our inversion attack not only shows the potential risks but can be used to quantify and inform appro-

appropriate defense strategies against it. In this work, we explore both dropout [25] and differential privacy (DP) [4] to prevent data leakage during SL. Our contributions can be summarized as follows.

- We propose “Split-U-Net” and successfully apply SL for biomedical image segmentation for multi-institutional collaboration.
- We develop a successful inversion attack to measure and quantify data leakage in SL.
- We propose and evaluate defense measures (dropout and DP) to prevent data leakage.

2 Methods

2.1 Split-U-Net

The basis of our network is a common implementation of U-Net [21] with $L = 4$ down- and up-sampling levels. The default number of output features at each level are configured as shown in Table 1. To turn this network $F(x)$ into Split-U-Net $F(x) = g(f(x))$ used for multi-modal collaborative SL, we divided the number of encoder features by the number of sites/modalities K involved in training. The number of features in the decoder stays the same as in the default network. Fig. 1 shows an example setup with $K = 4$. In this example, the “split” is done at the bottleneck. All encoder layers participate in SL, and only the site with label images has the decoder for segmentation. During training, each site k computes the forward pass $\{x_0^k, x_1^k, \dots, x_L^k\} = f^k(I_k)$ where I_k is a mini-batch of size B of input images with modality k and x_i^k is the feature map of layer i of L . Corresponding batch indices are communicated to each site before each training step. The site k with label images then takes the activation maps from all other participating sites and concatenates them at the appropriate feature levels (see Fig. 1). It then computes the loss and backward pass to obtain a gradient

$$\nabla \leftarrow \mathcal{L}_{seg} \left(g(\{x_0^k, x_1^k, \dots, x_L^k\}), labels \right) \quad (1)$$

and performs an optimizer update on its part of the network $g(x)$. The gradient ∇ at the split level is then communicated back to all $f^k(x)$ to complete the backward pass and update their parts of the model (the encoder branches $f^k(x)$), and the process is iterated until convergence. Note that in SL, a larger batch size can be used to reduce the total number of communication steps needed [24]. We assume that at each iteration, a random set of batch indices b_i is selected such that each

Table 1: U-Net and Split-U-Net features for brain tumor segmentation.

Level i	In	0	1	2	3	4	5	6	7	8	Out
U-Net (default) $F(x)$	4	32	32	64	128	256	128	64	32	32	4
Split-U-Net Encoder (per site) $f^k(x)$	1	8	8	16	32	64	-	-	-	-	-
Split-U-Net Decoder $g(x)$	-	-	-	-	-	-	128	64	32	32	4

client uses the same patients’ data and augmentation to build their mini-batch I_k . Furthermore, each site might apply additional spatial normalization steps, e.g., using non-linear image registration to bring their images from different modalities into a common data space to help the network better encode common anatomical features across modalities [3,30].

2.2 Measuring data leakage by inversion attack

In SL, activation maps are shared to complete each iteration step [7]. Therefore, a potential malicious actor, e.g., Site-1 in Fig. 1, receiving the activation maps from other sites may invert them to recover the underlying private data. Such attacks used to recover the data are called “inversion attacks” [27]. In the following, we explain how our inversion attack is executed and how its result can be used to measure the data leakage in order to inform an appropriate defense strategy. Our attack tries to optimize a randomly initialized C -channel input $\tilde{I}_i$ such that the activations $\tilde{x}_i$ at the forward layer of the attacker model $\tilde{f}_i(x)$ become the same as the intercepted activations x_i . In this work, we assume the attacker has access to the current state of the model used by the client to generate the forward pass. Therefore $\tilde{f}_i^k(x) \equiv f_i^k(x)$ given the same input x . This setting is typically referred to as a “white-box” attack [11]. In practical implementations of SL, this could be the case if a common network is used to initialize $f^k(x)$ on each participating client. The main loss used to align both activation maps is a L^2 -norm. Furthermore, we employ two common image prior losses often used in inversion attacks [6,33], namely total variation [22] (TV) and L^2 -norm of the recovered image $\tilde{I}$. The main loss for the inversion attack hence becomes

$$\mathcal{L}_{\text{inv}}(x_i^k, \tilde{x}_i^k, \tilde{I}_i^k) = \alpha_{\text{act}} \|x_i^k - \tilde{x}_i^k\|_2 + \alpha_{\text{tv}} \text{TV}(\tilde{I}_i^k) + \alpha_{l_2} \|\tilde{I}_i^k\|_2. \quad (2)$$

Therefore, the final inversion attack to recover an image $\tilde{I}_i$ from activation x_i at level i of Split-U-Net can be formulated as

$$\tilde{I}_i = \underset{\tilde{I}}{\text{argmin}} \mathcal{L}_{\text{inv}}(x_i^k, \tilde{x}_i^k, \tilde{I}_i^k), \quad (3)$$

where $\tilde{I}_i \in \mathbf{R}^{B,C,H,W}$ with B, C, H, W being the batch size, number of channels, height and width of the image, respectively. Note that the inversion can be run on large batch sizes B or independently for each activation in a mini-batch, depending on the compute resources of the attacker. The data inversions from intercepted activation maps can be seen in Fig. 2. To measure the amount of data leakage, we compute a common similarity metric between the recovered image $\tilde{I}_i$ and the original image I used to produce the activation x_i . Structural Similarity index (SSIM) [29] aims to provide a more intuitive and interpretable metric compared to other commonly used metrics like root-mean-squared error or peak signal-to-noise ratio. We, therefore, use SSIM in our analysis, but including other metrics would be possible.

2.3 Defenses

A straightforward defense strategy is to not send feature activation maps from early layers (x_0 , x_1 , and x_2) which are likely to leak more data (see Fig. 2 and

Fig. 4). We also investigate dropout [25] as an effective tool against inversion attacks. Each layer of the encoder can randomly drop the activations from neurons of the network with a probability of p_{dropout} . Another effective tool often used in the FL literature [15, 13, 10, 31], is differential privacy (DP). DP in its simplest form adds some calibrated random noise to any shared values in order to preserve the privacy of individual data entries. Here, we use a Gaussian mechanism [31] to add random noise sampled from a normal distribution $N(0, \sigma^2)$ to each activation mask x_i^k before sharing it with the next participant.

3 Experiments & Results

Data: In our study, we assume a collaborative model training setup where four institutes, here referred to as “sites”, jointly train a multi-modal image segmentation model using split learning. We use the *Medical Segmentation Decathlon* MSD¹ brain tumor segmentation dataset (Task 1) to simulate this setup. Each 3D volume in the dataset contains four MRI modalities, namely T1-weighted, post-Gadolinium contrast T1-weighted, T2-weighted, and T2 Fluid-Attenuated Inversion Recovery volumes [23]. For the purpose of this study, we extract one axial slice from each volume through the center of the tumor and formulate the task as a 2D semantic segmentation problem, resulting in a total of 484 images with ground truth annotation masks. We randomly split the data into 338 training, 49 validation, and 97 testing images, corresponding to 70%, 10%, 20% of the data, respectively. The segmentation task is to predict the brain tumor sub-regions, i.e., edema, enhancing, and non-enhancing tumor. Therefore, our network predicts four output classes, including the background, using a final softmax activation. Given the $K=4$ MRI input modalities, we simulate the Split-U-Net to be trained collaboratively among four sites, as shown in Fig. 1. Each site possesses the images for all patients but for just one modality. Site-1 is assumed to also have the annotation masks and can therefore compute the objective function using a combined Dice loss and cross-entropy loss.

Data leakage of shared activation maps: First, we investigate how much data the activation map at each layer can leak when sharing them during Split-U-Net training. We invert all activation maps of a mini-batch from layers x_0, x_1, x_2, x_3, x_4 , respectively. One can observe that the amount of data leakage reduces with the depth of the network and the resolution of the share activation map. The first level x_0 with a resolution similar to the input image is practically non-distinguishable from the original augmented images fed to the encoder networks during training. All inversions computed in this work used $\alpha_{\text{act}} = 1e-3$, $\alpha_{\text{tv}} = 1e-4$, and $\alpha_{l_2} = 1e-5$ (see Eq. 2). We used the Adam optimizer to solve Eq. 3 using a cosine learning rate decay with an initial rate of 0.1².

¹ <http://medicaldecathlon.com>

² **Implementation:** We utilize components from MONAI³ and NVIDIA FLARE⁴ to implement our SL simulation. In particular, we utilize MONAI’s *BasicUNet* as basis for Split-U-Net. All experiments were run on NVIDIA V100 GPUs with 16 GB memory.

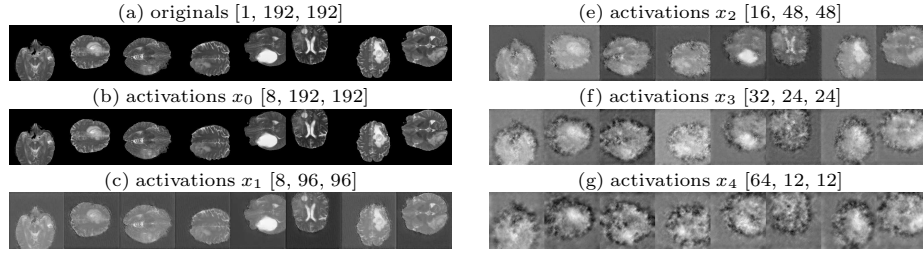


Fig. 2: Inversions from activations sent from different layers of the Split-U-Net encoder of Site-4 possessing one MRI modality when training with a mini-batch size of 8. Activations from earlier layers from the encoder are more likely to leak data, i.e., $x_0 \sim x_2$. The inversions from other sites and modalities are of the same quality.

Collaborative multi-modal image segmentation: To evaluate the effectiveness of Split-U-Net, we compare it to a baseline U-Net model taking the four MRI modalities directly as input (see the default setup in Table 1). In Table 2, we show Split-U-Net performs on par with its centralized counterpart (U-Net). The performance is comparable⁵ to the MSD challenge results reported for 3D tumor segmentation [2].

Table 2: Comparison of a centralized U-Net and different Split-U-Net settings with different privacy-preserving measures (dropout and differential privacy (DP)). The setting “w” and “w/o” indicates the performance of Split U-Net with and without skip connections, respectively; “ x_3, x_4 only” indicates the performances when only activations from later layers are being shared. The best Dice score achieved with Split-U-Net for each data subset is highlighted in **bold**.

Dice	Training (n=338)	Validation (n=49)	Testing (n=97)
U-Net	0.732	0.701	0.698
Split-U-Net (w/o skip)	0.743	0.619	0.599
Split-U-Net (w skip)	0.882	0.663	0.693
Split-U-Net (x_3, x_4 only)	0.821	0.675	0.650
Split-U-Net ($p_{\text{dropout}}=0.1$)	0.818	0.648	0.681
Split-U-Net ($p_{\text{dropout}}=0.2$)	0.843	0.658	0.683
Split-U-Net ($p_{\text{dropout}}=0.5$)	0.766	0.637	0.665
Split-U-Net ($p_{\text{dropout}}=0.8$)	0.719	0.643	0.650
Split-U-Net (DP $\sigma=1$)	0.865	0.671	0.691
Split-U-Net (DP $\sigma=2$)	0.797	0.669	0.695
Split-U-Net (DP $\sigma=3$)	0.821	0.658	0.666
Split-U-Net (DP $\sigma=5$)	0.811	0.684	0.687
Split-U-Net (DP $\sigma=50$)	0.543	0.394	0.393

Effectiveness of defenses: Adding dropout and Gaussian noise during training can be an effective defense (Fig. 3). The SSIM scores between originals and

⁵ The current **leading entry - SwinUNETR** [9] achieves an average Dice score of 0.647 for the three foreground tumor classes.

inversions go down with higher p_{dropout} or σ as shown in Fig. 4. The model performance is less affected when adding DP and even benefits from it during training, as seen in Table 1 for $\sigma=2.0$ in contrast to using dropout as a defense.

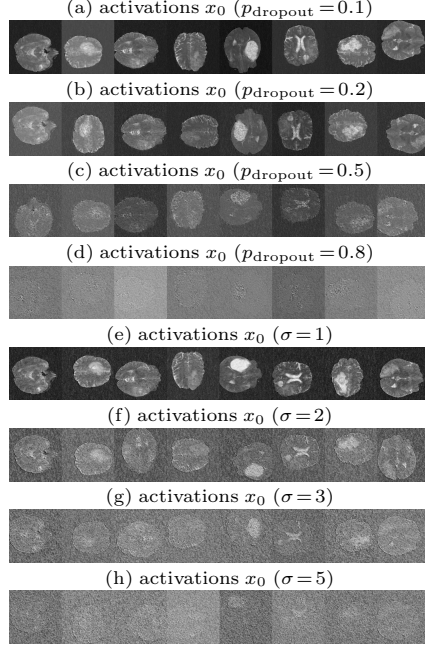


Fig. 3: Dropout (a-d) and differential privacy (e-h) as a defense against inversion attacks.

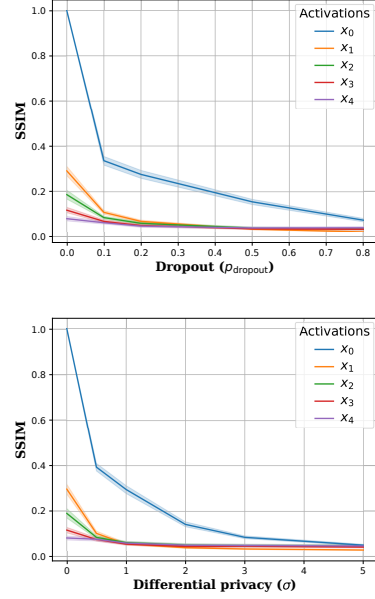


Fig. 4: Structural SIMilarity index (SSIM) [29] between the original images and inversions of each activation.

4 Discussion

To the best of our knowledge, our work was the first to apply SL to a multi-modal image segmentation task. We showed competitive results of Split-U-Net for 2D brain tumor segmentation on a relatively small dataset (only one slice per original volume). Further hyperparameter tuning and data augmentation might improve the performance. It should be investigated if weight sharing between the encoder branches could allow for further performance boosts [26]. An extension of Split-U-Net to 3D semantic segmentation tasks would be straightforward.

A major focus of this work is on the security aspect when applying SL. As shown in our results, depending on the depth of activation layers inside Split-U-Net, the data inversion attack can be successful, generating inversions that are visually indistinguishable from the original images (SSIM close to 1.0). This is the

case, especially for the first layer (x_0). Finding an appropriate defense strategy against such inversion attacks is very important. It can be assumed that the same defense settings are effective for each modality used in Split-U-Net training. Therefore, a recommendation would be for the site possessing both images and labels to study the data leakage vulnerabilities using our proposed data inversion and data leakage metrics to establish a secure setting that each collaborator can use. Of course, this assumes a level of trust in this site but might help protect against a potentially malicious server that coordinates the split learning. At the same time, each site could utilize public datasets with images and labels, as we have done in this study, to measure the data leakage risks of the network architecture they would like to train in real-world SL. Our results indicate that the dangers come from the architecture itself rather than the particular dataset used for training (see Fig. 2 where the inversion quality is not affected by different samples in the batch). This is in contrast to other studies in horizontal FL, where certain images in the batch are more likely to leak data [33,10].

Our study also has some limitations. For example, the inversion attack assumes to have access to the current state of the model that the data site uses to compute its forward pass ($f^k(X)$). This setting is typically referred to as a “white-box” attack [11] and assumes the attacker has knowledge about the state of the model during training. This could be true in some implementations of SL where one of the participants sends an initialization for all participants. As the training continues, this initial model will become less and less useful to the attacker. At the same time, our finding shows that a potential avenue for more secure implementations of SL is to not use a common initialization but let each participant randomly initialize their part of the model. A “black-box” attack [11] where the inversion needs to optimize for both the inputs and the current state of the model could be implemented next to better measure the data leakage risks in such a scenario. Furthermore, we assumed the participating sites to have a common anonymous identifier used to build mini-batches with images of corresponding patients. In real-world scenarios, a pre-processing step to securely compute the intersecting set of patients between sites has to be performed [1]. Also, some synchronization of data augmentation across different modalities should be incorporated in the communication protocols. An additional privacy risk in SL is the inversion of label sets from the shared model gradients. A similar attack to the one presented in this work could be applied to match gradients during SL to recover the label masks. However, we assumed that tumor segmentation masks are less likely to leak patient-identifiable information and therefore focused on the data/image recovery in this work. In this work, we simulated a multi-site FL study using pre-registered multi-modal MRI scans. In reality, more variations that are potentially critical to model performance would need to be considered before performing similar collaborative model training, including temporal and spatial misalignment across images of the same patient and mismatch between image and annotation masks. Finally, cryptographic techniques like homomorphic encryption [34] or secure multi-party computation [14] could be employed to reduce the risk of data leakage in SL. Those techniques typically

come with higher computation costs but should be explored, especially for medical image analysis tasks where patient privacy is of utmost concern.

In conclusion, we provided strong evidence that SL can be useful for biomedical image segmentation tasks when taking the appropriate security considerations into account. A real-world implementation of SL will provide clinical collaborators the chance to jointly leverage all available data to train more robust and generalizable AI models.

References

1. Angelou, N., Benaissa, A., Cebere, B., Clark, W., Hall, A.J., Hoeh, M.A., Liu, D., Papadopoulos, P., Roehm, R., Sandmann, R., et al.: Asymmetric private set intersection with applications to contact tracing and private vertical federated machine learning. arXiv preprint arXiv:2011.09350 (2020)
2. Antonelli, M., Reinke, A., Bakas, S., Farahani, K., Landman, B.A., Litjens, G., Menze, B., Ronneberger, O., Summers, R.M., van Ginneken, B., et al.: The medical segmentation decathlon. arXiv preprint arXiv:2106.05735 (2021)
3. C. Studholme, D. L. G. Hill, D.J.H.: Automated 3d registration of mr and pet brain images by multi-resolution optimisation of voxel similarity measures. *Medical Physics* **24**(1), 25–35 (1997)
4. Dwork, C., McSherry, F., Nissim, K., Smith, A.: Calibrating noise to sensitivity in private data analysis. In: *Theory of cryptography conference*. pp. 265–284. Springer (2006)
5. Erdogan, E., Kupcu, A., Cicek, A.E.: Unsplit: Data-oblivious model inversion, model stealing, and label inference attacks against split learning. arXiv preprint arXiv:2108.09033 (2021)
6. Geiping, J., Bauermeister, H., Dröge, H., Moeller, M.: Inverting gradients-how easy is it to break privacy in federated learning? *Advances in Neural Information Processing Systems* **33**, 16937–16947 (2020)
7. Gupta, O., Raskar, R.: Distributed learning of deep neural network over multiple agents. *Journal of Network and Computer Applications* **116**, 1–8 (2018)
8. Ha, Y.J., Lee, G., Yoo, M., Jung, S., Yoo, S., Kim, J.: Feasibility study of multi-site split learning for privacy-preserving medical systems under data imbalance constraints in covid-19, x-ray, and cholesterol dataset. *Scientific Reports* **12**(1), 1–11 (2022)
9. Hatamizadeh, A., Nath, V., Tang, Y., Yang, D., Roth, H., Xu, D.: Swin UNETR: Swin transformers for semantic segmentation of brain tumors in mri images. arXiv preprint arXiv:2201.01266 (2022)
10. Hatamizadeh, A., Yin, H., Molchanov, P., Myronenko, A., Li, W., Dogra, P., Feng, A., Flores, M.G., Kautz, J., Xu, D., et al.: Do gradient inversion attacks make federated learning unsafe? arXiv preprint arXiv:2202.06924 (2022)
11. He, Z., Zhang, T., Lee, R.B.: Model inversion attacks against collaborative inference. In: *Proceedings of the 35th Annual Computer Security Applications Conference*. pp. 148–162 (2019)
12. Jin, X., Chen, P.Y., Hsu, C.Y., Yu, C.M., Chen, T.: Catastrophic data leakage in vertical federated learning. *Advances in Neural Information Processing Systems* **34** (2021)
13. Kaissis, G., Ziller, A., Passerat-Palmbach, J., Ryffel, T., Usynin, D., Trask, A., Lima, I., Mancuso, J., Jungmann, F., Steinborn, M.M., et al.: End-to-end privacy preserving deep learning on multi-institutional medical imaging. *Nature Machine Intelligence* **3**(6), 473–484 (2021)
14. Kaissis, G.A., Makowski, M.R., Rückert, D., Braren, R.F.: Secure, privacy-preserving and federated machine learning in medical imaging. *Nature Machine Intelligence* **2**(6), 305–311 (2020)
15. Li, W., Milletari, F., Xu, D., Rieke, N., Hancox, J., Zhu, W., Baust, M., Cheng, Y., Ourselin, S., Cardoso, M.J., et al.: Privacy-preserving federated brain tumour segmentation. In: *International workshop on machine learning in medical imaging*. pp. 133–141. Springer (2019)

16. McMahan, B., Moore, E., Ramage, D., Hampson, S., y Arcas, B.A.: Communication-efficient learning of deep networks from decentralized data. In: Artificial intelligence and statistics. pp. 1273–1282. PMLR (2017)
17. Pasquini, D., Ateniese, G., Bernaschi, M.: Unleashing the tiger: Inference attacks on split learning. In: Proceedings of the 2021 ACM SIGSAC Conference on Computer and Communications Security. pp. 2113–2129 (2021)
18. Poirot, M.G.: Split Learning in Health Care: Multi-center Deep Learning without sharing patient data. Master’s thesis, University of Twente (2020)
19. Poirot, M.G., Vepakomma, P., Chang, K., Kalpathy-Cramer, J., Gupta, R., Raskar, R.: Split learning for collaborative deep learning in healthcare. arXiv preprint arXiv:1912.12115 (2019)
20. Rieke, N., Hancox, J., Li, W., Milletari, F., Roth, H.R., Albarqouni, S., Bakas, S., Galtier, M.N., Landman, B.A., Maier-Hein, K., et al.: The future of digital health with federated learning. NPJ digital medicine **3**(1), 1–7 (2020)
21. Ronneberger, O., Fischer, P., Brox, T.: U-Net: Convolutional networks for biomedical image segmentation. In: International Conference on Medical image computing and computer-assisted intervention. pp. 234–241. Springer (2015)
22. Rudin, L.I., Osher, S., Fatemi, E.: Nonlinear total variation based noise removal algorithms. Physica D: nonlinear phenomena **60**(1-4), 259–268 (1992)
23. Simpson, A.L., Antonelli, M., Bakas, S., Bilello, M., Farahani, K., Van Ginneken, B., Kopp-Schneider, A., Landman, B.A., Litjens, G., Menze, B., et al.: A large annotated medical image dataset for the development and evaluation of segmentation algorithms. arXiv preprint arXiv:1902.09063 (2019)
24. Singh, A., Vepakomma, P., Gupta, O., Raskar, R.: Detailed comparison of communication efficiency of split learning and federated learning. arXiv preprint arXiv:1909.09145 (2019)
25. Srivastava, N., Hinton, G., Krizhevsky, A., Sutskever, I., Salakhutdinov, R.: Dropout: a simple way to prevent neural networks from overfitting. The journal of machine learning research **15**(1), 1929–1958 (2014)
26. Thapa, C., Chamikara, M.A.P., Camtepe, S., Sun, L.: Splitfed: When federated learning meets split learning. arXiv preprint arXiv:2004.12088 (2020)
27. Usynin, D., Ziller, A., Makowski, M., Braren, R., Rueckert, D., Glocker, B., Kaissis, G., Passerat-Palmbach, J.: Adversarial interference and its mitigations in privacy-preserving collaborative machine learning. Nature Machine Intelligence **3**(9), 749–758 (2021)
28. Vepakomma, P., Gupta, O., Swedish, T., Raskar, R.: Split learning for health: Distributed deep learning without sharing raw patient data. arXiv preprint arXiv:1812.00564 (2018)
29. Wang, Z., Bovik, A.C., Sheikh, H.R., Simoncelli, E.P.: Image quality assessment: from error visibility to structural similarity. IEEE transactions on image processing **13**(4), 600–612 (2004)
30. Xu, Z., Burke, R.P., Lee, C.P., Baucom, R.B., Poulouse, B.K., Abramson, R.G., Landman, B.A.: Efficient multi-atlas abdominal segmentation on clinically acquired ct with simple context learning. Medical image analysis **24**(1), 18–27 (2015)
31. Yang, M., Lyu, L., Zhao, J., Zhu, T., Lam, K.Y.: Local differential privacy and its applications: A comprehensive survey. arXiv preprint arXiv:2008.03686 (2020)
32. Yang, Q., Liu, Y., Chen, T., Tong, Y.: Federated machine learning: Concept and applications. ACM Transactions on Intelligent Systems and Technology (TIST) **10**(2), 1–19 (2019)

33. Yin, H., Mallya, A., Vahdat, A., Alvarez, J.M., Kautz, J., Molchanov, P.: See through gradients: Image batch recovery via gradinversion. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 16337–16346 (2021)
34. Zhang, C., Li, S., Xia, J., Wang, W., Yan, F., Liu, Y.: Batchcrypt: Efficient homomorphic encryption for {Cross-Silo} federated learning. In: 2020 USENIX Annual Technical Conference (USENIX ATC 20). pp. 493–506 (2020)