

Adversarial Machine-Learning-Enabled Anonymization of OpenWiFi Data

Sambita Kuili¹, Kareem Dabbour¹, Irtiza Hasan¹, Andrea Herscovich¹, Burak Kantarci^{1,}, Marcel Chenier² and Melike Erol-Kantarci¹*

Abstract: Data privacy and protection through anonymization is a critical issue for network operators or data owners before it is forwarded for other possible use of data. With the adoption of Artificial Intelligence (AI), data anonymization augments the likelihood of covering up necessary sensitive information; preventing data leakage and information loss. OpenWiFi networks are vulnerable to any adversary who is trying to gain access or knowledge on traffic regardless of the knowledge possessed by data owners. The odds for discovery of actual traffic information is addressed by applied conditional tabular generative adversarial network (CTGAN). CTGAN yields synthetic data; which disguises as actual data but fostering hidden acute information of actual data. In this paper, the similarity assessment of synthetic with actual data is showcased in terms of clustering algorithms followed by a comparison of performance for unsupervised cluster validation metrics. A well-known algorithm, K-means outperforms other algorithms in terms of similarity assessment of synthetic data over real data while achieving nearest scores 0.634, 23714.57, and 0.598 as Silhouette, Calinski and Harabasz and Davies Bouldin metric respectively. On exploiting a comparative analysis in validation scores among several algorithms, K-means forms the epitome of unsupervised clustering algorithms ensuring explicit usage of synthetic data at the same time a replacement for real data. Hence, the experimental results aim to show the viability of using CTGAN-generated synthetic data in lieu of publishing anonymized data to be utilized in various applications.

Keywords: Anonymization, clustering techniques, cluster validation, generative CTGAN.

1. Introduction

WiFi-enabled services have been booming with the advent of widespread adoption of wirelessly connected devices. Simultaneously, a potential vulnerability persists in ensuring the privacy of

usage information. Previous studies employed data privacy techniques such as k-anonymity [1], l-diversity [2], t-closeness [3] and differential privacy [4]. However, those techniques can be complemented by considering data correlation, which possesses utmost significance while exploiting big data [5]. The heterogeneous characteristics underlying in OpenWiFi data possess identical traits of Big Data. Hence, adoption of novel machine learning techniques can eventually enable data privacy through anonymization.

Data anonymization is an essential task prior to making the data public for numerous domain applications. Apart from assuring privacy, it is also challenging to develop anonymized traces, which will resemble the original data. Therefore, retaining the resemblance of original in anonymized data by incorporating abundant noise alters the variance among features of data. This in turn raises issue on usability of data; which is where a trade-off exists while ensuring both data privacy and utility simultaneously on published data [6]. Previously, anonymization has been carried out mostly on healthcare-related domain pertaining to privacy for medical records of patients [7]. To the best of our knowledge, this paper paves the way for anonymization of WiFi usage streams for the first time.

A collaboration of unsupervised and generative adversarial network (GAN) is considered a holistic approach to initiate privacy preservation. Several clustering algorithms, namely K-means, Density Based Spatial Clustering of Application with Noise (DBSCAN) [8], Gaussian Hidden Markov Model [9] and Agglomerative are employed to analyze the actual pattern distribution of the usage traffic. While examining real-time data, there are no prior information on labels corresponding to each stream of data per timestamps. Therefore, employing clustering techniques provides useful insights to comprehend the behavioural patterns among samples in a multivariate data set. On the other hand, adopting a deep neural generative model GAN [10, 11] harbors generation of synthetic samples.

Commonly GAN is quite well-known in computer vision [12, 13]. There are other robust architectures for GANs which can yield remarkable performance on tabular data sets or non-image data sets. In light of these, this work aims to generate synthetic samples by leveraging a conditional tabular GAN (CTGAN) [14]; which is a modified version of architecture to traditional tabular GAN (TGAN) [15]. Hence, leveraging both clustering algorithms and computational power of CTGAN, we narrow down our work to anonymization by testing and validating in terms of quality of clusters. Therefore with support of cluster metrics it becomes easier to validate the performance of CTGAN; which aids us in assessing the quality of generated synthetic

¹School of Electrical Engineering and Computer Science, University of Ottawa, Ottawa, Canada

²NetExperience Inc., Ottawa, Canada

E-mail: skuil016@uottawa.ca; kdabb095@uottawa.ca;

ihasa074@uottawa.ca; ahers063@uottawa.ca; burak.kantarci@uottawa.ca;

marcel@netexperience.com; melike.erolkantarci@uottawa.ca

*Corresponding Author

Manuscript received 20 July 2023, accepted 02 January 2024, and ready for publication 21 March 2024.

© 2024 River Publishers

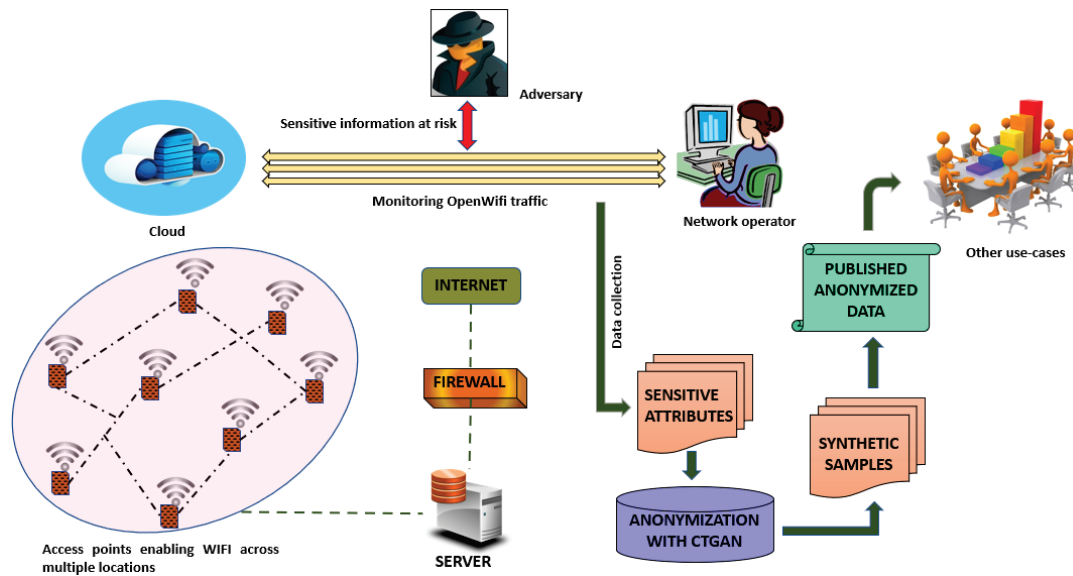


Figure 1.

Anonymization of OpenWiFi network traffic data.

samples. At the same time, it is crucial to test the amount of distortion being carried out in original samples. Simultaneously, we want to ensure that the synthetic samples resembles statistical properties of the original data. Figure 1 showcases topology for demonstration of data anonymization.

Key contributions of this paper are listed below with the ultimate goal to successfully build data anonymization while maintaining data utility and data privacy:

- The initiation of anonymization is being undergone by leveraging heuristic clustering algorithms: K-means, DBSCAN, Gaussian Hidden Markov Model (GHMM) and Agglomerative on a normalized data to showcase the potential cluster labels from the standpoint of distance, density and probability function.
- Leveraging unsupervised algorithms to present the formation of cluster labels and considering clustering membership information as a discrete variable and one of the pivotal parameters to train CTGAN.
- In addition, production of synthetic samples by training CTGAN on the real normalized samples and adopting effective measures to validate similarity performance of synthetic over real samples in terms of cluster metrics, for instance, Silhouette, Calinski and Harabasz (CH) and Davies-Bouldin (DB) scores.

The aforementioned state-of-the-art aligns to aim for high-quality data anonymization technique while advocating for the reliability and trustworthiness of using synthetic data from the standpoint of best-unsupervised algorithm and a number of cluster validation scores regardless of the distorted distribution of data. Over the course of this study, meticulous background information is provided in Section II highlighting the importance of data privacy algorithms and details of use-cases for generative adversarial networks (GANs) in terms of the generation of synthetic data. Moreover, Section III discusses the coherency of the data anonymization undertaken gradually on a multivariate time series

OpenWiFi data. Furthermore, an intensive evaluation showcases the significant performance comparison of adopted unsupervised clustering algorithms with regard to several validation metrics in Section IV, evaluating the quality of anonymized data in comparison to that of the real data. Finally, Section V provides the concluding remarks summarizing the research undergone in this study.

2. Related Work

Data anonymization is defined as the strategy of encrypting sensitive information specifically personally identifiable data in a dataset. In addition, the goal is to assure privacy protection of any information belonging to an attribute in the data by preventing information leakage when being exploited by others performing numerous use-cases. Among various methods of anonymization, for instance, generalization, masking, suppression, perturbation, and usage of synthetic data, the latter is given the major importance by advocating a machine learning model such as a Conditional Tabular Generative Adversarial Network (CTGAN). Consequently, prominent divergence infused on the real data reveals statistical alteration of the data while ensuring commonalities between synthetic and real data.

2.1. Data Privacy Algorithms

Previously data privacy preserving methods such as k -anonymity, l -diversity, t -closeness and δ -disclosure are leveraged to initiate privacy by compromising certain amendments to the attributes present in data set. Before exploring the process of those methods, it is crucial to understand following key terminologies identifiers, quasi-identifiers (QIDs) and sensitive attributes [16]. Generally, key identifiers signify those attributes with a unique code or

number; which are lacking by default in our data. However, quasi-identifier signifies attributes with discrete value denoting multiple parameters; such as MAC address, location IDs, equipment IDs, and timestamps in our work. Lastly, sensitive attributes are those consisting of other attributes which can neither be termed as identifiers or quasi-identifiers (QIDs). Sweeney [17] proposed an algorithm k-anonymity by elaborating on data; associated with person specific information to ensure privacy of identity and simultaneously considering the issue of re-identification attack into account. In fact, this algorithm works suitably well for selective attributes. Therefore, there is always a way for launching an adversarial attack using other attributes which are not being anonymized using k-anonymity. k-anonymity performs poorly in term of guaranteeing privacy due to presence of unequivocal homogeneity and background knowledge attacks. This is exactly when l-diversity plays a humongous role in ensuring privacy. On the contrary, the adversary seems to develop a thorough knowledge of background distribution of data even after relying on l-diversity algorithm. In order to avoid this consequence, Li et al. [2] introduced a novel algorithm t-closeness, which aligns the distribution of sensitive attributes with other attributes and therefore diminishing plausible background knowledge. Differential privacy is another well-known technique which shows improvement while incorporating privacy measures related to health care domains. Jain et al. [4] discusses on the amount of distortion incorporated by database into the data; which is determined in maintaining privacy and consequently can be found useful by data analysts.

2.2. Conditional Tabular Generative Adversarial Network

The above mentioned techniques still lack non trivial knowledge to convince data privacy for big data. As discussed earlier, Open-WiFi follows 5th Vs of big data and continual deployment of diverse data streams in high volume, therefore, those techniques are not sufficient to ensure data anonymization. Generative adversarial networks [11] are receiving notability lately for creation of fake synthetic samples and are employed for anonymization application by researchers. The architecture of this generative model comprises of two components: generator and discriminator. Those components can be designed by adopting different types of neural networks based on the applicability. Due to the complex architecture of GANs, classifying a sample into fake or original simply eases the entire process. This is the adversarial minmax game where the terminology of GANs turns up. The idea behind minmax game demonstrates that generator will minimize its performance of generating new synthetic training samples, on the contrary, discriminator will aim for maximizing its performance in classifying the generated and original samples into accurate scalar label.

$$\min_G \max_D V(D, G) = \mathbb{E}_{x \sim p_{data}(x)} [\log(D(x))] + \mathbb{E}_{z \sim p_z(z)} [\log(1 - D(G(z)))] \quad (1)$$

$p_z(z)$ is input noise distribution variable representing generator's distribution p_g over data samples x . A data mapping $G(z; \theta_g)$ denotes differential function G with parameters θ_g . Similarly, there is another differential function D with θ_d as parameters, which represents a mapping space for discriminator $D(x; \theta_d)$. In addition,

$D(x)$ signifies likelihood of x coming from real samples and not from distribution p_g . While training GANs, the motive is to train discriminator in assigning accurate labels for samples coming from each distribution mapping space. At the same time, train the generator to minimize the variable $\log(1 - D(G(z)))$. More clarification about the variables referred to (1) can be understood in [11].

There are other variant of generative models which are exploited for use case of generation of synthetic samples. Park et al. [16] introduce the adoption of tableGAN; which includes the architecture of deep convolutional generative adversarial networks (DCGAN) [12]. The main focus is being shown on generating relational synthetic tables based on existence of data attributes associated with convolutional neural networks [16].

Another interesting study undergone by Hajihassani et al. [18] which highlights their major contribution on anonymizing time series sensor data obtained from IoT devices with the help of Variational Autoencoder (VAE). Additionally, their work shares knowledge on understanding the anonymization into two aspects: algorithmic and systemic solutions. Considering algorithmic solutions involving machine learning, anonymization is being previously explored on the factors of security and protection with the assistance of deep neural networks and PrivacyNet [19]. Further insights on systemic solutions can be referred on [18] since those are not related with this work. As Xu and Veeramachaneni [15] propose the novel architecture of conditional tabular GAN emphasizing on the marginal analysis of data with an adoption of recurrent neural network, their work also shares other ways to obtaining synthetic data.

Adoption of other GANs, for instance, RGAN and RCGAN [20] are popular for generation of time series data. McCoy in [21] share the importance of CTGAN while training recommender systems which are useful to preserve user's privacy. Furthermore, this work also focuses on the reduction of actual raw training samples and in fact, expect an augmentation of actual samples in disguise of synthetic samples. Therefore, reliance on large retrieval of raw data is significantly reduced while still ensuring user privacy.

3. Methodology

Prior to the experiment on anonymization, processing the raw data forms primary step of pipeline. This dataset comprises of 20,000 service records collected from a small network operated over 4 weeks. Moreover, this network is composed of 2 access points averaging 2 connected clients a day. Furthermore, these access points are up and running efficiently, with clients in ON/OFF mode. By setting the desired duration and timestamps, 20K records are collected via REST APIs Cloud Service Portal for Open-WiFi until sufficient information on the active access points and clients connected to them is retrieved. In terms of processing and setting raw data ready for CTGAN, removal of any specific quasi-identifiers; including but not limited to any equipment IDs, location IDs, or customer IDs are implemented. In addition, scaling raw data samples address concerns with mode-specific normalization of CTGAN. Thus, modelling a CTGAN builds on existing numerical or sensitive attributes that contain at most 15,000 traffic records forming a multivariate time series data set.

Moreover, collected data does not contain any prior information on labels. Therefore, adoption of unsupervised learning

blends perfectly to comprehend underlying behavioural distribution of sensitive attributes. Those distributions are analyzed over clustering techniques such as K-means, Density Based Spatial Clustering of Application with Noise (DBSCAN) [22], Gaussian Hidden Markov Model (GHMM) [9]; establishing three different outlooks with respect to distance, density and probability factors. Moreover, an aggregation of clusters is visualized by selecting Agglomerative clustering algorithm. Quality of cluster validation for each unsupervised algorithm is assessed via Silhouette, CH and DB scores. At each stage of execution, potential clusters are demonstrated via Principal Component Analysis (PCA)-based dimensional reduction for more coherent understanding of original observable samples as well as synthetic samples.

Our processed multivariate data set consists of 15 feature variables with (11,900+) instances. While selection of optimal number of clusters for k-means clustering is challenging, setting a range from 2 to 10 is chosen. A trial and error execution is implemented to determine a best optimal number of cluster out of the arbitrary selection the aforementioned range. To obtain the optimal value k , a cluster validation metric Silhouette score is determined for each value k . Consequently, the value k for k-means turns out to be 2 with the highest Silhouette score ensuring grouping of observations based on distance heuristic with the focus on statistical similarity in data distribution of a multivariate dataset. Furthermore, the Silhouette score is an unsupervised clustering metric which quantifies the quality of the clustering of samples into the chosen optimal number of clusters i.e. k . The higher the Silhouette score (with a maximum score of 1) indicates a higher quality of grouping within each cluster and sufficient separability between them. On the other hand, estimating optimal parametric values for DBSCAN (ϵ and minPts) undergoes an exhaustive search approach highlighting highest Silhouette score [22]. The parameters ϵ and minPts are data dependant where ϵ stands for the radius to be followed with neighboring points from a particular data point in a cluster. In addition, minPts signifies the minimum neighboring data points required to be at a distance ϵ forming a dense distribution of data in a cluster. DBSCAN further identifies some samples as outliers which is bound to occur on high dimensional heterogeneous real data.

Therefore, minimum neighbouring points (minPts) 5 and a distance threshold epsilon (ϵ) 0.038 are selected. The number of states obtained for GHMM is domain-dependent. Kullback Leibler (KL) divergence has been used to obtain empirical probability distribution functions on HMM models with states 2, 3, and 4 on 1000 real RSSI WiFi signal samples [23]. Since there is a small difference between the KL divergence when using 2, 3, or 4 hidden states [24], a GHMM with 3 states is used to model this time series data. Additionally, in order to obtain a more cohesive view of the data, using agglomerative clustering with Ward linkage is also carried out in this work. Generally, agglomerative clustering follows an iterative hierarchical process starting from a single data point as a cluster and eventually merges into another cluster by identifying the closest pair of clusters.

4. Experimental Results

Tuning of parameters to achieve the highest Silhouette score enables analysis of the points exhibiting similar statistical property grouped into clusters. Initially, the original behaviour of real

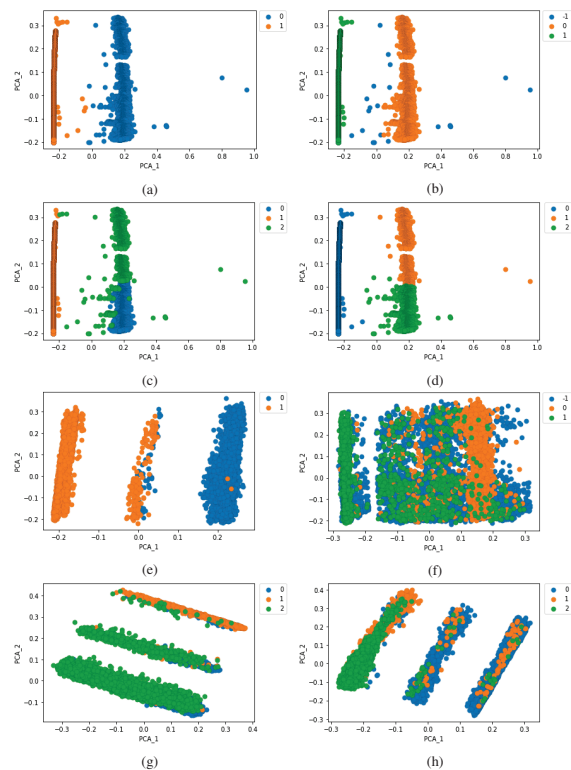


Figure 2.

Clustering visualization of real: (a–d) and synthetic data: (e–h).

data set right after preprocessing is learnt through clustering techniques as demonstrated in Figure 2(a–d); representing K-means, DBSCAN, GHMM and agglomerative algorithms respectively. On examining the Silhouette scores obtained for each technique, it is observed K-means acquires the highest Silhouette score with optimal cluster value 2 (k). Similarly, optimal cluster values for DBSCAN, GHMM and agglomerative obtains $k = 3$ respectively. At the same time, training a complex deep neural network architecture like CTGAN takes longer than expected while searching for best parameters suited for trained model. It is also challenging to validate the performance of synthetic samples generated by CTGAN. A well trained CTGAN is analyzed by understanding the performance loss convergence corresponding to two neural networks: generator and discriminator respectively as shown in Figure 3.

Most common problems which thrive while training GANs are mode collapse and loss convergence. This incorporates adversarial game between two neural networks where either there is a decrement in performance of generator loss and simultaneously an increment on discriminator loss. In addition, initial performance on acquiring the cluster validation scores for real data is analyzed as in Table 1. Furthermore, CTGAN is being trained to generate synthetic samples on being conditioned to discrete column of a data set; which contains cluster labels from each clustering algorithms at a time. The main motive is not limited to simply generate synthetic samples by exploiting generative model; but to check whether synthetic data can be relied upon as compared to real data. In order to assess the quality of synthetic data, clustering

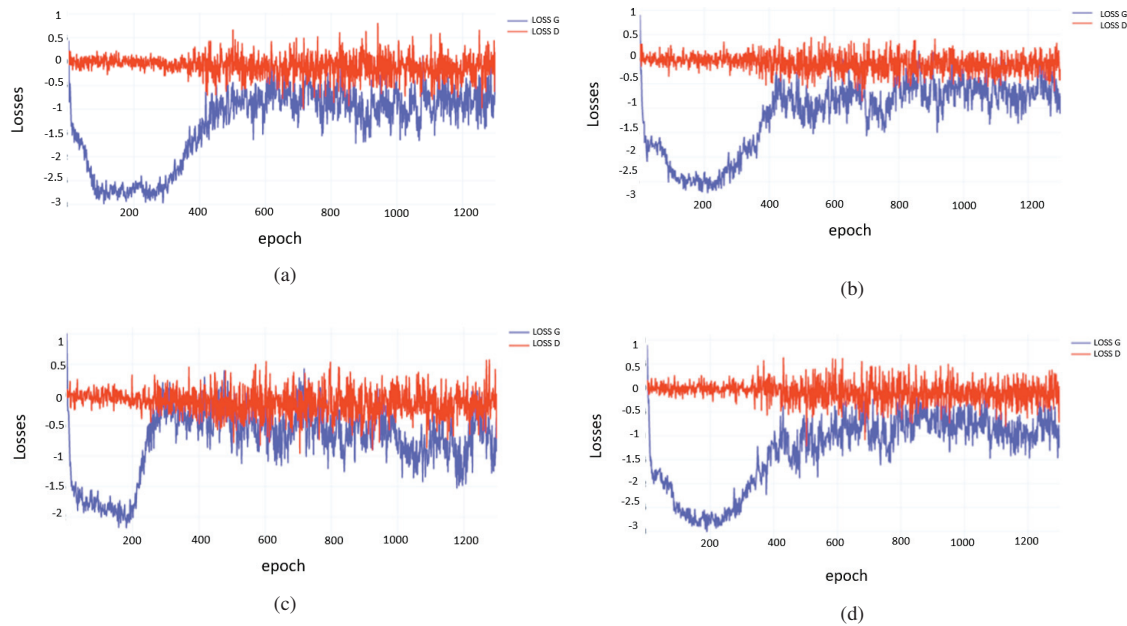


Figure 3.
Convergence of loss curves for CTGAN.

Table 1.
Cluster validation scores of unsupervised algorithms on real data

Unsupervised Algorithms	Real Data		
	Silhouette	Calinski Harabasz	Davies Bouldin
<i>K-means</i>	0.642	23342.92	0.599
<i>DBSCAN</i>	0.552	11659.34	5.559
<i>GHMM</i>	0.629	24788.87	0.529
<i>Agglomerative</i>	0.635	26109.21	0.501

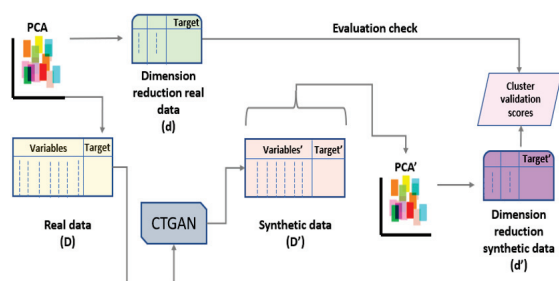


Figure 4.
Initiation of data anonymization leveraging clustering algorithms.

algorithms are employed which reflect the performance score for each synthetic and real data. As a result, more similarity or close the value for each metric, higher is the indication of replacing real with synthetic data. The initial real data (D) in Figure 4 is set with target attribute; which consists of labels obtained from K-means.

Moreover, we retain dimensionality reduction of real data (D) through PCA obtaining low dimensional data (d); which also

mitigates the problem collinearity among features. The core task for training generator of CTGAN requires conditional vectors. However, target attribute with labels 0 and 1 of K-means is considered as discrete variable while continuous variables as rest of the attributes or features existing on real data. Therefore, labels in discrete variable are utilized as conditional vectors to generate synthetic samples. A synthetic data (D') is obtained after training CTGAN while tuning for convergence with specific hyper-parameters demonstrated in the form of a loss curve in Figure 3(a).

This synthetic data (D') is further visualized with the help of PCA; forming another low dimensional data set (d'). Accordingly, evaluation of cluster metrics on (d') are computed. This exact method is carried out followed by altering the target variable subsequently with labels of DBSCAN, GHMM and agglomerative clustering algorithms, and correspondingly scores are calculated. Consequently, loss curves demonstrated in Figure 3(b), (c) and (d) are also checked for convergence after setting the labels for DBSCAN, GHMM and agglomerative algorithms as conditional vectors for generator. The cluster validation scores computed on synthetic data can be understood from Table 2.

While comparing the validation scores of K-means on both Tables 1 and 2, there is decent clarity which shows resemblance in scores for all three metrics. However, there is significant amount of deviation in scores for DBSCAN with respect to Silhouette, CH and DB. Similar to the Silhouette score, CH and DB scores aim to quantify quality of cluster formation. For instance, a higher CH score signifies better quality of clustering from the perspective of factors inclusive of within-cluster variance and between-cluster variance. Moreover, DB is quite the opposite of Silhouette and CH scores as it measures the average computation of inter-cluster distance and within-cluster dispersion of a cluster. More importantly, a lower DB score denotes greater cluster formation

Table 2.

Cluster validation scores of unsupervised algorithms on synthetic data			
Unsupervised Algorithms	Silhouette	Synthetic Data	
		Calinski Harabasz	Davies Bouldin
<i>K-means</i>	0.634	23714.57	0.598
<i>DBSCAN</i>	0.171	3742.45	1.836
<i>GHMM</i>	0.409	8865.44	0.770
<i>Agglomerative</i>	0.507	12188.80	0.769

quality. While comparing scores for GHMM, CH score reflects high deviation as compared to Silhouette and DB. On the contrary, agglomerative achieves reasonable scores for all metrics of both real and synthetic data. In this study, we intend to look for similarity between the scores for both real and synthetic data. For instance, K-Means preserves the statistical properties in synthetic data even after incorporating noise or conditional vectors by CTGAN into real data.

K-means is compatible while analyzing the performance of cluster validation scores from the perspective of data anonymization [25]. Additionally, a baseline algorithm K-Means on real data is selected to compare the cluster validation scores with each algorithm from synthetic data. It is observed that K-means exhibits resemblance in scores while comparing cluster validation metrics of real and synthetic data as shown in Tables 1 and 2, respectively. This coherently indicates that there is no significant statistical implication in the application context of data anonymization. However, DBSCAN, GHMM, and agglomerative algorithms result in the variation of both low Silhouette and Calinski Harabasz scores, and higher Davies Bouldin scores. In the context of application data anonymization and deviation in scores, it is understandable that DBSCAN is sensitive to outliers which results in performance downgrade under synthetic data when compared with the baseline algorithm. Similarly, the poor performance of GHMM and agglomerative algorithms highlights presence of noise in the synthetic data and involvement of model complexity with parameter initialization based on the data distribution. To comprehend deviation, we establish a visual analysis forming clusters of synthetic data for each algorithm shown in Fig. 2 (e-h); denoting K-means, DBSCAN, GHMM and agglomerative algorithms. On examine those figures, we encounter a skewed distribution of cluster labels as compared to Figure 2(a-d); which is rational to occur on a synthetic data. The reason of such distortion depends on the conditional vectors provided to train the generator of CTGAN. The resemblances in validation scores as seen on both real and synthetic data also provide assurance of preserving statistical properties similar to samples of real data; which is an essential factor as discussed in the earlier sections. Furthermore, this ensures that the synthetic data can be relied upon and effectively utilized for numerous applications.

5. Open Issues, Challenges, and Opportunities

The usage of multivariate time series OpenWiFi data where this study aims to adopt data anonymization and forecast the generation of the skewed distribution of synthetic samples already lack prior information of ground truth. Simultaneously, the analysis

focuses on mimicking cluster labels / ground truth based on the data distribution present in real and synthetic data setting an example to comprehend the problem only from the standpoint of the statistical context. However, other ways to augment the trustworthiness of synthetic data include privacy risk assessment, evaluating downstream tasks, and incorporating the review by domain experts who are knowledgeable about real data and will provide valuable feedback on synthetic data.

Every new data anonymization technique needs to cope with several challenges. Re-identifying instances of anonymized data, and safeguarding data utility is just one to mention. The latter can be understood if there is a loss of pivotal information in the real data, which can pave the way for an effective decision capability while exploiting the anonymized data. In addition, the dynamicity of multivariate time series data needs to be taken into account since it consists of temporal dependency and may impact anonymized data. Moreover, the scalability of the proposed framework is also of paramount importance which needs to be addressed while dealing with large volume of data. However, consistent performance of data anonymization is expected while safeguarding the data utility regardless of exploiting big data. Furthermore, scalability may also be studied from the use of CTGAN, which expects a large volume of data to train itself, therefore would ensure better prediction ability by studying the distinctions between synthetic and real data.

The idea of adversarial machine learning-enabled anonymization imparts the efficacy of data sharing across different domains such as ethical use of data by adopting a measure of privacy protection of sensitive information in the data.

6. Conclusion

Data anonymization is a crucial task for operators to augment privacy of information; which furthermore meets required user demands. This work summarizes the generation of synthetic samples by deploying conditional tabular GAN to replace real data instances with synthetic samples. Incorporation of clustering mechanism highlighted heterogeneous distribution of OpenWiFi traffic on the grounds of underlying factors such as distance, density and probability of samples existing in both real and synthetic data. In addition, computation of cluster validation metrics by well-known Silhouette, Calinski and Harabasz and Davies-Bouldin scores have been undergone to comprehend reliability of synthetic data. A wise similarity comparison of those scores has been checked while assessing synthetic with real data. Moreover, two dimensional visualization at every stage of implementation enables us to understand the original and the skewed distribution of synthetic data. This work has particularly focused on resemblance in the unsupervised validation metrics, which invokes the preservation of statistically correlated properties among samples. By leveraging this knowledge, extension of this work includes incorporation of other machine learning and privacy-based metrics to improve data anonymization.

Acknowledgment

This work was supported in part by Mathematics of Information Technology and Complex Systems (MITACS) Accelerate Program under Project Number IT24702.

References

- [1] A. Machanavajjhala, D. Kifer, J. Gehrke, and M. Venkatasubramanian, "l-diversity: Privacy beyond k-anonymity," *ACM Transactions on Knowledge Discovery from Data (TKDD)*, vol. 1, no. 1, pp. 3–es, 2007.
- [2] N. Li, T. Li, and S. Venkatasubramanian, "t-closeness: Privacy beyond k-anonymity and l-diversity," in *2007 IEEE 23rd international conference on data engineering*. IEEE, 2007, pp. 106–115.
- [3] X. Yang, T. Wang, X. Ren, and W. Yu, "Survey on improving data utility in differentially private sequential data publishing," *IEEE Transactions on Big Data*, vol. 7, no. 4, pp. 729–749, 2017.
- [4] P. Jain, M. Gyanchandani, and N. Khare, "Big data privacy: a technological perspective and review," *Journal of Big Data*, vol. 3, no. 1, pp. 1–25, 2016.
- [5] S. Biswas, N. Khare, P. Agrawal, and P. Jain, "Machine learning concepts for correlated big data privacy," *Journal of Big Data*, vol. 8, no. 1, pp. 1–32, 2021.
- [6] W. Liao, J. He, S. Zhu, C. Chen, and X. Guan, "On the tradeoff between data-privacy and utility for data publishing," in *2018 IEEE 24th International Conference on Parallel and Distributed Systems (ICPADS)*. IEEE, 2018, pp. 779–786.
- [7] J. Coutinho-Almeida, P. P. Rodrigues, and R. J. Cruz-Correia, "Gans for tabular healthcare data generation: A review on utility and privacy," in *International Conference on Discovery Science*. Springer, 2021, pp. 282–291.
- [8] E. Schubert, J. Sander, M. Ester, H. P. Kriegel, and X. Xu, "DbSCAN revisited, revisited: why and how you should (still) use dbSCAN," *ACM Transactions on Database Systems (TODS)*, vol. 42, no. 3, pp. 1–21, 2017.
- [9] L. R. Rabiner, "A tutorial on hidden markov models and selected applications in speech recognition," *Proceedings of the IEEE*, vol. 77, no. 2, pp. 257–286, 1989.
- [10] J. Yoon, L. N. Drumright, and M. Van Der Schaar, "Anonymization through data synthesis using generative adversarial networks (adsgan)," *IEEE journal of biomedical and health informatics*, vol. 24, no. 8, pp. 2378–2388, 2020.
- [11] I. Goodfellow, J. Pouget-Abadie, M. Mirza, B. Xu, D. Warde-Farley, S. Ozair, A. Courville, and Y. Bengio, "Generative adversarial nets," *Advances in neural information processing systems*, vol. 27, 2014.
- [12] A. Radford, L. Metz, and S. Chintala, "Unsupervised representation learning with deep convolutional generative adversarial networks," *CoRR*, vol. abs/1511.06434, 2015. [Online]. Available: <https://api.semanticscholar.org/CorpusID:11758569>.
- [13] J.-Y. Zhu, T. Park, P. Isola, and A. A. Efros, "Unpaired image-to-image translation using cycle-consistent adversarial networks," in *Proceedings of the IEEE international conference on computer vision*, 2017, pp. 2223–2232.
- [14] L. Xu, M. Skoularidou, A. Cuesta-Infante, and K. Veeramachaneni, "Modeling tabular data using conditional gan," *Advances in Neural Information Processing Systems*, vol. 32, 2019.
- [15] L. Xu and K. Veeramachaneni, "Synthesizing tabular data using generative adversarial networks," *arXiv preprint arXiv:1811.11264*, 2018.
- [16] N. Park, M. Mohammadi, K. Gorde, S. Jajodia, H. Park, and Y. Kim, "Data synthesis based on generative adversarial networks," *Proc. VLDB Endow.*, vol. 11, pp. 1071–1083, 2018. [Online]. Available: <https://api.semanticscholar.org/CorpusID:47017667>.
- [17] L. Sweeney, "k-anonymity: A model for protecting privacy," *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems*, vol. 10, no. 05, pp. 557–570, 2002.
- [18] O. Hajihassani, O. Ardakanian, and H. Khazaei, "Anonymizing sensor data on the edge: a representation learning and transformation approach," *ACM Transactions on Internet of Things*, vol. 3, no. 1, pp. 1–26, 2021.
- [19] V. Mirjalili, S. Raschka, and A. Ross, "Privacynet: semi-adversarial networks for multi-attribute face privacy," *IEEE Transactions on Image Processing*, vol. 29, pp. 9400–9412, 2020.
- [20] C. Esteban, S. L. Hyland, and G. Rätsch, "Real-valued (medical) time series generation with recurrent conditional gans," *arXiv preprint arXiv:1706.02633*, 2017.
- [21] S. V. McCoy, "Exploration of user privacy preservation via ctgan data synthesis for deep recommenders."
- [22] S. Kuili, B. Kantarci, M. Erol-Kantarci, M. Chenier, and B. Herscovici, "A holistic machine learning approach to identify performance anomalies in enterprise wifi deployments," in *Big Data IV: Learning, Analytics, and Applications*, vol. 12097. SPIE, 2022, pp. 105–117.
- [23] Ó. Belmonte-Fernández, E. Sansano-Sansano, A. Caballer-Miedes, R. Montoliu, R. García-Vidal, and A. Gascó-Compte, "A generative method for indoor localization using wi-fi fingerprinting," *Sensors*, vol. 21, no. 7, p. 2392, 2021.
- [24] M. Vidyasagar, "Bounds on the kullback-leibler divergence rate between hidden markov models," in *2007 46th IEEE Conference on Decision and Control*. IEEE, 2007, pp. 6160–6165.
- [25] M. E. Ferrão, P. Prata, and P. Fazendeiro, "Utility-driven assessment of anonymized data via clustering," *Scientific Data*, vol. 9, no. 1, p. 456, 2022.

Biographies



Samhita Kuili is a PhD candidate and a researcher in Smart Connected Vehicles Innovation Centre (SCVIC) at the University of Ottawa. Her research direction includes AI-enabled anomaly detection in 5G and local wireless network connection WiFi 6. She collaborated on her research by working alongside an enterprise in the field of OpenWiFi networks. She is working on a 5G wireless network to detect security attacks and anomalies. She is one of the founding members of ACM-Women Student Chapter at University of Ottawa. She is an active student member of the IEEE Communication Society and IEEE Communication Systems Integration and Modeling Technical Committee.



Kareem Dabbour graduated with an honours Bachelor of Science in Computer Science specializing in Data Science from the University of Ottawa. He is a computer science and math

enthusiast that is committed to leveraging the power of both disciplines for innovative problem-solving.



Irtiza Hasan holds a Bachelor of Science degree in Computer Science from the University of Ottawa, specializing in Data Science. He is interested in exploring solutions that prioritize ethical considerations while advancing the development of artificial intelligence at a graduate level.

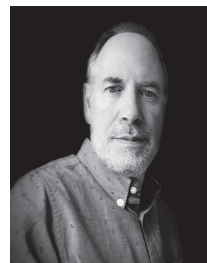


Andrea Herscovich works as a software developer at CoLab software. She has a Bachelor of Science in Computer Science with Honours from the University of Ottawa.



Burak Kantarci (S'05, M'09, SM'12) received the Ph.D degree in computer engineering in 2009. He is a Full Professor and the Founding Director of Smart Connected Vehicles Innovation Centre (SCVIC), and the Founding Director of the Next Generation Communications and Computing Networks (NEXTCON) Research Lab, University of Ottawa. He has coauthored over 250 publications in established journals and conferences, and 15 book chapters. He is well known for his contributions to the quantification of data trustworthiness in mobile crowd-sensing (MCS) systems, and game theoretic incentives to promote user participation in MCS campaigns with

high value data; as well as AI-backed access control, authentication and machine learning-backed intrusion detection solutions in sensing environments. In 2022, he has been awarded a Minister's Award of Excellence in Innovation and Entrepreneurship from Ontario Ministry of Colleges and Universities. He served as the Chair of IEEE Communications Systems Integration and Modeling Technical Committee, and has served as the Technical Program Co-Chair/Symposium Co-chair of more than twenty international conferences/symposia/workshops, including IEEE Global Communications Conference (GLOBECOM) – Communications Systems QoS, Reliability and Modeling (CQRM) symposium. In 2021, he has been elected as the new Secretary of IEEE Social Networks Technical Committee. He is an Editor of the IEEE Communications Surveys & Tutorials, IEEE Internet of Things Journal, Vehicular Communications (Elsevier), and an Associate Editor for IEEE Networking Letters, and Journal of Cybersecurity and Privacy. He is Editor for several IEEE and Elsevier journals. He was an ACM Distinguished Speaker in 2019–2021, currently an ACM Senior Member. He is IEEE Systems Council and IEEE Communications Society Distinguished Lecturer.



Marcel Chenier is CTO and co-founder of NetExperience where he brings three decades of technology leadership experience in the field of public and enterprise wireline, wireless and cloud networking. In his role, he is responsible for technology direction and ensuring product feedback from Net Experience's Service Provider customers.

Prior to NetExperience, Marcel took the role of CTO and VP of Engineering for KodaCloud building a Cloud Network control solution for MSP WLAN networks. Previous to KodaCloud, Marcel lead product architecture for the WLAN portfolio and Small Cell Plug&Play Solutions at Ericsson. Prior to Ericsson, Marcel was Vice-President of engineering at BelAir Networks, a WLAN carrier-class small-cell equipment provider which was acquired by Ericsson. Prior to BelAir, Marcel was head of engineering at Zhone Technologies Canada from early company stage to IPO responsible for voice and video packet products. As VP of R&D for Premisis, a smart DLC equipment provider, Marcel was part of the leadership team that lead the acquisition by Zhone. Marcel started his career with Nortel Networks working on silicon, hardware and software research and development and as an R&D Executive for the Signaling System Product Line. Marcel is a co-inventor on several patents in the area of wireless networking architecture.

Marcel holds a Bachelor of Engineering from the University of Ottawa. He was a member of the Intel Technical Advisory board from 2011 to 2014.



Melike Erol-Kantarci is Chief Cloud RAN AI/ML Data Scientist at Ericsson and Canada Research Chair in AI-enabled Next-Generation Wireless Networks and Full Professor at the School of Electrical Engineering and Computer Science at the University of Ottawa. She is the founding director of the Networked Systems and Communications Research (NETCORE) laboratory. She has received numerous awards and recognitions. Dr. Erol-Kantarci is the co-editor of three books on smart grids, smart cities and intelligent transportation. She has over 200 peer-reviewed publications. She has delivered 70+ keynotes, plenary talks and tutorials around the globe. Dr. Erol-Kantarci is on the editorial board of the IEEE Transactions on Communications, IEEE Transactions on Cognitive Communications and Networking and IEEE Networking Letters. She has acted as the general chair and technical program chair for many international conferences and workshops. Her main research interests are AI-enabled wireless networks, 5G and 6G wireless communications, smart grid and Internet of Things. Dr. Erol-Kantarci is an IEEE ComSoc Distinguished Lecturer, IEEE Senior member and ACM Senior Member.

