



Politecnico di Bari

Repository Istituzionale dei Prodotti della Ricerca del Politecnico di Bari

Frequency Matters: On the Impact of Carrier Frequency on Privacy in Radio Fingerprinting

This is a post print of the following article

Original Citation:

Frequency Matters: On the Impact of Carrier Frequency on Privacy in Radio Fingerprinting / Huso, I., Sciancalepore, S., Oligeri, G., Piro, G., Boggia, G.. - In: IEEE WIRELESS COMMUNICATIONS LETTERS. - ISSN 2162-2337. - STAMPA. - 14:7(2025), pp. 1904-1908. [10.1109/LWC.2025.3558440]

Availability:

This version is available at <http://hdl.handle.net/11589/286040> since: 2025-08-17

Published version

DOI:10.1109/LWC.2025.3558440

Publisher:

Terms of use:

(Article begins on next page)

Frequency Matters: On the Impact of Carrier Frequency on Privacy in Radio Fingerprinting

Ingrid Huso , Savio Sciancalepore , Gabriele Oliveri , Giuseppe Piro , and Gennaro Boggia 

Abstract—Radio Frequency Fingerprinting (RFF) relies on unique inherent imperfections in radios’ hardware to authenticate devices based on Radio Frequency emissions. In this work, we consider that fingerprints collected for multi-channel transmitters on certain frequencies get partially leaked to an adversary willing to track them, without information about the frequency used for training. In this scenario, we evaluate the performance of various state-of-the-art Convolutional Neural Networks for image-based RFF when the testing and training frequencies do not match. We demonstrate that RFF performances degrade significantly when training and testing frequencies differ, down to a random guess when they are sufficiently apart.

Index Terms—Physical Layer Security, Internet of Things (IoT), Authentication

I. INTRODUCTION

Ensuring the authenticity of communications and devices within a wireless network is a major challenge. Traditional cryptographic methods, while effective, often require significant computational resources, and request the possibility to modify the devices’ firmware and software. In this context, Radio Frequency Fingerprinting (RFF) techniques are emerging as a highly promising means of identification and authentication by taking advantage of the unique and inherent characteristics of Radio Frequency (RF) signals emitted by wireless devices [1]. RFF operates under the assumption that each wireless device, due to small manufacturing variances, exhibits distinct signal differences or “fingerprints” that can be exploited to enforce security properties such as the authentication of the end-user device. Specifically, unavoidable variances in electronic components like oscillators, amplifiers, and modulators may impact the signal phase and frequency without affecting the quality of the received signal. Thus, by analyzing these characteristics, RFF can uniquely identify (i.e.,

fingerprint) a device, like a biometric identifier in humans. In detail, the RFF process involves collecting RF emissions at a receiver, extracting relevant features, and training a model to recognize specific transmitters in the wild [2]. On the one hand, RFF is emerging as a promising technique to guarantee authenticity in wireless communications; on the other hand, it raises significant challenges, especially concerning privacy. When devices operate on a certain frequency, a dedicated receiver, which is tuned on the same frequency, collects their RF emissions to create radio fingerprints and then saves them on a server (e.g., a database in a local security service) [3]. Thus, unauthorized access to servers responsible for storing radio fingerprints poses a significant risk of data leakage [4]. Moreover, since each device features a unique fingerprint, a leakage of the RFF model (or portion of it) raises serious privacy concerns in its usage for inferring an individual’s location, habits, and even his social interactions, thus allowing for surveillance and infringing personal privacy [5]. In this context, to the best of the author’s knowledge, there has been no attempt to evaluate the impact of a partial leakage of the RFF model on the devices’ anonymity and privacy. Current research mostly assumes that the RFF model of a device is unique and remains unaffected by operational factors, such as the frequency at which the receiver gathers the RF emissions. Thus, assuming the availability of the RFF model to an adversary, no research evaluated the performance of Deep Learning (DL) image-based RFF when no information on the carrier frequency on which the model was collected and trained is known. To bridge this gap, our work studies (i) the consistency of the RFF model of a device across various carrier frequencies and (ii) the consequences of tracking attacks carried out via RFF with a mismatch between the carrier frequency at training time and the one used at testing time (i.e., partial leakage of the RFF model of a device).

Contribution. In detail, we investigate the impact of partial information leakage about the RFF model of an RF device on the overall RFF procedure. Through several tailored controlled experiments using Software Defined Radios (SDRs) and state-of-the-art DL image-based RFF models, we demonstrate that the RFF model of an RF device strongly depends on the carrier frequency used for the communication. While very close carrier frequencies are generally characterized by similar RFF models, differences increase with larger frequency shifts. When attackers obtain a stolen leaked RFF model without prior knowledge of the carrier frequency, these differences lead to a notable drop in the RFF accuracy, potentially reducing RFF performance to random guessing, even in scenarios with a small number of devices in the network.

Ingrid Huso, Giuseppe Piro, and Gennaro Boggia are with the Dept. of Electrical and Information Engineering, Politecnico di Bari, Bari, Italy; {ingrid.huso, giuseppe.piro, gennaro.boggia}@poliba.it.

Savio Sciancalepore is with the Eindhoven University of Technology and Eindhoven Artificial Intelligence Systems Institute (EAISI), Eindhoven, Netherlands; s.sciancalepore@tue.nl.

Gabriele Oliveri is with College of Science and Engineering (CSE), Hamad Bin Khalifa University (HBKU), Doha, Qatar; goligeri@hbku.edu.qa.

This work was supported by the European Union under the Italian National Recovery and Resilience Plan (NRRP) of NextGenerationEU, in the context of partnership on “Telecommunications of the Future” (PE00000001 - program “RESTART”, CUP: D93C22000910001), national center on “Sustainable Mobility” (CN00000023 - program “MOST”, CUP: D93C22000410001), and partnership on “Cybersecurity” (PE00000007 - program “SERICS”, CUP: D33C22001300002, project ISP5G+). It was also supported by the PRIN 2022 projects INSPIRE (grant no. 2022BEXMXN_01) and HORUS (grant no. 2022P44KA8) funded by the Italian MUR, by the HORIZON MSCA project BRIDGITISE (grant no. 101119554), and by “The house of emerging technologies of Matera (CTEMT)” project funded by the Italian MIMIT.

Our manuscript is the first to show that the lack of knowledge about the frequency used at training time makes the leaked RFF model unusable to the attackers and forces them to acquire a new profile, with significant time loss and cost. At the same time, frequency hopping represents a viable option for the RF device to make tracking attacks via RFF harder for the attacker, so contributing to keeping anonymity and location privacy longer. We also release our data as open source at [6], to foster reproducibility and extension of our results.

Paper organization. The rest of this paper is organized as follows. Sec. II provides the preliminaries, Sec. III introduces the scenario and adversary model, Sec. IV discusses our methodology, Sec. V summarizes our results and, finally, Sec. VI concludes the paper and outlines future work.

II. BACKGROUND

Digital Modulation. Wireless communication systems use digital modulation techniques to preprocess baseband signals before transmitting them at high carrier frequencies [7]. Specifically, a digital modulation scheme creates a modulated signal which consists of an in-phase and a quadrature component, typically coupled in a complex value $I + jQ$, where the I and Q vectors are the real and imaginary component, respectively. Given a bit sequence, the transmitter uses the modulation scheme to translate bits into IQ samples, while the receiver uses it to retrieve the original bit value from the received IQ samples.

Radio Frequency Fingerprinting. Research on RFF in recent years focused on efficient and effective solutions to extract meaningful features from received RF signals. Approaches working on raw samples acquire a given number of raw samples of the signals at the Physical (PHY) layer, namely IQ samples, and feed them directly into a DL classifier [8], [9]. Although this strategy is successfully deployed in a wide range of scenarios, it generates sensitive DL models that hardly adapt to different channel conditions, mobility, and power cycling of RF devices [10]. Conversely, image-based RFF systems assemble raw IQ samples into 2-D or 3-D images, transforming the RFF problem into an image classification problem. Several contributions shown that image-based RFF systems are more successful than the previous ones when it comes to identifying RF devices in challenging channel conditions [11] and across various power cycles of devices [10].

Deep Learning. DL techniques have been recently adapted to the wireless communication and image recovery domains [12]. In this context, Convolutional Neural Network (CNN) models demonstrate to achieve better performances, especially for image processing tasks [13]. Indeed, by exploiting their ability to learn and extract features, CNNs are deployed into a wide range of scenarios, including image classification [14]. In detail, CNNs comprise several neurons that apply convolutional operations and improve their performances via a learning process [15]. The architecture of a CNN consists of a sequence of three layers. Convolutional layers are primarily responsible for feature extraction by generating a feature map consisting of the convolution of the input layer. Consequently, pooling layers minimize the spatial dimensions of the feature

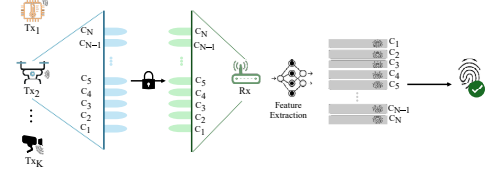


Fig. 1. Reference scenario: the transmitter communicates with the receiver on a pseudo-random channel through encrypted communication.



Fig. 2. Adversary model: the adversary is challenged to identify the transmitter by resorting to RFF while exploiting (leaked) information associated with the DL-based RFF model of the transmitter on channel c_x .

map. Finally, fully connected layers perform complex feature extraction tasks and generate predictions [15].

III. REFERENCE SCENARIO AND ADVERSARY MODEL

Reference Scenario. Fig. 1 depicts our reference scenario, including one or more RF transmitters and one RF receiver. The RF devices communicate wirelessly on a pseudo-random channel $c(t)$ chosen by the transmitter and the receiver as a function of previously established secrets. For sake of clarity, when referring to a channel, we specifically mean a RF communication channel characterized by a designated carrier frequency. Herein, the main objective of a transmitter is to stay anonymous to all devices in the network except to the receiver, which in turn performs the transmitter's physical-layer authentication via RFF. To this aim, the receiver resorts to a dataset of pre-trained models $\mathcal{M} = \{M_1, \dots, M_N\}$ constituted by N RFF models, each of them referring to a specific channel, which is stored on a server. It is noteworthy to mention that, the legitimate receiver already knows the carrier frequency required to identify the transmitter due to synchronization for communication. Therefore, the receiver tunes to the pre-defined channel c_x , select M_x , with $x \in [1, N]$, initiates the reception of the signals from the transmitter, and finally validates the transmitter applying M_x . Specifically, following the contribution in [16], the receiver deploys image-based RFF techniques using CNN pre-trained models for image classification achieving robustness against channel fluctuations.

Adversary model. Fig. 2 shows our adversary model. We consider that one of the RFF models $M_a \in \{M_1, \dots, M_N\}$ with $a \in [1, N]$ has been leaked to the adversary, as a result of the adversary gaining unauthorized access to the server where the RFF model is stored. Specifically, the adversary is aware of the RFF model, but they do not know which channels and carrier frequencies are associated with that specific RFF model. Thus, the adversary aims to identify the transmitter (i.e., fingerprint) by leveraging the leaked model M_a . Note that, given the high-security constraints, i.e., a random selection of the communication channel, carrier frequency, transmitter anonymity, and unawareness of the association of a model to a particular channel, the adversary

task is particularly challenging. In this work, we investigate the impact of M_a leakage on transmitter anonymity and the influencing configuration parameters.

IV. METHODOLOGY

This section reports the methodology used by both the legitimate receiver and the adversary to assess the impact of partial leakage of the RFF model and carrier frequency on the RFF. It exploits DL image-based RFF models, which are more robust to environmental changes compared to those using raw IQ samples [10], [16]. Specifically, our approach involves converting raw samples of the signal taken from the radio spectrum into images, due to its remarkable robustness to noise and other side effects. In line with such a methodology, we transform the RFF problem into an image classification problem. The main steps involve: (i) *IQ sample collection*, (ii) *images generation*, and finally (iii) *multi-class classification*.

IQ sample collection. We collect IQ samples by tuning both the transmitter and the receiver on the same channel. We consider the Binary Phase-Shift Keying (BPSK) modulation scheme, where the in-phase component assumes a value either -1 or +1, while the quadrature component is always zero. By mapping the I component and the Q component to the real and imaginary parts of a complex number, respectively, the BPSK decoding process follows Eq. 1:

$$x(t) = \begin{cases} -1 \cos(2\pi f_0 t), & \text{if } b = 0, \\ +1 \cos(2\pi f_0 t), & \text{if } b = 1, \end{cases} \quad (1)$$

where $x(t)$ is the transmitted modulated signal, f_0 is the carrier frequency, and b is the bit value. Thus, given a carrier frequency f_0 , the couples $[-1, 0]$ and $[1, 0]$ represent the theoretical position of the received IQ samples in the IQ plane. However, due to radio imperfections, the collected IQ samples are distributed in the IQ plane, generating a specific pattern that identifies the device's fingerprint.

Image Generation. The image generation phase processes the collected raw IQ samples and generates grayscale images following the baseline procedure described in [10]. Specifically, our procedure involves collecting K IQ samples and then dividing the IQ plane and the clouds of points created by such IQ samples into $N \times M$ tiles, with the values of N and M determining the image dimensions. Afterward, for each tile (m, n) , we count the IQ samples that fall in the tile (bivariate histogram). To guarantee that such value maps to a correct pixel value in the generated image, if the count exceeds 255 (i.e., the maximum possible value of the pixel of an image), the value is truncated to 255. To this aim, it is fundamental to calibrate the number of IQ samples per image to minimize the loss of information, i.e., too many tiles exceeding 255 samples.

Multi-class Classification. This task involves classifying images generated during the previous phase. Specifically, in line with the adversary model described earlier, the aim of the proposed multi-class classification problem is to identify the transmitter device. We first divided the collected IQ samples into three subsets, i.e., training, validation, and testing. We then considered state-of-the-art CNNs pre-trained on the ImageNet database [17], i.e., Alexnet, Resnet-18, Resnet-50,

Resnet-101, Inceptionv3. We used their implementation in MatLab2023b, where we adapted the input and output layers to fit our problem. The input layers are re-sized to fit the size of the images generated from raw IQ samples, while the output layers are re-designed to accommodate the number of classes in the specific experiment.

We run two main experiments. We first investigate the impact of the carrier frequency on the RFF model of a device. Specifically, it aims to assess to what extent the fingerprint of a particular device changes when changing the carrier frequency. Thus, we run a multi-class classification task where the number of classes is coincidental with the number of tested channels. Secondly, we evaluate the impact of a partial information leakage of the RFF model by analyzing the mismatch in the training and testing channel used for the RFF, in line with the adversary model discussed in Sec. III. Herein, we consider a multi-class classification problem where the number of classes is coincidental with the number of distinct devices in our setup. For each test and device, we train the RFF model on the IQ samples acquired at a specific channel, and we test it using IQ samples acquired at a different channel. Thus, we denote δ as the absolute value of the difference between the channel considered for training and the one used for testing, i.e., $\delta = |ch_{train} - ch_{test}|$, and we evaluate the performance of the RFF models for increasing values of δ .

V. PERFORMANCE EVALUATION

Experimental Testbed. We collected the IQ samples by resorting to four LimeSDRs devices. The considered SDRs feature the *LMS7002M* RF Transceiver, capable of running any wireless standard and mobile communication, including WiFi and 4G [18]. We connect the SDRs to an Ubuntu 22.04 workstation, equipped with a 12th Gen Intel(R) Core(TM) i7 @2.70 GHz processor. Moreover, investigating carrier frequency mismatches in real-world environments is challenging due to multipath, shadowing, and interference, which can significantly impact RFF accuracy. Thus, in line with several previous studies [10], [16], [19]–[21], we employ a direct wired connection between the transmitter and the receiver via a coaxial RF cable, ensuring that the observed RFF characteristics primarily reflect the radio itself rather than environmental variations. At the same time, note that this setup constitutes an extreme advantage for the adversary: indeed, as a legitimate receiver, the adversary can also focus only on the RFF task, without worrying about the mentioned noise figures.

To drive the behavior of the SDR, we use the *GNU Radio 3.10* software, offering the possibility to configure the radios with the desired communication parameters. Thus, we set the transmitter and the receiver gain to 50 dB and 70 dB, respectively. We consider the communication frequencies and channels defined by the IEEE 802.15.4 communication technology [22] and used by devices compliant with the Zigbee specification [23]. Thus, we consider 16 channels in the frequency range 2405 – 2480 MHz, each characterized by a bandwidth of 2 MHz and an inter-channel spacing of 5 MHz, according to the IEEE 802.15.4 standard specification. The transmission chain defined on GNURadio consists of

four blocks: i) a *File Source*, used to generate a message consisting of a string of 256 bytes with incremental values; ii) a *Constellation Modulator*, configured to handle the BPSK modulation scheme; iii) a *Multiply Constant*, used to adjust the amplitude of the signal to avoid saturation, and iv) the *LimeSuite Sink*, where the radio signals are up-converted to the selected carrier frequency, with a sample rate of 256K samples per second. On the receiver, we use six main blocks: i) the *LimeSDR Source*, used to receive the radio signal at the selected carrier frequency; ii) a *Rational Resampler*, acting as a filter; iii) an *AGC*, used for mitigating channel fluctuations; iv) a *Symbol Sync*, in charge of decoding the digital signal; v) a *Costas Loop*, deployed for phase and frequency mitigation offsets, and vi) a *File Sink*, storing the output of the whole reception chain into a ".iq" file.

Specifically, we deploy a selected LimeSDR as the receiver while we connect the other three radios alternately to act as transmitters. For the sake of clarity, to avoid the influence of power cycling on the RFF [10], the data collection campaign is conducted without switching off the receiver. Overall, we ran 60 tests lasting 90 seconds for each of the 16 carrier frequencies on each transmitter. Subsequently, all the collected data are uploaded to a centralized server for running our tests. The collected data are available open-source at [6]. For the data analysis, we use the High-Performance Computing (HPC) cluster available at TU/e in Eindhoven (NL), providing 2 GPUs Tesla V1000 with 256 GB of RAM. We consider $K = 10^6$ IQ samples per image, and we use images of size $M \times N = 255 \times 255$, in line with the size of the images of the ImageNet database in MATLAB2023b. We used 60%, 20%, and 20% of the data for training, validation, and testing, respectively and, due to the tested number of channels, we test the performance of the various RFF models by considering values of δ in the interval $\delta = [0, 15]$.

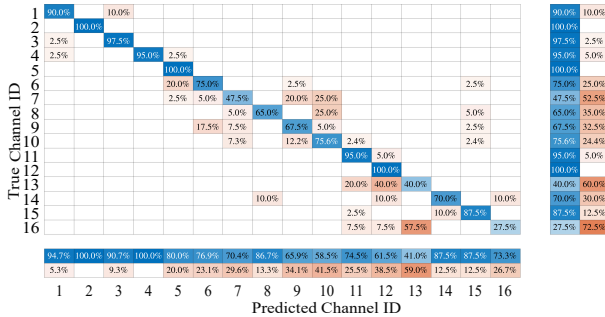


Fig. 3. Alexnet accuracy confusion chart. We train and test Alexnet on 16 channels, using a test set of 40 images per channel.

Fingerprint Robustness to Carrier Frequency. In this section, we investigate how a model trained on a specific channel performs when challenged with data coming from any other channel. We start with a preliminary example that takes into account one device and one CNN, i.e., Alexnet. We trained a model considering images generated from all the available 16 channels, and then, we tested the same model considering images from any channel. Our training set consists of 84 minutes of measurements for each channel, while the test set sums up to 6 minutes, 560 and 40 images, respectively.

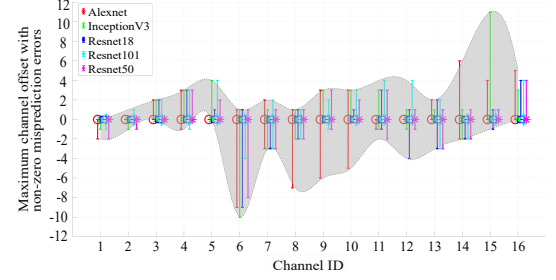


Fig. 4. Maximum channel offset with misprediction greater than zero, considering various CNNs.

Fig. 3 shows the resulting confusion chart. Firstly, we observe that the accuracy tends to be maximum when the data extracted from the channel considered for the test set is the same as that of the training set. In addition, the confusion chart highlights that miss predictions fall only on adjacent channels, e.g., channel 13 is predicted as 11, 12, and 13 in 20%, 40%, and 40% of the cases, respectively. Another example is channel 7, which has an accuracy of about 47% when testing on the same channel, while the remainder 53% is spread over channels 5, 6, 9, and 10 respectively. Since this pattern is consistent for all the channels, i.e., only cells close to the diagonal are likely to have an accuracy greater than zero, we claim that neighbor channels are better at preserving the features of the trained model. We also notice minor fluctuations in the classifier accuracy, mostly due to unpredictable fluctuations of real-world channel conditions. However, as shown in Fig. 3, the overall trend of the classification accuracy remains consistent across all channels. Once the model is trained on a particular channel, device fingerprint classification errors predominantly occur on adjacent channels, with misclassifications on distant channels being relatively rare. In fact, recalling our testbed and the associated assumptions, mispredictions occur when the features of the transmitter extracted from a given channel are detected in another channel, making the model associated with a specific channel (M_a) usable to identify a transmitter on another channel. We also observe some exceptions, e.g., channel 6 has been predicted as channel 15; we will investigate the phenomenon in the following. We now focus on the maximum offset between the channel considered for training and the one for testing with misprediction greater than zero. As an example, recalling Fig. 3, we observe that the maximum offset associated with channels 1, 2, and 6 is 2, 0, and 9, respectively. We stress that we take into account the maximum offset (distance between the channel adopted for training and the one considered for testing) independently of the misprediction error—this one being a conservative approach since, in the vast majority of the cases, the misprediction is characterized by a small accuracy. Considering different CNNs, i.e., Alexnet, Inceptionv3, Resnet-18, Resnet-101, and Resnet-50, Fig. 4 shows the maximum channel offset leading to a misprediction greater than zero as a function of the (reference) channel. Fig. 4 confirms that only adjacent channels tend to preserve the features of the transmitter. In fact, the grey-shaded area is mainly concentrated in a range of ± 5 (maximum offset). We observe a few exceptions; e.g., channels 5 and 14 are affected

by a large offset for all the considered networks. Moreover, we want to stress that our analysis took into account the conservative assumption of mispredictions greater than zero independently of their values. In fact, recalling Fig. 3, we note that channel 6 has a maximum offset of 9 (channel 15) but with a misprediction of 2.5%. Therefore, our analysis confirms that features extracted from a signal on a specific channel can be effectively used to test the model on the same channel or, with lower accuracy, on adjacent channels.

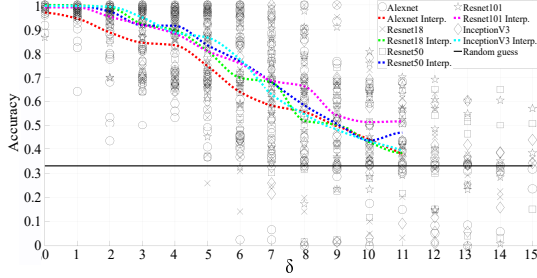


Fig. 5. RFF Accuracy at channel distances $\delta \in [0, 15]$, using various CNNs.

Tracking Attacks. We now focus on tracking attacks, i.e., the identification of the device from the leaked profile, when considering different channels when training and testing. We consider three transmitting devices and one receiver. We recall the configuration adopted for the previous test, i.e., a training set constituted of 252 minutes (84 minutes for each device) of measurements for each channel (for a total of 16 channels) and a test set of 18 minutes (6 minutes for each device) of measurements for each channel. We train various CNNs models on the measurements of a transmitter on a specific channel, and we challenge them to identify the same transmitter among the others with testing measurements acquired on channels at various distances $\delta \in [0, 15]$. Fig. 5 shows the accuracy of the CNN in correctly identifying the transmitter as a function of the offset δ between the channels used for training and testing. The best-case scenario is constituted by $\delta = 0$, i.e., the testing is performed on the same channel of the training. In such a case, the adversary maximizes the likelihood of guessing the transmitter, i.e., accuracy higher than 0.95 for all the considered CNN. Another example is the offset $\delta = 2$, where the average accuracy (dashed lines) for different networks is still higher than 0.9 but with higher variance (0.5 to 1). The accuracy drops when considering higher offset values up to the offset value equal to 11, approximating the random guess (0.33). On the one hand, these results confirm that the model acquired for RF devices serves to identify them only around the carrier frequency where it has been collected, while it becomes useless on other (farther) carrier frequencies. On the other hand, the RF device could keep anonymity longer by hopping among the available channels, maximizing the time necessary to build a reliable profile for RFF.

VI. CONCLUSIONS

In this paper, considering the use case of a partial leakage of the RFF model of an RF device, we explored the relation between the carrier frequency and the performance of

RFF models. Our real experiments, conducted using SDRs, highlight the critical dependence of the RFF accuracy on the matching between the training and testing frequencies. Our findings confirm that RFF models have high accuracy when the carrier frequency remains consistent throughout training and testing. Variations in frequency introduce significant challenges, resulting in lower accuracy. In turn, frequency discrepancies improve radio anonymity against adversaries that feature models trained on specific (and unknown) channels. Future research involves investigating RFF across multiple frequencies in low Signal-to-Noise Ratio (SNR) real-world scenarios characterized by purely wireless channels.

REFERENCES

- [1] A. Jagannath, et al., “A comprehensive survey on radio frequency (RF) fingerprinting: Traditional approaches, deep learning, and open challenges,” *Computer Networks (Elsevier)*, 2022.
- [2] N. Soltanieh, et al., “A Review of Radio Frequency Fingerprinting Techniques,” *IEEE Journal of Radio Frequency Identification*, 2020.
- [3] S. H. et al., “Radio fingerprinting for anomaly detection using federated learning in lora-enabled industrial internet of things,” *Future Generation Computer Systems (Elsevier)*, 2023.
- [4] Z. L. et al., “Non-inducible rf fingerprint hiding via feature perturbation,” in *IEEE International Conference on Communications*, 2023.
- [5] L. Abanto-Leon, et al., “Stay Connected, Leave no Trace: Enhancing Security and Privacy in WiFi via Obfuscating Radiometric Fingerprints,” *Proc. of ACM on Measur. and Analys. of Comput. Syst.*, 2020.
- [6] I. Huso, et al., “Open Source Data of SDRs on Different Channels,” <https://tinyurl.com/rxjv5wrc>, accessed: 23-May-2024.
- [7] T. S. Rappaport, *Wireless communications: principles and practice*. Cambridge University Press, 2024.
- [8] A. Al-Shawabka, et al., “Exposing the Fingerprint: Dissecting the Impact of the Wireless Channel on Radio Fingerprinting,” in *IEEE INFOCOM*, 2020.
- [9] B. Hamdaoui, et al., “Deep-learning-based device fingerprinting for increased lora-iot security: Sensitivity to network deployment changes,” *IEEE Network*, 2022.
- [10] S. Alhazbi, et al., “The Day-After-Tomorrow: On the performance of radio fingerprinting over time,” in *Proc. of ACSAC*, 2023.
- [11] G. Oliveri, et al., “PAST-AI: Physical-layer authentication of satellite transmitters via deep learning,” *IEEE Trans. on Inf. Forens. and Secur.*, 2023.
- [12] H. He, et al., “Deep learning-based channel estimation for beamspace mmwave massive mimo systems,” *IEEE Wirel. Commun. Lett.*, 2018.
- [13] A. Hermawan, et al., “CNN-Based Automatic Modulation Classification for Beyond 5G Communications,” *IEEE Commun. Lett.*, 2020.
- [14] J. Xie, et al., “Activity Pattern Aware Spectrum Sensing: A CNN-Based Deep Learning Approach,” *IEEE Commun. Lett.*, 2019.
- [15] A. Younesi, et al., “A Comprehensive Survey of Convolutions in Deep Learning: Applications, Challenges, and Future Trends,” *IEEE Access*, 2024.
- [16] L. Papangelo, et al., “Adversarial Machine Learning for Image-Based Radio Frequency Fingerprinting: Attacks and Defenses,” *IEEE Commun. Magaz.*, 2024.
- [17] O. Russakovsky, et al., “ImageNet Large Scale Visual Recognition Challenge,” *International journal of computer vision*, 2015.
- [18] Lime Microsystems, “LMS7002M – FPRF MIMO Transceiver IC,” Data Sheet.
- [19] Y. L. et al., “Radio frequency fingerprinting exploiting non-linear memory effect,” *IEEE Transactions on Cognitive Communications and Networking*, 2022.
- [20] R. Kong and H. Chen, “Csi-rff: Leveraging micro-signals on csi for rf fingerprinting of commodity wifi,” *IEEE Transactions on Information Forensics and Security*, 2024.
- [21] W. G. et al., “Radio frequency fingerprint acquisition and identification for small sample dmr signals under blind synchronization,” *IEEE Access*, 2025.
- [22] IEEE, “IEEE Standard for Low-Rate Wireless Networks,” *IEEE Std 802.15.4-2020 (Revision of IEEE Std 802.15.4-2015)*, 2020.
- [23] ZigBee Alliance, “Zigbee specification,” Specification, August 2015, ZigBee Document – 05-3474-21.