

Titre: Transfer learning-empowered physical layer security in aerial
Title: reconfigurable intelligent surfaces-based mobile networks

Auteurs: Yosefine Triwidyastuti, Tri Nhu Do, Ridho Hendra Yoga Perdana,
Authors: Kyusung Shim, & Beongku An

Date: 2025

Type: Article de revue / Article

Référence: Triwidyastuti, Y., Do, T. N., Perdana, R. H. Y., Shim, K., & An, B. (2025). Transfer
Citation: learning-empowered physical layer security in aerial reconfigurable intelligent
surfaces-based mobile networks. IEEE Access, 13, 5471-5490.
<https://doi.org/10.1109/access.2025.3526178>

Document en libre accès dans PolyPublie

Open Access document in PolyPublie

URL de PolyPublie:
PolyPublie URL: <https://publications.polymtl.ca/61781/>

Version: Version officielle de l'éditeur / Published version
Révisé par les pairs / Refereed

Conditions d'utilisation:
Terms of Use: Creative Commons Attribution 4.0 International (CC BY)

Document publié chez l'éditeur officiel

Document issued by the official publisher

Titre de la revue: IEEE Access (vol. 13)
Journal Title:

Maison d'édition: IEEE
Publisher:

URL officiel: <https://doi.org/10.1109/access.2025.3526178>
Official URL:

Mention légale: This work is licensed under a Creative Commons Attribution 4.0 License. For more
Legal notice: information, see <https://creativecommons.org/licenses/by/4.0/>

Received 9 December 2024, accepted 27 December 2024, date of publication 6 January 2025, date of current version 9 January 2025.

Digital Object Identifier 10.1109/ACCESS.2025.3526178

RESEARCH ARTICLE

Transfer Learning-Empowered Physical Layer Security in Aerial Reconfigurable Intelligent Surfaces-Based Mobile Networks

YOSEFINE TRIWIDYASTUTI¹, TRI NHU DO², (Member, IEEE),
RIDHO HENDRA YOGA PERDANA¹, (Graduate Student Member, IEEE),
KYUSUNG SHIM³, (Member, IEEE), AND BEONGKU AN⁴, (Senior Member, IEEE)

¹Department of Software and Communications Engineering, Graduate School, Hongik University, Sejong 30016, Republic of Korea

²Department of Electrical Engineering, Polytechnique Montréal, Montreal, QC H3T 1J4, Canada

³School of Computer Engineering and Applied Mathematics, Hankyong National University, Anseong 17579, Republic of Korea

⁴Department of Software and Communications Engineering, Hongik University, Sejong 30016, Republic of Korea

Corresponding author: Beongku An (beongku@hongik.ac.kr)

This work was supported by the National Research Foundation of Korea (NRF) grant funded by Korea Government Ministry of Science and ICT (MSIT) under Grant NRF-2022R1A2B5B01001190.

ABSTRACT This paper investigates the enhancement of physical layer security (PHY security) in Reconfigurable Intelligent Surfaces (RIS)-aided terrestrial and non-terrestrial networks (TN/NTN), focusing on the challenges posed by node mobility. In the context of next-generation mobile networks, ensuring secure communication is critical, especially under varying channel conditions caused by mobility. We explore different mobility models, including random walk, Gauss-Markov, and reference point group mobility, to assess their impact on key security metrics such as secrecy capacity and average secrecy rate. To address these challenges, we develop robust algorithms for optimizing the phase-shift configurations of RIS. Additionally, we employ Artificial Intelligence (AI) and Machine Learning (ML) techniques, specifically Deep Neural Networks (DNN), for performance prediction of PHY security metrics. We also leverage transfer learning to enhance model robustness across different mobility scenarios through domain adaptation. Our results demonstrate the effectiveness of our proposed methods in maintaining high levels of PHY security despite the dynamic nature of the channel conditions and the mobility of nodes. The proposed phase-shift configuration algorithms and ML-based solutions ensure secure and resilient communication in aerial RIS-aided TN/NTN, contributing to the advancement of secure mobile networks.

INDEX TERMS Physical layer security, reconfigurable intelligent surface, reference point group mobility, transfer learning, unmanned aerial vehicle.

I. INTRODUCTION

The next generation network toward 6G is envisioned to provide limitless connectivity through the integration of terrestrial networks (TNs) and non-terrestrial networks (NTNs). TNs primarily consist of ground-based infrastructure such as base stations and user equipment (UE). Conversely, NTNs incorporate aerial platforms like satellites, high-altitude platforms (HAPs), and unmanned aerial vehicles (UAVs), offering a more comprehensive coverage and enhanced

connectivity, especially in remote or underserved areas. Examples of UAV use case in TNs and NTNs for remote area are wildlife tracking, disaster recovery and public safety missions [1]. Looking forward to 6G, the integration of TNs and NTNs is expected to become more seamless, promising ubiquitous connectivity and improved network performance [2].

One of the persistent challenges in both TNs and NTNs is managing the issues of line-of-sight (LOS) and non-line-of-sight (NLOS) propagation. In LOS conditions, signals travel directly from the transmitter to the receiver, typically resulting in stronger and more reliable communication links.

The associate editor coordinating the review of this manuscript and approving it for publication was Cong Pu¹.

However, in NLOS conditions, where signals are obstructed by buildings, terrain, or other obstacles, signal quality can severely degrade. Additionally, the mobility of nodes in these networks—such as moving vehicles, UAVs, and users—introduces further variability and complexity to the channel conditions. To meet the increasing demand for mobility in TN/NTN, it is necessary to be aware of the advancements in innovative architectures, breakthrough technologies, and adaptive strategies [3].

Reconfigurable Intelligent Surfaces (RIS) have emerged as a promising solution to address the challenges associated with NLOS conditions and node mobility. RIS are artificial surfaces with electronically controllable elements that can manipulate electromagnetic waves [4]. By dynamically adjusting these elements, RIS can reflect signals in desired directions, effectively creating virtual LOS paths even in NLOS environments [5]. Integrating RIS increases the likelihood of having a LOS communication and improves the probability of LOS availability [6]. This capability not only enhances signal strength and reliability but also improves the overall network performance in dynamic and complex scenarios [7].

Physical layer security (PHY security) is a critical concern in both TNs and NTN, aiming to protect data from eavesdropping and unauthorized access by leveraging the physical characteristics of the communication medium [8]. However, the introduction of RIS poses new challenges to PHY security. While RIS can enhance signal propagation, it can also inadvertently reflect signals in directions that could be intercepted by eavesdroppers.

There are several researches concerning physical layer security in RIS deployment. Authors in [5], [9], [10], and [11] analyzed the secrecy performance of RIS-enabled communication in the presence of an eavesdropper. However, the analysis in the previous researches considered an ideal continuous phase-shift at the RIS. High resolution phase-shifters are costly. Therefore, in practice, only limited phase shifts are available to be configured [12], [13].

Other authors in [14], [15], [16], [17], [18], and [19] analyzed the secrecy performance in RIS-aided network with the assumption of perfect knowledge in both legitimate and illegitimate channel state information (CSI). In real world, this assumption is hard to fulfill, especially in the condition of passive eavesdropping. The direct CSI of an illegitimate node is available only if it is active or it is a licensed user that has legal access to the legitimate communication system [20]. However, when the illegitimate user only overhears the legitimate user transmission, it does not send its CSI feedback to transmitter, which causes one of challenges in this literature.

Node mobility in TN/NTN also adds another challenge in PHY security. One of the primary characteristics of mobile communications is the rapid time-variation of the channel coefficients induced by the Doppler spread. The fast time-variation in the channel makes mobile networks susceptible to channel aging, accordingly effective countermeasures

for PHY security in mobility scenarios becomes more difficult [12]. Hence, Doppler effect, channel selectivity, and shorter coherence time in mobile networks constitute more pronounced problem from a PHY security perspective [21].

Designing approaches with low complexity and efficient system performance to optimize aerial RIS-enhanced TN/NTN are challenging, especially when UAVs are deployed in a partially unknown environment. Artificial intelligence (AI) and machine learning (ML) approaches are powerful tools for designing and optimizing such networks. Deep neural networks (DNN) as a class of ML algorithms excels in handling non-linear models. However, several challenges still need to be investigated — for example, large computational processing power, high energy consumption, and latency [22].

Transfer learning is one of the potential candidates to make rapid decision-making with fast sampling efficiency and mitigate the issues in most ML methods [23]. Transfer learning is an advanced ML technique that involves transferring knowledge from one domain (source domain) to another (target domain). This approach is particularly beneficial in scenarios where the target domain has limited data but shares similarities with the source domain.

Khan et al. leveraged domain adaptation under transfer learning paradigm to deal with the outdated channel [24]. Domain adaptation refers to the process of using a DNN model to discover and transfer latent knowledge from the source domain to target domain. In [24], the transfer learning approach allows the designed detector to adapt itself properly into different channel environments and improve the system performance. In addition, authors in [25] established a deep transfer learning (DTL)-based framework to optimize the phase shifts at the RIS. However, both researches in [24] and [25] do not consider PHY security in the RIS-assisted networks.

Given these challenges, the key research focus is to develop strategies that ensure PHY security in aerial RIS-aided TNs and NTNs under various mobility scenarios. This includes leveraging advanced techniques such as transfer learning to adaptively enhance security measures in real-time, considering the dynamic nature of mobile networks. Addressing these issues is crucial for the next generation of secure, efficient, and resilient mobile communication systems.

In this paper, we focus on evaluating PHY security performance in aerial RIS-aided TN/NTN with emphasis on node mobility. Our goal is to design robust algorithms for discrete phase-shift configuration of RIS under limited resource and adapt transfer learning approach to deal with the varying channel conditions due to node mobility. The detailed contributions of our work are outlined as follows:

- We analyze the impact of different mobility models on PHY security performance in aerial RIS-aided TN/NTN. By considering the random walk, Gauss-Markov, and reference point group mobility models,

we aim to capture a broad spectrum of mobility patterns encountered in real-world scenarios.

- We develop robust algorithms for optimizing the phase-shift configuration of quantized RIS elements, namely optimal secrecy-oriented phase shift (OSP) and maximizing main channel real coefficient phase-shift configuration (MRC). OSP considers perfectly known CSI of the legitimate and eavesdropper channel, while MRC assumes without the knowledge of eavesdropper CSI. We also consider random phase-shift (RPS) as the performance benchmark.
- We build DNN model and apply transfer learning technique to address the regression problem of performance prediction in aerial RIS-aided TN/NTN. Our study is not only a mere application of transfer learning method, but also a thorough investigation of the transfer learning impacts from the model modifications and dataset variations on the secrecy performance of aerial RIS-aided TN/NTN under different mobility models.
- We comprehensively evaluate PHY security performance under various mobility models using metrics such as secrecy capacity and average secrecy rate. This analysis provides insights into the effectiveness of our proposed algorithms in protecting information from eavesdropper in aerial RIS-aided TN/NTN.

The organization of this paper is as follows: In Section I, we present an introduction to the topic, outlining the significance of enhancing PHY security in aerial RIS-aided TN/NTN and the challenges posed by node mobility. Section II provides a detailed description of the system model we proposed, including the integration of RIS and the leakage information in eavesdropper. Section III describes the details of various mobility models considered for every node in the proposed system model. In Section IV, we formulate the optimization problem aimed at designing an effective transmission protocol by optimizing the phase-shift configurations of the RIS. Section IV elaborates on the robust algorithms we developed to solve this optimization problem, including Optimal Secrecy-oriented Phase-shift (OSP), Maximizing Real Coefficient (MRC), and Random Phase Shift (RPS). Section V discusses the use of DNN model and transfer learning techniques for efficient performance prediction of PHY security metrics. In Section VI, we present our simulation results, providing a comprehensive analysis of the impact of different mobility models and the effectiveness of our proposed algorithms. Finally, Section VII concludes the paper by summarizing our key findings and insights, emphasizing the practical solutions provided for ensuring secure and resilient communication in aerial RIS-aided TN/NTN.

II. SYSTEM MODEL

In this paper, we consider an aerial RIS-aided secure transmission mobile system with one source node (S), one RIS mounted on UAV (U), one destination node (D), and

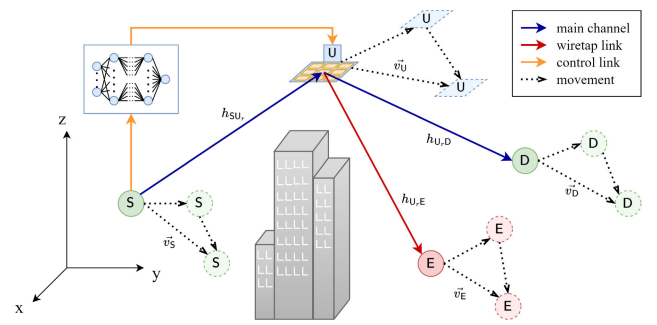


FIGURE 1. Illustration of the considered system model of the RIS-assisted communication.

an eavesdropper (E) as illustrated in Fig. 1. We assume that all nodes in the system (S, U, D and E) are mobile with the velocity vector $\vec{v}$ and has its own speed v and direction η . In addition, all nodes are assumed to be equipped with a single antenna. Due to far distance and the presence of obstacles, there is NLOS transmission from S to D and from S to E. Hence, the main channel transmission is assisted by the installation of a RIS with R elements that is attached to U. At the same time, E wiretaps the information from U that makes the system unsecured.

Let $\Theta = \text{diag}([\kappa_1 e^{j\theta_1}, \dots, \kappa_r e^{j\theta_r}, \dots, \kappa_R e^{j\theta_R}])$ denotes the phase-shift matrix of U, where $\kappa_r \in (0, 1]$ and $\theta_r \in [0, 2\pi)$ represent the amplitude reflection coefficient and the phase-shift of the r -th reflecting element at U, respectively. For the sake of convenience, we assume that the amplitude reflection coefficient for all R elements are the same ($\kappa = \kappa_r \forall 1 \leq r \leq R$), then the phase-shift matrix of U becomes $\Theta = \kappa \text{diag}([e^{j\theta_1}, \dots, e^{j\theta_R}])$.

The received signal at D that is transmitted from S and reflected by U can be written as

$$s_D = \sqrt{P_S} \sum_{r=1}^R \sqrt{F_{UD}} \tilde{h}_{U,D} \kappa e^{j\theta_r} \sqrt{F_{SU}} \tilde{h}_{S,U} u_S + n_D, \quad (1)$$

where P_S denotes the transmit power at S and u_S denotes the transmitted signal from S. $\tilde{h}_{S,U}$ and $\tilde{h}_{U,D}$ denote the complex channel coefficients from S to r -th element of U and from r -th element of U to D, respectively. Meanwhile, n_D is the additive white Gaussian noise (AWGN) at D with zero mean and variance σ_D^2 .

F_{AB} with $A \in \{S, U\}$ and $B \in \{U, D, E\}$ in (1) denotes the large-scale fading of the wireless channel from A to B that can be formulated as [26]

$$F_{AB} = 20 \log_{10} \left(\frac{4\pi f}{c} \right) + 10n \log_{10}(d_{AB}) + \chi_\sigma, \quad (2)$$

where f , c , n , and χ_σ denote the carrier frequency, light speed, path loss exponent, and shadow fading standard deviation, respectively. The distance between two nodes is denoted by d_{AB} that can be calculated in the xyz -plane as

$$d_{AB} = \sqrt{(x_A - x_B)^2 + (y_A - y_B)^2 + (z_A - z_B)^2}. \quad (3)$$

It is further assumed that within a coherence time block, the location of the mobile nodes and their connections do not change. Thus, the network topology is considered static but changes independently from one time block to another. By incorporating mobility model in a series of time blocks, we can capture the dynamic nature of nodes' mobility and characterize the time-series changing of nodes' locations and connections [27].

Using polar representation, the complex channel coefficients in (1) can be expressed as $\tilde{h}_{SU_r} = \hat{h}_{SU_r} e^{j\phi_{SU_r}}$ and $\tilde{h}_{U_rD} = \hat{h}_{U_rD} e^{j\phi_{U_rD}}$, where $\hat{h}_{SU_r}$ and $\hat{h}_{U_rD}$ are the magnitudes of the channel coefficients while ϕ_{SU_r} and ϕ_{U_rD} are their phases. Hence, the received signal-to-noise ratio (SNR) at D can be formulated as

$$\gamma_D = \rho_D F_{SU} F_{UD} \kappa^2 \left| \sum_{r=1}^R \hat{h}_{SU_r} \hat{h}_{U_rD} e^{j(\theta_r + \phi_{SU_r} + \phi_{U_rD})} \right|^2, \quad (4)$$

where $\rho_D = P_S / \sigma_D^2$ denotes the average SNR at D. Additionally, the achievable capacity of the received signal at D is given by

$$C_{SUD} = \log_2(1 + \gamma_D). \quad (5)$$

In the wiretap link, let $\tilde{h}_{U_rE}$ be the complex channel coefficient from the r -th element of U to E. The received signal at E can be written as

$$s_E = \sqrt{P_S} \sum_{r=1}^R \sqrt{F_{UE}} \tilde{h}_{U_rE} \kappa e^{j\theta_r} \sqrt{F_{SU}} \tilde{h}_{SU_r} u_S + n_E, \quad (6)$$

where n_E is the AWGN with zero mean and variance σ_E^2 at E. Incorporating more factors such as interference, environmental noise and imperfect CSI is also critical importance in PHY security of RIS-aided networks. In this work, the influence of those factors can be treated as noise effect. Detail analysis for different type of such factors can be dealt with in the future works. The resulting SNR at E based on (6) then can be expressed as

$$\gamma_E = \rho_E F_{SU} F_{UE} \kappa^2 \left| \sum_{r=1}^R \hat{h}_{SU_r} \hat{h}_{U_rE} e^{j(\theta_r + \phi_{SU_r} + \phi_{U_rE})} \right|^2, \quad (7)$$

where $\rho_E = P_S / \sigma_E^2$ denotes the average SNR at E. Without loss of generality, we assume that $\rho = \rho_D = \rho_E$. The magnitude and phase of the channel coefficient from the r -th element of U to E are denoted by $\hat{h}_{U_rE}$ and ϕ_{U_rE} , respectively. Thus, the achievable capacity of the overheard signal at E is given by

$$C_{SUE} = \log_2(1 + \gamma_E). \quad (8)$$

III. MOBILITY MODELS

In TN/NTN, node mobility introduces significant variability and complexity to the communication channels, impacting the performance and security of the network. To accurately characterize these effects, we focus on several mobility models: the random walk mobility (RWM) model, the

Gauss-Markov mobility (GMM) model, and the reference point group (RPG) mobility model. The random walk model captures the unpredictable and erratic movements of nodes, representing high randomness in mobility. The Gauss-Markov model provides a more realistic scenario where the current velocity and direction are correlated with past behavior, simulating smoother transitions. The reference point group mobility model is particularly relevant for scenarios involving coordinated movements, such as those of a group of nodes following a leader, which is common in applications involving UAVs and other aerial platforms. By studying these models, we aim to comprehensively understand and mitigate the impact of mobility on the performance and security of aerial RIS-aided TN/NTN systems, leading to robust and adaptive phase-shift optimization strategies.

A. LEGITIMATE USERS

Regarding to all nodes' mobility, we assume that all nodes during their movements maintain the same height through all the time. Specifically, we also assume that legitimate node movements in the xy -plane are modeled as the RPG mobility model, where they move toward the same direction. In this system, the role of U is to help or assist main channel transmission such that the information can be received securely at D. Then, the movements of other legitimate nodes (S and U) are directed by the movement of D.

The location of D as the reference point of RPG mobility model in the xy -plane at time t can be expressed as

$$x_g(t) = x_g(t-1) + v_g(t) \cos(\eta_g(t)), \quad (9)$$

$$y_g(t) = y_g(t-1) + v_g(t) \sin(\eta_g(t)), \quad (10)$$

where $v_g(t)$ and $\eta_g(t)$ represent the speed and direction of the group. We assume that the movement of D as the speed and direction of the group follows GMM model, which can be expressed as

$$v_g(t) = \alpha v_g(t-1) + (1-\alpha) \mathbb{E}[v_g] + \sigma_{v_g} \sqrt{(1-\alpha^2) m_{v_g}(t)}, \quad (11)$$

$$\eta_g(t) = \alpha \eta_g(t-1) + (1-\alpha) \mathbb{E}[\eta_g] + \sigma_{\eta_g} \sqrt{(1-\alpha^2) m_{\eta_g}(t)}, \quad (12)$$

where α indicates the memory level. $m_{v_g}(t)$ denotes the Gaussian distribution of group speed with mean $\mathbb{E}[v_g]$ and variance $\sigma_{v_g}^2$. Meanwhile, $m_{\eta_g}(t)$ denotes the Gaussian distribution of group direction with mean $\mathbb{E}[\eta_g]$ and variance $\sigma_{\eta_g}^2$ [28]. GMM model has smooth movement that is more similar with real mobile node movement. GMM model was originally proposed for the simulation of a personal communication system (PCS) [29].

According to the RPG mobility model, the speed and direction of other legitimate nodes at time t is a combination of group movement and random deviation, which can be

mathematically expressed as

$$v_{\text{node}}(t) = \min\{\max\{v_g(t) + \Delta v, 0\}, V_{\max}\}, \quad (13)$$

$$\eta_{\text{node}}(t) = \eta_g(t) + \Delta\eta, \quad (14)$$

where $\text{node} \in \{S, U\}$. Δv and $\Delta\eta$ represent random speed and direction deviation of each node which follows a uniform distribution [30]. $V_{\max}$ denotes the maximum speed of every node. Consequently, the position of each legitimate node at time t can be written as

$$x_{\text{node}}(t) = x_{\text{node}}(t-1) + v_{\text{node}}(t) \cos(\eta_{\text{node}}(t)), \quad (15)$$

$$y_{\text{node}}(t) = y_{\text{node}}(t-1) + v_{\text{node}}(t) \sin(\eta_{\text{node}}(t)). \quad (16)$$

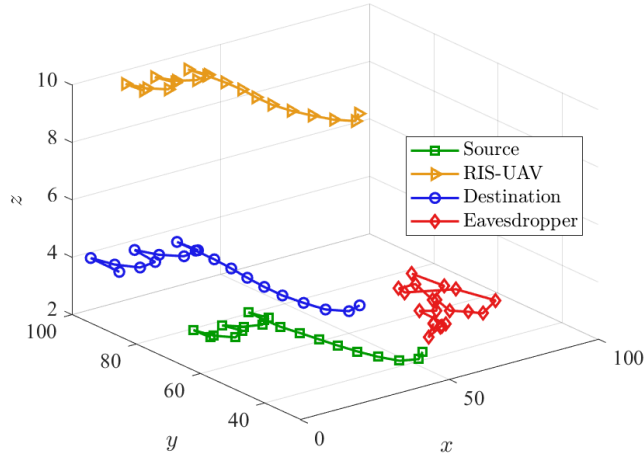


FIGURE 2. The track of 20 different locations for each node under RWM model.

Aside from group motion occurs frequently in ad hoc networks and there is relationship among mobile nodes [30], RPG mobility model can be used to describe users' movement for many scenarios, such as in military battlefield communications and during disaster recovery in search and rescue (SAR) operations [31]. With RPG and GMM model for node movement in aerial RIS-aided TN/NTN, the mobility scenarios in this research are general to diverse conditions as group movement in ad hoc networks and individual node movement in PCS.

B. EAVESDROPPER MOBILITY

In this paper, we consider three mobility models for eavesdropper movement. RWM and GMM scenarios represent independent eavesdropper. Meanwhile, RPG scenario represents smart and structured eavesdropper. Thus, the mobility scenarios in this study are sufficient enough to characterize eavesdropper behavior in mobile network. Further sophisticated eavesdropper with active attacks or adaptive eavesdropper mode can be dealt with in the future works.

1) RANDOM WALK MOBILITY MODEL

In RWM model, instead of following the speed and direction of D, the speed and direction of E at time t are randomly

selected from the uniform distribution, i.e., $v_E^{\text{RWM}}(t) \sim \mathcal{U}[V_{\min}, V_{\max}]$ and $\eta_E^{\text{RWM}}(t) \sim \mathcal{U}[0, 2\pi)$, respectively. The location of E at time t with RWM model can be further expressed as

$$x_E^{\text{RWM}}(t) = x_E^{\text{RWM}}(t-1) + v_E^{\text{RWM}}(t) \cos(\eta_E^{\text{RWM}}(t)), \quad (17)$$

$$y_E^{\text{RWM}}(t) = y_E^{\text{RWM}}(t-1) + v_E^{\text{RWM}}(t) \sin(\eta_E^{\text{RWM}}(t)). \quad (18)$$

Fig. 2 shows a snapshot of the node movements with E using RWM. The movement of S and U correlates with D's movement. Meanwhile, E moves according to RWM model that follows random uniform distribution.

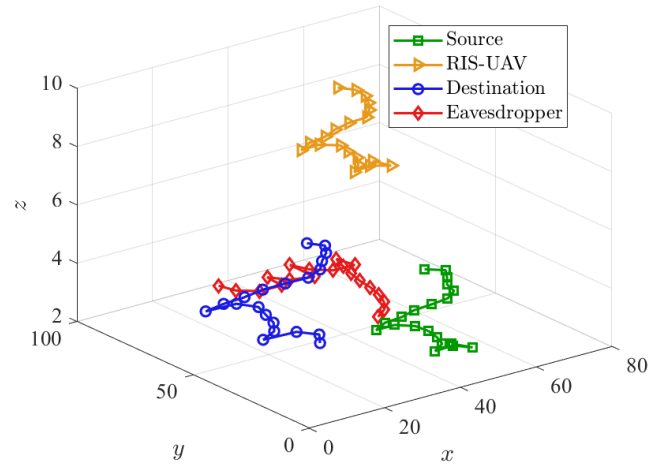


FIGURE 3. The track of 20 different locations for each node under GMM model.

2) GAUSS-MARKOV MOBILITY MODEL

In GMM model, E also does not follow the movement of D. Node E has its own unique speed and direction with some tendency not only continuing the current speed and direction (Markov property), but also randomly following Gaussian distribution. GMM model is known for its smooth movement and thereby more similar with real mobile node movement. The speed and direction of E at time t in GMM model can be expressed as

$$v_E^{\text{GMM}}(t) = \alpha_E v_E^{\text{GMM}}(t-1) + (1 - \alpha_E) \mathbb{E}[v_E^{\text{GMM}}] + \sigma_{v_E^{\text{GMM}}} \sqrt{(1 - \alpha_E^2) m_{v_E^{\text{GMM}}}(t)}, \quad (19)$$

$$\eta_E^{\text{GMM}}(t) = \alpha_E \eta_E^{\text{GMM}}(t-1) + (1 - \alpha_E) \mathbb{E}[\eta_E^{\text{GMM}}] + \sigma_{\eta_E^{\text{GMM}}} \sqrt{(1 - \alpha_E^2) m_{\eta_E^{\text{GMM}}}(t)}, \quad (20)$$

where α_E indicates the memory level of E's movement (Markov's property). $m_{v_E^{\text{GMM}}}(t)$ denotes the Gaussian distribution of E's speed with mean $\mathbb{E}[v_E^{\text{GMM}}]$ and variance $\sigma_{v_E^{\text{GMM}}}^2$. Meanwhile, $m_{\eta_E^{\text{GMM}}}(t)$ denotes the Gaussian distribution of E's direction with mean $\mathbb{E}[\eta_E^{\text{GMM}}]$ and variance $\sigma_{\eta_E^{\text{GMM}}}^2$. The

location of E then can be written as

$$x_E^{\text{GMM}}(t) = x_E^{\text{GMM}}(t-1) + v_E^{\text{GMM}}(t) \cos(\eta_E^{\text{GMM}}(t)), \quad (21)$$

$$y_E^{\text{GMM}}(t) = y_E^{\text{GMM}}(t-1) + v_E^{\text{GMM}}(t) \sin(\eta_E^{\text{GMM}}(t)). \quad (22)$$

Fig. 3 shows a snapshot of the node movements with E using GMM. As can be seen in Fig. 3, the movement of S and U correlate with D 's movement. Meanwhile, E moves according to GMM model that also has smooth movement as the legitimate node group but owns different Gauss-Markov properties with them.

3) REFERENCE POINT GROUP MOBILITY MODEL

In the assumption of E that can acquire the legitimate users mobility information, E can follow the legitimate nodes' mobility model. Therefore, the movement of eavesdropper can be modeled as the RPG mobility model, whose speed and direction at time t can be expressed as

$$v_E^{\text{RPG}}(t) = \min\{\max\{v_g(t) + \Delta v, 0\}, V_{\max}\}, \quad (23)$$

$$\eta_E^{\text{RPG}}(t) = \eta_g(t) + \Delta \eta. \quad (24)$$

Consequently, the location of E with RPG mobility model can be written as

$$x_E^{\text{RPG}}(t) = x_E^{\text{RPG}}(t-1) + v_E^{\text{RPG}}(t) \cos(\eta_E^{\text{RPG}}(t)), \quad (25)$$

$$y_E^{\text{RPG}}(t) = y_E^{\text{RPG}}(t-1) + v_E^{\text{RPG}}(t) \sin(\eta_E^{\text{RPG}}(t)). \quad (26)$$

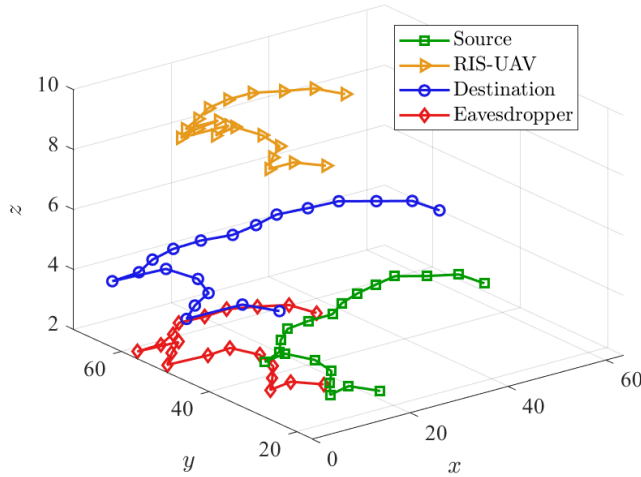


FIGURE 4. The track of 20 different locations for each node under RPG mobility model.

Fig. 4 shows a snapshot of the node movements with RPG mobility model. The movement of S , U , and E correlates with D 's movement. Thus, although all nodes do not have exactly the same speed and direction, they have similar pattern of movement track as shown in Fig. 4.

IV. PHASE SHIFT OPTIMIZATION

Optimizing the configuration of the phase-shifts of RIS is essential to control the propagation environment dynamically and enhance the overall communication performance.

By adjusting the phase-shifts based on the instantaneous CSI, we can steer the reflected signals toward the intended receiver, thereby maximizing the signal strength and quality at the destination. This targeted signal steering helps in creating virtual line-of-sight paths in non-line-of-sight scenarios, significantly improving the link reliability. Additionally, optimizing the phase-shifts is crucial for minimizing the signal leakage toward potential eavesdroppers, thereby enhancing physical layer security. This dual objective of enhancing legitimate communication while suppressing eavesdropping necessitates precise and adaptive phase-shift configurations to respond to the dynamic nature of wireless channels, especially under varying conditions such as node mobility in TNs and NTN. This motivation leads to the need of our proposed algorithms, which dynamically optimize RIS phase-shifts by leveraging instantaneous CSI to enhance signal strength at the destination and minimize leakage to eavesdroppers, thereby ensuring robust and secure communication in dynamic wireless environments.

A. PERFORMANCE METRICS

Evaluating the secrecy performance in a wireless system involves several key metrics. In this study, we calculate the instantaneous secrecy capacity (SC) that is defined as the difference between the main channel rate and the eavesdropper channel rate, and can be formulated as

$$SC = \max\{C_{\text{SUD}} - C_{\text{SUE}}, 0\}. \quad (27)$$

Consequently, the average secrecy rate (ASR) of the system can be mathematically calculated as

$$ASR = \mathbb{E}[SC]. \quad (28)$$

B. PROBLEM STATEMENT

In the proposed system model, we aim to maximize SC by considering quantized phase shift instead of high phase-shift resolution [32] and due to hardware limitation. The phase-shift configuration at U is modeled as linear quantization with Q levels. Hence, the objective function of our maximization problem is

$$\underset{\theta}{\text{maximize}} \quad SC(\theta) = C_{\text{SUD}} - C_{\text{SUE}} \quad (29a)$$

$$\text{subject to} \quad \theta_r \in \left\{0, \frac{2\pi}{Q}, \frac{4\pi}{Q}, \dots, \frac{2\pi(Q-1)}{Q}\right\}, \quad (29b)$$

$$\rho \leq P_S, \quad (29c)$$

where $\theta = [\theta_1, \dots, \theta_R]$ represents the phase-shift vector of all R elements at U .

C. OPTIMAL SECRECY-ORIENTED PHASE-SHIFT

In the condition that U can obtain the CSIs of S , D , and E , U can use these CSIs to obtain the optimal phase shift configuration that maximizes the instantaneous secrecy

capacity of the system, which is mathematically expressed as

$$\begin{aligned} & \text{(OSP problem)} \quad \underset{\theta}{\text{maximize}} \quad C_{\text{SUD}} - C_{\text{SUE}} \\ & \text{subject to} \quad (29b), (29c). \end{aligned} \quad (30)$$

Hence, the instantaneous SNR of the main and eavesdropper channels with optimal secrecy-oriented phase shift (OSP) can be expressed as

$$\gamma_{\text{D}}^{\text{OSP}} = \rho \kappa^2 F_{\text{SU}} F_{\text{UD}} \left| \sum_{r=1}^R \hat{h}_{\text{SU},r} \hat{h}_{\text{U},\text{D}} e^{j(\theta_r^{\text{OSP}} + \phi_{\text{SU},r} + \phi_{\text{U},\text{D}})} \right|^2, \quad (31)$$

$$\gamma_{\text{E}}^{\text{OSP}} = \rho \kappa^2 F_{\text{SU}} F_{\text{UE}} \left| \sum_{r=1}^R \hat{h}_{\text{SU},r} \hat{h}_{\text{U},\text{E}} e^{j(\theta_r^{\text{OSP}} + \phi_{\text{SU},r} + \phi_{\text{U},\text{E}})} \right|^2. \quad (32)$$

Algorithm 1 OSP Exhaustive Search

Input: $R, Q, \kappa, \rho, \mathbf{h}_{\text{SU}}, \mathbf{h}_{\text{UD}},$ and $\mathbf{h}_{\text{UE}}$;

Output: $\theta^{\text{OSP}}, \text{SC}^{\text{OSP}};$

```

1:  $\text{SC}^{\text{OSP}} = 0 \leftarrow$  secrecy capacity initialization;
2: for  $i \in [1, \dots, Q^R]$  do
3:   for  $r \in [1, \dots, R-1]$  do
4:      $\theta_r \leftarrow \frac{2\pi}{Q} \lceil \frac{i}{Q^{R-r}} \rceil;$ 
5:   end for
6:    $\theta_R \leftarrow \frac{2\pi i}{Q};$ 
7:    $\Theta \leftarrow \kappa \text{diag}([e^{j\theta_1}, \dots, e^{j\theta_r}, \dots, e^{j\theta_R}]);$ 
8:    $\gamma_{\text{D}}^{\text{OSP}} \leftarrow \rho |\mathbf{h}_{\text{UD}} \Theta \mathbf{h}_{\text{SU}}'|^2;$ 
9:    $\gamma_{\text{E}}^{\text{OSP}} \leftarrow \rho |\mathbf{h}_{\text{UE}} \Theta \mathbf{h}_{\text{SU}}'|^2;$ 
10:   $C_{\text{SUD}} \leftarrow \log_2(1 + \gamma_{\text{D}}^{\text{OSP}});$ 
11:   $C_{\text{SUE}} \leftarrow \log_2(1 + \gamma_{\text{E}}^{\text{OSP}});$ 
12:   $\text{SC}_{\text{temp}} \leftarrow \max[C_{\text{SUD}} - C_{\text{SUE}}, 0];$ 
13:  if  $\text{SC}_{\text{temp}} > \text{SC}^{\text{OSP}}$  then
14:     $i^{\text{OSP}} \leftarrow i; \text{SC}^{\text{OSP}} \leftarrow \text{SC}_{\text{temp}};$ 
15:  end if
16: end for
17: for  $r \in [1, \dots, R-1]$  do
18:    $\theta_r^{\text{OSP}} \leftarrow \frac{2\pi}{Q} \lceil \frac{i^{\text{OSP}}}{Q^{R-r}} \rceil;$ 
19: end for
20:  $\theta_R^{\text{OSP}} \leftarrow \frac{2\pi i^{\text{OSP}}}{Q};$ 
21: return
```

Let SC^{OSP} denote the instantaneous secrecy capacity for OSP. Meanwhile, channel coefficient vector $\mathbf{h}_{\text{AB}} \in \mathbb{C}^{R \times 1}$ is the combination of large scale and small scale fading for the channel from A to B, where $h_{\text{AB},r} = \sqrt{F_{\text{AB}}} \hat{h}_{\text{AB},r} e^{j\phi_{\text{AB},r}}$. The exhaustive search algorithm for OSP is shown in Algorithm 1. OSP algorithm obtains the optimal phase-shift by calculating the secrecy capacity (SC) from all Q^R possible phase-shift configurations, then searching the maximum SC value among all the resulted Q^R SCs. Therefore, OSP algorithm always obtains the global optimum phase-shift and OSP complexity is $\mathcal{O}(Q^R)$. Because there are Q^R combinations of phase-shift configuration that should be calculated to obtain the optimal

Algorithm 2 Maximizing Real Coefficient

Input: $R, Q, \mathbf{h}_{\text{SU}},$ and $\mathbf{h}_{\text{UD}};$

Output: $\theta^{\text{MRC}};$

```

1: for  $r \in [1, \dots, R]$  do
2:    $\text{val} \leftarrow 0;$ 
3:   for  $q \in [1, \dots, Q]$  do
4:      $\text{temp} \leftarrow \Re \left\{ \hat{h}_{\text{SU},r} \hat{h}_{\text{U},\text{D}} e^{j(\frac{2\pi q}{Q} + \phi_{\text{SU},r} + \phi_{\text{U},\text{D}})} \right\};$ 
5:     if  $\text{val} < \text{temp}$  then
6:        $\text{val} \leftarrow \text{temp};$ 
7:        $\theta_r^{\text{MRC}} \leftarrow \frac{2\pi q}{Q};$ 
8:     end if
9:   end for
10: end for
```

Calculating Instantaneous SC of MRC

Input: $\kappa, \rho, \mathbf{h}_{\text{SU}}, \mathbf{h}_{\text{UD}}, \mathbf{h}_{\text{UE}},$ and $\theta^{\text{MRC}};$

Output: $\text{SC}^{\text{MRC}};$

```

11:  $\Theta \leftarrow \kappa \text{diag}([e^{j\theta_1^{\text{MRC}}}, \dots, e^{j\theta_R^{\text{MRC}}}]);$ 
12:  $\gamma_{\text{D}}^{\text{MRC}} \leftarrow \rho |\mathbf{h}_{\text{UD}} \Theta \mathbf{h}_{\text{SU}}'|^2;$ 
13:  $\gamma_{\text{E}}^{\text{MRC}} \leftarrow \rho |\mathbf{h}_{\text{UE}} \Theta \mathbf{h}_{\text{SU}}'|^2;$ 
14:  $C_{\text{SUD}} \leftarrow \log_2(1 + \gamma_{\text{D}}^{\text{MRC}});$ 
15:  $C_{\text{SUE}} \leftarrow \log_2(1 + \gamma_{\text{E}}^{\text{MRC}});$ 
16:  $\text{SC}^{\text{MRC}} \leftarrow \max[C_{\text{SUD}} - C_{\text{SUE}}, 0];$ 
17: return
```

phase shift in OSP, this exhaustive search makes the OSP not feasible to be implemented within limited time.

D. MAXIMIZING REAL COEFFICIENT

In the condition of U that cannot obtain the CSI of E, U is assumed that it can still acquire the CSIs of S and D to maximize the received SNR at D. Thus, we propose phase-shift optimization based on the maximum real component of the main channel coefficient multiplication, which can be expressed as [33]

$$\begin{aligned} & \text{(MRC problem)} \quad \underset{\theta}{\text{maximize}} \quad \Re \left\{ \hat{h}_{\text{SU},r} \hat{h}_{\text{U},\text{D}} e^{j(\theta_r + \phi_{\text{SU},r} + \phi_{\text{U},\text{D}})} \right\} \\ & \text{subject to} \quad (29b). \end{aligned} \quad (33)$$

where $\Re\{\cdot\}$ denotes the real component of the complex number. Maximizing real coefficient (MRC) is considered as theoretical-based optimization algorithm to obtain the optimal phase-shift configuration in discrete-valued phase-shift RIS. Let θ^{MRC} denotes the phase shift vector for all R elements of RIS with MRC optimization, the resulting instantaneous SNR at D and E then becomes

$$\gamma_{\text{D}}^{\text{MRC}} = \rho \kappa^2 F_{\text{SU}} F_{\text{UD}} \left| \sum_{r=1}^R \hat{h}_{\text{SU},r} \hat{h}_{\text{U},\text{D}} e^{j(\theta_r^{\text{MRC}} + \phi_{\text{SU},r} + \phi_{\text{U},\text{D}})} \right|^2, \quad (34)$$

$$\gamma_{\text{E}}^{\text{MRC}} = \rho \kappa^2 F_{\text{SU}} F_{\text{UE}} \left| \sum_{r=1}^R \hat{h}_{\text{SU},r} \hat{h}_{\text{U},\text{E}} e^{j(\theta_r^{\text{MRC}} + \phi_{\text{SU},r} + \phi_{\text{U},\text{E}})} \right|^2. \quad (35)$$

TABLE 1. Computational Complexity of Phase-Shift Optimization Methods.

Optimization Method	Complexity
Optimal Secrecy-oriented Phase-shift (OSP)	$\mathcal{O}(Q^R)$
Maximizing Real Coefficient (MRC)	$\mathcal{O}(Q \times R)$
Random Phase Shift (RPS)	$\mathcal{O}(R)$

The algorithm of MRC is shown in Algorithm 2. The process of MRC algorithm is first in every RIS element it calculates a composite coefficient that is the product of the channel response from S to U, the reflection coefficient at U, and the channel response from U to D, i.e., $\hat{h}_{SU,r} \kappa e^{j\theta_r} \hat{h}_{U,D} = \hat{h}_{SU,r} e^{j\phi_{SU,r}} \kappa e^{j\theta_r} \hat{h}_{U,D} e^{j\phi_{U,D}} \approx \hat{h}_{SU,r} \hat{h}_{U,D} e^{j(\theta_r + \phi_{SU,r} + \phi_{U,D})}$ when we consider amplitude reflection $\kappa = 1$. Then, MRC searches the maximum real coefficients from all Q possible composite coefficients as the number of possible phase-shifts at every element is Q .

The computational complexity of maximum real searching at every RIS element is $\mathcal{O}(Q)$. Additionally, because the maximum real searching is computed repeatedly R times, then MRC complexity becomes $\mathcal{O}(Q \times R)$. MRC algorithm has linearly increasing complexity, while OSP algorithm has exponentially increasing complexity. As a result, MRC has lower complexity than OSP.

E. RANDOM PHASE SHIFT

The random phase shift (RPS) algorithm is considered as the baseline configuration for comparing the phase-shift optimization performance. The RPS algorithm randomly tunes the phase of the passive reflecting elements at RIS. It means that RPS algorithm does not require any CSI to tune the passive reflecting elements, which can be mathematically expressed as

$$\theta_r^{\text{RPS}} \sim \mathcal{U}(0, 2\pi), \quad (36)$$

where $\mathcal{U}(\cdot)$ represents the uniform distribution of the phase shifts in (29b). Thus, the instantaneous SNR of the main and eavesdropper channels with RPS can be expressed as

$$\gamma_D^{\text{RPS}} = \rho \kappa^2 F_{SU} F_{UD} \left| \sum_{r=1}^R \hat{h}_{SU,r} \hat{h}_{U,D} e^{j(\theta_r^{\text{RPS}} + \phi_{SU,r} + \phi_{U,D})} \right|^2, \quad (37)$$

$$\gamma_E^{\text{RPS}} = \rho \kappa^2 F_{SU} F_{UE} \left| \sum_{r=1}^R \hat{h}_{SU,r} \hat{h}_{U,E} e^{j(\theta_r^{\text{RPS}} + \phi_{SU,r} + \phi_{U,E})} \right|^2. \quad (38)$$

F. PRELIMINARY RESULTS

Computational complexity of the phase-shift optimization methods using OSP, MRC, and RPS is presented in Table 1. By considering the number of RIS elements R and the quantization levels in each element Q , OSP needs to calculate all Q^R combinations to obtain the optimum phase-shift configuration. On the other hand, MRC needs $Q \times R$

calculations since it only calculates the main-channel real coefficient multiplication at every element, whereas RPS only needs R computations due to its random configuration.

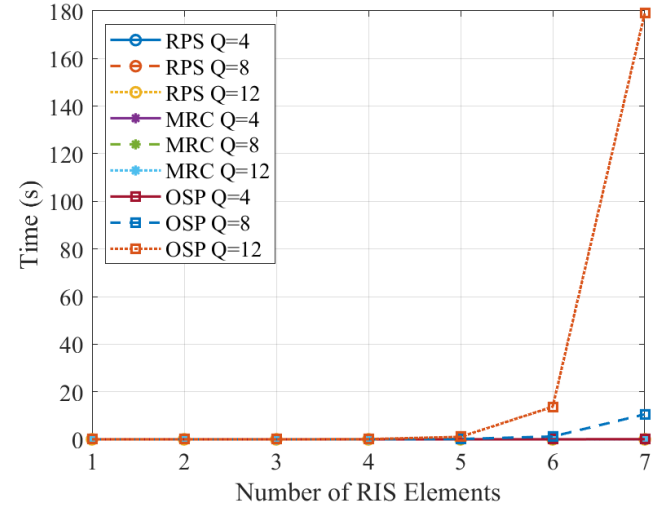
**FIGURE 5. Execution time of phase-shift optimization methods.**

Fig. 5 shows the execution time of the OSP, MRC, and RPS with different number of RIS elements and quantization levels. MRC and RPS have feasible execution time regardless of the number of RIS elements and quantization levels because the computational complexity only increases linearly with the number of RIS elements. However, OSP execution time increases exponentially as the number of RIS elements and quantization levels also increase. The reason is that because OSP needs to calculate all Q^R combinations to obtain the optimum phase-shift configuration. Thus, OSP is not feasible to be applied within limited time using high number of RIS elements.

Considering the high execution time of the phase-shift optimization methods, we propose a DNN model to predict the secrecy performance. DNN arises as a well-suited method for real-time system performance evaluation. DNN model has the shortest execution time compared to the Monte-Carlo simulation and mathematical analysis [34].

V. DEEP TRANSFER LEARNING

Addressing the environment changes due to node mobility, we propose a transfer learning method that has less training time than re-training the DNN model with the new environment data. In transfer learning approach, there are a source model and a target model, which are denoted by subscript S and T, respectively. Analogues to traditional machine learning, source model corresponds to the model in the training process, whereas target model is related to the model used in testing [35].

Fig. 6 illustrates the source and target model in transfer learning method to predict SC. The source model is a DNN model with several hidden layers. On the other hand, target model layers are transferred from the source model layers

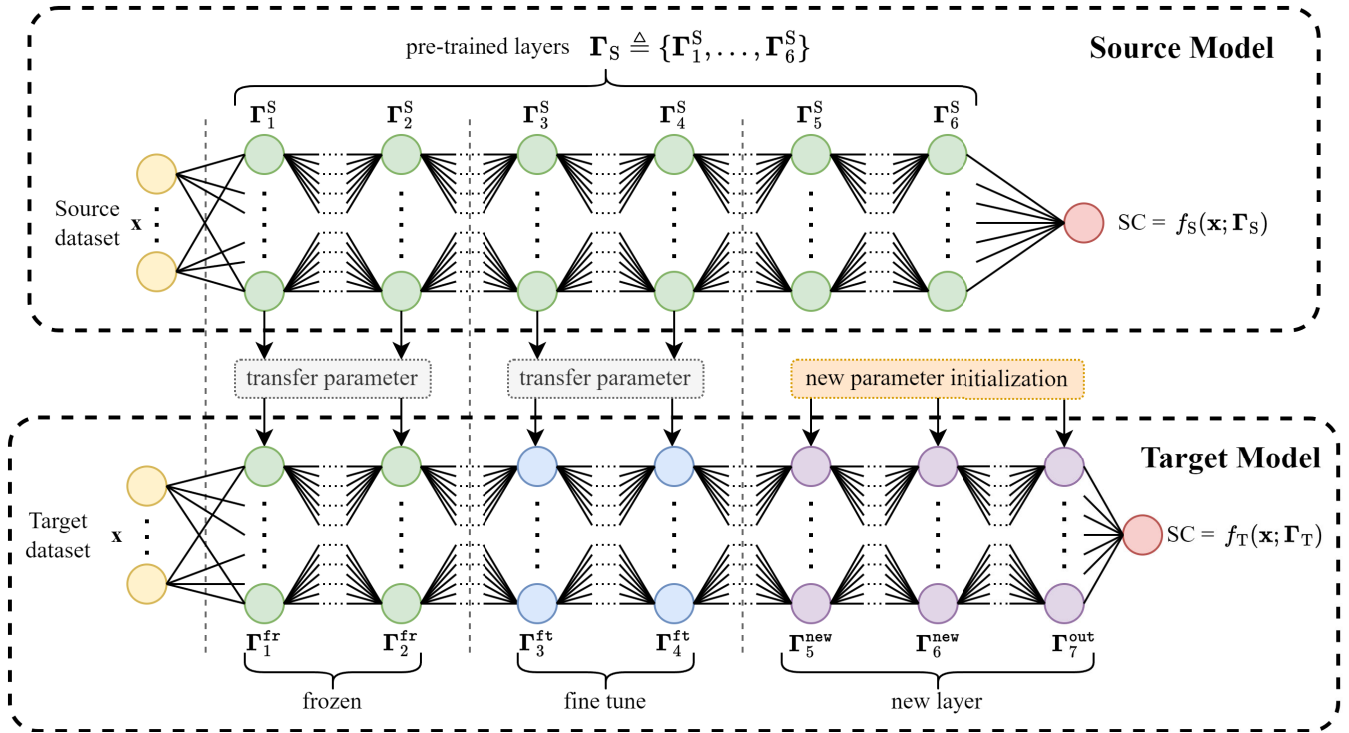


FIGURE 6. Source model and target model with transfer learning-based regression.

with specific arrangement. First, frozen layer and fine-tuned layer parameters are copied from the source model. These layers are indicated by the gray boxes of “transfer parameter” in Fig. 6. The input layer and frozen layers are settled into fixed parameters and do not join the backward propagation in target model training.

Second, new layers are located before the output layer. The new layers are small number of new layers where their parameter initialization are given randomly, no relation with the final trained parameter values in source model. These layers are denoted by the “new parameter initialization” box in Fig. 6. Adding new layers is intuitive method to transfer knowledge in deep transfer learning [36]. Because there are only small number of target dataset, we transfer or copy first several layers from source model into the first several layers of target model and add only small number of new layers in target model. The transferred layers correspond to general characteristics of features that can be applicable to source and target datasets [37]. However, because target dataset is different from source dataset, we need to add several new layers located directly after the transferred layers to learn the specific characteristics of features in target dataset.

Lastly, fine-tuned layers along with the new layer are trained with the new environment data. Fine-tuning aims to train just few layers of the source model by using few labeled samples in the target environment. Fine-tuning can dramatically reduce training costs because only parameters from few layers will be trained [38]. The proposed transfer learning approach in this paper is considered using a

fine-tuned model to learn new environment data efficiently. In this case, the optimal parameters in the trained source model are used and transferred into the target model to enhance the new training although the target model is only trained with insufficient amount of data [39].

A. DOMAIN AND TASK

Both of the source model and target model in transfer learning method have their own domain and task. A domain is defined by $\mathcal{D} = \{\mathcal{X}, p(\mathbf{x})\}$, where $\mathcal{X}$ is a feature space, $p(\mathbf{x})$ is a marginal probability distribution, and $\mathbf{x} \in \mathcal{X}$. For a given domain $\mathcal{D}$, a task is defined by $\mathcal{T} = \{\mathcal{Y}, f(\cdot)\}$, where $\mathcal{Y}$ is a label space and $f(\cdot)$ is a predictive function $f: \mathcal{X} \rightarrow \mathcal{Y}$ which is learned from the feature vector and label pairs $(\mathbf{x}, y)$ with $y \in \mathcal{Y}$ [40]. From a probabilistic perspective, we can also rewrite the task as $\mathcal{T} = \{\mathcal{Y}, p(y|\mathbf{x})\}$, where $p(y|\mathbf{x})$ denotes the posterior probability of label y for a given feature vector $\mathbf{x}$.

Assuming a source domain $\mathcal{D}_S$ with a corresponding source task $\mathcal{T}_S$ and a target domain $\mathcal{D}_T$ with a corresponding task $\mathcal{T}_T$, transfer learning goal is to improve the target predictive function $f_T(\cdot)$ by using the related information from $\mathcal{D}_S$ and $\mathcal{T}_S$, where $\mathcal{D}_S \neq \mathcal{D}_T$ and/or $\mathcal{T}_S \neq \mathcal{T}_T$. In addition, in most cases, the size of $\mathcal{D}_S$ is much larger than the size of $\mathcal{D}_T$ [41]. In nonstationary wireless networks, it is difficult to obtain a large number of new training samples within a short time in the new target environment, then transfer learning should be applied and fine-tuned with small data set in new target domain [42].

In the special case of $\mathcal{D}_S \neq \mathcal{D}_T$ and $\mathcal{T}_S = \mathcal{T}_T$ as in this study, this case is categorized as transductive transfer learning [43]. In this paper, we consider a source model with one source domain $\mathcal{D}_S$ and one source task $\mathcal{T}_S$, and a target model with one target domain $\mathcal{D}_T$ and one target task $\mathcal{T}_T$. We define $\mathcal{D}_S$, $\mathcal{T}_S$, $\mathcal{D}_T$, and $\mathcal{T}_T$ as

$$\begin{aligned}\mathcal{D}_S &= \{\mathcal{X}_S, p(\mathbf{x}_S)\}, \mathbf{x}_S \in \mathcal{X}_S, \\ \mathcal{T}_S &= \{\mathcal{Y}_S, f_S\}, f_S : \mathcal{X}_S \rightarrow \mathcal{Y}_S, y_S = f_S(\mathbf{x}_S) \in \mathcal{Y}_S, \\ \mathcal{D}_T &= \{\mathcal{X}_T, p(\mathbf{x}_T)\}, \mathbf{x}_T \in \mathcal{X}_T, \\ \mathcal{T}_T &= \{\mathcal{Y}_T, f_T\}, f_T : \mathcal{X}_T \rightarrow \mathcal{Y}_T, y_T = f_T(\mathbf{x}_T) \in \mathcal{Y}_T,\end{aligned}\quad (39)$$

respectively. The source and target model have the same task, which is to predict SC. However, they have different domain since the node mobility causes environment changing. Considering that the data sets in new environment follow different distributions, deep transfer learning has the ability to fine-tune DNN model with new training samples as target domain [44]. On the other hand, the case of traditional machine learning is $\mathcal{D}_S = \mathcal{D}_T$ and $\mathcal{T}_S = \mathcal{T}_T$. Conventional machine learning analyzes the same task under the same domains and assumes there is no difference in source and target environment.

B. SOURCE MODEL

In order to obtain the system secrecy performance in a short time, we build a source model with DNN architecture that is capable of predicting the system performance with supervised learning [45]. The developed deep learning model utilizes regression analysis to estimate the relationship between the dependent variable SC that denotes the instantaneous secrecy capacity in (27) and an input feature vector $\mathbf{x}$ that can be expressed as

$$\begin{aligned}\mathbf{x} &= [\rho, \mathbf{p}_S, \mathbf{p}_U, \mathbf{p}_D, \mathbf{p}_E, d_{SU}, d_{UD}, d_{UE}, \phi_{SU_1}, \dots, \phi_{SU_R}, \\ &\phi_{U_1D}, \dots, \phi_{U_RD}, \phi_{U_1E}, \dots, \phi_{U_RE}, \Re\{h_{SU_1}\}, \Im\{h_{SU_1}\}, \dots, \\ &\Re\{h_{SU_R}\}, \Im\{h_{SU_R}\}, \Re\{h_{U_1D}\}, \Im\{h_{U_1D}\}, \dots, \Re\{h_{U_RD}\}, \\ &\Im\{h_{U_RD}\}, \Re\{h_{U_1E}\}, \Im\{h_{U_1E}\}, \dots, \Re\{h_{U_RE}\}, \Im\{h_{U_RE}\}]^T,\end{aligned}\quad (41)$$

where h_{U_rE} denotes the channel coefficient of eavesdropper CSI from the r^{th} element of RIS to E, $\mathbf{p}_X \triangleq [x_X, y_X, z_X]$ denotes the 3D Cartesian coordinates of node X, and $\Im\{\cdot\}$ denotes the imaginer component of the complex number. It is noted that the total number of input features is $16 + 9R$.

The source model in our proposed transfer learning approach consists of five fully connected (FC) layers and an output layer that can be written as

$$\begin{aligned}f_S(\mathbf{x}; \mathbf{\Gamma}_S) &= \left[f_6(\mathbf{\Gamma}_6^S) \circ f_5(\mathbf{\Gamma}_5^S) \circ f_4(\mathbf{\Gamma}_4^S) \circ f_3(\mathbf{\Gamma}_3^S) \right. \\ &\quad \left. \circ f_2(\mathbf{\Gamma}_2^S) \circ f_1(\mathbf{\Gamma}_1^S) \right](\mathbf{x}),\end{aligned}\quad (42)$$

where $(f_2(\mathbf{\Gamma}_2^S) \circ f_1(\mathbf{\Gamma}_1^S))(\mathbf{x}) \triangleq f_2(f_1(\mathbf{x}; \mathbf{\Gamma}_1^S); \mathbf{\Gamma}_2^S)$ and $\mathbf{\Gamma}_S \triangleq \{\mathbf{\Gamma}_1^S, \dots, \mathbf{\Gamma}_6^S\}$ denotes the parameters (weights and biases of all layers) in the source model that should be optimized during deep learning process.

Let $\mathbf{o}_l \in \mathbb{R}^{1 \times a_l}$ be the output vector of the l^{th} FC layer with a_l neurons, the output vector of every layer can be formulated as [46]

$$\mathbf{o}_l = \varphi_l(\mathbf{o}_{l-1} \mathbf{W}_l + \mathbf{b}_l), \quad (43)$$

where $\varphi_l(\cdot)$ is the activation function used at the l^{th} layer, $\mathbf{W}_l \in \mathbb{R}^{a_{l-1} \times a_l}$ is the weight matrix from the $(l-1)^{\text{th}}$ FC layer to the l^{th} FC layer, and $\mathbf{b}_l \in \mathbb{R}^{1 \times a_l}$ is the bias vector for all of the a_l neurons at the l^{th} layer. In this paper, we consider all FC layers use rectified linear unit (ReLU) $\varphi(o) = \max\{0, o\}$ as the activation function. The parameters of the l^{th} FC layer then can be denoted as $\mathbf{\Gamma}_l = \{\mathbf{W}_l, \mathbf{b}_l\}$.

Using the assumption of the same number of neurons in all five FC layers ($a_l = a \forall 1 \leq l \leq L$), the number of trainable parameters that are updated in the first FC layer is $(16 + 9R)a + a$. Furthermore, the number of trainable parameters in the rest of FC layer is $(L - 1)(a^2 + a)$. When there is only one predicting neuron in the output layer representing the SC value, then the number of trainable parameters in the output layer is $a + 1$. Thus, total number of parameters in our proposed source model is $(18 + 9R)a + (L - 1)(a^2 + a) + 1$.

C. TRAINING PHASE

In the source model and target model training phase, dataset is prepared with the input vector in (41) and SC values as the output. During the training step, an optimizer is used to obtain optimal parameter $\mathbf{\Gamma}$ by minimizing the loss function of the predicted secrecy capacity $\widehat{SC}$ from the DNN model and the actual secrecy capacity SC from (27). The optimizer used in this study is adaptive moment estimation (Adam). The loss function that calculates the squared error for each k -th sample in the j -th mini-batch of the i -th epoch can be written as

$$\mathcal{L}^{(i,j,k)}(\mathbf{\Gamma}) = \frac{1}{2} \left(SC^{(i,j,k)} - \widehat{SC}^{(i,j,k)} \right)^2, \quad (44)$$

where j is the index of the mini-batch in the i -th epoch.

Let ζ be the neural network learning rate, that is a predetermined small positive value to scale the parameters update in the direction of the negative loss gradient. The updated parameter in every j -th mini-batch of the i -th epoch can be expressed as

$$\mathbf{\Gamma} = \mathbf{\Gamma} - \nabla \mathcal{L}_{\text{batch}}^{(i,j)}(\mathbf{\Gamma}) = \mathbf{\Gamma} - \frac{\zeta}{S_{\text{batch}}} \sum_{k=1}^{S_{\text{batch}}} \nabla \mathcal{L}^{(i,j,k)}(\mathbf{\Gamma}), \quad (45)$$

where S_{batch} is the size of the mini-batch and $\nabla \mathcal{L}_{\text{batch}}^{(i,j)}(\mathbf{\Gamma})$ is the loss gradient value of the j -th mini-batch in the i -th epoch.

It is noted that $\mathcal{D}_S \neq \mathcal{D}_T$ at mobile TN/NTN due to environment changes from node mobility. In most cases, the size of $\mathcal{D}_S$ is much larger than the size of $\mathcal{D}_T$ [41]. Then, the number of mini-batches in source model training (N_{batch}^S) is much larger than the number of mini-batches in target model training (N_{batch}^T) or we can say that $N_{\text{batch}}^S \gg N_{\text{batch}}^T$ when we consider the same batch size S_{batch} to train both the source and target model.

TABLE 2. Number of Parameters in Fine-tuning.

DNN Model	Number of Non-trainable Parameters	Number of Trainable Parameters
Source Model	0	$(18 + 9R)a + (L - 1)(a^2 + a) + 1$
Target Model	$(17 + 9R)a + (M - 1)(a^2 + a)$	$(L - M)(a^2 + a) + (a + 1)$

Algorithm 3 Transfer Learning Algorithm**Input:** $S_{\text{batch}}, N_{\text{epoch}} = 150, \text{SPLIT_SIZE} = 0.1$;1: **Load:** source dataset;2: **Load:** target dataset;**Phase 1:** Source model training phase**Input:** source training dataset;**Output:** trained parameters Γ_S and trained source model $f_S(\mathbf{x}; \Gamma_S)$;3: **build** source model using (42);4: $i \leftarrow 0$;5: **while** $i < N_{\text{epoch}}$ **do**6: $i \leftarrow i + 1; j \leftarrow 0$;7: **while** $j < N_{\text{batch}}^S$ **do**8: $j \leftarrow j + 1$;9: $\Gamma_S \leftarrow \text{Optimizer}(\Gamma_S, \nabla \mathcal{L}_{\text{batch}}^{(i,j)})$;10: **end while**11: **end while**12: **return**
Phase 2: Target model training phase**Input:** target training dataset, trained parameters Γ_S ;**Output:** trained parameters Γ_T and trained target model $f_T(\mathbf{x}; \Gamma_T)$;13: **transfer** Γ_S into Γ^{fr} and Γ^{ft} using (47) and (48);14: **build** target model using (46);15: $i \leftarrow 0$;16: **while** $i < N_{\text{epoch}}$ **do**17: $i \leftarrow i + 1; j \leftarrow 0$;18: **while** $j < N_{\text{batch}}^T$ **do**19: $j \leftarrow j + 1$;20: $\Gamma_T \leftarrow \text{Optimizer}(\Gamma_T, \nabla \mathcal{L}_{\text{batch}}^{(i,j)})$;21: **end while**22: **end while**23: **return****Phase 3:** Inference phase**Input:** testing dataset;**Output:** predicted SC and RMSE values;24: **compute** $\widehat{SC}^{(k)} = f_T(\mathbf{x}_{\text{test}}; \Gamma_T)$ using the trained target model $f_T(\mathbf{x}; \Gamma_T)$;25: **compute** RMSE using (52);26: **return****D. TARGET MODEL**

The target model in our proposed transfer learning approach also consists of several FC layers and an output layer with trained layers that are obtained from the source model and can be written as

$$f_T(\mathbf{x}; \Gamma_T) = \left[\overbrace{f_7(\Gamma_7^{\text{out}}) \circ f_6(\Gamma_6^{\text{new}}) \circ f_5(\Gamma_5^{\text{new}})}^{\text{new layers}} \right]$$

$$\left[\overbrace{f_4(\Gamma_4^{\text{ft}}) \circ f_3(\Gamma_3^{\text{ft}})}^{\text{fine-tuned layers}} \circ \overbrace{f_2(\Gamma_2^{\text{fr}}) \circ f_1(\Gamma_1^{\text{fr}})}^{\text{non-trainable layers}} \right](\mathbf{x}), \quad (46)$$

where $\Gamma_T \triangleq \{\Gamma_1^{\text{fr}}, \Gamma_2^{\text{fr}}, \Gamma_3^{\text{fr}}, \Gamma_4^{\text{ft}}, \Gamma_5^{\text{new}}, \Gamma_6^{\text{new}}, \Gamma_7^{\text{out}}\}$ denotes the parameters (weights and biases) of all layers in the target model. Γ^{fr} indicates the frozen layer parameters where the initial values are transferred from the source model and not updated in the target model training, which can be written as

$$\Gamma_l^{\text{fr}}(i, j) = \Gamma_l^S(i = N_{\text{epoch}}, j = N_{\text{batch}}^S). \quad (47)$$

On the contrary, Γ^{ft} indicates fine-tuning parameters where the initial values are also transferred from the source model but updated in the target model training based on the loss gradient value that can be calculated as

$$\Gamma_l^{\text{ft}}(i = 1, j = 1) = \Gamma_l^S(i = N_{\text{epoch}}, j = N_{\text{batch}}^S), \quad (48)$$

$$\Gamma_l^{\text{ft}}(i, j) = \Gamma_l^{\text{ft}}(i, j - 1) - \frac{\zeta}{S_{\text{batch}}} \sum_{k=1}^{S_{\text{batch}}} \frac{\partial \mathcal{L}^{(k)}(i, j - 1)}{\partial \Gamma_l^{\text{ft}}(i, j - 1)}. \quad (49)$$

In (46), Γ^{new} and Γ^{out} indicate the new additional layer and the output layer parameters respectively, where the values are randomly initialized and iteratively updated in the target model training, which can be expressed as

$$\Gamma_l^{\text{new}} = \Gamma_l^{\text{new}} - \frac{\zeta}{S_{\text{batch}}} \sum_{k=1}^{S_{\text{batch}}} \frac{\partial \mathcal{L}^{(k)}}{\partial \Gamma_l^{\text{new}}}, \quad (50)$$

$$\Gamma^{\text{out}} = \Gamma^{\text{out}} - \frac{\zeta}{S_{\text{batch}}} \sum_{k=1}^{S_{\text{batch}}} \frac{\partial \mathcal{L}^{(k)}}{\partial \Gamma^{\text{out}}}. \quad (51)$$

These layers are denoted by the “new parameter initialization” box in Fig. 6.

Considering there are M frozen layers in target model with total L layers, parameters in the target model can be divided into non-trainable parameters for the frozen layers (Γ^{fr}) and trainable parameters for the fine-tuning layers (Γ^{ft}), new layers (Γ^{new}), and output layer (Γ^{out}). Using the same assumption in the source model that each FC layer has a neurons, the number of non-trainable parameters in the target model is $(17 + 9R)a + (M - 1)(a^2 + a)$ and the number of the trainable parameters is $(L - M)(a^2 + a) + (a + 1)$. The total number of parameters in target model is the same with the total number of parameters in the source model. The number of parameters in the source model and target model are summarized in Table 2. The target model has less number of trainable parameters than the source model, because non-trainable parameters in target model is

TABLE 3. DNN Model Architecture.

Parameters	Values
Number of RIS elements, R	4
Number of input layer neurons	52
Number of hidden layers in the source model	5
Number of neurons at each hidden layer, a	512
Batch size, S_{batch}	512
Number of epochs, N_{epoch}	150
Number of transferred layers from the source model	4
Number of new hidden layers in the target model	2

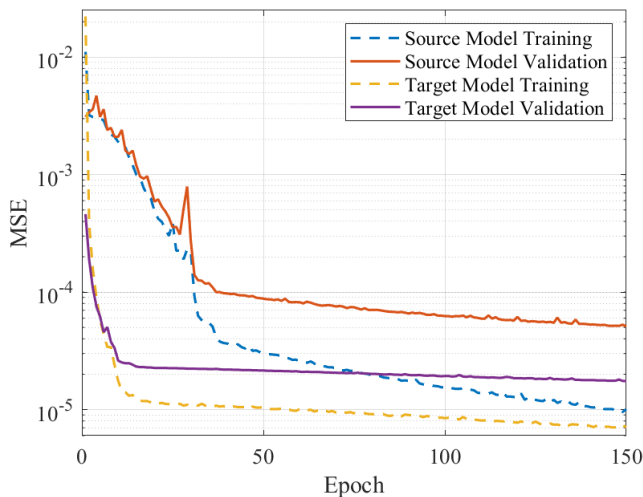
transferred from the source model training. The justification for selecting the number of frozen layers at the target model is explained by Fig. 13 in Section VI that illustrates the impact of the number of non-trainable parameters or number of frozen layers over its achieved ASR value.

E. TESTING PHASE

In the testing step, the trained DNN model predicts the secrecy capacity as the output value. We assess the performance of DNN model in terms of root mean squared error (RMSE) function. The RMSE is used to determine the discrepancy between the predicted value and the actual value over all testing dataset that can be expressed as

$$\text{RMSE} = \sqrt{\frac{\sum_{k=1}^{N_{\text{test}}} (\text{SC}^{(k)} - \widehat{\text{SC}}^{(k)})^2}{N_{\text{test}}}}, \quad (52)$$

where N_{test} denotes the number of data in the testing dataset. In the last inference phase, the trained DNN model is used to predict different dataset and compared with the simulation results. Overall, the algorithm of transfer learning is shown in Algorithm 3. Furthermore, the source model and target model architectures used in this study specifically their number of neurons and layers are summarized in Table 3. Both source and target model have the same batch size and number of epochs.

**FIGURE 7.** The loss MSE in source model and target model.

The mean squared error (MSE) value of the source model and target model with three frozen layers during training and validation is presented in Fig. 7. Target model in training and validation process has lower value of MSE than source model since target model has transfer knowledge from the source model training that makes the target model learning better and produces less error. Target model testing has RMSE value of 0.004004, while source model testing obtains 0.006943. The low value of RMSE means that the DNN model prediction has low discrepancy with the actual instantaneous SC value in simulation. The lower RMSE value in target model is the result of target model training that takes the advantage of parameter transfer and learns feature characteristics better than the source model.

F. DNN COMPLEXITY ANALYSIS

The computational complexity of the DNN model can be calculated as

$$\begin{aligned} & \mathcal{O}\left(\underbrace{(16 + 9R)a + a}_{\text{first FC layer}} + \underbrace{(L - 1)(a^2 + a)}_{\text{remaining FC layer}} + \underbrace{a + 1}_{\text{output layer}}\right) \\ &= \mathcal{O}((18 + 9R)a + (L - 1)(a^2 + a) + 1) \\ &\approx \mathcal{O}(aR + (L - 1)a^2) \\ &\approx \mathcal{O}(a^2) \end{aligned} \quad (53)$$

since $R \ll a$ and $L \ll a$. The resulted computational complexity shows that the DNN model calculates its prediction without considering the number of quantization levels Q for every RIS elements.

For source model training, when there are $S_{\text{batch}} \times N_{\text{batch}}^S$ samples to train the DNN model with number of epochs N_{epoch} , the training complexity of source DNN model from (53) can be recalculated as

$$\mathcal{O}(N_{\text{epoch}} \times S_{\text{batch}} \times N_{\text{batch}}^S \times a^2) \approx \mathcal{O}(a^4) \quad (54)$$

regarding of $N_{\text{epoch}} \ll a$, $S_{\text{batch}} \approx a$, and $N_{\text{batch}}^S \approx a$. For target model training with the same batch size and number of epochs, its complexity is given by

$$\mathcal{O}(N_{\text{epoch}} \times S_{\text{batch}} \times N_{\text{batch}}^T \times a^2) \approx \mathcal{O}(a^3) \quad (55)$$

because $N_{\text{batch}}^T \ll a$. The training complexity neglects the number of frozen layers (M) since $M < L \ll a$.

While the training complexity is $\mathcal{O}(a^4)$ for source model and $\mathcal{O}(a^3)$ for target model, the inference complexity is $\mathcal{O}(a^2)$ from (53) because inference phase computes only one sample to obtain the prediction and it becomes the main computational complexity in DNN approach with the assumption that training phase can be computed offline [47], [48]. Hence, DNN model with complexity $\mathcal{O}(a^2)$ is considered as having low complexity that neither grows with the increasing number of RIS elements nor quantization levels. Therefore, the proposed transfer learning-empowered PHY security is practical to optimize the secrecy performance in aerial RIS-aided network.

TABLE 4. Simulation Parameters.

Parameters	Values
Carrier frequency, f_c (GHz)	3
Noise figure, NF (dBm)	10
Thermal noise power density (dBm/Hz)	-174
Noise bandwidth, BW (MHz)	10
Reference distance, d_0 (m)	1
Gauss-Markov memory level of legitimate nodes group, α	0.6
Maximum speed of every node, $V_{\max}$	3
Mean of group speed, $\mathbb{E}[v_g]$	2
Standard deviation of group speed, σ_{v_g}	0.1
Mean of group direction, $\mathbb{E}[\eta_g]$	π
Standard deviation of group direction, σ_{η_g}	$\pi/2$
Gauss-Markov memory level at E, α_E	0.6
Mean of E's speed, $\mathbb{E}[v_E^{\text{GMM}}]$	4
Standard deviation of E's speed, $\sigma_{v_E^{\text{GMM}}}$	0.5
Mean of E's direction, $\mathbb{E}[\eta_E^{\text{GMM}}]$	$\pi/2$
Standard deviation of E's direction, $\sigma_{\eta_E^{\text{GMM}}}$	$\pi/4$
Number of RIS elements, R	4
Amplitude reflection coefficient, κ	1
Quantization levels, Q	16
Transmit Power at S, P_S (dBm)	[-40,80]
Width of the Gaussian kernel, γ_G	1

G. MAXIMUM MEAN DISCREPANCY

Since $\mathcal{D}_S \neq \mathcal{D}_T$ in transductive transfer learning, consequently how to measure the distribution difference or the similarity between $\mathcal{D}_S$ and $\mathcal{D}_T$ is an important issue. The measurement metric termed maximum mean discrepancy (MMD) is widely applied in the field of transfer learning [49]. MMD is a nonparametric distribution discrepancy measure used to compare the distributions of two different domains.

In order to calculate MMD, there are two samples as $\{\mathbf{x}_S^i\}_{i=1}^{n_S}$ from $\mathcal{D}_S$ and $\{\mathbf{x}_T^j\}_{j=1}^{n_T}$ from $\mathcal{D}_T$ in (39) and (40), respectively. The squared value of MMD between the two distributions $p(\mathbf{x}_S)$ and $p(\mathbf{x}_T)$ is formulated as [50]

$$\begin{aligned} \text{MMD}^2(\mathcal{D}_S, \mathcal{D}_T) &= \frac{1}{n_S^2} \sum_{i=1}^{n_S} \sum_{i'=1}^{n_S} k(\mathbf{x}_S^i, \mathbf{x}_S^{i'}) \\ &+ \frac{1}{n_T^2} \sum_{j=1}^{n_T} \sum_{j'=1}^{n_T} k(\mathbf{x}_T^j, \mathbf{x}_T^{j'}) \\ &- \frac{2}{n_S n_T} \sum_{i=1}^{n_S} \sum_{j=1}^{n_T} k(\mathbf{x}_S^i, \mathbf{x}_T^j), \end{aligned} \quad (56)$$

where $k(\mathbf{x}, \mathbf{x}')$ is a Gaussian kernel function given by

$$k(\mathbf{x}, \mathbf{x}') = \exp\left(-\frac{\|\mathbf{x} - \mathbf{x}'\|^2}{2\sigma_G^2}\right). \quad (57)$$

$\|\mathbf{x} - \mathbf{x}'\|^2$ is the Euclidean distance between two vectors $\mathbf{x}$ and $\mathbf{x}'$ that can be written as

$$\|\mathbf{x} - \mathbf{x}'\|^2 = \sum_{i=1}^{n_F} (x_i - x'_i)^2, \quad (58)$$

where $\mathbf{x}$ is the feature vector and n_F is the number of feature in (41), which is $16 + 9R$. Meanwhile, x_i and x'_i are the i -th

feature of the vector $\mathbf{x}$ and $\mathbf{x}'$, respectively. σ_G^2 is the variance of the Gaussian kernel function with $\gamma_G = \frac{1}{2\sigma_G^2}$ that controls the width of the Gaussian kernel.

In (56), the squared MMD value contains three parts of kernel computations $k(\mathbf{x}, \mathbf{x}')$, which are the pairwise kernel calculations within the source domain samples $k(\mathbf{x}_S^i, \mathbf{x}_S^{i'})$, the pairwise kernel calculations within the target domain samples $k(\mathbf{x}_T^j, \mathbf{x}_T^{j'})$, and the cross-domain kernel calculations between the source and target samples $k(\mathbf{x}_S^i, \mathbf{x}_T^j)$. Therefore, the more similar $\mathcal{D}_S$ and $\mathcal{D}_T$ results in the higher value of the cross-domain computation in the third part calculation of (56) and yields lower value of the squared MMD as the final calculation of the similarity metric between source and target domain.

VI. RESULTS AND DISCUSSION

In this section, we present our simulation results of the system secrecy performance.¹ The parameter setting that we used in the simulation is presented in Table 4. Considering NLOS condition, the path-loss calculated in this research is integrated in the channel gain of the Rayleigh fading as $G_A + G_B - 22.7 - 26 \log(f_c) - 36.7 \log(d_{AB}/d_0)$, where G is node antenna gain and d_0 is the reference distance.

A. SIMULATION RESULTS

There are three major approaches in RIS optimization research field, which are model-based methods, meta-heuristic techniques, and ML algorithms [51]. In this research, OSP and MRC algorithms are our proposed model-based solutions that represent the first category of RIS optimization approaches. Second category of RIS optimization approaches is meta-heuristic technique. However, meta-heuristic algorithms have possibility for early convergence in local optima and sensitive to parameter settings [51]. Therefore, we are not considering this second category of RIS optimization approaches in this study. The third category of RIS optimization approaches is machine learning (ML) technique. Our proposed DNN model with transfer learning approach is the representative of this third category of RIS optimization approaches. By implementing at PHY security field, all of our methods in the paper has acknowledged major optimization approaches in RIS researches.

Fig. 8 displays the instantaneous secrecy capacity of the system with transmit power 45 dBm at 40 different locations. OSP shows the highest secrecy performance among others since OSP obtains the optimal phase-shift configuration in RIS by calculating the maximum secrecy capacity and using the CSIs of main channel and wiretap link, whereas MRC only utilizes the main channel information to obtain the phase-shift configuration without considering the wiretap channel. Thus, the secrecy capacity of MRC is lower than OSP, yet higher than RPS. Similar results are also obtained for the scenario of E with RWM and GMM mobility model.

¹For the sake of reproducible results, our code can be found at <https://github.com/YTriwidayastuti/RIS-PLS-TL>

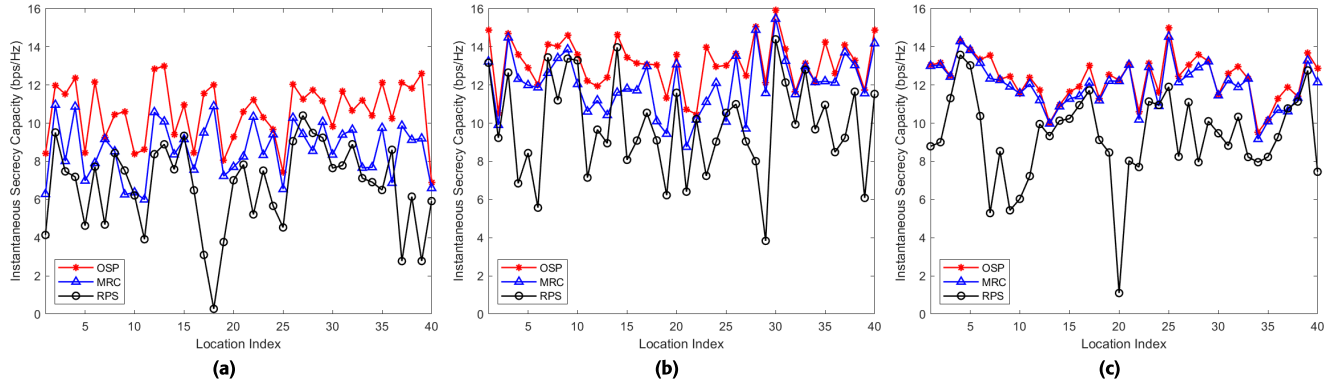


FIGURE 8. Instantaneous secrecy capacity of the three phase-shift methods for E with different mobility models. (a) RPG. (b) RWM. (c) GMM.

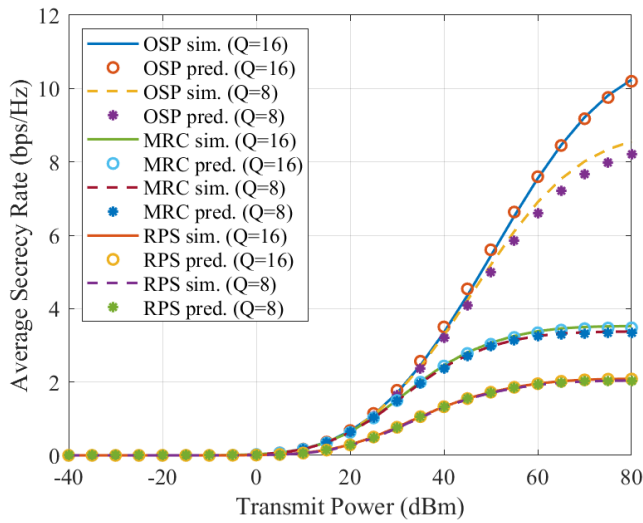


FIGURE 9. The impact of transmit power on the ASR with $R = 4$ and RPG eavesdropper.

Fig. 9 depicts the effect of transmit power and number of quantization levels on the ASR. As can be seen in Fig. 9, the ASR increases as the transmit power increases. The reason is that high transmit power increases the main channel and wiretap received signal rate, but gives more impact on the higher main channel data rate because the signal transmission from RIS elements more focuses at D rather than at E. Fig. 9 also shows that ASR increases as the number of RIS quantization levels increases because higher quantization levels at RIS elements makes signal reflections more focus on D. The prediction results from the DNN model show good correlation with the simulation results. Deep learning can achieve similar performance with OSP, MRC, and RPS as the DNN model obtains the relationship between the secrecy performance and the input parameters through deep learning from all dataset.

Let us assume that the total circuit power in the system is defined by [7]

$$P_0 \triangleq P_S^{\text{dyn}} + P_S^{\text{sta}} + P_D^{\text{sta}} + RP_{U_b}, \quad (59)$$

where $P_S^{\text{dyn}} = 10$ dBm represents the dynamic power consumption, related to the power radiation across all circuit blocks at S. $P_S^{\text{sta}} = 15$ dBm denotes the static power used by the cooling system at S. $P_D^{\text{sta}} = 5$ dBm is the hardware static power dissipation at D. P_{U_b} is the power dissipation per element at U, which is caused by the circuitry required for adaptive phase-shifting with b -bit resolution, where $b = \log_2 Q$. Typical power consumption values of each phase shifter are 1.5 and 4.5 mW for 3- and 4-bit resolution phase shifting, respectively. The system secrecy energy efficiency then can be mathematically expressed as

$$\text{SEE} = \frac{\text{BW} \times \text{SC}}{\frac{1}{\epsilon} P_S + P_0}, \quad (60)$$

where $\epsilon = 0.5$ is the transmit power efficiency [52].

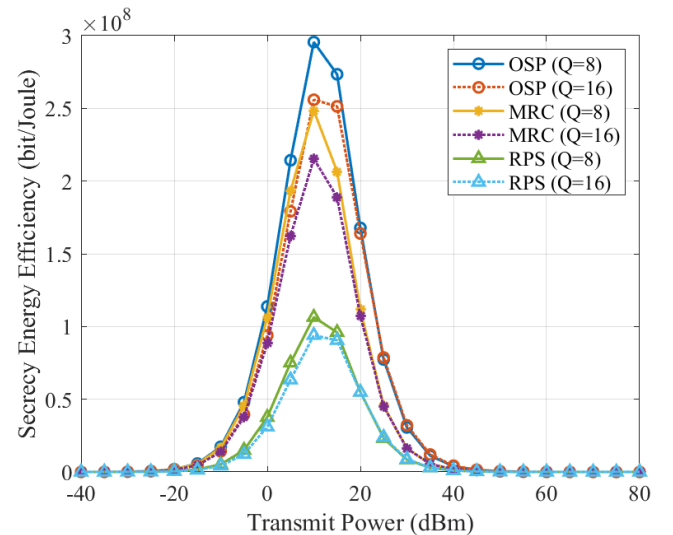


FIGURE 10. Average secrecy energy efficiency for eavesdropper with RPG mobility model.

The result of the secrecy energy efficiency (SEE) for RPG eavesdropper is presented in Fig. 10. At low transmit power, SEE is proportional to the ASR. However, SEE is decreasing at high transmit power due to the increasing

power consumption. As shown in Fig. 10, higher number of quantization levels has lower SEE because it has higher power dissipation per element, even though it can achieve higher SC.

B. TRANSFER LEARNING ANALYSIS

The proposed transfer learning-empowered PHY security is general enough for various scenarios. Deep transfer learning can handle model mismatch [53], nonstationary environment [42], and diverse distribution in device data sets [44], while still addressing PHY security performance in short time. Thus, the proposed transfer learning-empowered PHY security is feasible for practical applications in aerial RIS-aided TN/NTN.

The instantaneous secrecy capacity prediction of the source and target model with OSP phase-shift configuration in the scenario of RPG, RWM, and GMM eavesdropper mobility models are shown in Fig. 11. The source and target model can closely predict the instantaneous secrecy capacity in the shown 100 different locations with various mobility models because the source and target model can learn the features of the node positions and channel coefficients during their training process. The averages of source model and target model prediction also have similar pattern with the simulation average. The similar pattern means that source and target model successfully captures the feature characteristics from the input data in their training.

Fig. 12 presents the ASR result of the source model prediction and target model prediction with the transfer learning technique for E with RWM. As can be seen in Fig. 12, target model prediction has more similar ASR result than the source model prediction. Target model training takes the benefit of parameter transfer that has been trained in the source model. Thus, target model learns better than the source model, even though there are only small number of dataset in the target domain. Similar results are also obtained for different phase-shift configurations, where the source model and target model prediction for OSP has the highest ASR.

In addition to the investigation of transfer learning secrecy performance that is impacted by eavesdropper mobility scenarios, we also explore the effectiveness of our DNN model architecture that is determined by how well it is trained using data. It is not easy to train a model with a large amount of data. Minor adjustments in parameter values can have a significant impact on the performance of the DNN model [54].

Due to insufficient number of target dataset, we freeze M transferred layers in target model to avoid over-fitting. It is well documented that the lower layers of a deep learning model extract generic features that are common across multiple tasks, and the upper layers extract task-specific features [36]. If the good and general enough transferred layers from source model are fine-tuned by small new dataset, they can results in over-fitting [37]. We fine-tune only a small number of last layers because we only need to learn

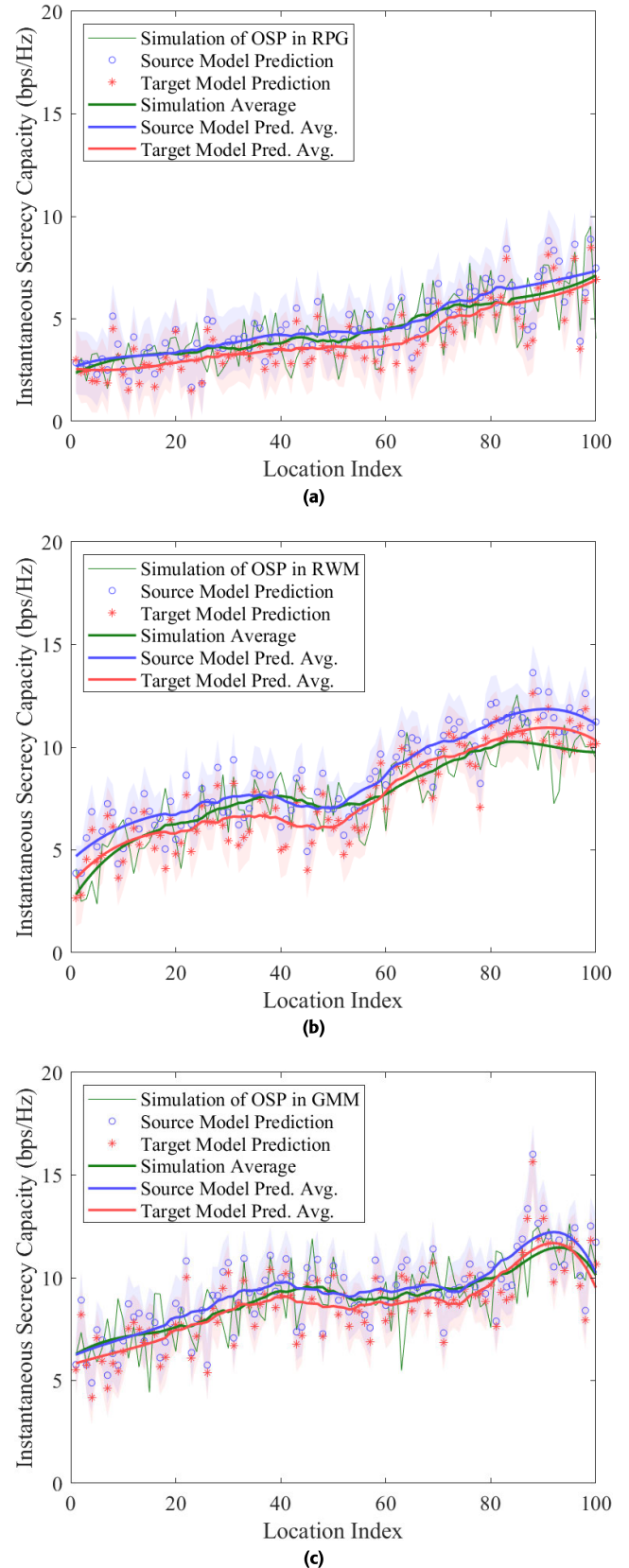


FIGURE 11. Instantaneous secrecy capacity of OSP and transfer learning prediction with $P_S = 40$ dBm, $R = 4$ and $Q = 16$ on different mobility models. (a) RPG. (b) RWM. (c) GMM.

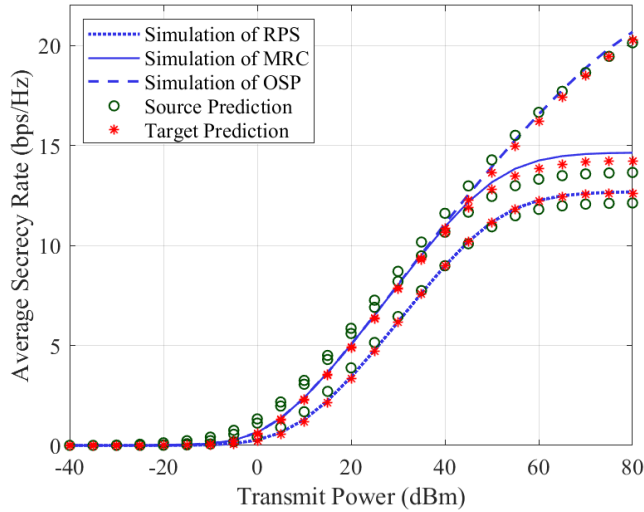


FIGURE 12. The ASR result of the source and target model with $R = 4$, $Q = 16$ and RWM eavesdropper.

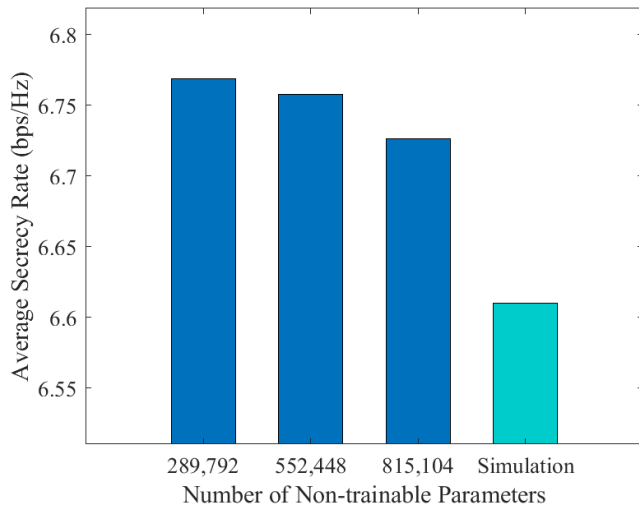
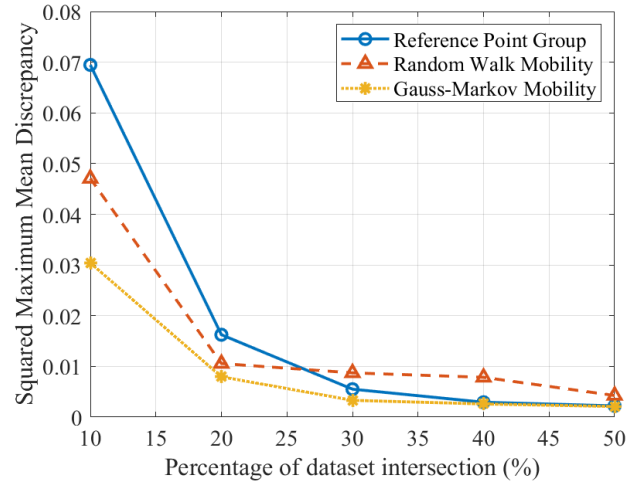


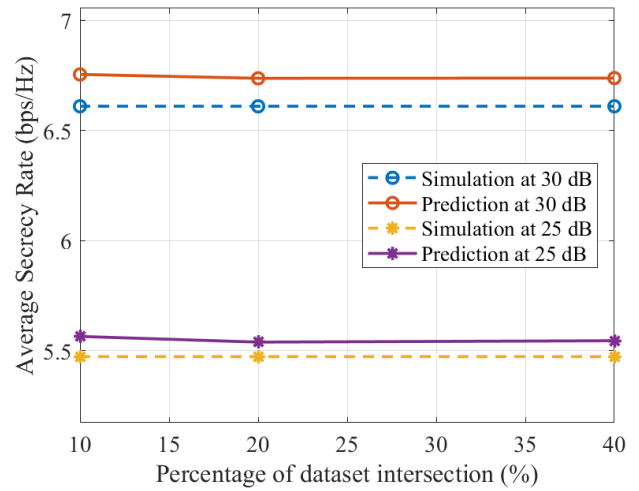
FIGURE 13. ASR of the target model prediction with transmit power 30 dBm as the function of the number of non-trainable parameters.

the specific characteristics of features from the small target dataset.

Fig. 13 shows the ASR result from the target model prediction as a function of the number of non-trainable parameters which are the number of parameters in the frozen layers. The source model has a total of 1,078,273 trainable parameters with 815,104 parameters that are transferred into the target model. As the more number of non-trainable parameters in the target model, the better target model prediction is. Since the source model parameters are trained with large number of dataset, the trained parameters in the source model captures more accurate features of the source domain. When there are only small number of dataset in the target model training, the trained parameters from the source model already represent the general features of the target domain well enough and help the small number of trainable parameters in the target model to learn sufficiently only the specific features of the target domain.



(a)



(b)

FIGURE 14. Impact of intersection percentage between source and target domain on (a) squared MMD and (b) ASR.

It is also difficult to learn the data pattern of nodes because of the mobile and dynamic natures of the traffic obtained from the network. Moreover, the employed parameters cannot be defined as a constant value owing to the limitless quantity of mobility [54]. Therefore, in our proposed transfer learning approach we also conduct a process for assessing and understanding the dataset distribution in mobile TN/NTN.

Fig. 14(a) depicts the impact of intersection percentage between source and target domain on the squared MMD value. As the intersection percentage between source and target domain increases, the squared MMD decreases. This is the result of the increasing similarity between source and target domain, such that the cross-domain kernel computations between the source and target domain increases. Hence, it reduces the MMD of source and target domain as the final result in (56). The same decreasing trend also happens in different mobility models with their own source and target domains.

Fig. 14(b) displays the ASR result of the target model prediction with variation of intersection percentage between the

TABLE 5. Computational Complexity for Multi-RIS.

Method or Technique	Complexity
Optimal Secrecy-oriented Phase-Shift (OSP)	$\mathcal{O}(Q^{R_{\text{tot}}})$
Maximizing Real Coefficient (MRC)	$\mathcal{O}(Q \times R_{\text{tot}})$
Random Phase Shift (RPS)	$\mathcal{O}(R_{\text{tot}})$
Transfer Learning	$\mathcal{O}(a^2)$

source and target domains. When the percentage of domain intersection increases, target model prediction becomes more similar with the simulation result. Increasing the similarity between the source and target domain enhances the target model training because target model can learn from more number of similar dataset that had been learned in the source model training. This tendency also happens in different value of transmit power as presented in Fig. 14(b).

C. SCALABILITY ANALYSIS

In the scenario of multi-RIS, let N be the number of RISs with R elements for each RIS, the total number of RIS elements that should be optimized is $R_{\text{tot}} = N \times R$. The computational complexity of all methods with Q quantization levels and a neurons in hidden layer then can be presented in Table 5. Transfer learning in multi-RIS has the same computational complexity with single-RIS's, that is $\mathcal{O}(a^2)$ since its complexity does not grow with the number of RIS elements. Thus, transfer learning-empowered PHY security can scale efficiently to larger networks with multiple RIS units.

In addition, to conquer both the computation and communication scalability problem, especially to handle larger-scale deployment, the GPU-accelerated hardware, e.g., scalable GPU server [55] can be considered as the next generation hardware in future mobile TN/NTN and as one of potential way to handle larger networks. Hence, the proposed transfer learning-empowered PHY security can overcome the challenge to be applied in real-time, especially for high-dimensional environment.

D. EXTENSION POSSIBILITY FOR ACTIVE RIS

Active RIS is equipped with phase-shift circuits and reflection-type amplifiers, such that it is not only able to adjust the phase shifts but also amplify the received signal attenuated from the first hop to a normal strength level. Accordingly, active RIS can overcome the product/double path loss attenuation in source-RIS and RIS-destination paths [56].

The implementation of active RIS introduces additional noise power from its components, the received signal at D that is transmitted from S and reflected by U then can be written as

$$s_D^{\text{act}} = \sqrt{P_S} \sum_{r=1}^R \sqrt{F_{UD}} \tilde{h}_{U,D} \kappa_{\text{act}} e^{j\theta_r} \sqrt{F_{SU}} \tilde{h}_{SU} u_S + \sum_{r=1}^R \sqrt{F_{UD}} \tilde{h}_{U,D} \kappa_{\text{act}} e^{j\theta_r} n_{U_r} + n_D, \quad (61)$$

where κ_{act} denotes the amplification gain for RIS active element. Meanwhile, n_{U_r} is the thermal noise introduced by active RIS components at the r -th element and modeled as additive white Gaussian noise (AWGN) with zero mean and variance $\sigma_{U_r}^2$.

The received signal-to-noise ratio (SNR) at D with active RIS can be formulated as

$$\gamma_D^{\text{act}} = \frac{P_S F_{SU} F_{UD} (\kappa_{\text{act}})^2 \left| \sum_{r=1}^R \hat{h}_{SU_r} \hat{h}_{U,D} e^{j(\theta_r + \phi_{SU_r} + \phi_{U,D})} \right|^2}{F_{UD} (\kappa_{\text{act}})^2 \sigma_U^2 \sum_{r=1}^R \left| \hat{h}_{U,D} e^{j(\theta_r + \phi_{U,D})} \right|^2 + \sigma_D^2}, \quad (62)$$

when active noise variance for all R elements are the same ($\sigma_{U_r}^2 = \sigma_U^2 \forall 1 \leq r \leq R$). Additionally, the achievable capacity of the received signal at D is given by

$$C_{\text{SUD}}^{\text{act}} = \log_2 (1 + \gamma_D^{\text{act}}). \quad (63)$$

In the wiretap link, the received signal at E can be formulated as

$$s_E^{\text{act}} = \sqrt{P_S} \sum_{r=1}^R \sqrt{F_{UE}} \tilde{h}_{U,E} \kappa_{\text{act}} e^{j\theta_r} \sqrt{F_{SU}} \tilde{h}_{SU} u_S + \sum_{r=1}^R \sqrt{F_{UE}} \tilde{h}_{U,E} \kappa_{\text{act}} e^{j\theta_r} n_{U_r} + n_E. \quad (64)$$

The resulting SNR at E then can be calculated as

$$\gamma_E^{\text{act}} = \frac{P_S F_{SU} F_{UE} (\kappa_{\text{act}})^2 \left| \sum_{r=1}^R \hat{h}_{SU_r} \hat{h}_{U,E} e^{j(\theta_r + \phi_{SU_r} + \phi_{U,E})} \right|^2}{F_{UE} (\kappa_{\text{act}})^2 \sigma_U^2 \sum_{r=1}^R \left| \hat{h}_{U,E} e^{j(\theta_r + \phi_{U,E})} \right|^2 + \sigma_E^2}. \quad (65)$$

Thus, the achievable capacity of the overheard signal at E is given by

$$C_{\text{SUE}}^{\text{act}} = \log_2 (1 + \gamma_E^{\text{act}}). \quad (66)$$

The OSP problem in active RIS then can be mathematically expressed as

$$\underset{\theta}{\text{maximize}} \quad \text{SC}^{\text{act}}(\theta) = C_{\text{SUD}}^{\text{act}} - C_{\text{SUE}}^{\text{act}}, \quad (67a)$$

$$\text{subject to} \quad \theta_r \in \left\{ 0, \frac{2\pi}{Q}, \frac{4\pi}{Q}, \dots, \frac{2\pi(Q-1)}{Q} \right\}, \quad (67b)$$

$$\rho \leq P_S. \quad (67c)$$

For simplicity, we assume $\sigma^2 = \sigma_U^2 = \sigma_D^2 = \sigma_E^2$ and define $\rho = P_S/\sigma^2$ as the average SNR. The amplification gain of the active RIS is given by [57]

$$(\kappa_{\text{act}})^2 = \frac{P_S}{P_S F_{SU} \sum_{r=1}^R \left| \tilde{h}_{SU_r} e^{j\theta_{\text{act}}^{\text{OSP}}} \right|^2 + \sigma^2 \sum_{r=1}^R \left| e^{j\theta_{\text{act}}^{\text{OSP}}} \right|^2}. \quad (68)$$

The instantaneous SNR of the main and eavesdropper channels with OSP can be expressed as

$$\gamma_{D_{act}}^{OSP} = \frac{\rho F_{SU} F_{UD} (\kappa_{act})^2 \left| \sum_{r=1}^R \tilde{h}_{SU,r} \tilde{h}_{U,D} e^{j\theta_r^{OSP}} \right|^2}{F_{UD} (\kappa_{act})^2 \sum_{r=1}^R \left| \hat{h}_{U,D} e^{j(\theta_r^{OSP} + \phi_{U,D})} \right|^2 + 1}, \quad (69)$$

$$\gamma_{E_{act}}^{OSP} = \frac{\rho F_{SU} F_{UE} (\kappa_{act})^2 \left| \sum_{r=1}^R \tilde{h}_{SU,r} \tilde{h}_{U,E} e^{j\theta_r^{OSP}} \right|^2}{F_{UE} (\kappa_{act})^2 \sum_{r=1}^R \left| \hat{h}_{U,E} e^{j(\theta_r^{OSP} + \phi_{U,E})} \right|^2 + 1}. \quad (70)$$

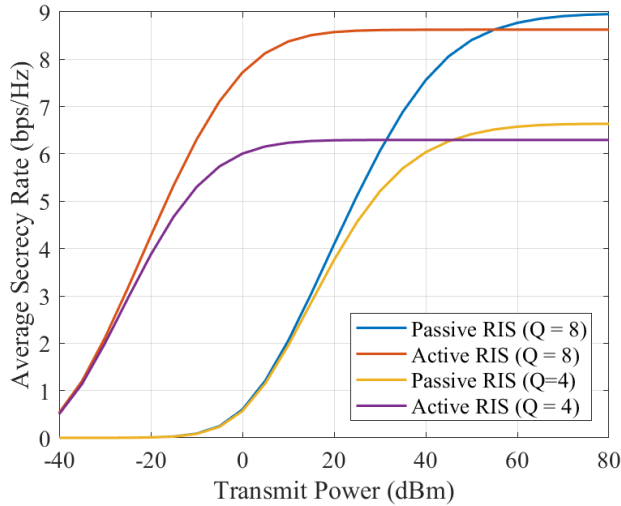


FIGURE 15. ASR performance of active and passive RIS.

The ASR result of the active and passive RIS for OSP optimization method and eavesdropper with RPG mobility model is shown in Fig. 15. Active RIS has higher ASR than passive RIS for low transmit power because element amplification in active RIS can overcome multiplicative or double path loss fading at D. However, active RIS has lower ASR than passive RIS for high source transmit power. The possible reason is that the amplification in active RIS also amplifies the channel gain at eavesdropper's link. Hereafter, active RIS amplification has possibility to make the transmission less secure.

Considering the total power consumption, active RIS consumes more power than passive RIS due to the amplification in every element. The overall power consumption of active and passive RIS-aided systems are respectively given by [56]

$$Q_{act} = P_S + R P_{SW} + R P_{DC} + P_{RIS} \quad (71)$$

$$Q_{pas} = P_S + R P_{SW} \quad (72)$$

where P_{SW} is the power consumed by the phase shift switch and control circuit in each RIS elements, P_{DC} is the direct current biasing power used by the amplifier in each active RIS element, and P_{RIS} is the power of amplified signal reflected by the active RIS.

Regarding to the extra power consumption in (71) and the possibility of signal leaking due to amplification gain, we need to reconsider when applying active RIS in our

system. Since UAVs are powered by batteries, the energy available for flight and communication is very limited. It is important to reasonably allocate the power resource [58], even though the proposed methods can be extended using active RIS. Thus, passive RIS is more suitable for aerial RIS in TN/NTN.

VII. CONCLUSION

In this paper, we focused on enhancing PHY security in aerial RIS-aided TN/NTN, specifically addressing the challenges introduced by node mobility. We proposed a system model integrating RIS into both terrestrial and non-terrestrial networks, investigating the impact of various mobility models—random walk, Gauss-Markov, and reference point group mobility—on key security metrics like secrecy capacity and average secrecy rate. To design an effective transmission protocol, we formulated an optimization problem focused on the phase-shift configurations of the RIS, aiming to maximize the secrecy capacity and average secrecy rate under different mobility scenarios. We developed several robust algorithms: Optimal Secrecy-oriented Phase-shift (OSP), which uses perfect knowledge of the CSI for both legitimate and eavesdropper channels to achieve optimal secrecy; Maximizing Real Coefficient (MRC), which optimizes phase shifts without requiring eavesdropper CSI; and Random Phase Shift (RPS), used as a performance benchmark. We employed both simulation and deep learning methods to validate our results, utilizing DNNs and transfer learning techniques to predict PHY security metrics efficiently. Our key insights include the superior performance of OSP and MRC algorithms, significant effects of different mobility models on PHY security, benefits of transfer learning in enhancing model robustness and reducing computational costs, and the close correlation of our DNN models with simulation results. This study provides valuable insights and practical solutions for ensuring secure and resilient communication in aerial RIS-aided TN/NTN, significantly contributing to the advancement of secure mobile networks by addressing the critical challenges posed by node mobility in next-generation communication systems.

REFERENCES

- [1] G. Geraci, D. López-Pérez, M. Benzaghta, and S. Chatzinotas, "Integrating terrestrial and non-terrestrial networks: 3D opportunities and challenges," *IEEE Commun. Mag.*, vol. 61, no. 4, pp. 42–48, Apr. 2023.
- [2] T. N. Do and G. Kaddoum, "Distributed machine learning for terrestrial and non-terrestrial Internet of Things networks," *IEEE Internet Things Mag.*, vol. 6, no. 4, pp. 54–61, Dec. 2023.
- [3] J. Zheng, J. Zhang, H. Du, D. Niyato, B. Ai, M. Debbah, and K. B. Letaief, "Mobile cell-free massive MIMO: Challenges, solutions, and future directions," *IEEE Wireless Commun.*, vol. 31, no. 3, pp. 140–147, Jun. 2024.
- [4] R. H. Y. Perdana, T.-V. Nguyen, and B. An, "Adaptive user pairing in multi-IRS-aided massive MIMO-NOMA networks: Spectral efficiency maximization and deep learning design," *IEEE Trans. Commun.*, vol. 71, no. 7, pp. 4377–4390, Jul. 2023.
- [5] J. Zhang, H. Du, Q. Sun, B. Ai, and D. W. K. Ng, "Physical layer security enhancement with reconfigurable intelligent surface-aided networks," *IEEE Trans. Inf. Forensics Security*, vol. 16, pp. 3480–3495, 2021.

- [6] L. Bariah, L. Mohjazi, H. Abumarshoud, B. Selim, S. Muhaidat, M. Tatipamula, M. A. Imran, and H. Haas, "RIS-assisted space-air-ground integrated networks: New horizons for flexible access and connectivity," *IEEE Netw.*, vol. 37, no. 3, pp. 118–125, May 2022.
- [7] R. H. Y. Perdana, T.-V. Nguyen, Y. Pramitarini, and B. An, "Deep learning-based energy efficiency maximization in massive MIMO-NOMA networks with multiple RISs," in *Proc. Int. Conf. Artif. Intell. Inf. Commun. (ICAIIIC)*, Osaka, Japan: IEEE, Feb. 2024, pp. 382–387.
- [8] Y. Pramitarini, R. H. Y. Perdana, K. Shim, and B. An, "Opportunistic scheduling scheme to improve physical-layer security in cooperative NOMA system: Performance analysis and deep learning design," *IEEE Access*, vol. 12, pp. 58454–58472, 2024.
- [9] W. Wang, H. Tian, and W. Ni, "Secrecy performance analysis of IRS-aided UAV relay system," *IEEE Wireless Commun. Lett.*, vol. 10, no. 12, pp. 2693–2697, Dec. 2021.
- [10] Z. Zhang, C. Zhang, C. Jiang, F. Jia, J. Ge, and F. Gong, "Improving physical layer security for reconfigurable intelligent surface aided NOMA 6G networks," *IEEE Trans. Veh. Technol.*, vol. 70, no. 5, pp. 4451–4463, May 2021.
- [11] A. K. Yadav, S. Yadav, A. Pandey, and A. Silva, "On the secrecy performance of RIS-enabled wireless communications over Nakagami-m fading channels," *ICT Exp.*, vol. 9, no. 3, pp. 452–458, Jun. 2023.
- [12] C. Pan, G. Zhou, K. Zhi, S. Hong, T. Wu, Y. Pan, H. Ren, M. D. Renzo, A. L. Swindlehurst, R. Zhang, and A. Y. Zhang, "An overview of signal processing techniques for RIS/IRS-aided wireless systems," *IEEE J. Sel. Topics Signal Process.*, vol. 16, no. 5, pp. 883–917, Aug. 2022.
- [13] P. Ramezani, B. Lyu, and A. Jamalipour, "Toward RIS-enhanced integrated terrestrial/non-terrestrial connectivity in 6G," *IEEE Netw.*, vol. 37, no. 3, pp. 178–185, May 2022.
- [14] K. Feng, X. Li, Y. Han, S. Jin, and Y. Chen, "Physical layer security enhancement exploiting intelligent reflecting surface," *IEEE Commun. Lett.*, vol. 25, no. 3, pp. 734–738, Mar. 2021.
- [15] M. H. Khoshafa, T. M. N. Ngatched, and M. H. Ahmed, "Reconfigurable intelligent surface-aided physical layer security enhancement in D2D underlay communications," *IEEE Commun. Lett.*, vol. 25, no. 5, pp. 1443–1447, May 2021.
- [16] J. Li, S. Xu, J. Liu, Y. Cao, and W. Gao, "Reconfigurable intelligent surface enhanced secure aerial-ground communication," *IEEE Trans. Commun.*, vol. 69, no. 9, pp. 6185–6197, Sep. 2021.
- [17] S. Sun, F. Yang, J. Song, and Z. Han, "Optimization on multiuser physical layer security of intelligent reflecting surface-aided VLC," *IEEE Wireless Commun. Lett.*, vol. 11, no. 7, pp. 1344–1348, Jul. 2022.
- [18] X.-T. Dang, H. V. Nguyen, and O.-S. Shin, "Physical layer security for IRS-UAV-assisted cell-free massive MIMO systems," *IEEE Access*, vol. 12, pp. 89520–89537, 2024.
- [19] Y. Zhang, S. Zhao, Y. Shen, X. Jiang, and N. Shiratori, "Enhancing the physical layer security of two-way relay systems with RIS and beamforming," *IEEE Trans. Inf. Forensics Security*, vol. 19, pp. 5696–5711, 2024.
- [20] A. Almohamad, A. M. Tahir, A. Al-Kababji, H. M. Furqan, T. Khattab, M. O. Hasna, and H. Arslan, "Smart and secure wireless communications via reflecting intelligent surfaces: A short survey," *IEEE Open J. Commun. Soc.*, vol. 1, pp. 1442–1456, 2020.
- [21] M. S. J. Solaija, H. Salman, and H. Arslan, "Towards a unified framework for physical layer security in 5G and beyond networks," *IEEE Open J. Veh. Technol.*, vol. 3, pp. 321–343, 2022.
- [22] W. U. Khan, E. Lagunas, Z. Ali, M. A. Javed, M. Ahmed, S. Chatzinotas, B. Ottersten, and P. Popovski, "Opportunities for physical layer security in UAV communication enhanced with intelligent reflective surfaces," *IEEE Wireless Commun.*, vol. 29, no. 6, pp. 22–28, Dec. 2022.
- [23] K. Shafique and M. Alhassoun, "Going beyond a simple RIS: Trends and techniques paving the path of future RIS," *IEEE Open J. Antennas Propag.*, vol. 5, pp. 256–276, 2024.
- [24] S. Khan, S. Durrani, and X. Zhou, "Transfer learning based detection for intelligent reflecting surface aided communications," in *Proc. IEEE 32nd Annu. Int. Symp. Pers. Indoor Mobile Radio Commun. (PIMRC)*, Helsinki, Finland: IEEE, Sep. 2021, pp. 555–560.
- [25] Y. Ge and J. Fan, "Beamforming optimization for intelligent reflecting surface assisted MISO: A deep transfer learning approach," *IEEE Trans. Veh. Technol.*, vol. 70, no. 4, pp. 3902–3907, Apr. 2021.
- [26] S. Sun, T. S. Rappaport, S. Rangan, T. A. Thomas, A. Ghosh, I. Z. Kovacs, I. Rodriguez, O. Koymen, A. Partyka, and J. Jarvelainen, "Propagation path loss models for 5G urban micro- and macro-cellular scenarios," in *Proc. IEEE 83rd Veh. Technol. Conf. (VTC Spring)*, Nanjing, China, May 2016, pp. 1–6.
- [27] Q. Zhu, X. Zhang, Y. Xiao, Y. Gao, X. Lei, and Z. Xiong, "Optimization of intelligent reflecting surface aided wireless networks with user mobility," in *Proc. Int. Symp. Netw., Comput. Commun. (ISNCC)*, Jul. 2022, pp. 1–6.
- [28] T. Camp, J. Boleng, and V. Davies, "A survey of mobility models for ad hoc network research," *Wireless Commun. Mobile Comput.*, vol. 2, no. 5, pp. 483–502, Aug. 2002.
- [29] B. Liang and Z. J. Haas, "Predictive distance-based mobility management for multidimensional pcs networks," *IEEE/ACM Trans. Netw.*, vol. 11, no. 5, pp. 718–732, Oct. 2003.
- [30] X. Hong, M. Gerla, G. Pei, and C. Chiang, "A group mobility model for ad hoc wireless networks," in *Proc. 2nd ACM Int. Conf. (MSWiM)*, New York, NY, USA: ACM, Aug. 1999, pp. 53–60.
- [31] A. H. Sawalmeh, N. S. Othman, H. Shakhatreh, and A. Khreishah, "Wireless coverage for mobile users in dynamic environments using UAV," *IEEE Access*, vol. 7, pp. 126376–126390, 2019.
- [32] T. N. Do, G. Kaddoum, T. L. Nguyen, D. B. da Costa, and Z. J. Haas, "Multi-RIS-aided wireless systems: Statistical characterization and performance analysis," *IEEE Trans. Commun.*, vol. 69, no. 12, pp. 8641–8658, Dec. 2021.
- [33] Y. Triwidyastuti, R. H. Y. Perdana, K. Shim, T. N. Do, and B. An, "Exploiting secrecy performance of RIS-assisted networks: Deep learning-based evaluation," in *Proc. ICGHIT*, Hanoi, Vietnam, Jan. 2024, pp. 183–188.
- [34] K. Shim, T. N. Do, T.-V. Nguyen, D. B. d. Costa, and B. An, "Enhancing PHY-security of FD-enabled NOMA systems using jamming and user selection: Performance analysis and DNN evaluation," *IEEE Internet Things J.*, vol. 8, no. 24, pp. 17476–17494, Dec. 2021.
- [35] A. Tuan Nguyen, T. Tran, Y. Gal, P. H. S. Torr, and A. G. Baydin, "KL guided domain adaptation," 2021, *arXiv:2106.07780*.
- [36] H. Venkateswara, S. Chakraborty, and S. Panchanathan, "Deep-learning systems for domain adaptation in computer vision: Learning transferable feature representations," *IEEE Signal Process. Mag.*, vol. 34, no. 6, pp. 117–129, Nov. 2017.
- [37] J. Yosinski, J. Clune, Y. Bengio, and H. Lipson, "How transferable are features in deep neural networks," in *Proc. 27th NIPS*, vol. 27. Cambridge, MA, USA: MIT Press, Dec. 2014, pp. 3320–3328.
- [38] Z. Chen, M. Wu, A. Chan, X. Li, and Y.-S. Ong, "Survey on AI sustainability: Emerging trends on learning algorithms and research challenges," *IEEE Comput. Intell. Mag.*, vol. 18, no. 2, pp. 60–77, May 2023.
- [39] M. Ghous, T. L. Nguyen, T. N. Do, and G. Kaddoum, "Deep transfer learning-based performance prediction of URLLC in independent and not necessarily identically distributed interference networks," *IEEE Access*, vol. 12, pp. 99071–99093, 2024.
- [40] K. R. Weiss, T. M. Khoshgoftaar, and D. Wang, "A survey of transfer learning," *J. Big Data*, vol. 3, no. 1, pp. 1–40, May 2016.
- [41] C. Tan, F. Sun, T. Kong, W. Zhang, C. Yang, and C. Liu, "A survey on deep transfer learning," in *Proc. ICANN*, vol. 11141. Cham, Switzerland: Springer, Jan. 2018, pp. 270–279.
- [42] C. She, C. Sun, Z. Gu, Y. Li, C. Yang, H. V. Poor, and B. Vucetic, "A tutorial on ultrareliable and low-latency communications in 6G: Integrating domain knowledge into deep learning," *Proc. IEEE*, vol. 109, no. 3, pp. 204–246, Mar. 2021.
- [43] S. J. Pan and Q. Yang, "A survey on transfer learning," *IEEE Trans. Knowl. Data Eng.*, vol. 22, no. 10, pp. 1345–1359, Oct. 2009.
- [44] R. Dong, C. She, W. Hardjawana, Y. Li, and B. Vucetic, "Deep learning for radio resource allocation with diverse quality-of-service requirements in 5G," *IEEE Trans. Wireless Commun.*, vol. 20, no. 4, pp. 2309–2324, Apr. 2021.
- [45] T. N. Do, G. Kaddoum, T. L. Nguyen, D. B. da Costa, and Z. J. Haas, "Aerial reconfigurable intelligent surface-aided wireless communication systems," in *Proc. IEEE 32nd Annu. Int. Symp. Pers., Indoor Mobile Radio Commun. (PIMRC)*, Helsinki, Finland: IEEE, Sep. 2021, pp. 525–530.
- [46] A. Zhang, Z. C. Lipton, M. Li, and A. J. Smola, *Dive into Deep Learning*. Cambridge, U.K.: Cambridge Univ. Press, 2023. [Online]. Available: <https://D2L.ai>
- [47] T.-V. Nguyen, T.-N. Tran, K. Shim, T. Huynh-The, and B. An, "A deep-neural-network-based relay selection scheme in wireless-powered cognitive IoT networks," *IEEE Internet Things J.*, vol. 8, no. 9, pp. 7423–7436, May 2021.
- [48] Y. Liu, Z. Su, H. Peng, X. Luo, and H.-H. Chen, "Intelligent reflecting surface assisted physical layer security: A deep learning approach," *IEEE Wireless Commun.*, vol. 31, no. 5, pp. 52–60, Oct. 2024.
- [49] F. Zhuang, Z. Qi, K. Duan, D. Xi, Y. Zhu, H. Zhu, H. Xiong, and Q. He, "A comprehensive survey on transfer learning," *Proc. IEEE*, vol. 109, no. 1, pp. 43–76, Jan. 2021.

- [50] J. Zhu, N. Chen, and C. Shen, "A new deep transfer learning method for bearing fault diagnosis under different working conditions," *IEEE Sensors J.*, vol. 20, no. 15, pp. 8394–8402, Aug. 2020.
- [51] H. Zhou, M. Erol-Kantarci, Y. Liu, and H. V. Poor, "A survey on model-based, heuristic, and machine learning optimization approaches in RIS-aided wireless networks," *IEEE Commun. Surveys Tuts.*, vol. 26, no. 2, pp. 781–823, 2nd Quart., 2024.
- [52] E. Björnson, Ö. Özdogan, and E. G. Larsson, "Intelligent reflecting surface versus decode-and-forward: How large surfaces are needed to beat relaying?" *IEEE Wireless Commun. Lett.*, vol. 9, no. 2, pp. 244–248, Feb. 2020.
- [53] R. Dong, C. She, W. Hardjawana, Y. Li, and B. Vucetic, "Deep learning for hybrid 5G services in mobile edge computing systems: Learn from a digital twin," *IEEE Trans. Wireless Commun.*, vol. 18, no. 10, pp. 4692–4707, Oct. 2019.
- [54] K. M. Faisal and W. Choi, "Machine learning approaches for reconfigurable intelligent surfaces: A survey," *IEEE Access*, vol. 10, pp. 27343–27367, 2022.
- [55] H. Cui, H. Zhang, G. R. Ganger, P. B. Gibbons, and E. P. Xing, "GeePS: Scalable deep learning on distributed GPUs with a GPU-specialized parameter server," in *Proc. 11th Eur. Conf. Comput. Syst.* New York, NY, USA: ACM, Apr. 2016, pp. 1–16.
- [56] K. Zhi, C. Pan, H. Ren, K. K. Chai, and M. Elksashan, "Active RIS versus passive RIS: Which is superior with the same power budget?" *IEEE Commun. Lett.*, vol. 26, no. 5, pp. 1150–1154, May 2022.
- [57] Z. Zhang, L. Dai, X. Chen, C. Liu, F. Yang, R. Schober, and H. V. Poor, "Active RIS vs. passive RIS: Which will prevail in 6G?" *IEEE Trans. Commun.*, vol. 71, no. 3, pp. 1707–1725, Mar. 2023.
- [58] Y. Zhu, B. Mao, and N. Kato, "Intelligent reflecting surface in 6G vehicular communications: A survey," *IEEE Open J. Veh. Technol.*, vol. 3, pp. 266–277, 2022.



YOSEFINE TRIWIDYASTUTI received the B.Eng. and M.Eng. degrees in electrical engineering from the Sepuluh Nopember Institute of Technology (ITS), Indonesia, in 2007 and 2013, respectively. She is currently pursuing the Ph.D. degree with the Department of Software and Communications Engineering, Graduate School, Hongik University, Republic of Korea.

Her work experiences are as a RF and Optimization Engineer with PT Samsung Telecommunication Indonesia, from 2007 to 2011, and a Computer Engineering Lecturer with Dinamika University, Indonesia, from 2013 to 2022. Her research interests include wireless communication, physical layer security, and reconfigurable intelligent surface.



TRI NHU DO (Member, IEEE) was born in Da Nang, Vietnam. He received the bachelor's degree in electronics and telecommunications engineering from the Posts and Telecommunications Institute of Technology, Vietnam, in 2012, and the master's and Ph.D. degrees in electronics and computer engineering from Hongik University, South Korea, in 2015 and 2018, respectively.

From 2017 to 2018, he was a Teaching Associate with the Department of Software and Communications Engineering, Hongik University. In 2019, he was a Research Associate with the Department of Computer Science, The University of Texas at Dallas, Dallas, TX, USA. From 2020 to 2023, he was a Postdoctoral Research Fellow with the Resilient Machine Learning Institute, École de technologie supérieure, Montreal, Canada. He is currently an Assistant Professor with the Department of Electrical Engineering, Polytechnique Montréal, Canada.



RIDHO HENDRA YOGA PERDANA (Graduate Student Member, IEEE) received the B.S. degree in telecommunications engineering from the Electronic Engineering Polytechnic Institute of Surabaya, Surabaya, Indonesia, in 2012, and the M.S. degree in electrical engineering from the Sepuluh Nopember Institute of Technology (ITS), Surabaya, in 2014. He is currently pursuing the Ph.D. degree with the Department of Software and Communications Engineering, Graduate School,

Hongik University, Sejong, South Korea.

His research interests include mathematical modeling of 5G networks and optimization/machine learning for wireless communications.

Mr. Perdana received the Best Paper Award of the IEEE International Conference on Artificial Intelligence in Information and Communication 2024.



KYUSUNG SHIM (Member, IEEE) received the B.S. degree in computer and information communications engineering, the M.S. degree in information system, and the Ph.D. degree in electronics and computer engineering from Hongik University, Sejong, South Korea, in 2012, 2017, and 2021, respectively.

After graduation, he joined the Faculty of the School of Computer Engineering and Applied Mathematics, Hankyong National University, Anseong, South Korea, where he is currently an Assistant Professor. His main research interests include AI-based wireless communications and AI-based multicast routing protocol.



BEONGKU AN (Senior Member, IEEE) received the B.S. degree in electronic engineering from Kyungpook National University, Republic of Korea, in 1988, the M.S. degree in electrical engineering from Polytechnic University (NYU), NY, USA, in 1996, and the Ph.D. degree in electrical engineering from New Jersey Institute of Technology (NJIT), NJ, USA, in 2002.

After graduation, he joined the Faculty of the Department of Software and Communications Engineering, Hongik University, Republic of Korea, where he is currently a Professor. From 1989 to 1993, he was a Senior Researcher with RIST, Republic of Korea. His current research interests include mobile wireless networks and communications, such as ad-hoc and sensor networks, cognitive radio networks, cellular networks, and the IoT. In particular, he is interested in cooperative transmission, QoS routing and QoS multicast routing, energy harvesting, physical layer security (PLS), visible light communication (VLC), cross-layer technology, 5G/beyond 5G, and machine learning/deep learning applications.

Prof. An was listed in Marquis Who's Who in Science and Engineering and Marquis Whos Who in the World. He was the President of the IEIE Computer Society in 2012 and also the General Chair of the International Conference (ICGHIT) from 2013 to 2017.

...