

Spatial-Temporal Attention-based mmWave Link Quality Prediction under Dynamic Blockages

Zhizhen Li^{*}, Mingzhe Chen[†], Gaolei Li[‡], Yuchen Liu^{*}

^{*}North Carolina State University, USA, [†]University of Miami, USA

[‡]Shanghai Jiao Tong University, China

Abstract—Millimeter-wave (mmWave) communication is a promising technology that has become a key component of next-generation wireless networks due to its large available bandwidth. However, the susceptibility of mmWave link to dynamic blockages makes it challenging to maintain consistently high rate performance. Hence, it is imperative to have the knowledge of link quality in advance at the location of interest to proactively optimize the use of network resources. In this work, we propose a Spatial-Temporal Attention-based Prediction (STAP) framework to predict the link quality at arbitrary locations in the presence of dynamic blockages. Specifically, our STAP model is built to capture the spatial correlation and temporal dependency of mmWave wireless characteristics in an integrated module, followed by an attention mechanism to complement the link quality prediction task. On top of that, we also design a regional training approach with a weighted loss function to address the data imbalance problem of map-based prediction. Extensive evaluation results show that our framework effectively captures comprehensive spatial-temporal knowledge and achieves significantly higher accuracy than other baseline prediction methods.

I. INTRODUCTION

Millimeter-wave (mmWave) technology has been gaining increasing attention nowadays due to its ability to provide high-bandwidth and low-latency wireless communication, which enables many attractive scenarios and applications including 5G/6G cellular networks, wireless backhaul, Wi-Fi networks, and virtual reality [1]. However, one critical issue of mmWave communication is the high sensitivity to both static and dynamic blockages due to its short propagation distance and poor penetration ability. This problem is exacerbated in obstacle-rich environments, where radio propagation phenomena can be more complex and unpredictable. For instance, in the context of an indoor environment, various factors can affect the quality of mmWave links, including fixed obstacles and moving humans, leading to prominent multi-path effects, shadowing, and blockages.

The susceptibility of links to blockage effects particularly makes it difficult to maintain continuously high link quality, as small changes in the distribution of obstacles or the location of a client device can have constructive or destructive impact on the quality of a mmWave link. Thus, having the knowledge of link quality at locations of interest will significantly enhance network management. To be specific, when a mobile user is moving in an indoor environment, the quality of service experienced by mobile users may be significantly enhanced if information about future link quality along the users' routes is used for proactive resource allocation [2]. Furthermore, link

quality prediction can optimize the use of network resources, allowing for better network management and improved system performance. Thus, it is necessary to make a map-based link quality prediction to guarantee reliable communication and improve network resource scheduling.

Currently, most works focus solely on static scenarios with a long-term link quality prediction, whereas scant attention was paid on short-term prediction under dynamic blockages, e.g., caused by moving humans. In a static network scenario, a high-quality mmWave link always uses a line-of-sight (LoS) path between sender and receiver [3]. When objects made of highly reflective materials such as metal are present in the environment, reflected paths can be also found to maintain high link quality even when no LoS path exists between the two endpoints. However, the dynamic blockages due to moving humans may frequently break this steady state by disrupting the well-established links at different locations, resulting in fluctuations of received signal strength and making it challenging to estimate the link quality in space and time. Although a few prior works have addressed the problem considered herein, they either perform the prediction tasks in a relatively plain scenario with few obstacles [4] or only predict the link quality between a transmitter and a few receivers [5]. Particularly, [5], [6] designed long short-term memory (LSTM) models to predict multi-link quality under dynamic blockages. However, they focused on the link quality prediction at several dedicated locations with considering only temporal-domain information. None of them can dynamically predict a complete link quality map covering any location of interest, which is the subject of this work herein.

In our prior work [7], [8], we developed a machine learning and regression-based approach to link quality prediction for the static scenarios, i.e., addressing blockages due to static obstacles to permit link quality prediction several seconds into the future to facilitate proactive resource allocation. Augmented with dynamic prediction techniques, this paper addresses short-term predictions due to dynamic blockages caused by human obstacles, paving the way for real-time proactive networking configurations. Particularly, since moving obstacles may affect both nearby LoS paths and non-LoS reflection paths to more distant receivers, it is necessary to capture both temporal dependency and spatial correlation for an accurate map-based prediction.

In this paper, we propose a spatial-temporal attention-based framework to make link quality predictions under dynamic

blockages. First, the ray-tracing analysis is performed to synthetically generate sufficient high-quality training data covering a wide range of fine-grained mmWave network scenarios. Then, we investigate the data imbalance problem that exists in the map-based prediction and address it by introducing a novel regional learning mechanism, which strategically captures the critical information from the neighboring areas of moving obstacles. We also design a new loss function with the penalty weights to the minority of data samples collected from potentially affected areas. Next, a Spatial-Temporal Attention-based Prediction (STAP) framework is developed to capture spatial correlation and learn temporal dependency for predictions in both space and time. Although the spatial-temporal based deep learning models have been applied in the area of traffic flow prediction [9]–[12], but few attention is made to mmWave link quality prediction. We also add a soft attention mechanism to improve the prediction accuracy by learning the importance of the link quality variance at every moment. Extensive performance evaluations are conducted to validate the stability, effectiveness, generalization capability, and stretchable time-window prediction ability of the STAP model. The proposed scheme is also shown to outperform the baseline prediction approaches by up to 61% on the prediction accuracy. The main contributions of this work are summarized as follows.

- We propose a spatial-temporal learning framework to predict the link quality under dynamic blockages. This framework can efficiently construct a complete link quality map within a given environment. To the best of our knowledge, this is the first work that attempts to integrate temporal dependency and spatial correlation into the map-based mmWave link quality prediction.
- We investigate the data imbalance problem of dynamic link quality prediction, and propose a novel regional learning mechanism with a weighted loss function to strategically capture the critical information from the specific areas close to moving objects for training.
- We perform extensive performance evaluations of the proposed STAP framework, where the results show very good agreement with the ground truths. This demonstrates that mmWave link quality under dynamic blockages can be accurately predicted through exploiting spatial and temporal characteristics of wireless environment.

II. PRELIMINARIES AND PROBLEM FORMULATION

In this section, we first introduce the ray-tracing analysis to generate large volumes of data that will be used in our STAP model. Then, we formulate the problem through a space gridding scheme for the map-based link quality prediction. Next, we address a data imbalance problem with a novel regional learning mechanism, paving the way for developing our STAP framework as in the subsequent section.

A. Ray-Tracing Analysis

Although machine learning techniques could be used to make link quality predictions, it is challenging to collect a

sufficient volume of training data in real environment covering a complex range of network scenarios. To address this challenge, we adopt a ray-tracing based approach to synthetically generate high-quality training data covering a wide range of fine-grained mmWave network scenarios, which is then used to develop our regression-based approach to dynamic link quality prediction. As a complementary approach to the experimental measurements, the ray-tracing analysis can capture the geometrical properties of the wireless channel for each transceiver and generate the profile of delay τ , path gain, angle of departure (AoD) θ_t , angle of arrival (AOA) θ_r , etc, and then the received power (i.e., link quality) can be obtained by accumulating all signal profiles including LoS, reflection, and diffraction paths [13].

As depicted in Fig. 1 (a), we consider the 3-D layout of an office scenario with a size of $25\text{m} \times 25\text{m} \times 3\text{m}$, consisting of wooden tables, wooden chairs, metal cabinets, and several moving humans to simulate the dynamic obstacles. The transmitter (i.e., a mmWave access point) is placed at the center of the room with a height of 2.9m, and the receivers are evenly distributed with a spacing of 0.4m and at a height of 1m. In this work, we adopt a commercial ray tracer called *Wireless Insite*[®] to generate the above network environment and mmWave signal profiles. The dataset is publicly available at [14]. Specifically, we choose the 3-D ray-tracing model which has no restrictions on geometry shape or transceiver's height. For a cost-effective ray tracing analysis, the maximum order of reflection paths between a transmitter and a receiver is set to 4, which is a reasonable number in mmWave wireless contexts as the large-order reflection rays have negligible impacts on the overall link quality due to the cumulative reflection loss. Similarly, considering the significant signal strength drop after the first-order diffraction, we set the maximum order of diffraction to reach the receiver as 1. The corresponding link quality map is shown in Fig. 1 (b).

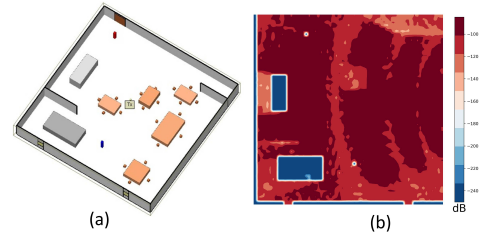


Fig. 1. (a) 3-D scenario layout; (b) The corresponding link quality map.

B. Problem Formulation

To perform a map-based link quality prediction, the geographical area of the environment space is partitioned into $M = N \times N$ grids, in which each grid represents a spatial region r_n ($1 \leq n \leq M$). At each region r_n , a receiver is placed to collect the link quality values in real time. X_v^t is defined as the link quality for all the regions in the set v at the t -th time slot. We aim to predict X_v^{t+1} based on previous T observed values $\mathcal{X}_v^{t-T:t} = (X_v^{t-T}, \dots, X_v^t)$. Since

we consider the information from both spatial and temporal domains, the prediction problem can be formulated as

$$\hat{X}_v^{t+1} = \mathcal{F}(\mathcal{X}_v^{t-T:t}, \mathcal{G}_v^{t-T:t}), \quad (1)$$

where $\mathcal{G}_v^{t-T:t} = \{\mathcal{X}_\mu^{t-T:t} | \mu \in \mathcal{NB}(v)\}$ denotes the link quality at v 's neighbor $\mathcal{NB}(v)$ during the same time period.

C. Regional Learning Mechanism

Typically, to predict the future link quality of the entire space, the input of prediction model should be the link quality values at any locations across the previous time steps. However, this straightforward method causes a data imbalance problem, making the prediction model fail to learn any effect brought by the dynamic blockages. This is because the link quality in majority of the areas is not likely to be affected by the moving objects. As a result, the valid information obtained from the blocked areas is much less than the redundant information retrieved from those unaffected areas, which implies that data training model should shift more attention to the link quality variance in the areas around the dynamic obstacles. [15]

To address this problem, we propose a regional learning mechanism that only considers the link quality status of adjacent regions of the moving obstacles as input to prediction model during the training process. This can be viewed as a data under-sampling method that decreases the samples from those unaffected areas. Specifically, the selected area can be a rectangle region with the same length as the whole room, but with a smaller width, only covering the neighboring area of the obstacle. Next, during the backpropagation process, a weighted loss function is designed to further address this data imbalance issue. Traditional loss functions using basic mean squared error (MAE) are inappropriate for our problem since the error is always small as long as the link quality is well predicted in those unblocked areas. To resolve this problem, we use the loss function with a penalty parameter α as follow:

$$\mathcal{L}_\delta = \frac{\sum_{i=1}^{n_1} |y_i - \hat{y}_i| + \sum_{j=1}^{n_2} \alpha |y_j - \hat{y}_j|}{n_1 + n_2}, \quad (2)$$

where y_i (y_j) and $\hat{y}_i$ ($\hat{y}_j$) represent the ground-truth value and the predicted value of link quality in the unblocked (blocked) areas. n_1 (n_2) represents the number of link quality in the unblocked (blocked) areas, respectively.

III. DYNAMIC LINK QUALITY PREDICTION

In this section, we present the proposed STAP framework for link quality predictions. In general, we first design a graph convolutional network (GCN) to extract the spatial-domain features of mmWave wireless environment, and then a long short-term memory (LSTM) based module is used to capture the temporal dependency for predicting link quality in future time steps. We also add a soft attention mechanism by assigning weights to the past time-series data to further improve the prediction accuracy.

A. Spatial-domain Correlation

In a dynamic mmWave wireless environment, the presence of moving obstacles can easily affect the link quality between

transceivers at arbitrary locations. Thus, it is necessary to capture the spatial correlation between link quality variance and environment details. To this end, we first partition the space (as shown in Fig. 1 (a)) into many grids and place a receiver at each grid to record the received signal strength based on our ray tracing analysis. That way, each receiver can be regarded as a vertex and assuming that the neighboring vertices of the receiver are highly correlated, we then add the edges between these neighbouring vertices to further construct a connected graph which contains detailed spatial information.

Next, we use two layers of GCN model to extract spatial-domain features, taking into account the graph node and the adjacent links of the node to capture the correlation between link quality and environment details. A multi-layer GCN can be expressed as:

$$H^{(l+1)} = \sigma(\tilde{D}^{-\frac{1}{2}} \tilde{A} \tilde{D}^{-\frac{1}{2}} H^{(l)} \theta^{(l)}), \quad (3)$$

where $\hat{A} = A + I$, A is the adjacency matrix of the constructed graph, and I is the identity matrix. $\tilde{D}$ is the degree matrix with $\tilde{D}_{ii} = \sum_j \hat{A}_{ij}$. $H^{(l)}$ is the output of the layer l . $\theta^{(l)}$ is the parameter of the layer l , and σ is the activation function.

In the stage of graph convolution, each node will combine the information received from its neighbouring nodes and then share the learned knowledge with each other. In this way, our GCN model encodes the topological structure of the graph and captures the spatial correlations among all nodes and links.

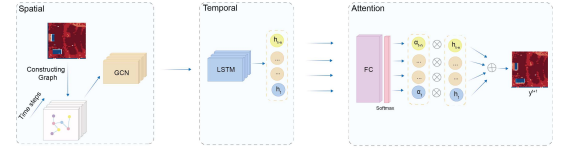


Fig. 2. Overview of the STAP framework.

B. Temporal-domain Dependency

To learn the temporal dependency of link quality variances caused by blockages and multi-path effects along the timeline, we use a LSTM layer in the framework to predict the received signal strength at any locations of a future time step. As a variant of recurrent neural network, LSTM is designed to circumvent the vanishing gradient problem and make use of the gate mechanism to capture long- and short-term dependencies. The model can be formulated as:

$$f_t = \sigma(W_f x_t + W_{hf} h_{t-1} + b_f) \quad (4)$$

$$i_t = \sigma(W_i x_t + W_{if} h_{t-1} + b_i) \quad (5)$$

$$o_t = \sigma(W_o x_t + W_{of} h_{t-1} + b_o) \quad (6)$$

$$c_t = f_t \circ c_{t-1} + i_t \sigma(W_c x_t + W_{cf} h_{t-1} + b_c) \quad (7)$$

$$h_t = o_t \circ \sigma(c_t), \quad (8)$$

where x_t is the input vector at time t . f_t, i_t, o_t represent the forget gate, input gate and output gate respectively. c_t denotes the memory cell in the unit and h_t stands for the hidden unit. All the W are the weight vectors in each gate. b stands for

the bias vector and σ is the activation function. In particular, the memory cell c_t combines the previous cell states at c_{t-1} , current input and previous output, to update hidden states h_t . The forget gate f determines whether the information in the previous memory should be discarded or not. The output gate learns how the memory cell should affect the hidden states. As such, the LSTM layer can well predict the link quality of the future time step based on the previous hidden state information and the input at the current time step, which captures the dynamic temporal variations with this gated mechanism.

C. Attention-based Enhancement

As the last component in Fig. 2, we add a soft attention layer in the STAP framework to learn the importance of the link quality at every moment. Since each past data in both space and time will have a different degree of impact on the link quality in future time steps, it is critical to strategically assign different weights to those historical data pieces for a more accurate prediction.

To be specific, suppose that the input time series is $X = \{x_1, x_2, \dots, x_n\}$, then for every single time step x_k in X , there is a corresponding hidden state h_k from the LSTM output. Typically, the hidden state h_n of the last input time step is used as the output for prediction. However, the information from much earlier time steps might not be totally ignored or addressed as it may also contain some important knowledge that contributes to the prediction at next time steps. In this way, the output of the attention layer is calculated in a weighted average way as:

$$\hat{h} = \sum_{i=1}^n \alpha_i h_i, \quad (9)$$

where α_i is the weight of each hidden layer. To calculate the weights, we train a fully connected layer on the hidden states to get a score for each state as follow:

$$s_i = \text{sigmoid}(w^T h_i + b_i), \quad (10)$$

where s_i is the calculated score. Then, we use a softmax function in Eq. (11) to normalize this score and get the weight for each hidden state.

$$\alpha_i = \frac{\exp(s_i)}{\sum_{k=1}^n \exp(s_k)}. \quad (11)$$

IV. PERFORMANCE EVALUATIONS

In this section, we evaluate the performance of our proposed STAP framework to predict the mmWave link quality under dynamic blockages, including the comparison with the state-of-the-art models, the stability of the proposed model under multi-human scenarios, and its transfer capability to various mobility patterns of obstacles.

A. Network Settings

a) *Scenario and model configurations:* We consider a mmWave network scenario as described in Sec. III, where several fixed objects are randomly placed on the ground, and some human objects are randomly moving at a regular speed of 1.4m/s. Then, we perform the space gridding to partition

the entire space into 3,969 (63×63) small grids with the equal sizes of 0.16m^2 . Next, we use our ray tracer to generate over 1,000,000 data samples, which include received signal strengths (RSS) accounting for the locations of transmitters and receivers, fixed and dynamic blockages, and multi-path effects including reflection and diffraction. The data is collected with a sample rate of 30ms and scaled by max-min normalization. We split the dataset into training and testing sets with a ratio of 70% and 30%. For the learning model configurations, we set the input length of time step to the STAP framework as 8 and predict the RSS for the next time step. Inside the model, the GCN is developed to learn the spatial representation with the number of hidden units of 64, and we set 40 hidden units and 40 hidden nodes in LSTM part and the attention layer, respectively. Lastly, the penalty parameter α in our designed loss function is set as 10.

b) *Evaluation metrics:* We evaluate the performance of our STAP model using the mean absolute error (MAE) and the performance difference ratio (PDR), measuring the difference between the predicted values and ground truths, where

$$\text{MAE} = \frac{\sum_{i=1}^n |\mathcal{S}_{pred} - \mathcal{S}_{truth}|}{n}, \quad (12)$$

$$\text{PDR} = |\mathcal{S}_{pred} - \mathcal{S}_{truth}| / (\mathcal{S}_{\max} - \mathcal{S}_{\min}). \quad (13)$$

In Eq. (12), we consider both global MAE and local MAE, where the former indicates the aggregated error across the entire area of interest and the latter focuses only on the error of those locations that affected by moving objects. Additionally, the PDR in Eq. (13) is used to evaluate the prediction accuracy with varying *error tolerance rate* (ETR), where the predicted link quality $\mathcal{S}_{pred}$ is accepted as an accurate result when the PDR is less than the given ETR.

B. Impact of Region Selection

As described in the Sec. II.C, we exploit a regional learning mechanism to overcome the data imbalance issue. Intuitively, considering a large region size in the model may compromise the prediction personalization, resulting in the increase of the local MAE, while a small region size will fail to capture the sufficient spatial information for prediction due to the high environment dependency of mmWave links. Therefore, it is utmost of importance to choose an appropriate region size in our STAP framework.

In this part, we evaluate the performance of STAP model with different region sizes and the results are reported in Fig. 3. First, it is expected to see that the global MAE increases with the larger region size due to the data imbalance issue. Then, it is interesting to observe that the local MAE decreases at first, but then starts to increase as the considered region size becomes larger. The initial decrease is due to more spatial information being considered as the selected area is expanded. However, as the region size keeps increasing, the data imbalance begins to dominate and overwhelm the benefits brought by spatial information, resulting in higher local MAE. In what follows, we select the 11% of the space size for

regional learning because it strikes a good balance between the local MAE and the global MAE.

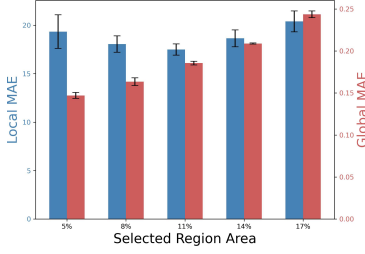


Fig. 3. Area percentage of selected region vs. MAE.

C. Model Comparison

Next, to validate the performance of our proposed STAP model, we compare with several baseline models including LSTM based model from [4], CNN-LSTM model from [10], GCN-LSTM model from [16], and the STAP model using the standard MSE based loss function (termed as STAP-STD), while our proposed STAP herein is trained with a modified loss function $\mathcal{L}_\delta$ in Eq. (2).

TABLE I
PERFORMANCE COMPARISONS.

Method	MAE	
	Local MAE	Global MAE
LSTM [4]	41.7542	0.2652
CNN-LSTM [10]	31.7453	0.2781
GCN-LSTM [16]	25.1281	0.2548
STAP-STD	18.1567	0.1922
STAP	16.1409	0.1902

Table. I shows the performance comparisons among all considered models. Obviously, the proposed STAP outperforms other baseline models in terms of both global MAE and local MAE. By capturing the spatial dependency information, our STAP, CNN-LSTM and GCN-LSTM can improve the prediction accuracy by up to 61%, 24%, and 39% compared to the pure LSTM, respectively, which demonstrates the importance of spatial correlations in mmWave link quality prediction. In addition, our STAP is superior to GCN-LSTM and CNN-LSTM by adding a soft attention mechanism, which considers the correlation between links in both space and time. We also find that the STAP shows the better performance than STAP-STD, and this validates the effectiveness of the modified loss function that well addresses the data imbalance issue.

Besides the quantitative results presented in Table. I, Fig. 4 depicts the visualized map-based prediction results. Specifically, we showcase the prediction error map (i.e., $\forall i \in L$, $|\hat{x}_i - x_i|/x_i$) for each model, where $\hat{x}_n$ and x_i are the predicted and ground-truth link quality at any location $i \in L$. The brighter pixel in the map indicates the larger prediction error, so the superiority of STAP model can be easily observed, which is consistent with the quantitative results in Table. I. Additionally, as discussed in Sec. II-C, we only predict the

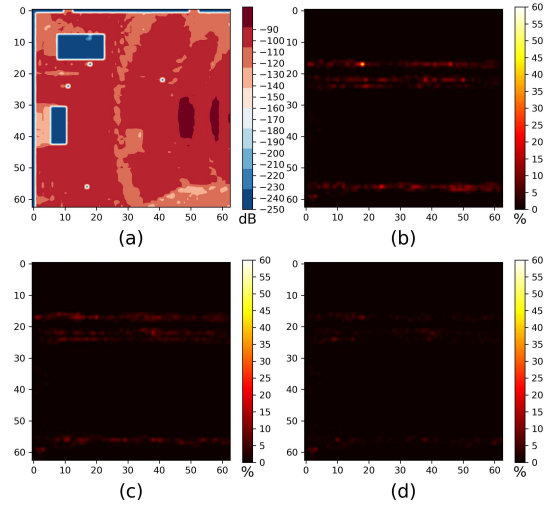


Fig. 4. Visualization of prediction results. (a) is the predicted link quality map of STAP; (b)-(d) are the error maps between predicted and ground-truth link quality maps from LSTM, GCN-LSTM, and our STAP models.

future link quality of the neighbouring area of obstacles. The link quality of the remaining area is the same as that of the last time step. As what we find from the error maps in Fig. 4(b)-(d), the majority of the error in those areas are around zero, which means the link quality from last time step is almost the same as the next time step. This result validates the effectiveness of our regional learning mechanism, namely achieving high prediction accuracy with less computational overhead.

D. Prediction on Stretchable Time Windows

In addition to predicting the link quality at only the next time step, our STAP model is capable of making predictions on a stretchable time window, i.e., generating link quality maps for next several time steps, where each time step is set as 30ms in this evaluated case. Here we first investigate the performance of our STAP model vs. the future time steps in Fig. 5(a). As expected, the prediction error increases when the

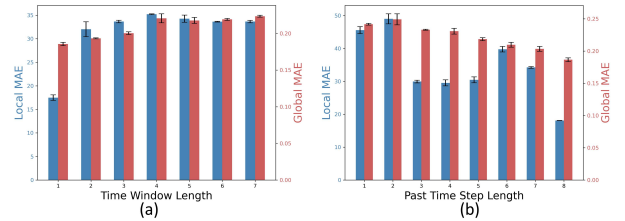


Fig. 5. (a) The length of future time window vs. MAE; (b) The length of past time window vs. MAE.

model becomes more farsighted. Additionally, we observe a significant increase of local MAE at first, but then it becomes marginal as the time step increases. Notably, both local and global MAE stay almost unchanged when the window length is larger than 4, where the global prediction error is maintained at around only 0.22. This result demonstrates the capability of our STAP model to predict link quality within a stretchable time window, exhibiting the potential use to allow for

proactive network configurations in different delay-sensitive applications.

Besides the study on the “lookahead” capability, here we use the term “lookback” to depict length of past time step needed for predicting the future link quality. Intuitively, a longer lookback period can encode more temporal information during the learning process, thus improving the prediction accuracy. This hypothesis is proved in Fig. 5(b), where we can see a decreasing trend in both local MAE and global MAE when more lookbacks are considered. Specifically, the prediction error becomes relatively small when the lookback period is more than 3 in the evaluated scenario. As a result, we conclude that the information from a few past time period might be sufficient to make an accurate link quality prediction.

E. Multi-human Scenario and Model Generalizability

In this part, we evaluate the performance of our STAP model with varying moving human density in the network scenario. Fig. 6(a). shows the PDR metric vs. the moving human density. We adopt different ETRs to evaluate the performance of the proposed prediction model, where the predicted link quality is accepted as an accurate result when the PDR is less than the given ETR. As expected, the increase of human density will cause a decrease in the percentage of accepted prediction results across all receiver locations in the scenario. However, our STAP model can still maintain around 85% and 97% prediction accuracy with a large dynamic blockage density when ETR is 0.01 and 0.03, respectively, which corresponds to the average link quality prediction error of just 1–3 dB across the entire scenario map. The results validate the stability of our proposed model, i.e., being able to predict the link quality variance within an acceptable accuracy as the density of dynamic blockages increases.

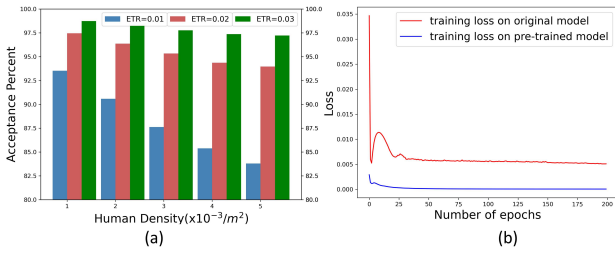


Fig. 6. (a) Prediction accuracy vs. moving obstacle densities; (b) Loss comparison on original model and pre-trained model.

Lastly, we investigate whether our STAP model is generalized to arbitrary mobility patterns of temporary obstacles. We evaluate the model performance in the case of humans moving in random directions, and the results are reported in Fig. 6(b). Specifically, the red line in Fig. 6(b). represents the learning loss vs. the used epochs when training a new model, while the blue line shows the convergence when new dataset consisting of a different moving pattern is used as input to a pre-trained model. In particular, it is observed that the initial loss on the pre-trained model is significantly lower than that of the newly trained model. Also, adding the new data to our pre-trained model converges faster and achieves the lower loss. This result

shows the generalizability of our model to mobility pattern of obstacles, which can be applied in various dynamic mmWave network scenarios, since only a few epochs are needed to train a link quality predictor based on the pre-trained model.

V. CONCLUSION

In this paper, we studied mmWave link quality prediction under dynamic blockages. We first investigated a data imbalance problem of the map-based prediction through a regional training mechanism. Then, an attention-based spatial-temporal learning framework was proposed to dynamically predict the link quality at any locations in the scenario. Extensive evaluation results showed that our approach can achieve fairly promising prediction accuracy and is robust to multiple dynamic obstacles with arbitrary mobility patterns.

ACKNOWLEDGEMENT

This research was supported by the National Science Foundation through Award CNS–2312138 and CNS–2312139.

REFERENCES

- [1] M. Chen, D. Gündüz, K. Huang, W. Saad, M. Bennis, A. V. Feljan, and H. V. Poor, “Distributed learning in wireless networks: Recent progress and future challenges,” *IEEE Journal on Selected Areas in Communications*, vol. 39, no. 12, pp. 3579–3605, 2021.
- [2] L. Shen, T. Wang, and S. Wang, “Proactive proportional fair: A novel scheduling algorithm based on future channel information in OFDMA systems,” in *IEEE/CIC Int’l Conf. on Communication in China*, 2019.
- [3] Y. Liu, Q. Hu, and D. M. Blough, “Joint link-level and network-level reconfiguration for mmWave backhaul survivability in urban environments,” in *Proceedings of International ACM Conference on Modeling, Analysis and Simulation of Wireless and Mobile Systems*, 2019.
- [4] S. H. A. Shah, M. Sharma, and S. Rangan, “LSTM-based multi-link prediction for mmWave and sub-THz wireless systems,” in *IEEE International Conference on Communications (ICC)*. IEEE, 2020.
- [5] S. H. A. Shah and S. Rangan, “Multi-cell multi-beam prediction using auto-encoder LSTM for mmWave systems,” *IEEE Transactions on Wireless Communications*, vol. 21, no. 12, pp. 10366–10380, 2022.
- [6] K. Ma, D. He, H. Sun, and Z. Wang, “Deep learning assisted mmWave beam prediction with prior low-frequency information,” in *ICC 2021-IEEE International Conference on Communications*. IEEE, 2021.
- [7] Y. Liu and D. M. Blough, “Environment-aware link quality prediction for millimeter-wave wireless lans,” in *Proceedings of ACM International Symposium on Mobility Management and Wireless Access*, 2022.
- [8] Y. Liu, “Location, Location, Location: Maximizing mmWave LAN Performance through Intelligent Wireless Networking Strategies,” 2022.
- [9] K. He, Y. Huang, X. Chen, Z. Zhou, and S. Yu, “Graph attention spatial-temporal network for deep learning based mobile traffic prediction,” in *IEEE Global Communications Conference (GLOBECOM)*. IEEE, 2019.
- [10] D. Wang, Y. Yang, and S. Ning, “DeepSTCL: A deep spatio-temporal ConvLSTM for travel demand prediction,” in *2018 international joint conference on neural networks (IJCNN)*. IEEE, 2018, pp. 1–8.
- [11] K. He, X. Chen, Q. Wu, S. Yu, and Z. Zhou, “Graph attention spatial-temporal network with collaborative global-local learning for citywide mobile traffic prediction,” *IEEE Trans on mobile computing*, 2020.
- [12] J. Bai, J. Zhu, Y. Song, L. Zhao, Z. Hou, R. Du, and H. Li, “A3t-gcn: Attention temporal graph convolutional network for traffic forecasting,” *ISPRS International Journal of Geo-Information*, vol. 10, no. 7, p. 485, 2021.
- [13] Y. Liu, S. K. Crisp, and D. M. Blough, “Performance study of statistical and deterministic channel models for mmWave wi-fi networks in ns-3,” in *Proceedings of the 2021 Workshop on ns-3*, 2021, pp. 33–40.
- [14] Z. Li, “Dataset for STAP framework under dynamic blockages.” [Online]. Available: <https://github.com/lzzz1998/moving-obj-data>
- [15] C. Zhang, J. Li, Y. Zhao, T. Li, Q. Chen, X. Zhang, and W. Qiu, “Problem of data imbalance in building energy load prediction: Concept, influence, and solution,” *Applied Energy*, vol. 297, p. 117139, 2021.
- [16] Z. Li, G. Xiong, Y. Chen, Y. Lv, B. Hu, F. Zhu, and F.-Y. Wang, “A hybrid deep learning approach with gcn and lstm for traffic flow prediction,” in *2019 IEEE intelligent transportation systems conference (ITSC)*. IEEE, 2019, pp. 1929–1933.