

Site-Specific Beamforming for Full-Duplex Massive MIMO Systems via Implicit Channel Estimation

Samuel H. Li and Ian P. Roberts

Abstract—Beamforming has proven to be valuable in enabling full-duplex massive MIMO base stations, but doing so effectively often requires knowledge of the self-interference channel matrix $\mathbf{H}$. Estimating this high-dimensional channel is costly in practice, however, since it requires a prohibitive number of measurements, especially in fast-fading conditions. In this work, we overcome this dilemma by designing full-duplex beams using *implicit* channel knowledge gathered from a relatively small number of measurements across $\mathbf{H}$. These measurements are collected by the base station using a sequence of beams tailored to both the deployment environment and the particular users being served. This is accomplished through site-specific training of a transformer-based deep learning model that learns to efficiently probe portions of $\mathbf{H}$ most relevant to the particular users being served by exploiting the underlying structure of the surrounding environment. The deep learning model then uses these probing measurements to design transmit and receive beams that couple low self-interference while delivering high gain to a pair of downlink and uplink users. For favorable multi-user scaling, a single set of probing measurements can be used by the model to serve several users throughout the coherence time of $\mathbf{H}$ by leveraging correlations across those users' channels. Simulation results using ray-tracing demonstrate that our proposed approach exceeds the best possible performance with explicit channel estimation across a wide range of scenarios, especially with large antenna arrays.

Index Terms—Full-duplex, beamforming, self-interference, phased arrays, massive MIMO, millimeter-wave, deep learning.

I. INTRODUCTION

Massive multiple-input multiple-output (MIMO) wireless transceivers are a cornerstone of 5G communication systems and are expected to be a mainstay in future 6G systems [1], [2]. To meet the demands of emerging applications, there has been recent interest in upgrading massive MIMO base stations with in-band full-duplex capability—the ability to transmit and receive at the same time and same frequency [3]. This has the potential to increase throughput, extend coverage, reduce latency, and enable integrated sensing and communication (ISAC) [4], but is challenged by the manifestation of *self-interference* that couples between the base station's transmitter and receiver [5]. More specifically, self-interference manifests between each *pair* of its transmit and receive antennas, the nature of which depends on near-field coupling as well as reflections off the environment [6], [7]. Enabling full-duplex operation requires some method to sufficiently cancel this self-interference [4], [5].

Prior work has shown that, if one has knowledge of all these coupling paths—i.e., the self-interference channel matrix $\mathbf{H}$ —then beamforming can be used to isolate the transmitter

and receiver from one another by forming spatial nulls around the channel [8]–[10]. Obtaining this self-interference channel state information (CSI) is particularly challenging in practice, however, since the number of required measurements can be prohibitively high, eroding the potential gains offered by full-duplex operation. Given N_t transmit antennas and N_r receive antennas at the base station—each often on the order of tens, hundreds, or even thousands—the number of measurements generally needed to estimate $\mathbf{H}$ is $\mathcal{O}(N_t N_r)$, as the majority of massive MIMO systems employ analog or hybrid beamforming architectures [11], [12]. Since reflections off surrounding objects can influence the self-interference channel, estimation of $\mathbf{H}$ must be repeated at a rate commensurate with the dynamics of the environment, further increasing measurement overhead. This dilemma marks our focus in this paper: how to best design beams to cancel self-interference and enable full-duplex massive MIMO systems without incurring the prohibitive measurement overhead associated with explicitly estimating $\mathbf{H}$?

A. Related Work

As mentioned, most existing full-duplex beamforming designs assume (often perfect) knowledge of the self-interference channel. In [13]–[18], for instance, the channel matrix $\mathbf{H}$ is a required input into optimization problems that solve for beams that maximize spectral efficiency or minimize self-interference. Other schemes use deep learning methods such as convolutional neural networks [19] and deep-unfolding [20] to design beams for full-duplex, yet also require explicit knowledge of $\mathbf{H}$. The overhead associated with estimating $\mathbf{H}$ and its implications on end performance have been largely overlooked in the vast majority of prior work, as they typically report signal-to-interference-plus-noise ratio (SINR) or raw spectral efficiency but do not account for the resources spent to estimate $\mathbf{H}$.

To our knowledge, the only existing designs that do not rely on *explicit* estimation of $\mathbf{H}$ are [21]–[23]. In [21], [22], this was accomplished through a heuristic search of beam candidates, measuring the self-interference coupled by each candidate in order to find those that offer best performance. This search typically requires 50–500 measurements and needs to be repeated anytime a new user pair is served or when the self-interference channel changes. A fundamental shortcoming of this approach is that it does not exploit all the spatial degrees-of-freedom offered by massive MIMO, since it searches over predefined beam candidates rather than the entire beam space. In [23], explicit estimation of $\mathbf{H}$ was circumvented by using a long short-term memory (LSTM)

architecture to design a sequence of *probing* beams to collect measurements of self-interference that are then used to design actual serving beams. Like [21], [22], the scheme in [23] must be repeated for each user pair and thus exhibits poor multi-user scaling. An additional drawback of [23] is that it couples measurement of $\mathbf{H}$ with measurement of the users' channels, which necessitates coordination and feedback from users and thus impedes its practical adoption.

It is worth noting that this problem of high-dimensional channel estimation also appears in *beam management*, where beams need to be steered between a massive MIMO base station and a user. To do so effectively, the two need some knowledge of the downlink/uplink channel, but estimating it directly would similarly incur excessive measurement overhead. This has led to codebook-based beam management becoming the *de facto* standard in wireless networks such as 5G [24], where a set of beams are swept to identify the one which maximizes signal-to-noise ratio (SNR). Within this beam-sweeping framework, researchers have proposed sparse probing techniques [24]–[30], which leverage deep learning to design small sets of probing beams tailored to particular environments in place of predefined beam codebooks; we take a similar approach in this paper. In other works [31]–[34], machine learning has been used to infer the optimal beam from side information such as GPS, camera, radar, lidar, or lower-frequency channel data. Also of note are compressive sensing techniques to estimate high-dimensional channels or identify angles-of-departure/arrival [35], [36]. The primary drawbacks of such approaches include their reliance on channel sparsity and their sensitivity to noise in practice; they can also require a prohibitively large number of measurements for practical deployments [37], [38].

B. Contributions

In this work, we introduce a novel approach to designing the transmit and receive beams of a full-duplex massive MIMO base station that circumvents explicit estimation of the self-interference channel $\mathbf{H}$. We employ a deep learning model to first design a set of *probing* beams that are swept by the base station to collect M measurements of self-interference, providing it with *implicit* knowledge of the underlying channel $\mathbf{H}$; these measurements are then processed by the model to synthesize transmit and receive beams which serve pairs of downlink and uplink users in a full-duplex fashion. Rather than repeat probing measurements for each user pair, we architect our solution such that a single set of M measurements can be reused across multiple user pairs for favorable scaling.¹

Our deep learning solution consists of multiple transformer-based modules that are trained end-to-end in a site-specific fashion, tailoring it to the particular environment in which the base station is deployed. This allows the model to learn to craft probing beams that exploit the underlying structure of the environment in order to efficiently probe the portions of $\mathbf{H}$ that are most relevant to the particular group of users

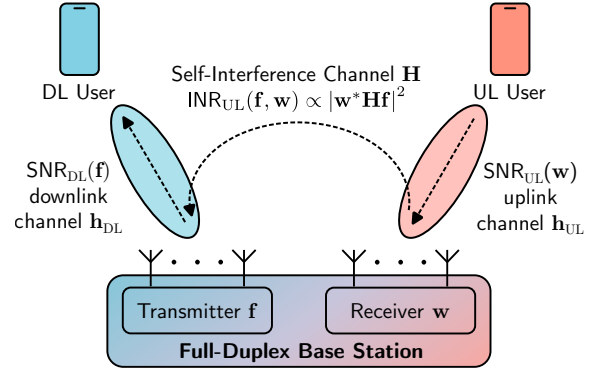


Fig. 1. An in-band full-duplex base station transmits to a downlink user while receiving from an uplink user across the same frequency band. This work aims to optimally design $\mathbf{f}$ and $\mathbf{w}$ without *explicitly* estimating $\mathbf{H}$.

to be served. The implicit channel knowledge contained in these measurements can then be interpreted by the model to synthesize serving beams to cancel self-interference while delivering high SNR to the users.

Evaluation of our proposed solution using a ray-tracing simulation demonstrates that our model is capable of designing near-optimal serving beams from a fraction of the measurements that would be required to explicitly estimate the self-interference channel $\mathbf{H}$. In many settings, our results outperform baselines in terms of spectral efficiency, especially when measurement overhead is accounted for. In one conservative setting, for instance, our solution can deliver over 80% of the full-duplex capacity to eight users with only $M = 16$ total probing measurements, whereas the *best possible* performance when explicitly estimating $\mathbf{H}$ is below 64% of the capacity. Our results also exhibit healthy scaling to an increased number of antennas; in fact, we see a net performance improvement thanks to the added degrees-of-freedom, without increases in measurement overhead or computational complexity.

C. Notation

The following notation is used throughout this paper: standard Roman font x is used to denote scalars; bold lowercase $\mathbf{x}$ is used to denote vectors; bold uppercase $\mathbf{X}$ is used to denote matrices; $[\mathbf{x}]_i$ denotes the i -th element of a vector $\mathbf{x}$; $\mathbb{R}$ and $\mathbb{C}$ denote the sets of real and complex numbers; $|\cdot|$, $\|\cdot\|_1$, $\|\cdot\|_2$, $\|\cdot\|_{\max}$, and $\|\cdot\|_F$ represent element-wise absolute value, ℓ_1 -norm, ℓ_2 -norm, max-norm, and Frobenius norm, respectively.

II. SYSTEM MODEL

As illustrated in Fig. 1, we consider an in-band full-duplex base station equipped with separate transmit and receive arrays, serving single-antenna downlink and uplink users. The base station's transmit and receive arrays consist of N_t and N_r antenna elements, respectively, each driven by a single radio frequency (RF) chain and equipped with analog beamforming. The base station transmits to the downlink user using transmit beam $\mathbf{f} \in \mathbb{C}^{N_t \times 1}$, while simultaneously receiving from the

¹This multi-user aspect is a key functional difference between the present paper and our conference paper [39]. As a result of this, the deep learning implementation herein is substantially more sophisticated than that in [39].

uplink user with receive beam $\mathbf{w} \in \mathbb{C}^{N_r \times 1}$. The resulting downlink and uplink SNRs are given by

$$\text{SNR}_{\text{DL}}(\mathbf{f}) = \frac{P_{\text{DL}} |\mathbf{h}_{\text{DL}}^* \mathbf{f}|^2}{N_t \sigma_{\text{DL}}^2}, \quad (1)$$

$$\text{SNR}_{\text{UL}}(\mathbf{w}) = \frac{P_{\text{UL}} |\mathbf{w}^* \mathbf{h}_{\text{UL}}|^2}{\|\mathbf{w}\|_2^2 \sigma_{\text{UL}}^2}. \quad (2)$$

Here, $\mathbf{h}_{\text{DL}} \in \mathbb{C}^{N_t \times 1}$ and $\mathbf{h}_{\text{UL}} \in \mathbb{C}^{N_r \times 1}$ are the downlink and uplink channel vectors; P_{DL} and P_{UL} are the transmit powers of the base station and of the uplink user; σ_{UL}^2 and σ_{DL}^2 are the noise variances at the base station and at the downlink user, respectively. The divisions by N_t in (1) and by $\|\mathbf{w}\|_2^2$ in (2) account for power splitting across the transmit antennas and noise combining at the receive array, respectively.

We assume analog beamforming is implemented via a network of analog phase shifters and variable attenuators, and this imposes a per-antenna power constraint on the transmit beam $\mathbf{f}$ as

$$|[\mathbf{f}]_i| \leq 1, \quad i = 1, \dots, N_t. \quad (3)$$

By transmitting and receiving in a full-duplex fashion, the base station inflicts interference upon itself across the MIMO channel matrix $\mathbf{H} \in \mathbb{C}^{N_r \times N_t}$. We model $\mathbf{H}$ as the sum of separate line-of-sight (LOS) and non-line-of-sight (NLOS) components by [7], [40]

$$\mathbf{H} = \sqrt{\frac{\kappa}{\kappa + 1}} \mathbf{H}_{\text{LOS}} + \sqrt{\frac{1}{\kappa + 1}} \mathbf{H}_{\text{NLOS}}, \quad (4)$$

where κ is the Rician factor that controls the relative strength of the LOS and NLOS components. The LOS component $\mathbf{H}_{\text{LOS}}$ captures near-field coupling between each pair of antennas in the arrays and is largely determined by the array geometry at the base station; it can therefore be assumed quasi-static [41]. The NLOS component $\mathbf{H}_{\text{NLOS}}$, on the other hand, captures multipath reflections from the surrounding environment and is assumed to vary more rapidly with time.

For a given transmit beam $\mathbf{f}$ and receive beam $\mathbf{w}$, the strength of self-interference incurred at the base station is proportional to $|\mathbf{w}^* \mathbf{H} \mathbf{f}|^2$, with the uplink interference-to-noise ratio (INR) putting this relative to noise as

$$\text{INR}_{\text{UL}}(\mathbf{f}, \mathbf{w}) = \frac{P_{\text{DL}} |\mathbf{w}^* \mathbf{H} \mathbf{f}|^2}{N_t \|\mathbf{w}\|_2^2 \sigma_{\text{UL}}^2}. \quad (5)$$

In this full-duplex setting, the downlink user may also experience cross-link interference from the uplink user's transmission, with the downlink INR given by

$$\text{INR}_{\text{DL}} = \frac{P_{\text{UL}} |h|^2}{\sigma_{\text{DL}}^2}, \quad (6)$$

where $h \in \mathbb{C}$ denotes the scalar interference channel between the uplink and downlink users. Cross-link interference is typically orders of magnitude weaker than that of self-interference, due to larger distances between users and lower uplink transmit powers [7], [42], making self-interference our primary focus.

Accounting for the effects of both additive noise and interference, we can formulate downlink and uplink SINRs as

$$\text{SINR}_{\text{DL}}(\mathbf{f}) = \frac{\text{SNR}_{\text{DL}}(\mathbf{f})}{1 + \text{INR}_{\text{DL}}}, \quad (7)$$

$$\text{SINR}_{\text{UL}}(\mathbf{f}, \mathbf{w}) = \frac{\text{SNR}_{\text{UL}}(\mathbf{w})}{1 + \text{INR}_{\text{UL}}(\mathbf{f}, \mathbf{w})}. \quad (8)$$

The SINRs dictate the achievable spectral efficiency on each link, which can be expressed as

$$R_{\text{DL}}(\mathbf{f}) = \log_2(1 + \text{SINR}_{\text{DL}}(\mathbf{f})), \quad (9)$$

$$R_{\text{UL}}(\mathbf{f}, \mathbf{w}) = \log_2(1 + \text{SINR}_{\text{UL}}(\mathbf{f}, \mathbf{w})). \quad (10)$$

Overall system performance can then be measured by the sum spectral efficiency (SSE)

$$R(\mathbf{f}, \mathbf{w}) = R_{\text{DL}}(\mathbf{f}) + R_{\text{UL}}(\mathbf{f}, \mathbf{w}), \quad (11)$$

which jointly captures the effects of downlink SNR, uplink SNR, self-interference, and cross-link interference as a function of the base station's transmit and receive beams.

III. PROBLEM STATEMENT

The goal of this work is to create a scheme which designs the transmit beam $\mathbf{f}$ and receive beam $\mathbf{w}$ of the full-duplex base station to maximize the SSE when serving pairs of downlink and uplink users over time. Within this context, we consider the case where the self-interference channel $\mathbf{H}$ varies with time according to a block-fading model. More specifically, the coherence time of $\mathbf{H}$ is defined such that K user pairs can be served, each for L time slots, before the channel changes. We term this a *coherent user group*. If the base station collects M measurements of self-interference to estimate $\mathbf{H}$ before serving these K user pairs, then the mean *effective sum spectral efficiency* among these users is equal to

$$\frac{KL}{M + KL} \left(\frac{1}{K} \sum_{k=1}^K R_k(\mathbf{f}_k^*, \mathbf{w}_k^*) \right), \quad (12)$$

where $(\mathbf{f}_k^*, \mathbf{w}_k^*)$ are the transmit and receive beams used to serve the k -th user pair and $R_k(\mathbf{f}_k^*, \mathbf{w}_k^*)$ is the resulting SSE delivered to that user pair. The principal objective of this work is to design near-optimal serving beams $\{(\mathbf{f}_k^*, \mathbf{w}_k^*)\}_{k=1}^K$ in order to maximize the SSE to those K user pairs.

From (12), we can observe the core trade-off at the center of this work. Designing $(\mathbf{f}_k^*, \mathbf{w}_k^*)$ to maximize spectral efficiency requires knowledge of self-interference, but if too many measurements M are used to gather this knowledge, the effective SSE will be degraded [43]. This motivates us to optimize these measurements through strategic design of the M *probing* beam pairs $\{(\mathbf{f}_m, \mathbf{w}_m)\}_{m=1}^M$ that are used to collect them. We aim to demonstrate that, if these M probing beam pairs are designed well, then many fewer measurements are needed compared to that for explicit estimation of $\mathbf{H}$, i.e., $M \ll N_t N_r$. As shall be seen, the key to fulfilling this aim is our use of deep learning to *jointly* design the M probing beam pairs and the final K serving beam pairs.

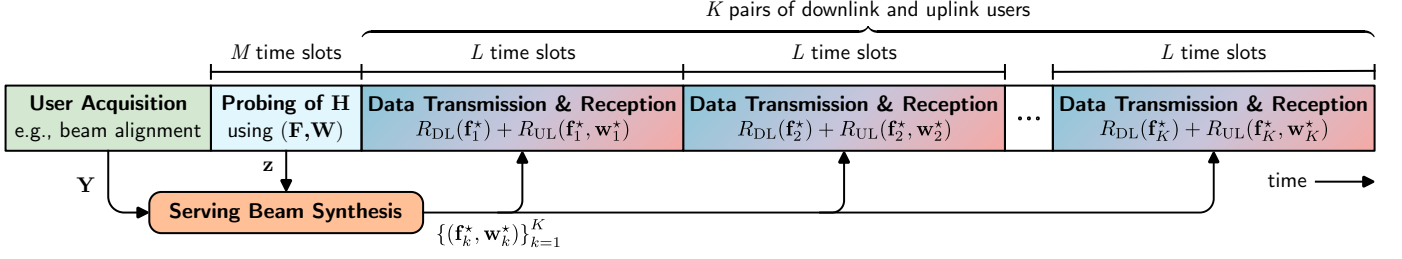


Fig. 2. Timeline of the envisioned use of our proposed scheme, with one time slot defined as the time to collect a single probing measurement across $\mathbf{H}$.

IV. SITE-SPECIFIC BEAMFORMING VIA IMPLICIT SELF-INTERFERENCE CHANNEL ESTIMATION

In this section, we detail our proposed full-duplex beamforming design, which aims to design the M probing beam pairs $\{(\mathbf{f}_m, \mathbf{w}_m)\}_{m=1}^M$ and the final K serving beam pairs $\{(\mathbf{f}_k^*, \mathbf{w}_k^*)\}_{k=1}^K$. As shown in Fig. 2 and Fig. 3, our envisioned scheme consists of three main components:

- 1) Conventional beam alignment (or a similar procedure) to obtain some knowledge of the users' channels.
- 2) Measurements of self-interference taken with M probing beam pairs to obtain *implicit* channel knowledge of $\mathbf{H}$.
- 3) Synthesis of K transmit and receive beams to serve K pairs of downlink and uplink users in a full-duplex fashion.

We overview each of these components, along with our final design problem, in the following subsections. Then, in the next section, we elaborate on the deep learning implementation used to realize this envisioned solution.

A. User Acquisition

Like in conventional half-duplex systems, the full-duplex base station requires knowledge of the users' downlink and uplink channels in order to serve them with high gain. Our proposed framework begins by assuming that existing mechanisms are used to obtain this information, such as beam alignment in 5G networks [26]. In general, this channel knowledge can either be explicit estimates of $\mathbf{h}_{\text{DL}}$ and $\mathbf{h}_{\text{UL}}$ for each user pair or implicit knowledge of some form, such as raw beam alignment measurements or angles-of-departure/arrival. We denote this user channel information across all K user pairs by $\{(\mathbf{y}_{\text{DL}}^{(k)}, \mathbf{y}_{\text{UL}}^{(k)})\}_{k=1}^K$. Naturally, gathering this channel knowledge would need to be repeated commensurate with the coherence time of the users' channels. We remark that, by not mandating a particular strategy for obtaining this user channel knowledge, our proposed scheme seeks to maintain backward compatibility with existing standards and facilitate its adoption in real-world networks.

B. Self-Interference Probing

The next stage of our proposed solution is to collect M measurements across the self-interference channel $\mathbf{H}$. To do so, we design M probing beam pairs $\{(\mathbf{f}_m, \mathbf{w}_m)\}_{m=1}^M$, where $(\mathbf{f}_m, \mathbf{w}_m)$ are the transmit and receive beams used by the base

station to collect the m -th measurement of self-interference, expressed as

$$z_m = \sqrt{\frac{P_{\text{DL}}}{N_t}} \mathbf{w}_m^* \mathbf{H} \mathbf{f}_m + \mathbf{w}_m^* \mathbf{n}_m, \quad (13)$$

where $\mathbf{n}_m \sim \mathcal{N}_{\mathbb{C}}(\mathbf{0}, \sigma_{\text{UL}}^2 \mathbf{I})$ denotes complex Gaussian noise at the receive array. All M measurements of self-interference can be written compactly in vector form as

$$\mathbf{z} = \sqrt{\frac{P_{\text{DL}}}{N_t}} \text{diag}(\mathbf{W}^* \mathbf{H} \mathbf{F}) + \text{diag}(\mathbf{W}^* \mathbf{N}) \in \mathbb{C}^{M \times 1}, \quad (14)$$

with $\mathbf{z} = [z_1 \cdots z_M]^T$, $\mathbf{N} = [\mathbf{n}_1 \cdots \mathbf{n}_M]$, and the transmit and receive probing beams stacked into the matrices

$$\mathbf{F} = [\mathbf{f}_1 \cdots \mathbf{f}_M] \in \mathbb{C}^{N_t \times M}, \quad (15)$$

$$\mathbf{W} = [\mathbf{w}_1 \cdots \mathbf{w}_M] \in \mathbb{C}^{N_r \times M}, \quad (16)$$

which we refer to as the transmit and receive *probing codebooks*. Note that these M measurements of self-interference will need to be recollected once $\mathbf{H}$ changes, which we assume is after the K user pairs have been served.

Core to our approach is its tailoring of the probing codebooks $\mathbf{F}$ and $\mathbf{W}$ to the particular K user pairs that the base station intends to serve. As detailed in the next section, this is carried out by a deep learning model which performs a mapping $\mathcal{P}(\cdot)$ from user channel knowledge to probing codebooks, i.e.,

$$(\mathbf{F}, \mathbf{W}) = \mathcal{P}(\mathbf{Y}), \quad (17)$$

where

$$\mathbf{Y} \triangleq \left\{ \left(\mathbf{y}_{\text{DL}}^{(k)}, \mathbf{y}_{\text{UL}}^{(k)} \right) \right\}_{k=1}^K. \quad (18)$$

Even though the self-interference channel $\mathbf{H}$ itself does not depend on the particular users being served, the amount of self-interference that is ultimately coupled does, since the base station will form its transmit and receive serving beams based on the users' channels. While explicitly estimating the entire matrix $\mathbf{H}$ would provide knowledge of the self-interference coupled by *any* possible serving beams $(\mathbf{f}, \mathbf{w})$, we hypothesize that focusing only on the portion of $\mathbf{H}$ that is relevant to these specific users will allow the base station to attain comparable performance with substantially fewer measurements compared to obtaining full knowledge of $\mathbf{H}$. We confirm this to indeed be true in our performance evaluation in Section VII.

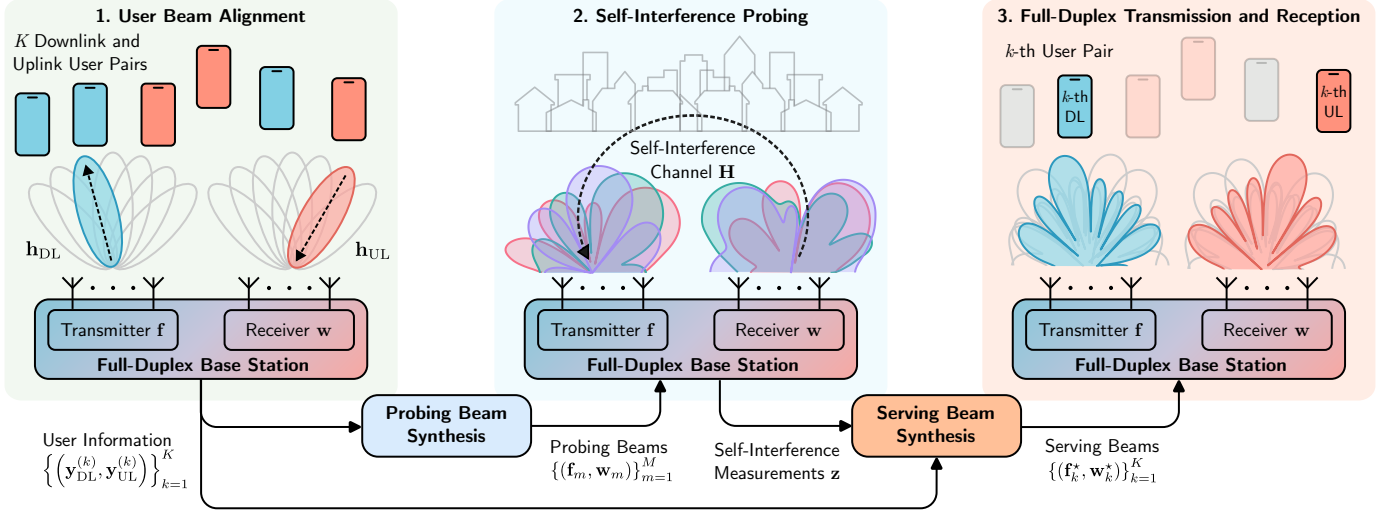


Fig. 3. Our proposed approach consists of three stages. First, beam alignment is used to obtain knowledge of the downlink and uplink users' channels. Then, using this user channel information, a sequence of M probing beam pairs are designed to collect M measurements across the self-interference channel. Finally, these measurements are used along with the user information to synthesize transmit and receive beams to serve user pairs in a full-duplex fashion.

C. Serving Beam Synthesis

The final step of our proposed approach lies in generating the transmit and receive beams used by the base station to serve each of the K user pairs in a full-duplex fashion. Like the probing beams, the serving beams will be designed using deep learning via a policy $\mathcal{S}(\cdot)$, which takes as input the self-interference probing measurements $\mathbf{z}$ and the user channel knowledge $\mathbf{Y}$, i.e.,

$$\{(\mathbf{f}_k^*, \mathbf{w}_k^*)\}_{k=1}^K = \mathcal{S}(\mathbf{z}, \mathbf{Y}). \quad (19)$$

As mentioned, these serving beams must deliver high SNR to the users while coupling low self-interference across $\mathbf{H}$ in order to enable full-duplex operation. To design such, we will jointly optimize $\mathcal{P}(\cdot)$ and $\mathcal{S}(\cdot)$ in order to design probing beams that gather implicit knowledge of $\mathbf{H}$ that is the most informative for synthesizing optimal serving beams.

D. Problem Formulation

To formalize the design of the probing beam policy $\mathcal{P}(\cdot)$ and serving beam policy $\mathcal{S}(\cdot)$ based on our stated goals, we assemble the following optimization problem.

$$\max_{\mathcal{P}, \mathcal{S}} \mathbb{E}_{\mathbf{H}, \mathbf{Y}} \left[\sum_{k=1}^K \bar{R}_k(\mathbf{f}_k^*, \mathbf{w}_k^*) \right] \quad (20a)$$

$$\text{s.t. } (\mathbf{F}, \mathbf{W}) = \mathcal{P}(\mathbf{Y}) \quad (20b)$$

$$\{(\mathbf{f}_k^*, \mathbf{w}_k^*)\}_{k=1}^K = \mathcal{S}(\mathbf{z}, \mathbf{Y}) \quad (20c)$$

$$\mathbf{z} = \sqrt{\frac{P_{\text{DL}}}{N_t}} \text{diag}(\mathbf{W}^* \mathbf{H} \mathbf{F}) + \text{diag}(\mathbf{W}^* \mathbf{N}) \quad (20d)$$

$$\|\mathbf{f}_k^*\|_{\max} \leq 1, \quad k = 1, \dots, K \quad (20e)$$

$$\|\mathbf{f}_m\|_{\max} \leq 1, \quad m = 1, \dots, M \quad (20f)$$

Here, our objective function has been defined as the expected sum of the *normalized* spectral efficiency across user pairs

rather than the raw spectral efficiency as in (12). This normalized SSE is defined as

$$\bar{R}_k(\mathbf{f}_k^*, \mathbf{w}_k^*) = \frac{R_k(\mathbf{f}_k^*, \mathbf{w}_k^*)}{C_k}, \quad (21)$$

with C_k the interference-free sum capacity of the k -th user pair, defined as

$$C_k = \log_2 \left(1 + \frac{P_{\text{DL}} \|\mathbf{h}_{\text{DL}}^{(k)}\|_2^2}{\sigma_{\text{DL}}^2} \right) + \log_2 \left(1 + \frac{P_{\text{UL}} \|\mathbf{h}_{\text{UL}}^{(k)}\|_2^2}{\sigma_{\text{UL}}^2} \right). \quad (22)$$

This normalization is performed to promote fairness, so that $\mathcal{P}(\cdot)$ and $\mathcal{S}(\cdot)$ are not biased toward favoring user pairs with stronger channels. Note that the expectation of the objective function is over the distributions of $\mathbf{H}$ and $\mathbf{Y}$. This will ensure that a single design of $\mathcal{P}(\cdot)$ and $\mathcal{S}(\cdot)$ will be useful regardless of the particular self-interference channel $\mathbf{H}$ and user grouping at any given time.

In problem (20), line (20b) tailors the M probing beams to the group of K user pairs to be served; line (20d) represents the M probing measurements of self-interference taken across $\mathbf{H}$; line (20c) synthesizes the K beam pairs used to serve the K user pairs; and finally lines (20e)–(20f) enforce per-antenna power constraints on the transmit beams. Recognize that this problem would be solved for a fixed number of measurements M and number of user pairs K , both of which may be tuned to optimize performance based on operating conditions, as we will see in our performance evaluation in Section VII. In the next section, we aim to solve problem (20) through a transformer-based deep learning implementation.

Note that, by limiting the scope of problem (20) to a single base station, $\mathcal{P}(\cdot)$ and $\mathcal{S}(\cdot)$ can be designed in a *site-specific* fashion, tailoring them to the particular self-interference distribution and user distribution seen by that base station—both of which depend on features of the surrounding environment, such as buildings, vehicles, foliage,

etc. In practice, this would be accomplished by training our model using channel data from measurements, a statistical model, a ray-tracing simulator, and/or digital twin; in the performance evaluation in Section VII, we use ray-tracing. While not strictly necessary, this site-specific approach has the potential to reduce the number of measurements M required for appreciable performance, by limiting the distributions of $\mathbf{H}$ and the user channels to what is seen within a particular environment, allowing the deep learning model to exploit their relations within that same environment.

V. DEEP LEARNING IMPLEMENTATION

In this section, we detail the deep learning implementation that we have developed to solve problem (20) in a site-specific fashion. Corresponding to the three components discussed in Section IV, our proposed deep learning framework consists of three modules: the *user encoder*, the *probing beam synthesizer*, and the *serving beam synthesizer*, as illustrated in Fig. 4. All three modules are transformer-based [44], as detailed next.

A. User Encoder

The first module is the *user encoder*, which uses a transformer encoder to act as a dedicated pre-processor that maps the user information $\mathbf{Y}$ to a *user embedding* tensor $\mathbf{E}_U$ used in subsequent tasks. This embedding can be thought of as a compact, real-valued representation of the information about the users' channels that is most relevant in designing the probing and final serving beams, with each user pair's embedding having a fixed dimension D_U .

Since each of the K entries of $\mathbf{Y}$ consists of two complex-valued vectors with dimensions equal to the number of transmit and receive antennas, they need to be converted to fixed-length, real-valued inputs $\mathbf{Y}_R = \{\mathbf{y}_R^{(k)}\}_{k=1}^K$ before being ingested by the encoder. To do so, the real and imaginary components of $\mathbf{y}_{DL}$ and $\mathbf{y}_{UL}$ are concatenated jointly along the array dimension to form a single real-valued vector of length $2(N_t + N_r)$, i.e.,

$$\mathbf{y}_R^{(k)} = \left[\Re\{\mathbf{y}_{DL}^{(k)}\}^T, \Im\{\mathbf{y}_{DL}^{(k)}\}^T, \Re\{\mathbf{y}_{UL}^{(k)}\}^T, \Im\{\mathbf{y}_{UL}^{(k)}\}^T \right]^T.$$

A linear projection layer subsequently reduces the dimension of each $\mathbf{y}_R^{(k)}$ from $2(N_t + N_r)$ to D_U to form the initial user embeddings. The transformer encoder then iteratively refines $\mathbf{E}_U$ through its self-attention layers, which selectively aggregate channel information across all K user pairs. The resulting $\mathbf{E}_U$ can thus contain valuable knowledge of not just the users' channels but the correlations across the coherent user group, which can be exploited when designing the probing and serving beams.

Note that, by operating on the fixed-dimension $\mathbf{E}_U$ rather than $\mathbf{Y}_R$ directly, the complexity of the transformer is independent of the number of transmit/receive antennas. Consequently, a single user encoder can accommodate varying antenna array sizes, in which separate, lightweight projection layers are used to map the user information from different array sizes to embeddings of fixed size D_U , which can be processed by the same transformer backbone and trained jointly across a

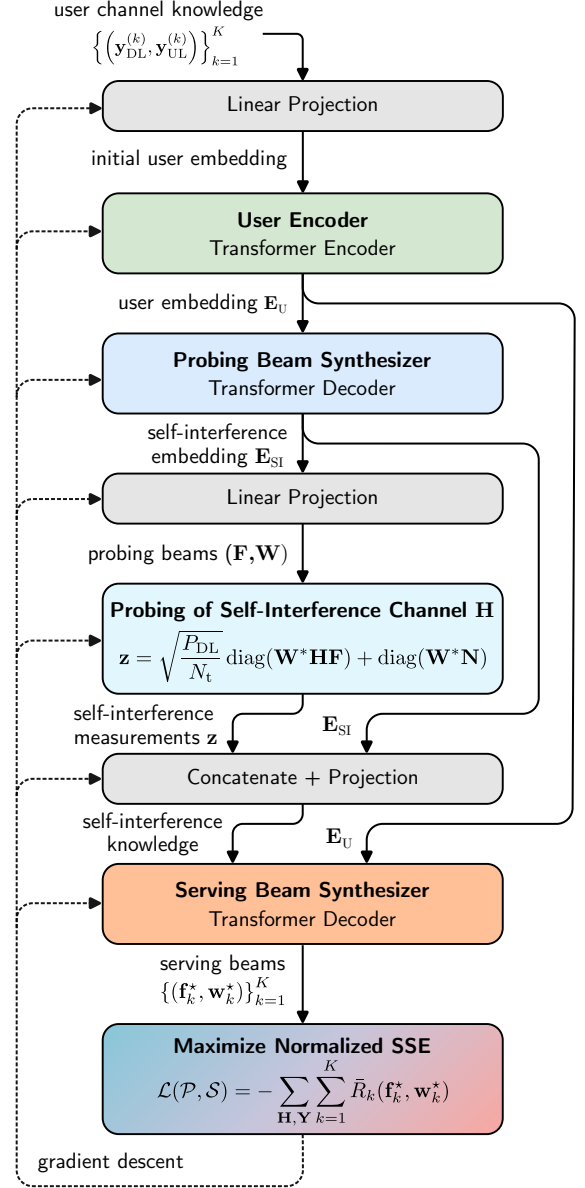


Fig. 4. Proposed transformer-based deep learning model to realize our envisioned site-specific beamforming solution. The entire model is trained in an end-to-end fashion to jointly design the probing beam synthesizer and serving beam synthesizer.

broad dataset of many different configurations under a variety of conditions. Additionally, no positional encoding or causal masking is applied within the self-attention layers, since all K pairs of user channels are available to the user encoder from the start, and their ordering does not carry semantic meaning in this context, so these embeddings should be attended to in parallel and in a permutation-invariant manner.

B. Probing Beam Synthesizer

After the user encoder, the *probing beam synthesizer* is responsible for realizing the probing policy $\mathcal{P}(\cdot)$. It conditions on the user embedding $\mathbf{E}_U$ to generate the self-interference probing beam codebooks $(\mathbf{F}, \mathbf{W})$ for the particular coherent user group. At the core of the probing synthesizer is a transformer decoder, which generates a set of M self-interference

embeddings $\mathbf{E}_{\text{SI}}$, each corresponding to a particular pair of probing beams $(\mathbf{f}_m, \mathbf{w}_m)$ to be generated.

Since the base station has no instantaneous knowledge of the particular $\mathbf{H}$ before probing, it must rely on its prior understanding of the channel's distribution, which is learned during training as will be discussed shortly. Alongside the model weights, this understanding is also captured through a learnable embedding table from which the initial M embedding vectors are drawn, thus providing a unique starting point for each of the M probing beam pairs.

The self-interference embeddings are subsequently passed through the transformer decoder with alternating self-attention and cross-attention layers. The cross-attention layer takes $\mathbf{E}_{\text{U}}$ as the key-value memory, pivoting each self-interference embedding to focus on the components of $\mathbf{H}$ most relevant to these particular users. The self-attention then interchanges information among $\mathbf{E}_{\text{SI}}$, so that probing beams yield *complementary* measurements rather than overlapping ones. Finally, a projection layer maps the m -th self-interference embedding to the real and imaginary components of the corresponding probing beams to construct $(\mathbf{f}_m, \mathbf{w}_m)$.

Since the output of the final projection layer does not necessarily satisfy the per-antenna power constraint in (20e), we explicitly normalize each transmit probing beam. To do so, each transmit probing beam vector is scaled by its maximum entry magnitude, e.g.,

$$[\mathbf{f}_m]_i = \frac{[\tilde{\mathbf{f}}_m]_i}{\max_j |[\tilde{\mathbf{f}}_m]_j|}, \quad (23)$$

where $\tilde{\mathbf{f}}_m$ is the m -th unnormalized transmit probing beam output by the model and $\mathbf{f}_m$ is its normalized version that populates the m -th column of $\mathbf{F}$. We observed that other normalizations, such as simple clipping, can also be used and offer similar performance. The normalized probing beam codebooks $\mathbf{F}$ and $\mathbf{W}$ are then used to collect M measurements of self-interference, populating $\mathbf{z}$ according to (14).

In a similar manner to the user encoder, the probing beam pairs in $\mathbf{F}$ and $\mathbf{W}$ are generated in parallel: the model generates the m -th beam pair $(\mathbf{f}_m, \mathbf{w}_m)$ conditioned only on $\mathbf{E}_{\text{U}}$ and not measurements from “preceding” beams, e.g., $(z_1, \dots, z_{m-1})$. Rather than generate probing beam pairs autoregressively, like in [23], this parallel approach was done deliberately to avoid the latency associated with processing and inference between measurements, thereby allowing for the probing measurements to be collected in rapid succession.

C. Serving Beam Synthesizer

The final module of our proposed solution is the *serving beam synthesizer*, which realizes the beam synthesis policy $\mathcal{S}(\cdot)$. It takes as input the user embedding $\mathbf{E}_{\text{U}}$, the self-interference embedding $\mathbf{E}_{\text{SI}}$, and the probing measurements $\mathbf{z}$, and generates the final serving beams $\{(\mathbf{f}_k^*, \mathbf{w}_k^*)\}_{k=1}^K$ for the coherent user group. It is implemented as a transformer decoder similar to the probing beam synthesizer, but with subtle differences to its input and output structure. Here, the user embedding $\mathbf{E}_{\text{U}}$ is supplied directly as the initial embedding into the transformer decoder. By attending to all K

user embeddings, the self-attention layers leverage both the individual channel information as well as the overall propagation environment shared across the coherent user group to provide high SNRs to the users. Meanwhile, the real and imaginary components of $\mathbf{z}$ are concatenated with $\mathbf{E}_{\text{SI}}$ and projected as the key-value memory to the cross-attention layers, which condition the beam embeddings on the implicit knowledge of $\mathbf{H}$ to steer the final serving beams away from self-interference. Finally, a projection layer maps the final beam embeddings to the real and imaginary components of the serving beam pairs, which are then normalized according to (23) to produce the final beams $\{(\mathbf{f}_k^*, \mathbf{w}_k^*)\}_{k=1}^K$. Like the probing beams, the serving beams are generated in a non-autoregressive manner, in parallel, and without positional encoding.

D. Site-Specific Unsupervised Training

Our entire deep learning model is trained end-to-end in an unsupervised manner to maximize the sum normalized spectral efficiency via the loss function

$$\mathcal{L}(\mathcal{P}, \mathcal{S}) = - \sum_{\mathbf{H}, \mathbf{Y}} \sum_{k=1}^K \bar{R}_k(\mathbf{f}_k^*, \mathbf{w}_k^*). \quad (24)$$

Here, batching is done over random user groupings and self-interference channel realizations to approximate the expectation in (20a).

A key aspect of our proposed approach is its *site-specific* design, which is accomplished by training on user channel data and self-interference channel data collected in a particular environment. During training, the model learns to exploit the user distribution and unique environmental features such as scatterers and blockage to optimize the probing and beam synthesis policies. At inference time, this site-specific knowledge is embedded in the trained model to produce effective probing beams and serving beams.

It is worth noting that our transformer-based architecture allows a single model to generalize across a variety of configurations and operating conditions, in addition to antenna count, as described in Section V-A. For instance, since K and M affect the sequence *lengths*, e.g. the number of vectors within their embeddings, rather than the *dimensions* of these vectors, a single trained model can be used for various sized coherent user groups and probing budgets. We take advantage of this via a two-stage training process where a foundational model is first pretrained broadly across diverse configurations, then fine-tuned to the specific deployment setting at a fraction of the cost of training from scratch.

VI. SIMULATION SETUP AND BASELINES

To evaluate the performance of our proposed approach, we simulate a full-duplex massive MIMO base station operating at a carrier frequency of 28 GHz. The base station employs separate transmit and receive antenna arrays equipped with analog beamforming. The arrays are placed on the same plane facing the center of the user distribution, with their centers separated horizontally by 10 wavelengths. Unless otherwise specified, we use 4×4 half-wavelength uniform planar arrays

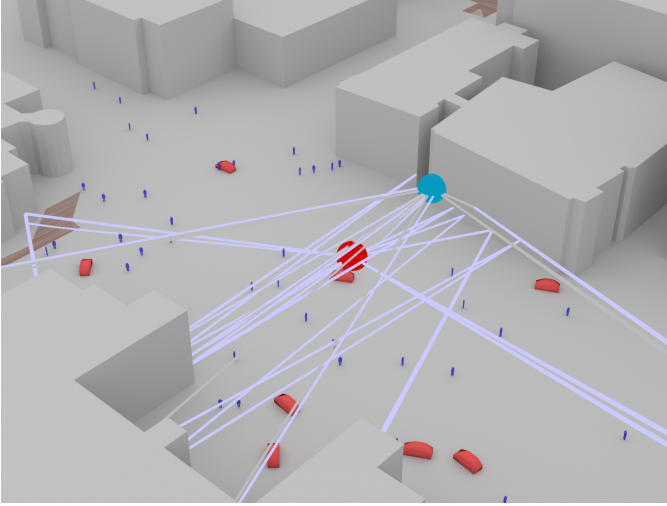


Fig. 5. A sample scene realization based on a reconstructed 3D model of Westwood Plaza on the UCLA campus, derived from OpenStreetMap [45] data. The cyan and red spheres indicate the locations of the base station atop a building and a downlink user in the courtyard, respectively. The blue figures represent pedestrians and the red objects represent vehicles, both randomly placed in the courtyard. The purple lines indicate representative ray-traced propagation paths between the base station and the user equipment. For readability, the rest of the downlink and uplink users are omitted from visualization but remain present in the simulation.

(UPAs) for both the transmitter and receiver, yielding $N_r = N_t = 16$ elements; we investigate performance with larger arrays shortly.

In an effort to provide realistic performance insights, we perform ray-traced channel modeling using NVIDIA's Sionna Ray Tracing library [46]. As shown in Fig. 5, we conduct simulations in a reconstructed 3D scene of Westwood Plaza on the UCLA campus. The base station is mounted in front of a building at a height of 20 meters. During each scene realization, the street level is populated with vehicles and pedestrians whose locations are sampled from a uniform distribution. These objects act as dynamic blockage and scattering bodies, while the surrounding buildings and streets act as a static backdrop of the propagation environment. A random subset of the pedestrians are designated as downlink or uplink users at a height uniformly sampled between 1 to 1.7 meters.

Given the geometry of the scene, Sionna computes the signal propagation paths through differentiable ray tracing. For added richness and realism, we randomly discard 10% of the output paths to mimic the effects of foliage blockage, small-scale fading, and other phenomena not captured by the ray-tracing simulator. The remaining propagation paths are summed together to generate narrowband downlink and uplink channels. This process is repeated multiple times, each with unique vehicle and user placements, in order to populate a dataset of channels that can be used to train and evaluate our deep learning model. To emulate beam management procedures in practical 5G networks [26], we assume that the base station only has knowledge of the most dominant (often LOS) path to each user, rather than the full channel. This comprises the user channel information $\mathbf{Y}$ that is input to our deep learning model.

TABLE I
MODEL HYPERPARAMETERS AND TRAINING DETAILS

Parameter	Value
Embedding dimensions D_U, D_{SI}	320
Attention heads	5
User encoder layers	3
Probing beam synthesizer layers	2
Serving beam synthesizer layers	2
Optimizer	AdamW [48]
Weight decay	0.01
Pretraining learning rate	5×10^{-5}
Pretraining LR schedule	Cosine-annealed to 5×10^{-6}
Fine-tuning LR schedule	1cycle policy [49]

Consistent with (4), our model of the self-interference channel $\mathbf{H}$ consists of a LOS component and a NLOS component. The NLOS component is produced by Sionna according to the environment's reflective components in each scene realization, and the LOS component is simulated separately using the near-field spherical-wave model [47]. These two components are then combined according to (4), with κ determining their relative strengths. We will use $\kappa = 0$ dB by default, but will sweep its value across a wide range to assess performance.

We use a lightweight implementation of the transformer-based architecture detailed in Section V, which consists of approximately 10 million parameters. The implementation details and training hyperparameters are summarized in Table I. During each experiment, the model is first pretrained across all combinations of K , M , and κ , and then fine-tuned on a specific configuration before evaluating performance. We found this fine-tuning to generally improve performance modestly, i.e., by a few percent. To standardize training and evaluation, the mean downlink and uplink SNR upper bounds are normalized to 10 dB, and the mean uplink INR upper bound is normalized to 40 dB. To focus exclusively on the effects of self-interference, we assume cross-link interference is negligible, as discussed in Section II.

We compare the performance of our proposed approach against three relevant baselines, two of which are upper bounds on performance under differing assumptions:

- *LMMSE*: In this baseline, the base station performs maximum ratio transmission (MRT) at its transmitter based on the same partial knowledge of the users' channels as our proposed scheme (from beam alignment). Then, the receiver scans a N_r -point discrete Fourier transform (DFT) codebook, and performs linear minimum mean square error (LMMSE) estimation of the *effective* self-interference channel vector $\mathbf{h}_{SI} = \mathbf{H}\mathbf{f} \in \mathbb{C}^{N_r \times 1}$. The receive beam is then designed as the LMMSE combiner in order to maximize uplink SINR. Because $\mathbf{h}_{SI}$ depends on the transmit beam $\mathbf{f}$, channel estimation associated with this LMMSE baseline must be repeated for each user pair, meaning it requires $K N_r$ measurements in total to serve all K user pairs.

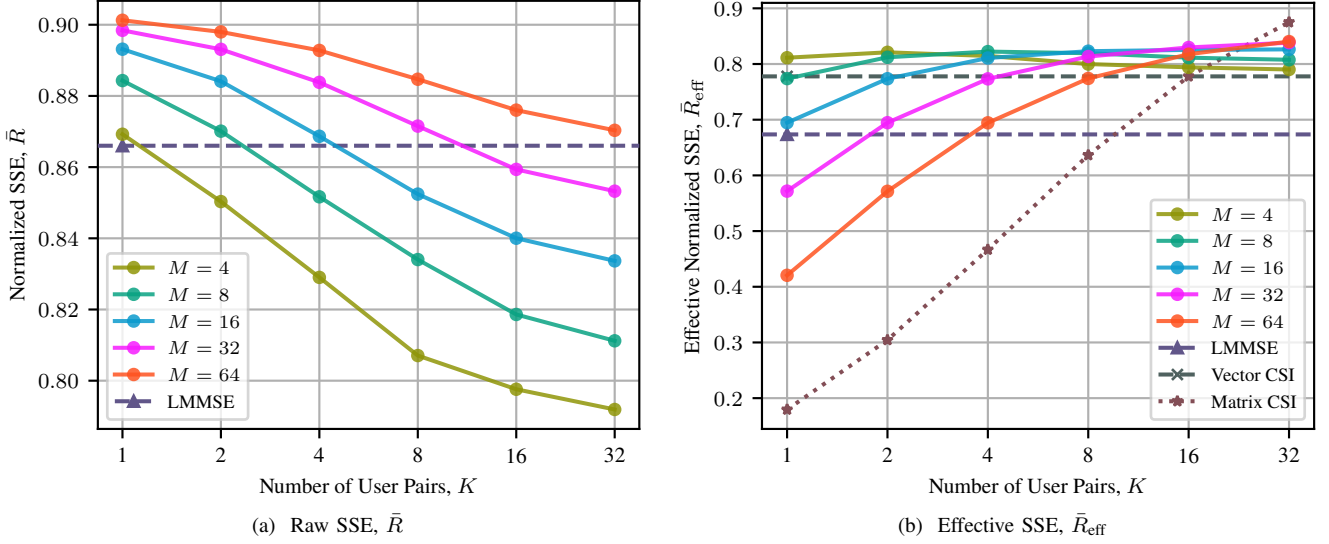


Fig. 6. Raw and effective normalized SSE as a function of the coherent user group size K for various probing budgets M , with $\kappa = 0$ dB. In (a), the Vector CSI and Matrix CSI baselines attain a raw SSE of $\bar{R} = 1$. In (b), each user is served for $L = 56$ time slots.

- **Vector CSI:** This baseline assumes the base station collects N_t measurements to estimate $\mathbf{h}_{\text{SI}}$ for each user pair, after which it attains the full-duplex capacity. By attaining the capacity, its raw normalized SSE is $\bar{R} = 1$, and its effective SSE is equal to $\bar{R}_{\text{eff}} = \frac{KL}{KN_t + KL}$. This represents *best-case* performance of any scheme which explicitly estimates $\mathbf{h}_{\text{SI}}$ for each user pair, including the LMMSE baseline.
- **Matrix CSI:** Analogous to the above Vector CSI baseline, this baseline assumes the base station attains the full-duplex capacity after collecting $N_t N_r$ measurements to estimate the entire matrix $\mathbf{H}$. Since this attains the capacity, its raw normalized SSE is $\bar{R} = 1$, and its effective SSE is equal to $\bar{R}_{\text{eff}} = \frac{KL}{N_t N_r + KL}$. This represents *best-case* performance of any scheme which explicitly estimates $\mathbf{H}$ from $N_t N_r$ measurements.

VII. PERFORMANCE EVALUATION

In the results that follow, we assess performance using the raw normalized SSE

$$\bar{R} = \frac{1}{K} \sum_{k=1}^K \bar{R}_k(\mathbf{f}_k^*, \mathbf{w}_k^*), \quad (25)$$

along with the *effective* normalized SSE

$$\bar{R}_{\text{eff}} = \frac{KL}{M + KL} \bar{R}, \quad (26)$$

which accounts for the measurement overhead associated with acquiring self-interference channel knowledge. Note that, in both, the spectral efficiency has been normalized to the true full-duplex capacity, based on perfect user channel information rather than the partial knowledge provided during initial beam alignment.

A. Coherent User Group Size, K

To begin our performance evaluation, we first investigate the effect of the coherent group size K ; recall, K represents the number of user pairs over which the self-interference channel $\mathbf{H}$ remains static and can thus be thought of as a proxy for its coherence time. In Fig. 6, we fix $\kappa = 0$ dB and plot both the raw and effective normalized SSE as a function of K for various numbers of probing measurements M . We also plot performance of the three baselines. Since $N_t = N_r = 16$ in this case, the LMMSE and Vector CSI baselines each require $KN_t = 16K$ measurements and the Matrix CSI baseline requires $N_t N_r = 256$ measurements.

In Fig. 6a, we see that when $K = 1$ user pair, our proposed scheme can outperform LMMSE in terms of raw SSE with only $M = 4$ probing measurements, representing a savings of 75% in measurement overhead. When M is fixed and K is increased, our scheme's performance naturally degrades, since the collected measurements are forced to generalize across more user pairs. The measurements thus become less informative for any individual user pair, degrading the model's ability to reliably cancel self-interference. Fixing K and increasing the number of measurements M , on the other hand, provides more informative implicit channel knowledge and in turn yields an increase in raw SSE. Note that the raw SSE of LMMSE does not vary as K is increased, since it is executed individually on each user pair; this comes at the cost of its measurement overhead scaling linearly with K .

To account for the overhead associated with more measurements, Fig. 6b repeats this evaluation in terms of the *effective* SSE with each user pair served for $L = 56$ time slots², where

²In 5G NR, a single slot is 14 OFDM symbols, and with analog beamforming, each probing measurement consumes an entire OFDM symbol. If each user is served for 2–8 slots, then this corresponds to $L = 28$ to $L = 112$ symbols. Considering operation at 28 GHz (FR2) with a subcarrier spacing of 120 kHz, then $L = 56$ symbols corresponds to around 500 μs . If the coherence time of $\mathbf{H}$ is around 4 ms, for instance, then $K = 8$ user pairs could be served from a single set of probing measurements $\mathbf{z}$.

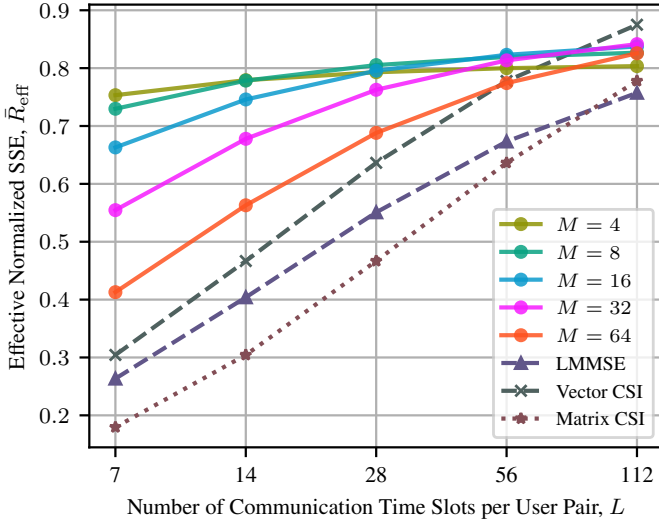


Fig. 7. Effective normalized SSE $\bar{R}_{\text{eff}}$ as a function of the number of communication time slots per user L for various probing budgets M , with $\kappa = 0$ dB and $K = 8$ user pairs.

a more nuanced trend emerges. It can be seen that increasing M can only be justified when K also increases, i.e., when the coherence time of $\mathbf{H}$ becomes longer, because this amortizes the measurement overhead across multiple user pairs. Our proposed approach outperforms all three baselines across the board, except for Matrix CSI when $K = 32$ which corresponds to a long coherence time; recall this is the upper bound with perfect knowledge of the downlink and uplink channels. This superiority indicates that our proposed scheme exhibits a more favorable trade-off between measurement overhead and raw SSE than the baselines. In other words, for a given number of measurements, our scheme tends to offer a higher raw SSE, thus yielding a higher effective SSE.

B. Communication Time Slots, L

The net cost of measurement overhead, captured by the fraction $\frac{KL}{M+KL}$, depends on the number of time slots L each user pair is served. This motivates our next evaluation in Fig. 7, which shows the effective SSE as a function of L for $K = 8$ users. With smaller L , the measurement overhead magnifies, and the best effective performance comes from fewer measurements. For this reason, the three baselines exhibit poorer performance overall, given LMMSE and Vector CSI require 16 measurements per user and Matrix CSI requires 256 total measurements. Taken together with Fig. 6b, these results suggest that optimal deployment of our proposed solution involves choosing M and K based on the coherence time of the self-interference channel relative to the duration in which users are served.

C. Average Probing Cost, M/K

Our approach collects M measurements at once and reuses the information contained in these measurements across all K user pairs. It is interesting to ask whether this is superior to an approach where the M measurements are divided up across

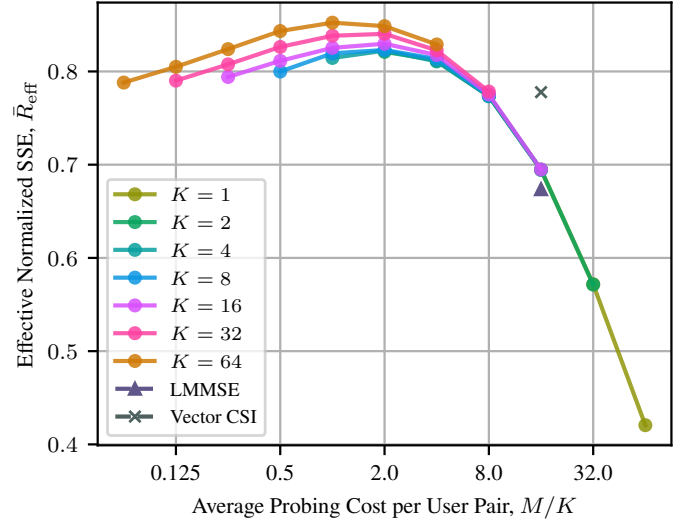


Fig. 8. Effective normalized SSE as a function of the average probing cost per user pair M/K for various coherence group sizes K . Each user is served for $L = 56$ time slots.

smaller groups of users, allowing each subset of measurements to focus on fewer user pairs. To investigate this, Fig. 8 plots the effective SSE as a function of M/K , the average probing cost per user pair.

Let us first consider the case of $M/K = 1$, meaning the number of measurements M equals the total number of user pairs K . If $M = K = 4$, for example, then 4 user pairs would be grouped together and 4 measurements would be spent per user group, for an effective cost of one probing measurement per user pair. If $M = K = 64$, this would yield the same per-user-pair probing cost, but would be accomplished with groupings of 64 user pairs and spending 64 measurements per group. While these two groupings exhibit the same relative overhead, Fig. 8 indicates that a larger grouping results in greater performance. This suggests that user pairs within a group stand to benefit from sharing with one another the implicit channel knowledge contained in $\mathbf{z}$. This can be explained by the fact that there exist correlations across the users' channels, and this means certain portions of the self-interference channel $\mathbf{H}$ may be of relevance to multiple user pairs within a given user group. This is most pronounced when M/K is low, since this is where performance is constrained by its limited self-interference channel knowledge. These results suggest that, in practice, it is best to group user pairs into larger groups and perform a single sweep of M probing measurements; in fact, it would be best to group as many user pairs as possible, provided the coherence time of $\mathbf{H}$ allows for such.

D. Structure of the Self-Interference Channel, κ

We now assess performance as a function of the self-interference channel structure. To do so, we vary its Rician factor κ from -20 dB to 20 dB in Fig. 9, where increasing κ strengthens the LOS component of the self-interference channel. We observe that, as κ is increased, the channel becomes less favorable structurally, as indicated by the decreas-

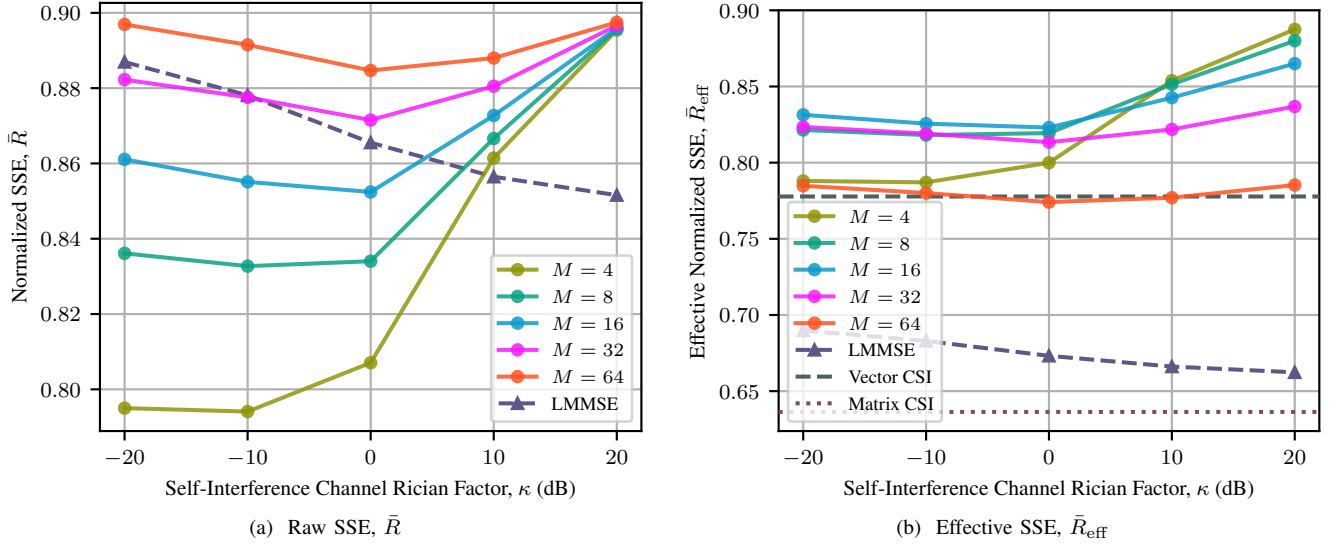


Fig. 9. Raw and effective normalized SSE as a function of the self-interference channel Rician factor κ for various probing budgets M , with $K = 8$ user pairs. In (b), each user is served for $L = 56$ time slots.

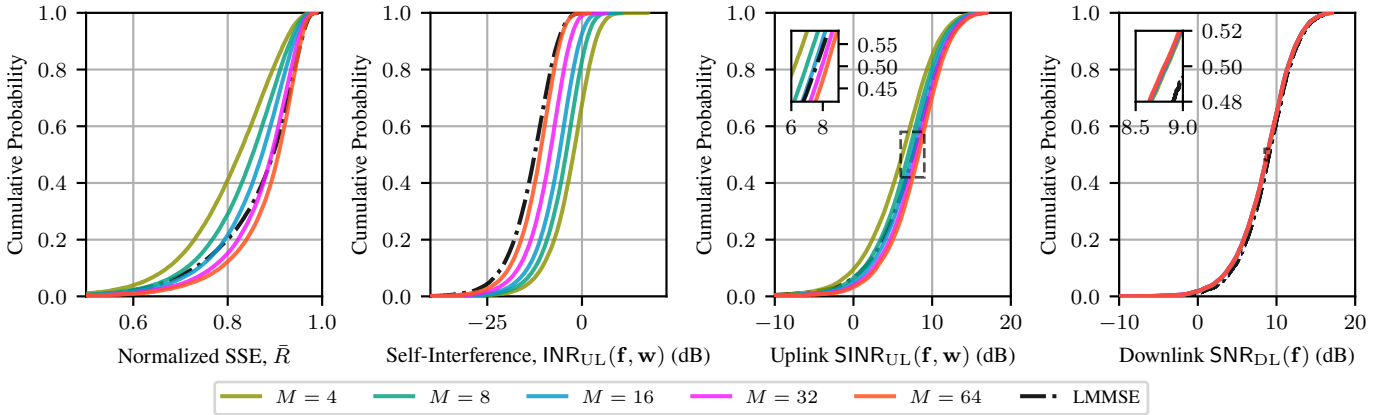


Fig. 10. CDFs of the normalized SSE $\bar{R}$, self-interference, uplink SINR, and downlink SNR for various probing budgets M , with $K = 8$ user pairs and $\kappa = 0$ dB.

ing performance in LMMSE. As the channel becomes more LOS-dominant, however, it also becomes more deterministic, allowing our proposed scheme to perform well with only a few probing measurements.

The net effect of this can be seen in the effective SSE of Fig. 9b. At low κ , a modest number of measurements is preferable, e.g., $M = 16$ and $M = 32$, since inferring a richer channel benefits from additional measurements. As κ grows, the raw performance gain of additional probing diminishes, since the channel becomes more deterministic, allowing performance with only a few measurements to climb and eventually overtake the case of heavy probing. An exception to this trend is the $M = 64$ case which offers the worst $\bar{R}_{\text{eff}}$ across the sweep, as its marginal gain in raw SSE does not justify its high overhead, in this particular case.

E. Performance Distribution

The CDFs in Fig. 10 provide a more detailed view of several relevant metrics in order to further explore the performance

distribution. As established before, increasing the number of measurements M reliably increases the raw SSE, and this is thanks to an improved ability to cancel self-interference. This results in an increase in uplink SINR that actually exceeds LMMSE, which can be justified by the fact that our proposed scheme makes minor sacrifices in downlink SNR, unlike LMMSE. Clearly, our proposed scheme is better at optimizing this trade-off between downlink and uplink to maximize SSE—a non-convex problem—compared to the LMMSE baseline.

F. Antenna Array Size, N_t and N_r

Finally, Fig. 11 examines how performance scales with the size of the base station's transmit and receive antenna arrays, which we assume are identical in size; while changing the arrays' size, we keep all other parameters constant, including the total transmit power. We begin with 4×4 UPAs at the transmitter and receiver and then scale up to 16×16 UPAs. This corresponds to self-interference channel matrices that range

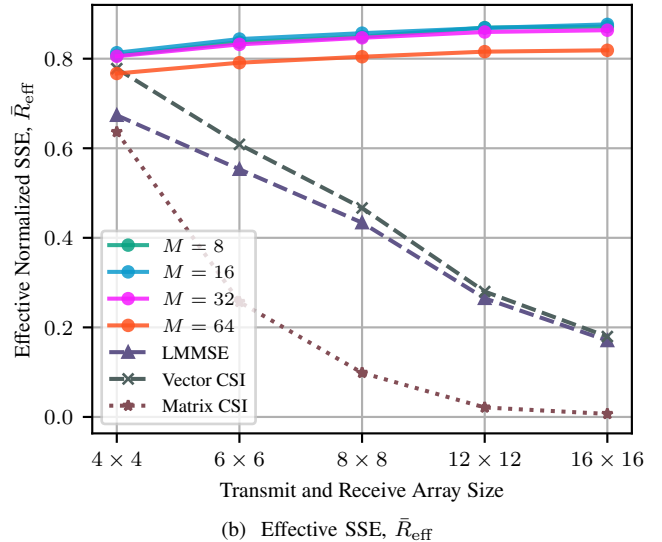
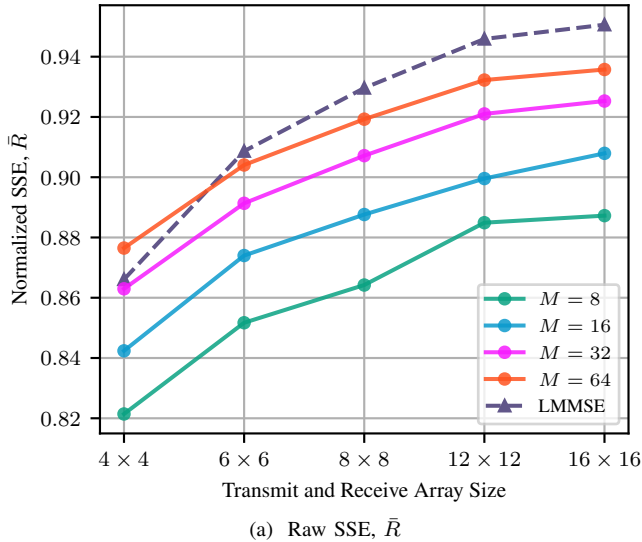


Fig. 11. Raw and effective normalized SSE as a function of the transmit and receive array size for various probing budgets M , with $K = 8$ user pairs and $\kappa = 0$ dB. In (b), each user is served for $L = 56$ time slots.

from $16^2 = 256$ entries to $256^2 = 65,536$ entries, which illustrates the high measurement overhead associated with the Matrix CSI baseline.

The raw SSE shown in Fig. 11a improves with the array size for the proposed scheme and the baselines, reflecting the beamforming gains afforded by larger arrays—both in terms of SNR gain and interference reduction. Once the probing overhead is accounted for in Fig. 11b, however, our proposed scheme diverges sharply from the baselines. The effective SSE offered by all three baselines plummet as the antenna count increases, since their overhead scales either linearly or quadratically with $N_t = N_r$. Performance with our proposed approach is relatively steady, consistent with the behavior in Fig. 11a, further demonstrating our model’s ability to efficiently gather and harness *implicit* channel knowledge. Since the makeup of the NLOS component of the self-interference channel—e.g., prominent reflectors and blockage—does not scale with the number of antennas, one would not expect the required number of measurements to necessarily increase, which is indeed the case with our approach. It is also worth noting that this increase in antenna count does not come at the cost of substantial increases in computational complexity, since our model’s transformer-based construction is based on a fixed embedding size. Overall, these results demonstrate that our proposed approach scales gracefully to large antenna arrays, making it an attractive solution for the coming 6G era as networks pursue base stations with an ever-increasing number of antennas.

VIII. CONCLUSION

This work introduces a novel approach to design the transmit and receive beams of a full-duplex massive MIMO base station. A particularly unique aspect of our proposed approach is that it designs such beams based on *implicit* knowledge of the self-interference channel $\mathbf{H}$, therefore avoiding the prohibitively high overhead associated with *explicit* estimation.

This implicit channel knowledge is obtained by sweeping relatively few *probing beams* that are tailored to the particular environment in which the base station is deployed. A transformer-based deep learning model is used to design these probing beams, as well as the final serving beams, through end-to-end training on channel data that can be obtained from ray-tracing, actual measurements, or a digital twin of a particular site. This site-specific training allows the model to exploit the underlying structure of the environment and how it influences the users’ channels and self-interference channel, ultimately obtaining superior performance with fewer measurements than existing baselines. Our simulation results using ray-tracing data demonstrate the effectiveness of our scheme across a variety of scenarios and conditions, making it a promising enabler of future full-duplex massive MIMO systems in 6G and beyond.

Valuable future directions include extending the proposed approach to frequency-selective self-interference channels or considering models other than block fading, perhaps through self-interference channel prediction. More broadly, our proposed approach could likely be extended to ISAC applications, where sparse, site-specific probing can also be used to estimate targets in a known environment.

ACKNOWLEDGMENTS

This work used computational and storage services associated with UCLA’s Hoffman2 Cluster, operated by the UCLA Office of Advanced Research Computing.

REFERENCES

- [1] J. G. Andrews, T. E. Humphreys, and T. Ji, “6G takes shape,” *IEEE BITS Inf. Theory Mag.*, vol. 4, no. 1, pp. 2–24, Mar. 2024.
- [2] Nokia, “Transforming the 6G vision to action,” Nokia, White Paper, Oct. 2025. [Online]. Available: <https://www.nokia.com/asset/t/214027/>
- [3] B. Smida *et al.*, “Full-duplex wireless for 6G: Progress brings new opportunities and challenges,” *IEEE J. Sel. Areas Commun.*, vol. 41, no. 9, pp. 2729–2750, Sep. 2023.

- [4] B. Smida *et al.*, "In-band full-duplex: The physical layer," *Proc. IEEE*, pp. 1–30, 2024.
- [5] K. E. Kolodziej, B. T. Perry, and J. S. Herd, "In-band full-duplex technology: Techniques and systems survey," *IEEE Trans. Microw. Theory Techn.*, vol. 67, no. 7, pp. 3025–3041, Feb. 2019.
- [6] E. Everett, A. Sahai, and A. Sabharwal, "Passive self-interference suppression for full-duplex infrastructure nodes," *IEEE Trans. Wireless Commun.*, vol. 13, no. 2, pp. 680–694, Feb. 2014.
- [7] I. P. Roberts, Y. Zhang, T. Osman, and A. Alkhateeb, "Real-world evaluation of full-duplex millimeter wave communication systems," *IEEE Wireless Commun. Lett.*, vol. 23, no. 9, pp. 10 803–10 819, Sep. 2024.
- [8] E. Everett, C. Shepard, L. Zhong, and A. Sabharwal, "SoftNull: Many-antenna full-duplex wireless via digital beamforming," *IEEE Trans. Wireless Commun.*, vol. 15, no. 12, pp. 8077–8092, Dec. 2016.
- [9] I. T. Cummings, J. P. Doane, T. J. Schulz, and T. C. Havens, "Aperture-level simultaneous transmit and receive with digital phased arrays," *IEEE Trans. Signal Process.*, vol. 68, pp. 1243–1258, 2020.
- [10] I. P. Roberts, J. G. Andrews, H. B. Jain, and S. Vishwanath, "Millimeter-wave full duplex radios: New challenges and techniques," *IEEE Wireless Commun.*, vol. 28, no. 1, pp. 36–43, Feb. 2021.
- [11] E. Björnson *et al.*, "Towards 6G MIMO: Massive spatial multiplexing, dense arrays, and interplay between electromagnetics and processing," Jan. 2024. [Online]. Available: <http://arxiv.org/abs/2401.02844>
- [12] A. F. Molisch *et al.*, "Hybrid beamforming for massive MIMO: A survey," *IEEE Commun. Mag.*, vol. 55, no. 9, pp. 134–141, Sep. 2017.
- [13] I. P. Roberts, J. G. Andrews, and S. Vishwanath, "Hybrid beamforming for millimeter wave full-duplex under limited receive dynamic range," *IEEE Trans. Wireless Commun.*, vol. 20, no. 12, pp. 7758–7772, Dec. 2021.
- [14] A. Koc and T. Le-Ngoc, "Intelligent non-orthogonal beamforming with large self-interference cancellation capability for full-duplex multiuser massive MIMO systems," *IEEE Access*, vol. 10, pp. 51 771–51 791, 2022.
- [15] R. López-Valcarce and M. Martínez-Cotelo, "Full-duplex mmWave MIMO with finite-resolution phase shifters," *IEEE Trans. Wireless Commun.*, vol. 21, no. 11, pp. 8979–8994, Nov. 2022.
- [16] I. P. Roberts, S. Vishwanath, and J. G. Andrews, "LoneSTAR: Analog beamforming codebooks for full-duplex millimeter wave systems," *IEEE Trans. Wireless Commun.*, vol. 22, no. 9, pp. 5754–5769, Sep. 2023.
- [17] Z. He *et al.*, "Full-duplex communication for ISAC: Joint beamforming and power optimization," *IEEE J. Sel. Areas Commun.*, vol. 41, no. 9, pp. 2920–2936, Sep. 2023.
- [18] R. Hernangómez, J. Fink, R. L. G. Cavalcante, and S. Stańczak, "CISSIR: Beam codebooks with self-interference reduction guarantees for integrated sensing and communication beyond 5G," *IEEE Trans. Wireless Commun.*, vol. 25, pp. 6523–6537, 2026.
- [19] S. Huang, Y. Ye, and M. Xiao, "Learning-based hybrid beamforming design for full-duplex millimeter wave systems," *IEEE Trans. on Cogn. Commun. Netw.*, vol. 7, no. 1, pp. 120–132, Mar. 2021.
- [20] I. Bilbao *et al.*, "Deep unfolding-powered analog beamforming for In-band full-duplex," *IEEE Open J. Commun. Society*, vol. 5, pp. 3753–3761, 2024.
- [21] I. P. Roberts, A. Chopra, T. Novlan, S. Vishwanath, and J. G. Andrews, "STEER: Beam selection for full-duplex millimeter wave communication systems," *IEEE Trans. Commun.*, vol. 70, no. 10, pp. 6902–6917, Oct. 2022.
- [22] I. P. Roberts, Y. Zhang, T. Osman, and A. Alkhateeb, "STEER+: Robust beam refinement for full-duplex millimeter wave communication systems," in *Proc. Asilomar Conf. Signals Sys. Comput.*, Oct. 2023, pp. 134–139.
- [23] J. M. Kong and I. P. Roberts, "Active beam learning for full-duplex wireless systems," in *Proc. Asilomar Conf. Signals Sys. Comput.*, Oct. 2024, pp. 865–869.
- [24] 3GPP, "5G; NR; physical layer procedures for data," 2024. [Online]. Available: <https://www.3gpp.org/DynaReport/38214.htm>
- [25] Y. Heng, Y. Zhang, A. Alkhateeb, and J. G. Andrews, "Site-specific beam alignment in 6G via deep learning," *IEEE Commun. Mag.*, vol. 62, no. 8, pp. 162–168, Aug. 2024.
- [26] Q. Xue *et al.*, "A survey of beam management for mmWave and THz communications towards 6G," *IEEE Commun. Surveys Tuts.*, pp. 1–1, 2024.
- [27] Y. Heng and J. G. Andrews, "Grid-free MIMO beam alignment through site-specific deep learning," *IEEE Trans. Wireless Commun.*, vol. 23, no. 2, pp. 908–921, Feb. 2024.
- [28] M. Alrabeiah, Y. Zhang, and A. Alkhateeb, "Neural networks based beam codebooks: Learning mmWave massive MIMO beams that adapt to deployment and hardware," *IEEE Trans. Commun.*, vol. 70, no. 6, pp. 3818–3833, Jun. 2022.
- [29] R. M. Dreifuerst and R. W. Heath, "Machine learning codebook design for initial access and CSI type-II feedback in sub-6-GHz 5G NR," *IEEE Trans. Wireless Commun.*, vol. 23, no. 6, pp. 6411–6424, Jun. 2024.
- [30] J. W. Kwak, H. Yoo, I. P. Roberts, and C.-B. Chae, "Site-specific beam alignment without explicit channel knowledge via deep learning," in *Proc. Asilomar Conf. Signals Sys. Comput.*, Oct. 2024, pp. 1139–1143.
- [31] V. Va, J. Choi, T. Shimizu, G. Bansal, and R. W. Heath, "Inverse multipath fingerprinting for millimeter wave V2I beam alignment," *IEEE Trans. Veh. Technol.*, vol. 67, no. 5, pp. 4042–4058, May 2018.
- [32] U. Demirhan and A. Alkhateeb, "Radar aided 6G beam prediction: Deep learning algorithms and real-world demonstration," in *Proc. IEEE WCNC*, 2022, pp. 2655–2660.
- [33] S. Jiang and A. Alkhateeb, "Computer vision aided beam tracking in a real-world millimeter wave deployment," in *Proc. IEEE GLOBECOM Wkshp.*, Dec. 2022, pp. 142–147.
- [34] B. Yang, W. Wang, and W. Zhang, "Radio map-based beamforming assisted with reduced pilots," *IEEE Trans. Wireless Commun.*, vol. 24, no. 10, pp. 8878–8891, Oct. 2025.
- [35] A. Alkhateeb, O. El Ayach, G. Leus, and R. W. Heath, "Channel estimation and hybrid precoding for millimeter wave cellular systems," *IEEE J. Sel. Topics Signal Process.*, vol. 8, no. 5, pp. 831–846, Oct. 2014.
- [36] J. Lee, G.-T. Gil, and Y. H. Lee, "Channel estimation via orthogonal matching pursuit for hybrid MIMO systems in millimeter wave communications," *IEEE Trans. Commun.*, vol. 64, no. 6, pp. 2370–2386, Jun. 2016.
- [37] A. Alkhateeb, G. Leus, and R. W. Heath, "Compressed sensing based multi-user millimeter wave systems: How many measurements are needed?" in *Proc. IEEE ICASSP*, Apr. 2015, pp. 2909–2913.
- [38] F. Sohrabi, T. Jiang, W. Cui, and W. Yu, "Active sensing for communications by learning," *IEEE J. Sel. Areas Commun.*, vol. 40, no. 6, pp. 1780–1794, Jun. 2022.
- [39] S. Li and I. P. Roberts, "Site-specific beam learning for full-duplex massive MIMO wireless systems," in *Proc. IEEE Mil. Commun. Conf.*, Oct. 2025, pp. 1236–1241.
- [40] M. Duarte, C. Dick, and A. Sabharwal, "Experiment-driven characterization of full-duplex wireless systems," *IEEE Trans. Wireless Commun.*, vol. 11, no. 12, pp. 4296–4307, Dec. 2012.
- [41] I. P. Roberts, A. Chopra, T. Novlan, S. Vishwanath, and J. G. Andrews, "Beamformed self-interference measurements at 28 GHz: Spatial insights and angular spread," *IEEE Trans. Wireless Commun.*, vol. 21, no. 11, pp. 9744–9760, Nov. 2022.
- [42] A. Sabharwal *et al.*, "In-band full-duplex wireless: Challenges and opportunities," *IEEE J. Sel. Areas Commun.*, vol. 32, no. 9, pp. 1637–1652, Sep. 2014.
- [43] B. Hassibi and B. Hochwald, "How much training is needed in multiple-antenna wireless links?" *IEEE Trans. Inf. Theory*, vol. 49, no. 4, pp. 951–963, Apr. 2003.
- [44] A. Vaswani *et al.*, "Attention is all you need," *Adv. Neural Inf. Process. Syst.*, vol. 30, pp. 1–11, 2017.
- [45] M. Haklay and P. Weber, "OpenStreetMap: User-generated street maps," *IEEE Pervasive Comput.*, vol. 7, no. 4, pp. 12–18, Oct. 2008.
- [46] J. Hoydis *et al.*, "Sionna," 2022. [Online]. Available: <https://nvlabs.github.io/sionna/>
- [47] J.-S. Jiang and M. Ingram, "Spherical-wave model for short-range MIMO," *IEEE Trans. Commun.*, vol. 53, no. 9, pp. 1534–1541, Sep. 2005.
- [48] I. Loshchilov and F. Hutter, "Decoupled weight decay regularization," in *Int. Conf. Learn. Represent. (ICLR)*, 2019. [Online]. Available: <https://arxiv.org/abs/1711.05101>
- [49] L. N. Smith and N. Topin, "Super-convergence: Very fast training of neural networks using large learning rates," in *Artif. Intell. Mach. Learn. Multi-Domain Oper. Appl.*, vol. 11006. SPIE, May 2019, pp. 369–386.