

Joint Communication and Parameter Estimation in MIMO Channels

Gökhan Yılmaz, Franz Lampel, Hamdi Joudeh and Giuseppe Caire

Abstract

We study a joint communication and sensing setting comprising a transmitter, a receiver, and a sensor, all equipped with multiple antennas. The transmitter sends an encoded signal over the channel with the dual purpose of communicating an information message to the receiver, and enabling the sensor to estimate a target parameter vector by generating back-scattered signals. We assume that the transmitter and sensor are co-located, or fully connected, giving the latter access to the transmitted signal. The target parameter vector is randomly drawn from a continuous distribution, yet remains fixed throughout the transmission block. We establish the fundamental performance trade-off between the communication and sensing tasks, captured in terms of a capacity-MSE function. In doing so, we identify optimal coding schemes for this multi-antenna joint communication and sensing setting. Moreover, we particularize our result to two practically-inspired scenarios where we showcase optimal schemes and trade-offs.

I. INTRODUCTION

Joint communication and sensing (JCAS), also known as integrated sensing and communication (ISAC), is expected to be a key feature of next-generation wireless technology. This emerging paradigm aims at designing efficient wireless systems with transceivers that can utilize the same physical resources (e.g., hardware, spectrum, and power) to carry out both communication and sensing tasks [1]–[5]. Envisaged use cases include cellular networks, where base stations will have the ability to communicate with active devices and simultaneously detect and track passive targets from back-scattered signals [6]; and automotive, where vehicles will be able to sense their

Gökhan Yılmaz, Franz Lampel, and Hamdi Joudeh are with the Department of Electrical Engineering, Eindhoven University of Technology, 5600 MB Eindhoven, The Netherlands (e-mail: g.yilmaz@tue.nl; f.lampel@tue.nl; h.joudeh@tue.nl).

Giuseppe Caire is with the Faculty of Electrical Engineering and Computer Science, Technical University of Berlin, 10623 Berlin, Germany (e-mail: caire@tu-berlin.de).

This work was supported by the European Research Council (ERC) under Grant 101116550.

surroundings to identify and track road obstacles, while communicating with other vehicles and the infrastructure [7], [8]. In both of these use cases, and several others, the underlying JCAS physical channel will very likely be a multiple-input and multiple-output (MIMO) channel, due to the employment of multi-antenna terminals, multiple sub-carriers, or both.

In recent years, a number of studies have emerged investigating various aspects of MIMO in JCAS; see, e.g., [2]–[5], [9]–[14]. With the exception of [13], [14], the vast majority of studies on MIMO JCAS focus mainly on signal processing and optimization aspects. In particular, a transmission scheme with a predetermined signal structure, often based on linear precoding and Gaussian signaling, is adopted; reducing the problem to the design of signal parameters (e.g., covariance matrices) to achieve different performance trade-offs. While these works have provided useful insights into the potential benefits of co-designed MIMO JCAS systems, without an information-theoretic treatment involving coding theorems and converse bounds, the optimality of these predetermined signal structures and schemes remains unclear. Moreover, we still lack performance benchmarks that indicate how much more can be gained.

Goal and considered setting: Our primary goal in this work is to shed light on the fundamental performance limits of MIMO JCAS systems by taking an information-theoretic approach. To this end, we focus on an elemental setting comprising a transmitter, a receiver, and a sensor, all equipped with multiple antennas (or any vector dimension, e.g., sub-carriers). The transmitter and sensor are co-located or strongly connected, enabling mono-static sensing, and henceforth can be thought of as components of the same cellular base station (BS). The receiver, on the other hand, belongs to an active user equipment (UE) served by the BS. The transmitter sends an encoded signal with the dual purpose of communicating a message to the receiver and, at the same time, probing the surrounding environment and generating a back-scatter signal to enable target sensing. We refer to this setting as a *MIMO JCAS channel*.

Information-theoretic formulations and performance limits for MIMO *communication* channels are well established and extensively studied [15]–[17]. To treat MIMO JCAS channels with the same degree of rigor, it is crucial to concretely formulate the *sensing* aspect of the problem. In the current work, we focus on the case where the sensing task involves estimating a *continuous* parameter vector of the back-scatter channel. The parameter vector is drawn at random according to a known prior distribution, yet remains *fixed* throughout the transmission block. This *abstraction* is meant to model scenarios where parameters of interest change at a much slower timescale compared to channel uses (i.e., symbols). Practical examples include direction-of-arrival (DoA),

range (delay), and velocity (Doppler) estimation, all scenarios in which only slight changes in parameters occur over the span of a typical transmission block (see Remark 2).

Fixed vs. varying parameter: Our choice to adopt a fixed random parameter model for sensing stands in sharp contrast to a prevalent approach in the information-theoretic literature on JCAS, where the parameter of interest is assumed to vary in an i.i.d. fashion from one channel symbol to the other. The i.i.d. model was proposed by Kobayashi *et al.* in [18], where the focus is on discrete memoryless and basic Gaussian channels; and has since then been widely adopted in several follow-up works [19]–[23]. In [13], Xiong *et al.* study a MIMO JCAS channel that is very similar to ours, with the exception that their sensing model also follows the i.i.d. philosophy in [18]. In particular, they consider the estimation of a continuous parameter vector that varies across channel symbols (or sub-blocks of symbols) in an i.i.d. fashion. An asymptotically large number of symbols (or sub-blocks) is considered, resembling ergodic fading in wireless communication. Moreover, to measure the sensing performance, a form of Bayesian Cramér-Rao bound (BCRB) [24] is used in [13] as a proxy for the mean square error (MSE) averaged over parameter states.

The BCRB is known to be generally not tight, even asymptotically, except in certain special cases [25]; hence, the performances characterized in [13] are not exact in general. Nevertheless, high signal-to-noise-ratio (SNR) analysis reveals two distinct types of trade-offs in the MIMO JCAS channel model of [13]. The first is referred to as the *deterministic–random trade-off (DRT)*, which arises when the family of Gaussian input distributions—optimal for communication alone—fails to achieve all Pareto-optimal performance trade-off points. In this case, “less random” input distributions are required to achieve trade-off points that favor sensing. The second is referred to as the *subspace trade-off (ST)*, which arises from the fact that the communication user and the sensing target often occupy distinct subspaces in the MIMO vector space. Hence, signals more aligned with the user (resp. target) favor communication (resp. sensing).

The i.i.d. model does not fully capture practical scenarios where the parameters of interest remain relatively constant within transmission blocks, as the ones highlighted earlier. The idea to study fundamental performance limits in JCAS channels with a fixed parameter originated from Joudeh and Willems in [26], where the focus is on the basic sensing task of target detection (i.e., a binary parameter). Extensions later followed in [14], [27]–[29], where the focus remained on discrete (mostly binary) parameters. The extension to MIMO JCAS channels in [14] is particularly relevant, and can be seen as a detection counterpart to the estimation problem we consider in the present paper. In a recent contribution, Seo [30] extended the fixed discrete

parameter JCAS setting in [26]–[28] to the continuous parameter regime. Our current paper follows in the footsteps of [30], and further extends the approach therein in several directions, as discussed in detail below in light of our contributions.

Overview and contribution: We consider a MIMO JCAS channel where the sensing task is to estimate a random yet fixed parameter vector, as described earlier. We adopt a Bayesian framework for parameter estimation, with an MSE-based risk function [25], [31]. For this setting, we establish the fundamental trade-off between communication and sensing efficiencies in the asymptotic regime of large block-length. Communication efficiency is measured in terms of the rate of reliable communication, in bits per channel use; whose ultimate asymptotic limit is the celebrated Shannon capacity [32]. Sensing efficiency, on the other hand, is measured in terms of the MSE risk scaled by the block-length, which we refer to as the *asymptotic MSE*. This choice of scaling is commonplace in asymptotic statistics, as pointed out in [33, Sec. I], since the MSE risk in such settings decays at a rate proportional to one over the block-length. Therefore, a larger asymptotic MSE implies a slower decay rate and vice versa.

In the main result of this paper, presented in Section III, we derive the exact optimal trade-off between the rate of reliable communication and the sensing asymptotic MSE, described in terms of a capacity-MSE function. The rate is characterized by a mutual information term, while the asymptotic MSE is characterized by a conditional variant of the Expected Cramér-Rao bound (ECRB) [25], adapted to our sensing setting. It is interesting to note that while the ECRB does not generally yield a valid lower bound at finite block-lengths, it is asymptotically tight [25], [34]. We extend ECRB’s asymptotic tightness to the case of independent non-identical measurements arising in JCAS, from which the conditional ECRB variant emerges. Unlike the BCRB used in [13], the ECRB enables an exact asymptotic characterization in our setting. Moreover, we showcase the established fundamental trade-off through two practically-inspired use cases, one involving DoA estimation and another involving OFDM signaling (see Section V).

An interesting feature revealed by our analysis is that all optimal rate-MSE trade-off points are characterized by Gaussian input distributions with different choices of covariance matrices. In the terminology of Xiong *et al.* [13], the MIMO JCAS setting considered here exhibits only an ST and no DRT, with Gaussian-like communication waveforms also being optimal for sensing. This is mainly because the sensing channel parameter is fixed. Arriving at this conclusion, however, is nontrivial, since i.i.d. Gaussian codebook ensembles widely used in standard achievability proofs do not directly apply in our case. Intuitively, i.i.d.-generated codebooks occasionally

contain codewords that are unfavorable for sensing. Alternatively, we rely on *almost-constant covariance codes* [14], which yield uniformly favorable sensing waveforms while asymptotically exhibiting Gaussian-like properties that preserve their communication optimality.

Related work: A capacity-MSE trade-off, similar to the one we present here, was recently derived by Seo [30], albeit in a much more restricted setting. A major difference is that the approach of Seo is only applicable to discrete memoryless channels (DMCs), where inputs and outputs come from finite sets (despite the parameter being continuous). In particular, Seo follows the approach developed by Wu and Joudeh in [27], [35], which relies on constant composition codes to prove achievability and type-by-type analysis to prove the converse. Such an approach is not applicable to Gaussian (or other continuous) channels, and cannot be used to analyze the MIMO JCAS setting at hand. In contrast, our achievability proof is based on Feinstein’s lemma, which is applicable to channels with general alphabets [36]. Moreover, our converse proof avoids the type-based argument in [30] and instead relies on a convexity property of the conditional ECRB, which we also prove in this paper. Another smaller difference is that Seo considers a scalar sensing channel parameter only, while we consider a general vector parameter.

The MIMO JCAS channel in the fixed-parameter regime has also been studied in the prior literature, e.g., [9]–[12]. With the exception of [12], however, these works adopt the classical CRB within a non-Bayesian framework. As noted earlier, they also focus primarily on signal processing and optimization aspects without providing information-theoretic optimality guarantees. Moreover, the reliance on single-letter models, prevalent in the wireless communication literature, typically involves an informal notion of empirical input covariance convergence under the generic pretext of “Gaussian signaling” (see, e.g., [9, eq. (3)] and [10]). We address this issue by adopting a multi-letter perspective and employing almost-constant covariance codes.

Finally, although Xiong *et al.* [13] consider the i.i.d. varying-parameter regime, the obtained rate-BCRB trade-off can be applied to the Bayesian fixed-parameter regime we consider in the current paper. As we show, this yields an outer bound that is strictly loose, and hence not achievable, in general. This has also been observed by Seo [30] in the context of the DMC.

Notation: We adopt basic notation, which is mostly clear from the context. Upper-case bold letters as $\mathbf{A}$ denote matrices; lower-case bold letters as $\mathbf{a}$ denote (column) vectors; while A and a are scalars. Calligraphic letters as $\mathcal{A}$ denote sets. Random variables are denoted as $\mathbf{A}$, $\mathbf{a}$ and a , while $\mathbf{A}$, $\mathbf{a}$ and a denote deterministic quantities (e.g., realizations).

II. PROBLEM SETTING

We consider a vector (i.e., MIMO) JCAS channel comprising an M -antenna transmitter, an L -antenna receiver (or user), and a T -antenna sensor.¹ The transmitter sends an encoded signal over the channel with the dual purpose of *communicating* an information message to the user, and enabling *sensing* of a target or the environment by generating back-scattered signals for the sensor. We assume that the sensor has access to the transmitted encoded signal and can perform coherent sensing. This is the case in mono-static sensing, where the transmitter and sensor are co-located; or in bi-static sensing with a strong link between the transmitter and sensor.

A. Signal model

We consider a base-band discrete-time model, in which each channel use can be interpreted as one symbol period. In the n -th use, the transmitter sends a vector channel symbol $\mathbf{x}_n \in \mathbb{C}^M$ through its M antennas. The user receives a linearly-transformed noisy vector $\mathbf{y}_n \in \mathbb{C}^L$ given as

$$\mathbf{y}_n = \mathbf{H}\mathbf{x}_n + \boldsymbol{\omega}_n. \quad (1)$$

$\mathbf{H} \in \mathbb{C}^{L \times M}$ is the channel matrix between the transmitter and the user and $\boldsymbol{\omega}_n \sim \mathcal{CN}(\mathbf{0}, \sigma_{\omega}^2 \mathbf{I})$ is a zero-mean circularly symmetric complex Gaussian noise vector, independent over channel uses n . We refer to $\mathbf{H}$ as the *user channel matrix*. This remains fixed over channel uses in the setting we consider in this paper, and it is assumed to be known to both the transmitter and the user. An underlying assumption here is that initial access, including channel training, between the transmitter and the user has already been performed.

The sensor receives a noisy echo signal $\mathbf{z}_n \in \mathbb{C}^T$ through a back-scatter channel modeled by

$$\mathbf{z}_n = \mathbf{G}(\boldsymbol{\theta})\mathbf{x}_n + \mathbf{v}_n \quad (2)$$

where $\mathbf{G}(\boldsymbol{\theta}) \in \mathbb{C}^{T \times M}$ is the channel response matrix from the transmitter to the target (or environment) and back to the sensor, and $\mathbf{v}_n \sim \mathcal{CN}(\mathbf{0}, \sigma_v^2 \mathbf{I})$ is the corresponding additive noise vector. We refer to $\mathbf{G}(\boldsymbol{\theta})$ as the *sensor channel matrix*. It is a function of a real *target parameter vector* $\boldsymbol{\theta} \in \mathbb{R}^K$, drawn from a continuous distribution with probability density function $f_{\boldsymbol{\theta}}(\boldsymbol{\theta})$ on $\mathbb{R}^K$. The sensor is interested in estimating the target parameter vector, as it contains the relevant information about the target or environment. We assume that the mapping $\mathbf{G} : \mathbb{R}^K \rightarrow \mathbb{C}^{T \times M}$ from

¹The vector dimension is interpreted as antennas, but can also represent time extension or sub-carriers (e.g., OFDM).

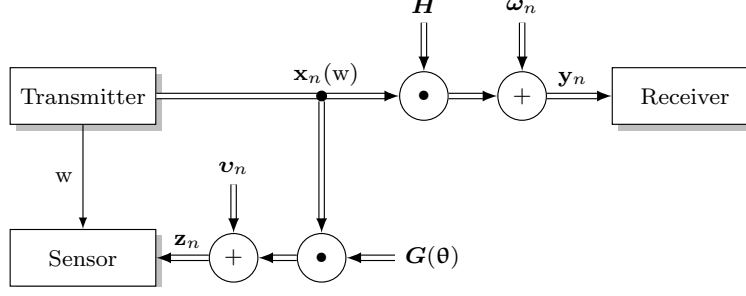


Fig. 1: The considered MIMO JCAS system model.

the target parameter vector to the sensor channel matrix is differentiable and invertible (almost surely on the support of Θ). As we shall see in Section V, where we discuss particularized examples of this model, these assumptions are justifiable in many settings of practical interest. An illustration of our model is shown in Fig. 1.

Note that while θ (and therefore $G(\theta)$) is drawn at random, it remains fixed throughout the transmission block in the setting we consider. This represents scenarios where the target parameter changes at a much slower time scale compared to channel uses within a single transmission block, which is the case in many practical scenarios, as explained earlier.

Transmission occurs over a block of N channel uses, where a signal (i.e., a sequence of vector channel symbols) given by $\mathbf{x}^N \triangleq \mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_N$ is sent. As we specify shortly, the transmitted signal $\mathbf{x}^N$ is a codeword selected from a codebook, and encodes a message intended for the user. It also serves the additional purpose of probing, to enable parameter estimation at the sensor. The transmitted signal is subject to an average (over symbols) power constraint expressed by

$$\frac{1}{N} \sum_{n=1}^N \|\mathbf{x}_n\|^2 \leq P \quad (3)$$

where P is the average transmit power budget.

B. Encoding, decoding, and estimation

We now formally describe a general JCAS scheme for the above setting, which involves encoding at the transmitter, decoding at the user, and estimation at the sensor.

1) *Encoding*: In a given transmission block of length N , the transmitter wishes to communicate a message w to the user, where w is drawn uniformly at random from a message set

of M_N indices given by $\mathcal{M}_N \triangleq \{1, 2, \dots, M_N\}$. To this end, an encoded signal (or codeword) $\mathbf{x}^N(w) = \mathbf{x}_1(w), \dots, \mathbf{x}_N(w)$ is sent. The set of all codewords, i.e., codebook, is given by

$$\mathcal{C}_N \triangleq \{\mathbf{x}^N(1), \mathbf{x}^N(2), \dots, \mathbf{x}^N(M_N)\}.$$

Each codeword in the codebook must satisfy the power constraint in (3). Moreover, since w is uniform, then the random transmitted signal $\mathbf{x}^N \triangleq \mathbf{x}^N(w)$ is uniform on $\mathcal{C}_N$. Note that the signals received by the user and the sensor depend on w through $\mathbf{x}^N(w)$, and hence conditioning on $w = w$ in what follows is often replaced with conditioning on $\mathbf{x}^N = \mathbf{x}^N(w)$.

2) *Decoding*: The user receives a random signal given by an N -sequence of vectors $\mathbf{y}^N \triangleq \mathbf{y}_1, \mathbf{y}_2, \dots, \mathbf{y}_N$ from which it produces a decoded message $\hat{w}(\mathbf{y}^N) \in \mathcal{M}_N$. Recall the assumption that the user has access to its channel matrix $\mathbf{H}$, on which the decoder may depend. Given that $w = w$ is sent, and hence $\mathbf{x}^N = \mathbf{x}^N(w)$ is transmitted, the probability of decoding error is

$$\varepsilon_N(w) \triangleq \mathbb{P} [\hat{w}(\mathbf{y}^N) \neq w \mid \mathbf{x}^N = \mathbf{x}^N(w)]. \quad (4)$$

The *communication reliability* is measured by the maximal probability of decoding error

$$\varepsilon_N \triangleq \max_{w \in \mathcal{M}_N} \varepsilon_N(w). \quad (5)$$

We shall refer to this as the *error probability* for short. The maximum over all messages in (5) reflects an underlying assumption that all messages are equally important, i.e., there is no reason to distinguish between different message reliabilities. It is sometimes technically more convenient to work with the average (over messages) decoding error probability $\varepsilon_N^{\text{av}} = \frac{1}{M_N} \sum_{w \in \mathcal{M}_N} \varepsilon_N(w)$. However, as is often the case in asymptotic analysis of point-to-point communication channels, capacity results remain unchanged regardless of whether we choose ε_N or $\varepsilon_N^{\text{av}}$ as the measure of decoding error probability. This is also the case for the results we present in this paper.

3) *Estimation*: The sensor receives the random signal given by $\mathbf{z}^N \triangleq \mathbf{z}_1, \mathbf{z}_2, \dots, \mathbf{z}_N$, and with knowledge of the codeword $\mathbf{x}^N(w)$, it produces an estimate $\hat{\boldsymbol{\theta}}_N(\mathbf{z}^N, \mathbf{x}^N(w))$ of the random target parameter vector $\boldsymbol{\theta}$. Note that the estimation function $\hat{\boldsymbol{\theta}}_N : \mathbb{C}^{T \times N} \times \mathcal{C}_N \rightarrow \mathbb{R}^K$ itself is a deterministic mapping, while the estimate $\hat{\boldsymbol{\theta}}_N(\mathbf{z}^N, \mathbf{x}^N(w))$ is a random variable due to the randomness of the observation $\mathbf{z}^N$. For brevity, we sometimes write $\hat{\boldsymbol{\theta}}_N(\mathbf{z}^N)$ as $\mathbf{x}^N(w)$ is given.

A quadratic loss criterion is used to measure the estimation error, leading to an MSE-based estimation risk described next. To formulate the estimation risk, we shall first define the MSE matrix. Given an estimator $\hat{\boldsymbol{\theta}}_N$ and that $\mathbf{x}^N(w)$ is sent, the MSE matrix is defined as

$$\boldsymbol{\Sigma}(\mathbf{x}^N(w), \hat{\boldsymbol{\theta}}_N) \triangleq \mathbb{E} \left[\left[\boldsymbol{\theta} - \hat{\boldsymbol{\theta}}_N(\mathbf{z}^N) \right] \left[\boldsymbol{\theta} - \hat{\boldsymbol{\theta}}_N(\mathbf{z}^N) \right]^T \mid \mathbf{x}^N = \mathbf{x}^N(w) \right] \quad (6)$$

where the expectation in (6) is with respect to $\boldsymbol{\theta}$ and $\mathbf{z}^N$, given $\mathbf{x}^N = \mathbf{x}^N(w)$.² The diagonal elements of $\boldsymbol{\Sigma}(\mathbf{x}^N(w), \hat{\boldsymbol{\theta}}_N)$ represent the MSEs of individual target parameters in $\boldsymbol{\theta}$. The error in estimating the target parameter vector, given codeword $\mathbf{x}^N(w)$, is measured in terms of the sum-MSE risk (or *MSE risk* for short), defined as

$$\xi_N(w) \triangleq \text{tr} \left(\boldsymbol{\Sigma} \left(\mathbf{x}^N(w), \hat{\boldsymbol{\theta}}_N \right) \right). \quad (7)$$

The *sensing risk* is chosen to be the *maximal MSE risk* (maximal over messages) defined as

$$\xi_N \triangleq \max_{w \in \mathcal{M}_N} \xi_N(w). \quad (8)$$

The maximum over messages in (8) ensures a uniform sensing risk across different messages, or codewords in the codebook. This reflects a natural operational requirement that the sensing performance should be independent of which message is being communicated, or equivalently, which codeword is being transmitted. Such a formulation gives rise to the interesting design problem of finding *good* channel codebooks with the important property that all codewords also constitute *equally good* probing signals for the sensing task.

A less stringent measure of sensing risk is the average (over messages) MSE risk, defined as $\xi_N^{\text{av}} \triangleq \frac{1}{M_N} \sum_{w \in \mathcal{M}_N} \xi_N(w)$. This, however, does not provide the same sensing performance guarantees. Interestingly, as we shall see further on, regardless of whether we choose ξ_N or ξ_N^{av} , the fundamental asymptotic trade-off we establish in the next section remains unchanged.

Remark 1. In many practical settings, individual target parameters may carry different importance. For such scenarios, we may introduce the weights $a_1, a_2, \dots, a_K$, where $a_i \geq 0$ signifies the “importance” of the i -th parameter in $\boldsymbol{\theta}$. Weights are assumed to be fixed and known beforehand. In this case, we modify the MSE in (7) into a weighted MSE (WMSE)

$$\xi_N(w) \triangleq \sum_{i=1}^K a_i \left[\boldsymbol{\Sigma} \left(\mathbf{x}^N(w), \hat{\boldsymbol{\theta}}_N \right) \right]_{ii} \quad (9)$$

$$= \text{tr} \left(\mathbf{A} \boldsymbol{\Sigma} \left(\mathbf{x}^N(w), \hat{\boldsymbol{\theta}}_N \right) \right) \quad (10)$$

where $\mathbf{A} = \text{diag}(a_1, a_2, \dots, a_K)$ is a diagonal matrix induced by the weights. The WMSE sensing risk is then obtained from (9) by taking the maximum or average over messages. For ease of exposition, we will present our results in terms of the unweighted MSE risk. Nevertheless, all results and proofs can be easily extended to account for the weighted case.

²Our choice of MSE is sometimes referred to as the Bayesian MSE [31], since it involves an expectation over the prior of $\boldsymbol{\theta}$.

C. Asymptotic rate-MSE trade-off

A JCAS scheme, as defined above, is referred to as an $(N, M_N, \varepsilon_N, \xi_N)$ -scheme if it has a codebook $\mathcal{C}_N$ of M_N length- N codewords, a decoder with maximal error probability ε_N , and an estimator with maximal MSE risk ξ_N . We are interested in characterizing the fundamental performance limits of such schemes in the asymptotic regime of $N \rightarrow \infty$.

As the block-length N grows, it is known from the channel coding theorem that the size of the message set M_N can be made to grow exponentially in N while guaranteeing reliable decoding, i.e., $\varepsilon_N \rightarrow 0$. Therefore, communication efficiency is often measured by the *rate* $\frac{1}{N} \log M_N$, which is understood as the number of communicated bits per channel use.

On the other hand, we expect the MSE sensing risk in (8) to decay to zero as N grows large at a rate of $O(1/N)$, since θ is fixed and the measurements are conditionally independent [25]. Therefore, an appropriate asymptotic measure of sensing risk is the scaled MSE $N\xi_N$. We will refer to the scaled MSE as the *asymptotic MSE*, or *MSE* for short, where it should be clear from the context whether we are referring to ξ_N or the scaled version $N\xi_N$. We are now in position to introduce the notion of asymptotically achievable rate-MSE pairs.

Definition 1. *The rate-MSE tuple (R, Δ) is said to be achievable if for every $\epsilon > 0$ and $N \geq N_\epsilon$ (possibly large and may depend on ϵ), there exists an $(N, M_N, \varepsilon_N, \xi_N)$ -scheme with*

$$\begin{aligned} \varepsilon_N &\leq \epsilon \\ \frac{1}{N} \log M_N &\geq R - \epsilon \\ N\xi_N &\leq \Delta + \epsilon. \end{aligned}$$

The capacity-MSE function $C(\Delta)$ is the maximum rate R for which (R, Δ) is achievable.

Stated differently, the above definition says that if a rate-MSE tuple (R, Δ) is achievable, then there exists an N -block scheme for the above setting that can reliably communicate any one of $M_N = 2^{N(R-\epsilon)}$ messages, with error probability $\varepsilon_N = \epsilon$; while simultaneously achieving a sensing risk of $\xi_N = (\Delta + \epsilon)/N$. Note that ϵ can be made as small as desired by choosing a large-enough block-length N . Our main result in this paper is to fully characterize $C(\Delta)$, and then showcase the trade-off through a couple of canonical examples.

Remark 2. Before we proceed, we justify some of our assumptions in this remark. First, for the fixed target parameter, this is practically motivated by the fact that many parameters of interest

in sensing remain unchanged or change very slightly over a typical transmission block. For instance, during a transmission period of 0.5 ms, a fast-traveling car at 100 km/h moves by just over 1.3 cm, and hence parameters of interest (e.g., range, velocity, angle) remain almost fixed.

Now moving on to the asymptotic regime ($N \rightarrow \infty$), we argue that this does not pose a contradiction when considered together with a fixed parameter. For instance, consider a wireless JCAS system operating with $B = 250$ MHz of bandwidth. A transmission block spanning $\tau = 0.5$ ms corresponds to roughly $N = 125000$ channel uses (i.e. $N \approx \tau B$ complex base-band samples). Even if we consider smaller bandwidths, shorter intervals, or lower bandwidth utilization due to pulse shaping, we will still obtain a large enough N that justifies asymptotic analysis.

III. MAIN RESULT

A. Preliminaries

To facilitate the presentation of our main result, we start by introducing common information measures, namely the mutual information and the Fisher information. We also introduce a function of the Fisher information matrix, related to the ECRB, which will be used to characterize the asymptotic MSE in the main theorem. In what follows, we use $f_{\mathbf{y}|\mathbf{x}}(\mathbf{y}|\mathbf{x})$ to denote the density of $\mathcal{CN}(\mathbf{H}\mathbf{x}, \sigma_\omega^2 \mathbf{I})$, i.e., the transition law of the user channel (1). Similarly, $f_{\mathbf{z}|\mathbf{x}, \boldsymbol{\theta}}(\mathbf{z}|\mathbf{x}, \boldsymbol{\theta})$ denotes the density of $\mathcal{CN}(\mathbf{G}(\boldsymbol{\theta})\mathbf{x}, \sigma_v^2 \mathbf{I})$, associated with the sensor channel (2).

Now suppose that the random vector $\mathbf{x} \in \mathbb{C}^M$ is used as input in a single use of the JCAS channel under consideration. The mutual information between $\mathbf{x}$ and the user output $\mathbf{y}$ is

$$I(\mathbf{x}; \mathbf{y}) \triangleq \mathbb{E} \left[\log \left(\frac{f_{\mathbf{y}|\mathbf{x}}(\mathbf{y}|\mathbf{x})}{f_{\mathbf{y}}(\mathbf{y})} \right) \right] \quad (11)$$

where $f_{\mathbf{y}}(\mathbf{y})$ is the probability density of the output distribution of $\mathbf{y}$ induced by the input $\mathbf{x}$. Note that the quantity inside the expectation in (11) is known as the information density, defined for any given input-output pair of realizations $\mathbf{x}$ and $\mathbf{y}$ as

$$\imath(\mathbf{x}; \mathbf{y}) \triangleq \log \left(\frac{f_{\mathbf{y}|\mathbf{x}}(\mathbf{y}|\mathbf{x})}{f_{\mathbf{y}}(\mathbf{y})} \right). \quad (12)$$

For $\tilde{\mathbf{x}} \sim \mathcal{CN}(\mathbf{0}, \mathbf{Q})$, the mutual information $I(\tilde{\mathbf{x}}; \mathbf{y})$ evaluates to

$$I(\tilde{\mathbf{x}}; \mathbf{y}) = \log \det \left(\mathbf{I} + \frac{1}{\sigma_\omega^2} \mathbf{H} \mathbf{Q} \mathbf{H}^H \right). \quad (13)$$

It is well known that (13) is non-decreasing, and concave in $\mathbf{Q} \in \mathbb{S}_+^M$, where $\mathbb{S}_+^M$ denotes the cone of $M \times M$ complex, symmetric, positive semi-definite matrices.

Now let us condition on $\boldsymbol{\theta} = \boldsymbol{\theta}$ and $\mathbf{x} = \mathbf{x}$. The Fisher information matrix (FIM) associated with estimating $\boldsymbol{\theta}$ from the random sensor observation $\mathbf{z} \sim f_{\mathbf{z}|\mathbf{x},\boldsymbol{\theta}}(\mathbf{z}|\mathbf{x},\boldsymbol{\theta})$ is given by

$$\mathbf{J}(\boldsymbol{\theta}|\mathbf{x}) \triangleq \mathbb{E} \left[\left[\nabla_{\boldsymbol{\theta}} \ln f_{\mathbf{z}|\mathbf{x},\boldsymbol{\theta}}(\mathbf{z}|\mathbf{x},\boldsymbol{\theta}) \right] \left[\nabla_{\boldsymbol{\theta}} \ln f_{\mathbf{z}|\mathbf{x},\boldsymbol{\theta}}(\mathbf{z}|\mathbf{x},\boldsymbol{\theta}) \right]^{\top} \right]. \quad (14)$$

For our particular case, the (i, j) -th element of $\mathbf{J}(\boldsymbol{\theta}|\mathbf{x})$ evaluates to

$$[\mathbf{J}(\boldsymbol{\theta}|\mathbf{x})]_{ij} = \frac{2}{\sigma_v^2} \operatorname{Re} \left\{ \mathbf{x}^{\text{H}} \frac{\partial \mathbf{G}(\boldsymbol{\theta})^{\text{H}}}{\partial \theta_i} \frac{\partial \mathbf{G}(\boldsymbol{\theta})}{\partial \theta_j} \mathbf{x} \right\} \quad (15)$$

$$= \frac{2}{\sigma_v^2} \operatorname{Re} \left\{ \operatorname{tr} \left(\frac{\partial \mathbf{G}(\boldsymbol{\theta})}{\partial \theta_j} \mathbf{x} \mathbf{x}^{\text{H}} \frac{\partial \mathbf{G}(\boldsymbol{\theta})^{\text{H}}}{\partial \theta_i} \right) \right\}. \quad (16)$$

A quantity derived from the FIM which is very relevant to us is the following:

$$\mathbf{K}(\boldsymbol{\theta}; \mathbf{z}|\mathbf{x}) \triangleq \mathbb{E} \left[\mathbb{E} [\mathbf{J}(\boldsymbol{\theta}|\mathbf{x})|\boldsymbol{\theta}]^{-1} \right]. \quad (17)$$

To parse (17), given $\boldsymbol{\theta} = \boldsymbol{\theta}$, we take the expectation $\mathbb{E} [\mathbf{J}(\boldsymbol{\theta}|\mathbf{x})]$ with respect to a random input $\mathbf{x}$, which we will refer to as the *conditional FIM*.³ We then invert the conditional FIM and take its expectation with respect to the random parameter $\boldsymbol{\theta}$. As shown in Section IV, $\mathbf{K}(\boldsymbol{\theta}; \mathbf{z}|\mathbf{x})$ is closely related to the ECRB for estimating $\boldsymbol{\theta}$ from the sensor observations with knowledge of the transmitted codeword, and exactly characterizes the asymptotic MSE. In fact, we can think of (17) as a form of *conditional ECRB* matrix.

Now suppose that we have an input $\tilde{\mathbf{x}} \sim \mathcal{CN}(\mathbf{0}, \mathbf{Q})$. We can see from (16) that the conditional FIM $\mathbb{E} [\mathbf{J}(\boldsymbol{\theta}|\tilde{\mathbf{x}})]$ in this case only depends on the input distribution through $\mathbf{Q}$. Therefore, we denote $\mathbb{E} [\mathbf{J}(\boldsymbol{\theta}|\tilde{\mathbf{x}})]$ by $\mathbf{J}(\boldsymbol{\theta}|\mathbf{Q})$, in which the (i, j) -th element is given by

$$[\mathbf{J}(\boldsymbol{\theta}|\mathbf{Q})]_{ij} = \frac{2}{\sigma_v^2} \operatorname{Re} \left\{ \operatorname{tr} \left(\frac{\partial \mathbf{G}(\boldsymbol{\theta})}{\partial \theta_j} \mathbf{Q} \frac{\partial \mathbf{G}(\boldsymbol{\theta})^{\text{H}}}{\partial \theta_i} \right) \right\}. \quad (18)$$

In this case, the conditional ECRB matrix in (17) becomes

$$\mathbf{K}(\boldsymbol{\theta}; \mathbf{z}|\tilde{\mathbf{x}}) = \mathbb{E} [\mathbf{J}(\boldsymbol{\theta}|\mathbf{Q})^{-1}]. \quad (19)$$

As we will see shortly, the asymptotic MSE is characterized in terms of the function

$$\operatorname{tr} (\mathbf{K}(\boldsymbol{\theta}; \mathbf{z}|\tilde{\mathbf{x}})) = \operatorname{tr} (\mathbb{E} [\mathbf{J}(\boldsymbol{\theta}|\mathbf{Q})^{-1}]) \quad (20)$$

which is perhaps not surprising in view of $\mathbf{K}(\boldsymbol{\theta}; \mathbf{z}|\tilde{\mathbf{x}})$ being effectively a conditional ECRB matrix. The exact relationship is discussed in more detail in Section IV.

³This is similar, in a sense, to the conditional entropy [32], which involves conditioning and then taking an expectation. Note that for fixed $\boldsymbol{\theta}$, the conditional FIM $\mathbb{E} [\mathbf{J}(\boldsymbol{\theta}|\mathbf{x})]$ appears in the *modified CRB* [37].

Lemma 1. $\text{tr}(\mathbb{E}[J(\boldsymbol{\theta}|\mathbf{Q})^{-1}])$ is monotonically non-increasing and convex in $\mathbf{Q} \in \mathbb{S}_+^M$.

The above Lemma will be useful for our analysis. Its proof is presented in Appendix A-A.

B. Result and insights

We are now in position to present the main result.

Theorem 1. *The capacity-MSE function $C(\Delta)$ is characterized by*

$$\begin{aligned} \max_{\mathbf{Q} \succeq \mathbf{0}} \quad & \log \det \left(\mathbf{I} + \frac{1}{\sigma_{\omega}^2} \mathbf{H} \mathbf{Q} \mathbf{H}^H \right) \\ \text{s.t.} \quad & \text{tr}(\mathbb{E}[J(\boldsymbol{\theta}|\mathbf{Q})^{-1}]) \leq \Delta \\ & \text{tr}(\mathbf{Q}) \leq P. \end{aligned} \quad (21)$$

A detailed proof of Theorem 1 is presented in Section IV. The rest of this section is dedicated to dissecting the above result and deriving some analytical insights.

1) *Trade-off:* First, we observe that any optimal rate-MSE pair $(R, \Delta) = (C(\Delta), \Delta)$, which must necessarily lie on the capacity-MSE curve $C(\Delta)$, is characterized by

$$R = I(\tilde{\mathbf{x}}; \mathbf{y}) \quad \text{and} \quad \Delta = \text{tr}(\mathbf{K}(\boldsymbol{\theta}; \mathbf{z}|\tilde{\mathbf{x}})) \quad (22)$$

for some Gaussian input vector $\tilde{\mathbf{x}} \sim \mathcal{CN}(\mathbf{0}, \mathbf{Q})$ with an optimized choice of input covariance matrix $\mathbf{Q}$ that satisfies the power constraint $\text{tr}(\mathbf{Q}) \leq P$. That is, Gaussian input distributions are sufficient for characterizing all optimal rate-MSE trade-off points. In the terminology of [13], the optimal characterization in Theorem 1 only exhibits a so-called ST, but no DRT. This is due to the fact that $\boldsymbol{\theta}$, albeit randomly chosen, remains fixed throughout the transmission.

2) *Almost-constant covariance codes:* Despite the fact that the capacity-MSE function $C(\Delta)$ is characterized using Gaussian input distributions, the common random coding approach based on i.i.d. Gaussian codebook ensembles [15], [32] seems to be insufficient for proving the achievability of $C(\Delta)$. This is mainly due to the worst-case nature of the sensing risk in (8). The intuition here is that an i.i.d. Gaussian process will occasionally generate codewords that are very bad for sensing, which will dominate the maximal MSE risk in (8). Alternatively, in our achievability proof, we rely on analyzing codebooks in which all codewords have approximately the same sample covariance matrix. In particular, we impose that $\frac{1}{N} \sum_{n=1}^N \mathbf{x}_n(w) \mathbf{x}_n(w)^H \approx \mathbf{Q}$ must hold for all $\mathbf{x}^N(w) \in \mathcal{C}_N$, in a sense that will be made precise in the proof. This restriction guarantees a uniform MSE risk performance across all codewords (or messages); and at the same

time yields an achievable rate equal to the mutual information $I(\tilde{\mathbf{x}}; \mathbf{y})$, which is identical to the one achieved with i.i.d. Gaussian codebook ensembles with covariance $\mathbf{Q}$.

Note that in previous works, it is commonly assumed that $\frac{1}{N} \sum_{n=1}^N \mathbf{x}_n(w) \mathbf{x}_n(w)^H \approx \mathbf{Q}$ holds under ‘‘Gaussian signaling’’ as N grows large [9], [10]. The almost-constant covariance codes framework adopted here can be viewed as a rigorous formalization of this assumption.

3) *Analytical properties:* (21) constitutes a convex optimization problem, since the objective to be maximized is concave in $\mathbf{Q}$, while the constraints are convex and linear. The convex constraint function $\text{tr}(\mathbb{E}[\mathbf{J}(\boldsymbol{\theta}|\mathbf{Q})^{-1}])$, however, involves an expectation with respect to the prior of $\boldsymbol{\theta}$. Hence, solving (21) may require approximating this expectation by its sample average. In the special case where $\mathbf{G}(\boldsymbol{\theta})$ is a linear map, the dependency of $\mathbf{J}(\boldsymbol{\theta}|\mathbf{Q})$ on the prior washes out, and the problem becomes simpler. This is seen in Section V, where we investigate case studies and present numerical examples.

It is also useful to note that the function $C(\Delta)$, as described in (21), is continuous, non-decreasing, and concave in Δ (see Appendix A-B). Consequently, the rate increase is expected to slow down as the MSE grows, and exhibit a ‘‘diminishing returns’’ behavior due to concavity. This trend is indeed confirmed by the examples in Section V.

4) *BCRB-based bound:* When proving a converse for $C(\Delta)$, instead of the ECRB-based approach we follow in Section IV-C, the BCRB (i.e., van Trees inequality) [25] can be used to lower bound the MSE. This leads to an upper bound on the capacity-MSE function given by

$$\begin{aligned} \max_{\mathbf{Q} \succeq \mathbf{0}} \quad & \log \det \left(\mathbf{I} + \frac{1}{\sigma_w^2} \mathbf{H} \mathbf{Q} \mathbf{H}^H \right) \\ \text{s.t.} \quad & \text{tr}(\mathbb{E}[\mathbf{J}(\boldsymbol{\theta}|\mathbf{Q})^{-1}]) \leq \Delta \\ & \text{tr}(\mathbf{Q}) \leq P \end{aligned} \quad (23)$$

which we denote by $C^{\text{ub}}(\Delta)$ (see Appendix C for details). Note that the prior information term in the BCRB vanishes since we scale by N and take $N \rightarrow \infty$. It is readily seen that

$$\text{tr}(\mathbb{E}[\mathbf{J}(\boldsymbol{\theta}|\mathbf{Q})^{-1}]) \leq \text{tr}(\mathbb{E}[\mathbf{J}(\boldsymbol{\theta}|\mathbf{Q})^{-1}]) \quad (24)$$

which holds due to Jensen’s inequality and convexity of the matrix inverse. In fact, due to strict convexity, equality in (24) holds only if $\mathbf{J}(\boldsymbol{\theta}|\mathbf{Q})$ is invariant under $\boldsymbol{\theta}$. From (24), it is clear that (23) gives an upper bound on $C(\Delta)$, which is not achievable in general. To illustrate this relation, we include the BCRB-based upper bound in our case studies in Section V.

Remark 3. In [13], Xiong *et al.* used the Miller-Chang type BCRB [24] as a proxy for the MSE, where the random input signal $\mathbf{x}^N$ is treated as a so-called “nuisance parameter”. Adapting this approach to our setting leads to the same upper bound in (23), see Appendix C.

IV. PROOF OF THEOREM 1

This section is dedicated to proving Theorem 1. We start by presenting essential preliminaries on Bayesian parameter estimation, which will be used further on in our proof.

A. Bayesian MMSE estimation and ECRB

Suppose that a codeword $\mathbf{x}^N \in \mathcal{C}^N$ is sent (we suppress the dependency on w in this part) and $\mathbf{z}^N$ is observed by the sensor. The optimal estimator that minimizes the MSE risk in (7) is the minimum MSE (MMSE) estimator, which is given by the conditional mean

$$\hat{\boldsymbol{\theta}}_N^*(\mathbf{z}^N, \mathbf{x}^N) = \mathbb{E} [\boldsymbol{\theta} | \mathbf{z}^N = \mathbf{z}^N, \mathbf{x}^N = \mathbf{x}^N]. \quad (25)$$

The corresponding MSE matrix satisfies

$$\boldsymbol{\Sigma}(\mathbf{x}^N, \hat{\boldsymbol{\theta}}_N^*) \preceq \boldsymbol{\Sigma}(\mathbf{x}^N, \hat{\boldsymbol{\theta}}_N) \quad (26)$$

compared to any estimator $\hat{\boldsymbol{\theta}}_N$, where $\boldsymbol{\Sigma}(\mathbf{x}^N, \hat{\boldsymbol{\theta}}_N^*)$ is known to be equal to the conditional covariance matrix $\mathbf{K}_{\boldsymbol{\theta}|\mathbf{z}^N, \mathbf{x}^N=\mathbf{x}^N}$ [25], [31]. Note that (26) directly implies

$$\text{tr}(\mathbf{A}\boldsymbol{\Sigma}(\mathbf{x}^N, \hat{\boldsymbol{\theta}}_N^*)) \leq \text{tr}(\mathbf{A}\boldsymbol{\Sigma}(\mathbf{x}^N, \hat{\boldsymbol{\theta}}_N)) \quad (27)$$

for any non-negative diagonal weighting matrix $\mathbf{A}$, and hence optimality is maintained for the WMSE risk in Remark 1. We are interested in the asymptotic MSE performance.

By extending (14), let $\mathbf{J}(\boldsymbol{\theta}|\mathbf{x}^N)$ be the FIM associated with estimating $\boldsymbol{\theta} = \boldsymbol{\theta}$ from the random sequence of sensor observations $\mathbf{z}^N$ given that $\mathbf{x}^N = \mathbf{x}^N$ is transmitted. From the additive property of the FIM under conditionally independent measurements, we have

$$\mathbf{J}(\boldsymbol{\theta}|\mathbf{x}^N) = \sum_{i=1}^N \mathbf{J}(\boldsymbol{\theta}|\mathbf{x}_i). \quad (28)$$

The classical CRB is obtained by inverting the FIM $\mathbf{J}(\boldsymbol{\theta}|\mathbf{x}^N)$, while the ECRB is obtained by taking the expectation of the CRB with respect to $\boldsymbol{\theta}$, and is given by

$$\mathbf{K}(\boldsymbol{\theta}; \mathbf{z}^N | \mathbf{x}^N = \mathbf{x}^N) = \mathbb{E} [\mathbf{J}(\boldsymbol{\theta}|\mathbf{x}^N)^{-1}]. \quad (29)$$

Unlike the BCRB [25], the ECRB does not always yield a lower bound (in the semi-definite sense) on the MSE matrix for every value of N . Nevertheless, in the i.i.d. measurement regime recovered by setting $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_N = \mathbf{x}$ in (28), the ECRB provides an exact characterization of the asymptotic MSE as $N \rightarrow \infty$. This was stated by Van Trees *et al.* in [25, Sec. 4.2–4.3] and studied more recently by Aharon and Tabrikian in [34]. In particular, [34] investigates a weighted variant of the BCRB, showing its validity as a lower bound for all N and its asymptotic tightness, with convergence to the ECRB as N grows large (under regularity conditions). Next, we adapt this conclusion to the sensing setting at hand, where measurements are independent but non-identical, arising from the fact that the codeword symbols are not identical in general.

To this end, we define the empirical covariance matrix of an input sequence $\mathbf{x}^N$ as

$$\hat{\mathbf{Q}}(\mathbf{x}^N) \triangleq \frac{1}{N} \sum_{i=1}^N \mathbf{x}_i \mathbf{x}_i^H. \quad (30)$$

From (18), it is clear that the N -letter FIM in (28) is equal to a scaled conditional FIM:

$$\mathbf{J}(\boldsymbol{\theta}|\mathbf{x}^N) = N\mathbf{J}(\boldsymbol{\theta}|\hat{\mathbf{Q}}(\mathbf{x}^N)). \quad (31)$$

This relation is key to formulating the conditional ECRB tightness result presented next.

Lemma 2. Fix $\epsilon > 0$. There exists a (possibly large) block-length N_ϵ such that

$$\text{tr} \left(\mathbb{E} \left[\mathbf{J}(\boldsymbol{\theta}|\hat{\mathbf{Q}}(\mathbf{x}^N))^{-1} \right] \right) - \epsilon \leq N \text{tr} \left(\boldsymbol{\Sigma} \left(\mathbf{x}^N, \hat{\boldsymbol{\theta}}_N^* \right) \right) \leq \text{tr} \left(\mathbb{E} \left[\mathbf{J}(\boldsymbol{\theta}|\hat{\mathbf{Q}}(\mathbf{x}^N))^{-1} \right] \right) + \epsilon \quad (32)$$

holds for every $N \geq N_\epsilon$, provided that some regularity conditions hold.

The above lemma follows by slightly adapting existing results for i.i.d. measurements, as shown by Seo [30] in the scalar parameter case. In particular, the lower bound in (32) can be obtained from the weighted BCRB in [34, Theorem 8], which is shown to approach the ECRB as N grows large (see [34, Prop. 3]). The upper bound in (32) can be obtained using the sub-optimal maximum likelihood (ML) estimator, whose asymptotic MSE is characterized exactly by the ECRB as discussed in [25, Sec. 4.2–4.3] (see also [34, eq. (12)]).

B. Achievability

In order to prove achievability, it suffices to show that for an arbitrary fixed covariance matrix $\mathbf{Q} \succeq 0$ that satisfies $\text{tr}(\mathbf{Q}) \leq P$, and for every $\epsilon > 0$ and N large enough, there exists an $(N, M_N, \varepsilon_N, \xi_N)$ -scheme with $\varepsilon_N \leq \epsilon$ and

$$\frac{1}{N} \log M_N \geq \log \det \left(\mathbf{I} + \frac{1}{\sigma_\omega^2} \mathbf{H} \mathbf{Q} \mathbf{H}^H \right) - \epsilon \quad (33)$$

$$N\xi_N \leq \text{tr} \left(\mathbb{E} \left[\mathbf{J}(\boldsymbol{\theta}|\mathbf{Q})^{-1} \right] \right) + \epsilon. \quad (34)$$

By optimizing $\mathbf{Q}$ and making $\epsilon \rightarrow 0$ through $N \rightarrow \infty$, we achieve all capacity-MSE trade-off points. We now break down the achievability proof into three parts as follows.

1) *Codebook constraint*: For fixed N and codebook $\mathcal{C}_N$, the sensing risk is dictated by the *worst* codeword, as clearly seen for (8). To guarantee a somewhat uniform MSE across codewords, we impose the constraint that codewords of $\mathcal{C}_N$ must be drawn from

$$\mathcal{B}_{N,\delta}(\mathbf{Q}) \triangleq \left\{ \mathbf{x}^N \in \mathbb{C}^N : \mathbf{Q} - \delta \mathbf{I} \preceq \hat{\mathbf{Q}}(\mathbf{x}^N) \preceq \mathbf{Q} + \delta \mathbf{I} \right\} \quad (35)$$

where $\delta > 0$ is sufficiently small. All sequences $\mathbf{x}^N \in \mathcal{B}_{N,\delta}(\mathbf{Q})$ have an empirical covariance matrix $\hat{\mathbf{Q}}(\mathbf{x}^N)$ that is sufficiently close to $\mathbf{Q}$, in the semi-definite sense. Next, we show that (i) by choosing all codewords from $\mathcal{B}_{N,\delta}(\mathbf{Q})$, the MSE in (34) is achieved; and (ii) there exists a codebook with codewords from $\mathcal{B}_{N,\delta}(\mathbf{Q})$ that achieves the rate in (33).

2) *MSE*: Suppose that $\mathcal{C}_N \subset \mathcal{B}_{N,\delta}(\mathbf{Q})$ and consider the codeword $\mathbf{x}^N(w) \in \mathcal{C}_N$. Assuming MMSE estimation, and starting from the upper bound in (32), we obtain

$$N\xi_N(w) \leq \text{tr} \left(\mathbb{E} \left[\mathbf{J}(\boldsymbol{\theta}|\hat{\mathbf{Q}}(\mathbf{x}^N(w)))^{-1} \right] \right) + \epsilon' \quad (36)$$

$$\leq \text{tr} \left(\mathbb{E} \left[\mathbf{J}(\boldsymbol{\theta}|\mathbf{Q} - \delta \mathbf{I})^{-1} \right] \right) + \epsilon' \quad (37)$$

$$\leq \text{tr} \left(\mathbb{E} \left[\mathbf{J}(\boldsymbol{\theta}|\mathbf{Q})^{-1} \right] \right) + \delta' + \epsilon' \quad (38)$$

where (37) is due to monotonicity while (38) is due to continuity (see Lemma 1). Note that $\delta' + \epsilon'$ can be made small by controlling δ and N . Since (38) holds for all $w \in \mathcal{M}_N$, it holds by taking the maximum, from which we obtain (34).

3) *Rate*: We now wish to show that there exists a codebook $\mathcal{C}_N \subset \mathcal{B}_{N,\delta}(\mathbf{Q})$ of size M_N that satisfies (33) and has a sufficiently small maximal decoding error probability.

To this end, we invoke a refined version of the channel coding theorem known in the literature as Feinstein's lemma (see [36] for a comprehensive treatment). The advantage of using Feinstein's lemma compared to the abundant typicality-based approach [32], [38] is that it provides guarantees on the maximal (as opposed to average) decoding error probability, and more importantly, it is better suited for handling multi-letter input distributions (non i.i.d. codebooks) and arbitrary codeword constraint sets. The version we adopt incorporates a constraint set, and is referred to in [36] as the extended Feinstein's lemma. To this end, recall the information density defined in (12), and let $\iota(\mathbf{x}^N; \mathbf{y}^N)$ be the corresponding N -letter extension. Moreover, in what follows,

let $\tilde{\mathbf{x}}^N = \tilde{\mathbf{x}}_1, \tilde{\mathbf{x}}_2, \dots, \tilde{\mathbf{x}}_n$ be an i.i.d. sequence of Gaussian vectors, in which every $\tilde{\mathbf{x}}_n$ is drawn independently from $\mathcal{CN}(\mathbf{0}, \mathbf{Q})$; and $\mathbf{y}^N$ be a random output sequence induced by the input $\tilde{\mathbf{x}}^N$.

Lemma 3 (Extended Feinstein's lemma [36, Th. 20.7]). *Fix N and $\gamma_N > 0$. There exists a codebook $\mathcal{C}_N \subset \mathcal{B}_{N,\delta}(\mathbf{Q})$ of size M_N and maximal decoding error probability that satisfies*

$$\varepsilon_N \times \mathbb{P}[\tilde{\mathbf{x}}^N \in \mathcal{B}_{N,\delta}(\mathbf{Q})] \leq \mathbb{P}[\iota(\tilde{\mathbf{x}}^N; \mathbf{y}^N) \leq \log \gamma_N] + \frac{M_N}{\gamma_N}. \quad (39)$$

Note that since $\tilde{\mathbf{x}}^N$ is i.i.d. and the channel is memoryless, the information density in our case is additive, and for every realization $(\tilde{\mathbf{x}}^N, \mathbf{y}^N)$ of $(\tilde{\mathbf{x}}^N, \mathbf{y}^N)$ we have

$$\iota(\tilde{\mathbf{x}}^N; \mathbf{y}^N) = \sum_{n=1}^N \iota(\tilde{\mathbf{x}}_n; \mathbf{y}_n) = \sum_{n=1}^N \log \frac{f_{\mathbf{y}|\mathbf{x}}(\mathbf{y}_n|\tilde{\mathbf{x}}_n)}{f_{\mathbf{y}}(\mathbf{y}_n)}. \quad (40)$$

Moreover, since the mutual information is the expected information density (see (11)), then

$$\mathbb{E}[\iota(\tilde{\mathbf{x}}^N; \mathbf{y}^N)] = I(\tilde{\mathbf{x}}^N; \mathbf{y}^N) = NI(\tilde{\mathbf{x}}; \mathbf{y}) \quad (41)$$

where $\tilde{\mathbf{x}} \sim \mathcal{CN}(\mathbf{0}, \mathbf{Q})$. Going back to Lemma 3, we set

$$\frac{1}{N} \log \gamma_N = I(\tilde{\mathbf{x}}; \mathbf{y}) - \frac{1}{2}\epsilon \quad \text{and} \quad \frac{1}{N} \log M_N = I(\tilde{\mathbf{x}}; \mathbf{y}) - \epsilon.$$

The error probability bound becomes

$$\varepsilon_N \times \mathbb{P}[\tilde{\mathbf{x}}^N \in \mathcal{B}_{N,\delta}(\mathbf{Q})] \leq \mathbb{P}\left[\frac{1}{N} \sum_{n=1}^N \iota(\tilde{\mathbf{x}}_n; \mathbf{y}_n) \leq I(\tilde{\mathbf{x}}; \mathbf{y}) - \frac{\epsilon}{2}\right] + 2^{-N\frac{\epsilon}{2}}. \quad (42)$$

Since the sequences involved are i.i.d., by the weak law of large numbers (WLLN), we get

$$\mathbb{P}\left[\frac{1}{N} \sum_{n=1}^N \iota(\tilde{\mathbf{x}}_n; \mathbf{y}_n) \leq I(\tilde{\mathbf{x}}; \mathbf{y}) - \frac{\epsilon}{2}\right] \rightarrow 0 \quad \text{as } N \rightarrow \infty. \quad (43)$$

Moreover, we have $2^{-N\frac{\epsilon}{2}} \rightarrow 0$ as $N \rightarrow \infty$. Therefore, the right-hand side of (42) approaches zero as N grows large. The following lemma is key for completing the proof.

Lemma 4. *Let $\tilde{\mathbf{x}}^N$ be an i.i.d. sequence of Gaussian vectors with distribution $\mathcal{CN}(\mathbf{0}, \mathbf{Q})$. Then*

$$\mathbb{P}[\tilde{\mathbf{x}}^N \in \mathcal{B}_{N,\delta}(\mathbf{Q})] \rightarrow 1 \quad \text{as } N \rightarrow \infty. \quad (44)$$

In words, the above lemma says that an i.i.d. $\mathcal{CN}(\mathbf{0}, \mathbf{Q})$ sequence will, with high probability, have an empirical covariance matrix that is close to $\mathbf{Q}$ in the positive semi-definite sense. A proof is presented in Appendix B. Putting everything together, we conclude that $\varepsilon_N \leq \epsilon$ is attainable by making $N \rightarrow \infty$ large enough, which concludes the proof of achievability.

C. Converse

Here we show that any achievable pair (R, Δ) must satisfy $R \leq C(\Delta)$, where $C(\Delta)$ is as characterized in Theorem 1. For this purpose, recall from Definition 1 that if (R, Δ) is achievable, then for large enough N , there exists an $(N, M_N, \varepsilon_N, \xi_N)$ -scheme with $\varepsilon_N \leq \epsilon$, $\frac{1}{N} \log M_N \geq R - \epsilon$ and $N\xi_N \leq \Delta + \epsilon$. We now proceed to bound the sensing MSE risk as follows

$$\Delta + \epsilon \geq N \max_{w \in \mathcal{M}_N} \xi_N(w) \quad (45)$$

$$\geq \frac{N}{M_N} \sum_{w \in \mathcal{M}_N} \xi_N(w) \quad (46)$$

$$\geq \frac{1}{M_N} \sum_{w \in \mathcal{M}_N} \text{tr} \left(\mathbb{E} \left[\mathbf{J}(\boldsymbol{\theta} | \hat{\mathbf{Q}}(\mathbf{x}^N(w)))^{-1} \right] \right) - \epsilon \quad (47)$$

$$\geq \text{tr} \left(\mathbb{E} \left[\mathbf{J} \left(\boldsymbol{\theta} \mid \frac{1}{M_N} \sum_{w \in \mathcal{M}_N} \hat{\mathbf{Q}}(\mathbf{x}^N(w)) \right)^{-1} \right] \right) - \epsilon \quad (48)$$

where (46) is since the maximum (over messages) is lower bounded by the average, (47) is from the lower bound in (32), and (48) is due to the convexity of (20) and Jensen's inequality (see Lemma 1). Next, we denote the empirical covariance matrix average over messages by

$$\bar{\mathbf{Q}} \triangleq \frac{1}{M_N} \sum_{w \in \mathcal{M}_N} \hat{\mathbf{Q}}(\mathbf{x}^N(w)) \quad (49)$$

which clearly must also satisfy $\bar{\mathbf{Q}} \succeq \mathbf{0}$ and $\text{tr}(\bar{\mathbf{Q}}) \leq P$. The bound in (48) is rewritten as

$$\Delta \geq \text{tr} \left(\mathbb{E} \left[\mathbf{J}(\boldsymbol{\theta} | \bar{\mathbf{Q}})^{-1} \right] \right) - 2\epsilon. \quad (50)$$

We proceed to bound the rate. To this end, we invoke Fano's inequality [32], [38] and obtain

$$R - \epsilon \leq \frac{1}{N} \log M_N \quad (51)$$

$$\leq \frac{1}{N} I(\mathbf{x}^N; \mathbf{y}^N) + \frac{1}{N} + \frac{\varepsilon_N^{\text{av}}}{N} \log M_N \quad (52)$$

where the term $\frac{1}{N} + \frac{\varepsilon_N^{\text{av}}}{N} \log M_N$ vanishes as N grows large, and hence we bound it by ϵ from now on. Moreover, it is worthwhile noting that Fano's inequality yields a bound in terms of $\varepsilon_N^{\text{av}}$. This is, however, valid for our purpose since $\varepsilon_N^{\text{av}} \leq \varepsilon_N$.

The next step is to bound the normalized multi-letter mutual information term $\frac{1}{N} I(\mathbf{x}^N; \mathbf{y}^N)$. To this end, it helps to remember that the random input sequence $\mathbf{x}^N = \mathbf{x}^N(w)$ is uniformly distributed on the codebook $\mathcal{C}_N$. Using the fact that the channel is memoryless, we proceed as

$$\frac{1}{N} I(\mathbf{x}^N; \mathbf{y}^N) \leq \frac{1}{N} \sum_{n=1}^N I(\mathbf{x}_n; \mathbf{y}_n) \quad (53)$$

$$\leq \frac{1}{N} \sum_{n=1}^N \log \det \left(\mathbf{I} + \frac{1}{\sigma_{\omega}^2} \mathbf{H} \mathbb{E} [\mathbf{x}_n \mathbf{x}_n^H] \mathbf{H}^H \right) \quad (54)$$

$$\leq \log \det \left(\mathbf{I} + \frac{1}{\sigma_{\omega}^2} \mathbf{H} \frac{1}{N} \sum_{n=1}^N \mathbb{E} [\mathbf{x}_n \mathbf{x}_n^H] \mathbf{H}^H \right). \quad (55)$$

The inequality in (54) follows from the extremal property of Gaussian inputs under a covariance (or correlation) matrix constraint [38, Sec. 2.2], which leads to $I(\mathbf{x}_n; \mathbf{y}_n)$ being maximized by an input distribution $\mathcal{CN}(\mathbf{0}, \mathbb{E} [\mathbf{x}_n \mathbf{x}_n^H])$. (55) is due to concavity and Jensen's inequality.

Using the fact that $\mathbf{x}^N = \mathbf{x}^N(w)$ is uniformly distributed on the codebook, we observe that

$$\frac{1}{N} \sum_{n=1}^N \mathbb{E} [\mathbf{x}_n \mathbf{x}_n^H] = \frac{1}{M_N} \sum_{w \in \mathcal{M}_N} \frac{1}{N} \sum_{n=1}^N \mathbf{x}_n(w) \mathbf{x}_n(w)^H = \bar{\mathbf{Q}}. \quad (56)$$

Plugging everything back into (52), we obtain

$$R \leq \log \det \left(\mathbf{I} + \frac{1}{\sigma_{\omega}^2} \mathbf{H} \bar{\mathbf{Q}} \mathbf{H}^H \right) + 2\epsilon \quad (57)$$

$$\leq C(\Delta + 2\epsilon) + 2\epsilon \quad (58)$$

where the last inequality is obtained by maximizing over $\bar{\mathbf{Q}}$ subject to the MSE constraint in (50), and the power constraint. Taking $\epsilon \rightarrow 0$ as $N \rightarrow \infty$, we obtain the desired result.

Remark 4. It can be seen from our converse proof that we bounded the maximal (over messages) MSE risk ξ_N in terms of its average counterpart ξ_N^{av} in (46); and the maximal decoding error probability ε_N in terms of its average counterpart $\varepsilon_N^{\text{av}}$ in (52). These are, of course, valid steps as they are applied in the converse part only. More importantly, they imply that the characterization obtained in Theorem 1 remains valid if we adopt less stringent average risk and error measures.

V. CASE STUDIES

A. JCAS with DoA Estimation

In our first case study, we consider scenarios where the transmitter is communicating with a single antenna user, i.e., $L = 1$; while estimating the DoA of a passive point target in its range. The transmitter and the sensor are equipped with uniform linear arrays (ULAs) of M and T antennas, respectively. Assuming half-wavelength element spacing for both arrays, the steering vectors at the transmitter and at the sensor are given by

$$\mathbf{v}_M(\theta) \triangleq \left[1 \ e^{i\pi \sin(\theta)} \ e^{i2\pi \sin(\theta)} \ \dots \ e^{i(M-1)\pi \sin(\theta)} \right]^T \quad (59)$$

$$\mathbf{v}_T(\theta) \triangleq \begin{bmatrix} 1 & e^{i\pi \sin(\theta)} & e^{i2\pi \sin(\theta)} & \dots & e^{i(T-1)\pi \sin(\theta)} \end{bmatrix}^T \quad (60)$$

respectively. We adopt a line-of-sight, single-ray model for the communication channel as

$$\mathbf{H} = \beta \mathbf{v}_M(\phi)^H \quad (61)$$

where $\phi \in (-\frac{\pi}{2}, \frac{\pi}{2})$ denotes the angle of departure (AoD) from the transmitter array towards the user's direction, and $\beta \in \mathbb{C}$ denotes the channel gain. As stated in Section II, we assume that the user channel matrix is fixed and known to both the transmitter and the user through, e.g., an initial access phase. For simplicity, we set $\beta = 1$. As for the sensing channel, this is given as

$$\mathbf{G}(\theta) = \lambda \mathbf{v}_T(\theta) \mathbf{v}_M(\theta)^H \quad (62)$$

where θ is the DoA parameter to be estimated, taking values from $(-\frac{\pi}{2}, \frac{\pi}{2})$, and $\lambda \in \mathbb{C}$ is the complex back-scatter channel gain (see, e.g., [9], among others). For simplicity, we assume that the channel gain is known *a priori* and set $\lambda = 1$. Furthermore, the angle of arrival (AoA) and the AoD of the sensing channel coincide in our monostatic setting, as the transmitter and the sensor arrays are co-located and parallel. θ is drawn at random according to a prior $f_\theta(\theta)$, yet remains fixed throughout the transmission. We examine two different prior distributions specified below. The procedure and analysis can be extended to the joint estimation of DoA and gain [39].

Capacity-MSE trade-off: To calculate the conditional FIM in accordance with (18), which reduces to a scalar for single-parameter estimation, the first-order derivative of the sensing channel matrix with respect to the target parameter θ is obtained as

$$\frac{\partial \mathbf{G}(\theta)}{\partial \theta} = i\pi \cos(\theta) (\mathbf{T}\mathbf{G}(\theta) - \mathbf{G}(\theta)\mathbf{M}) \quad (63)$$

where $\mathbf{T} \triangleq \text{diag}(0, \dots, T-1)$ and $\mathbf{M} \triangleq \text{diag}(0, \dots, M-1)$ represent element positions in the sensor and transmitter arrays, respectively. Defining $\tilde{\mathbf{G}}(\theta) \triangleq \mathbf{T}\mathbf{G}(\theta) - \mathbf{G}(\theta)\mathbf{M}$, we obtain

$$J(\theta|\mathbf{Q}) = \frac{2\pi^2}{\sigma_v^2} \cos^2(\theta) \text{tr} \left(\tilde{\mathbf{G}}(\theta) \mathbf{Q} \tilde{\mathbf{G}}(\theta)^H \right) \quad (64)$$

where the operator $\text{Re}\{\cdot\}$ is dropped since the trace of a Hermitian matrix is real.

From Theorem 1, the capacity-MSE function $C(\Delta)$ specialized for the present use case is

$$\begin{aligned} & \max_{\mathbf{Q} \succeq \mathbf{0}} \log \left(1 + \frac{1}{\sigma_\omega^2} \mathbf{v}_M(\phi)^H \mathbf{Q} \mathbf{v}_M(\phi) \right) \\ & \text{s.t. } \frac{\sigma_v^2}{2\pi^2} \mathbb{E} \left[\frac{1}{\cos^2(\theta) \text{tr}(\tilde{\mathbf{G}}(\theta) \mathbf{Q} \tilde{\mathbf{G}}(\theta))} \right] \leq \Delta \\ & \text{tr}(\mathbf{Q}) \leq P. \end{aligned} \quad (65)$$

For Δ sufficiently large, such that the sensing constraint is inactive, the capacity is given by $\log(1 + MP/\sigma_\omega^2)$, where M is the array gain. In this case, beamforming along the user's AoD is optimal, i.e. $\mathbf{Q} = \frac{P}{M} \mathbf{v}_M(\phi) \mathbf{v}_M(\phi)^H$. More generally, by optimizing $\mathbf{Q}$, which reflects the beamforming characteristics of the transmit array, a clear trade-off between the rate and the MSE is obtained, as demonstrated in our numerical results below.

In addition to $C(\Delta)$, we will also examine the BCRB-based upper bound $C^{\text{rub}}(\Delta)$ given in (23). This is obtained by replacing the sensing constraint in (65) with

$$\frac{\sigma_v^2}{2\pi^2} \frac{1}{\mathbb{E}[\cos^2(\theta) \text{tr}(\tilde{\mathbf{G}}(\theta) \mathbf{Q} \tilde{\mathbf{G}}(\theta))]} \leq \Delta \quad (66)$$

which is a looser constraint, as discussed in Section III-B4. In the numerical results presented below, we set $T = M = 16$ for the number of transmitter and sensor antennas. The communication SNR is defined as P/σ_ω^2 and set to 15 dB. Similarly, the sensing SNR is defined as P/σ_v^2 , and set to -25 dB. We examine two types of prior distribution for θ : uniform and non-uniform. The uniform case is meant to represent scenarios where no prior information is available; while the non-uniform case represents scenarios where prior information suggests that the target is more likely to be in a certain angular sector as in, e.g., tracking from a previous estimate.

Tapered uniform prior: Under a uniform prior, one would set $f_\theta(\theta) = 1/\pi$ for $\theta \in (-\frac{\pi}{2}, \frac{\pi}{2})$ and zero otherwise. However, as highlighted in [34], [40], Bayesian variants of the CRB (e.g., BCRB and ECRB) are not defined for such a prior as it violates the regularity conditions. Therefore, as suggested in [34], we adopt a tapered uniform distribution model with a raised-cosine

$$f_\theta(\theta) = \frac{1}{s(1+\kappa)} \begin{cases} 1, & |\theta| < s\kappa \\ \frac{1}{2} + \left(1 + \cos\left(\frac{\pi(|\theta|-s\kappa)}{s(1-\kappa)}\right)\right), & s\kappa \leq |\theta| \leq s \\ 0, & |\theta| > s \end{cases} \quad (67)$$

where $s \in \mathbb{R}^+$ determines the boundaries of the distribution and $\kappa \in [0, 1]$ is the roll-off factor that controls the sharpness at the edges. $\kappa = 0$ corresponds to a pure raised-cosine shaped distribution, whereas $\kappa = 1$ corresponds to a perfect uniform distribution. We set $s = \pi/2$ to cover the entire range and $\kappa = 0.7$ to have a tapered uniform distribution. Note that since ULAs have lower resolution in the endfire directions, i.e., $\theta \approx \pm \frac{\pi}{2}$, it is reasonable to bias the estimate away from these directions by setting a low prior value. We set $\phi = 0$ for the user channel. The tapered uniform prior and the user location are shown in Fig. 2a.

Under the above specified assumptions, the capacity-MSE function and its BCRB-based upper bound are shown in Fig. 2b. By varying Δ in (65), and solving the corresponding optimization

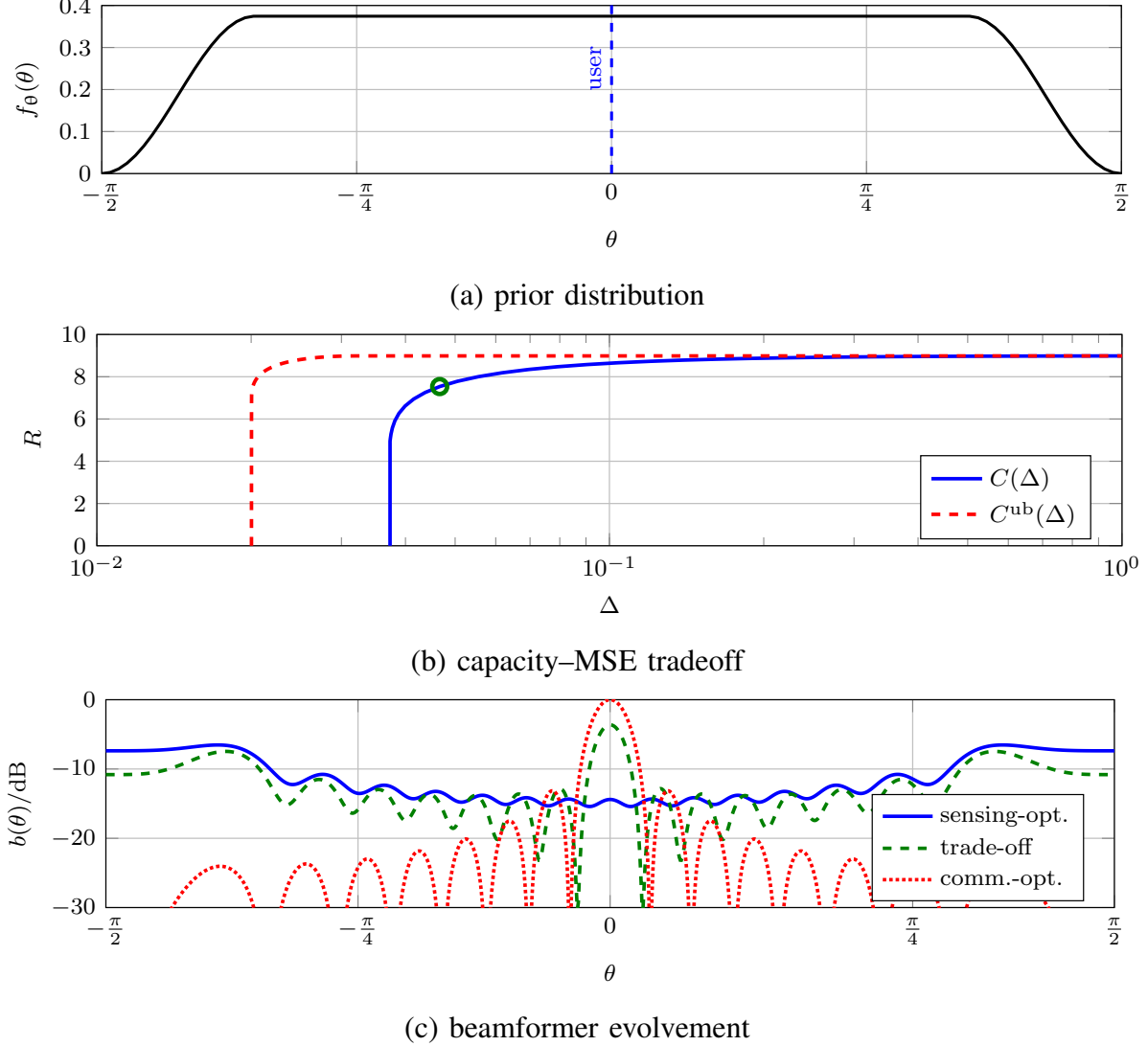


Fig. 2: Capacity-MSE trade-off with a uniform prior in JCAS with DoA estimation.

problem, different rate-MSE trade-off points are obtained. Moreover, the looseness of the BCRB-based upper bound compared to the ECRB-based exact characterization is clearly seen.

Next, we also utilize the projection $b(\theta) \triangleq \frac{1}{MP} \mathbf{v}_M(\theta)^H \mathbf{Q} \mathbf{v}_M(\theta)$ to calculate the normalized beamforming gain in any given angular direction θ , and show the result in Fig. 2c. To illustrate a specific trade-off choice, we select a point on the $C(\Delta)$ curve indicated with a marker in Fig. 2b, whose optimal beamforming characteristics are also illustrated in Fig. 2c. While the sensing-optimal beamforming favors a wide scan of the entire range, due to the almost uniform prior; beam steering towards the user's direction gradually turns into the optimal beamformer.

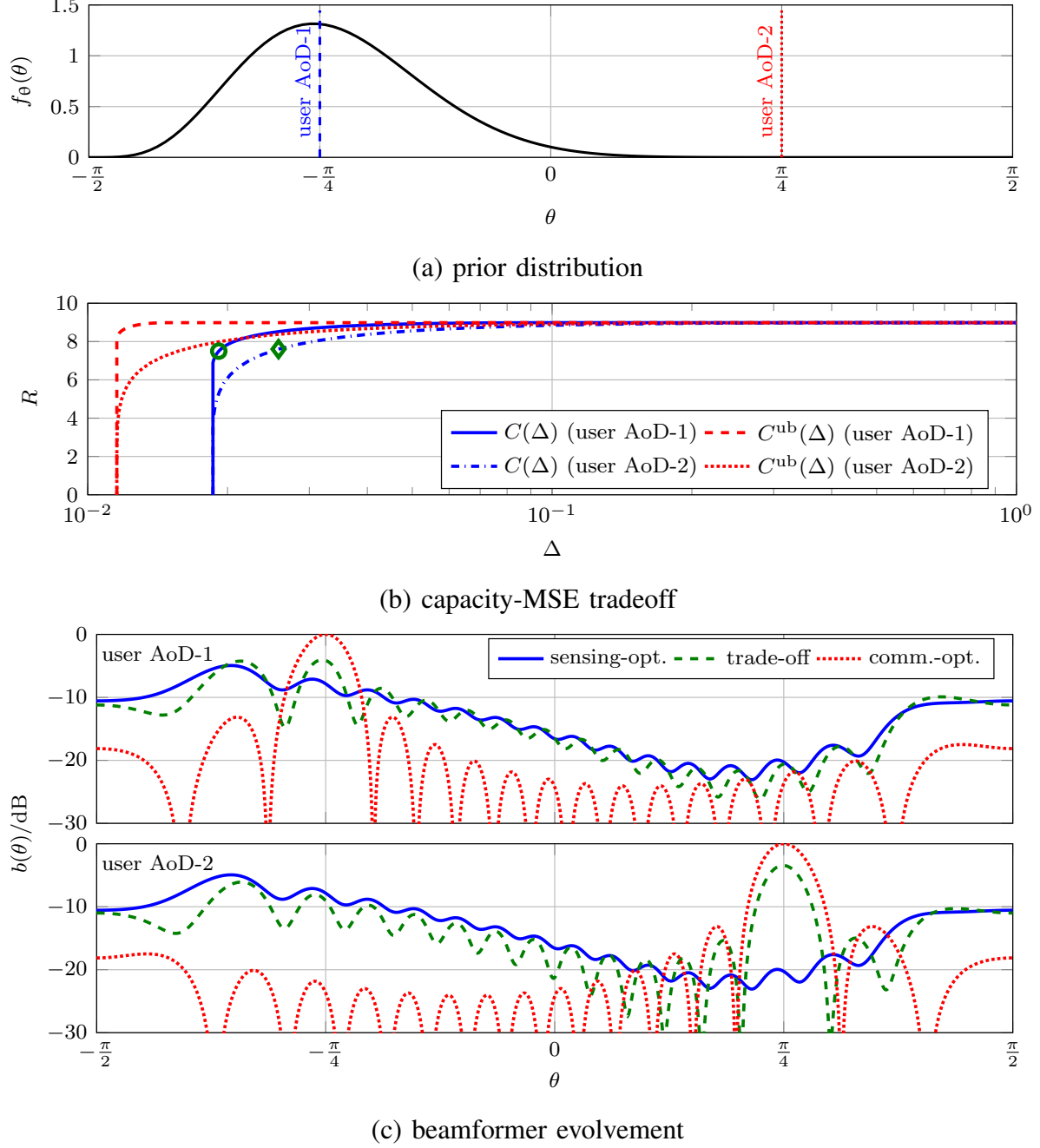


Fig. 3: Capacity-MSE trade-off with a non-uniform prior in JCAS with DoA estimation.

Non-uniform prior: For this case, we utilize the Beta distribution to model a concentrated prior for which $f_{\theta}(\theta) = 0$ when $\theta = \pm\pi/2$. In particular, we have

$$f_{\theta}(\theta) = \frac{(\theta - \theta_{\min})^{s_1-1}(\theta_{\max} - \theta)^{s_2-1}}{(\theta_{\max} - \theta_{\min})^{s_1+s_2-1}B(s_1, s_2)} \quad (68)$$

for $\theta \in [\theta_{\min}, \theta_{\max}]$ and zero otherwise, where $B(s_1, s_2) \triangleq \int_0^1 t^{s_1-1} (1-t)^{s_2-1} dt$ is the Beta function, and s_1 and s_2 are the shape parameters of the distribution. We set $\theta_{\min} = -\pi/2$ and $\theta_{\max} = \pi/2$ so that the support of the distribution spans the full angular range. The shape parameters are chosen as $s_1 = 5.5$ and $s_2 = 15$, which results in a distribution that is more concentrated around a specific region that peaks around $-\pi/4$. This selection ensures that the prior still satisfies the regularity conditions while favoring a particular region in the angular domain. With the non-uniform prior, we consider two scenarios for the communication channel: first, $\phi = -\pi/4$, i.e., the user is located in the region where the target is most likely present, and second, $\phi = \pi/4$, to represent the scenario where the target is well separated from the user in the angular domain. The prior and the user AoD in the two scenarios are shown in Fig. 3a.

The capacity-MSE function and its BCRB-based upper bound are depicted in Fig. 3b, for two different user AoD scenarios. The optimal beamforming gains for the two scenarios are also shown in Fig. 3c. For both user AoD scenarios, a point is selected on the capacity-MSE curve to illustrate the beamforming characteristics of an intermediate trade-off point.

For user AoD-1, since the prior already favors the user direction, the capacity-MSE function exhibits a limited trade-off. In this case, the user and the target are highly aligned, and therefore, a $\mathbf{Q}$ which is good for communication is also good for sensing. This is in contrast to the user AoD-2 scenario, where the target's prior concentrates further away from the user. The evolution of beam patterns in Fig. 3c is also a clear indication of the trade-off in the misaligned case.

B. JCAS with OFDM

We now consider a cyclic-prefix OFDM scenario, where the vector dimension represents sub-carriers instead of antennas. The system communicates with a user and simultaneously senses the back-scatter channel characterized by the frequency response vector $\boldsymbol{\theta} \in \mathbb{C}^K$, where K is the number of sub-carriers being utilized. Assuming that the effect of the cyclic prefix is removed at the receiver, the user channel matrix in this case is given as

$$\mathbf{H} = \text{diag}(\mathbf{F}\mathbf{h}) \quad (69)$$

where $\mathbf{F} \in \mathbb{C}^{K \times K}$ is the unitary discrete Fourier transform matrix (DFT) and $\mathbf{h} \in \mathbb{C}^K$ is the impulse response of the communication channel in the time-domain. We assume that $\mathbf{H}$ is known to the transmitter and receiver. Similarly, the sensor channel matrix in (2) can be written as

$$\mathbf{G}(\boldsymbol{\theta}) = \text{diag}(\mathbf{F}\boldsymbol{\theta}). \quad (70)$$

Note that in this use case, we have $M = L = T = K$. It is also clear that the sensing channel matrix is linear in the target parameter vector $\boldsymbol{\theta}$. To allow for concise representation, we allow the parameter vector, i.e. $\boldsymbol{\theta}$, to be complex-valued. The model can be equivalently expressed using a $2K$ -dimensional real parameter vector, consistent with (1) and (2). For more on the real-valued representation of complex parameters in estimation, see [31, Ch. 15].

As we shall see shortly, for this use case, the FIM is invariant under different realizations of the target parameter. Hence, the prior density $f_{\boldsymbol{\theta}}(\boldsymbol{\theta})$ does not influence the JCAS trade-off.

Capacity-MSE trade-off: Since the user channel matrix is diagonal, the mutual information in (13) is maximized by a diagonal input covariance matrix. Moreover, the conditional FIM and ECRB only depend on the input covariance matrix through its diagonal elements. To see this, we evaluate the conditional FIM, i.e., (18), for this setting. We first calculate the partial derivative of the sensor channel matrix with respect to the i -th entry of the parameter as

$$\frac{\partial \mathbf{G}(\boldsymbol{\theta})}{\partial \theta_i} = \text{diag}(\mathbf{f}_i) \quad (71)$$

where $\mathbf{f}_i$ denotes the i -th column of the unitary DFT matrix $\mathbf{F}$. From this, we obtain

$$[\mathbf{J}(\boldsymbol{\theta}|\mathbf{Q})]_{ij} = \frac{1}{\sigma_v^2} \text{tr} \left(\frac{\partial \mathbf{G}(\boldsymbol{\theta})}{\partial \theta_j} \mathbf{Q} \frac{\partial \mathbf{G}(\boldsymbol{\theta})}{\partial \theta_i} \right) \quad (72)$$

$$= \frac{1}{\sigma_v^2} \text{tr} (\text{diag}(\mathbf{f}_j) \mathbf{Q} \text{diag}(\mathbf{f}_i^H)) \quad (73)$$

$$= \frac{1}{\sigma_v^2} \mathbf{f}_i^H (\mathbf{I} \circ \mathbf{Q}) \mathbf{f}_j \quad (74)$$

where $\circ$ denotes the Hadamard (element-wise) product between two matrices. The resulting conditional FIM can be expressed more compactly as

$$\mathbf{J}(\boldsymbol{\theta}|\mathbf{Q}) = \frac{1}{\sigma_v^2} \mathbf{F}^H (\mathbf{I} \circ \mathbf{Q}) \mathbf{F}. \quad (75)$$

Note that this also directly follows from the FIM for linear models [31, Ex. 15.9]. From (75), we see that the conditional FIM does not depend on the off-diagonal terms of $\mathbf{Q}$. Hence, restriction to diagonal matrices does not influence the sensing performance either. Therefore, in what follows, we assume without loss of generality that $\mathbf{Q} = \text{diag}(\boldsymbol{\sigma}_x^2)$, where $\boldsymbol{\sigma}_x^2 \triangleq (\sigma_{x,0}^2, \dots, \sigma_{x,K-1}^2)$.

Continuing with a diagonal input covariance matrix, (75) becomes

$$\mathbf{J}(\boldsymbol{\theta}|\boldsymbol{\sigma}_x^2) = \frac{1}{\sigma_v^2} \mathbf{F}^H \text{diag}(\boldsymbol{\sigma}_x^2) \mathbf{F}. \quad (76)$$

and the asymptotic MSE is given by

$$\text{tr} (\mathbf{J}(\boldsymbol{\theta}|\boldsymbol{\sigma}_x^2)^{-1}) = \sigma_v^2 \sum_{k=0}^{K-1} \frac{1}{\sigma_{x,k}^2}. \quad (77)$$

On the other hand, the rate expression is easily obtained by specializing (13), leading to the well-known sum-rate expression of parallel Gaussian channels [32, Sec. 9.4]. Combining this with Theorem 1, we obtain the following capacity-MSE trade-off

$$\begin{aligned}
& \max_{\text{diag}(\sigma_{\mathbf{x}}^2) \succeq \mathbf{0}} \sum_{k=0}^{K-1} \log \left(1 + \frac{\sigma_{\mathbf{x},k}^2}{\sigma_{\omega}^2} |[\mathbf{H}]_{kk}|^2 \right) \\
& \text{s.t.} \quad \sigma_v^2 \sum_{k=0}^{K-1} \frac{1}{\sigma_{\mathbf{x},k}^2} \leq \Delta \\
& \quad \sum_{k=0}^{K-1} \sigma_{\mathbf{x},k}^2 \leq P
\end{aligned} \tag{78}$$

which boils down to power allocation across parallel channels under an ECRB constraint.

Remark 5. In this special case, where the sensor observations are a linear function of the parameter vector corrupted by Gaussian noise, the conditional FIM in (76) is invariant with respect to the parameter realization $\boldsymbol{\theta}$. Therefore, equality holds in (24) and the capacity-MSE function $C(\Delta)$ coincides with the BCRB-based upper bound $C^{\text{ub}}(\Delta)$. Moreover, $C(\Delta)$ remains the same if one adopts a non-Bayesian framework with the standard CRB instead of the ECRB, as in [10]. In that sense, the characterization in (78) can be seen as a special case of [10, eq. (18)], notwithstanding that the Bayesian framework in the present paper is more general.

The communication-optimal trade-off point—arising when Δ in (78) is sufficiently large so that the sensing constraint becomes inactive—is achieved by a water-filling sub-carrier power allocation [32, Sec. 9.6]. On the other hand, the sensing optimal trade-off point is achieved with an equal power allocation [10, Prop. 2]. This follows from the fact that the MSE in (77) is convex in the power allocation vector. Specifically, by Jensen’s inequality, we have

$$\sigma_v^2 \sum_{k=0}^{K-1} \frac{1}{\sigma_{\mathbf{x},k}^2} \geq \sigma_v^2 \frac{K}{\frac{1}{K} \sum_{k=0}^{K-1} \sigma_{\mathbf{x},k}^2} \tag{79}$$

$$= \sigma_v^2 K^2 / P \tag{80}$$

with equality when $\sigma_{\mathbf{x},k}^2 = P/K$ for all k . More generally, the structure of the optimal solution beyond these two extremes has been studied extensively in [10]. It is worth noting that the sensing constraint in (78) requires non-zero power allocation across all sub-carriers in order to yield a finite MSE Δ . This contrasts with the communication-optimal water-filling solution, which may deactivate some sub-carriers. This phenomenon is further explored next.

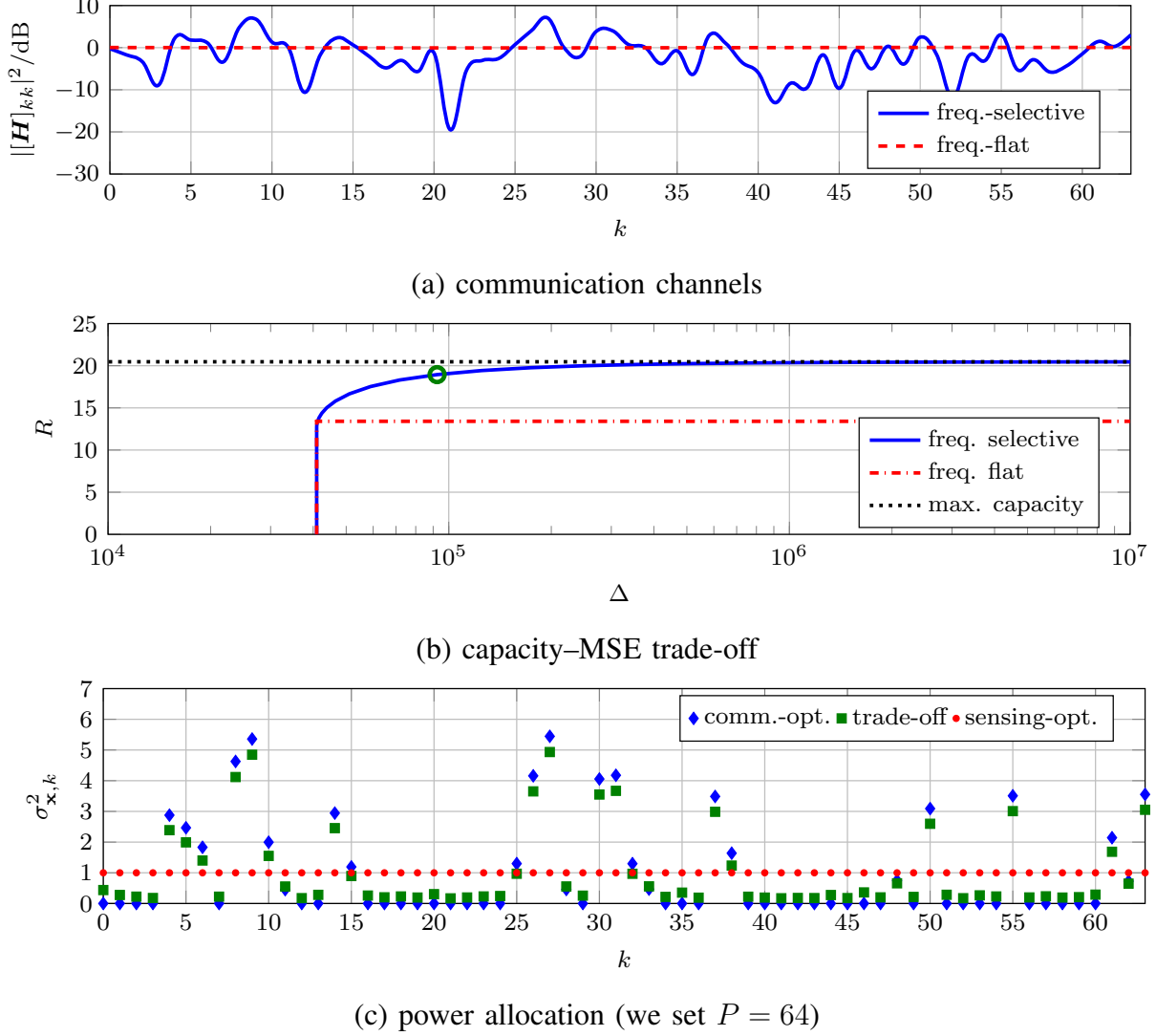


Fig. 4: Capacity-MSE trade-off in JCAS with OFDM.

Frequency-selective vs. frequency-flat communication channels: Here, we numerically evaluate the capacity-MSE function. We set $K = 64$ and study two cases: frequency-selective communication channel and frequency-flat communication channel. The former is meant to illustrate the trade-off between communication and sensing, while the latter will lead to no trade-off. To this end, we set the entries of the time-domain channel impulse response $\mathbf{h}$ as

$$h_k = \exp\left(-\frac{1}{2\alpha}k + i\varphi_k\right), \quad k = 0, 1, \dots, K-1 \quad (81)$$

where φ_k is a uniformly distributed phase term. The channel model has an exponential power delay profile with delay spread α . For the frequency-selective channel, we set $\alpha = 10$. To get an

almost frequency-flat channel, on the other hand, we chose $\alpha = 0.1$. In both cases, we normalize the channel impulse response to have unit energy per sub-carrier, i.e., $\|\mathbf{h}\|^2/K = 1$. We define the communication and sensing SNR as $P/\sigma_\omega^2 = 10$ dB and $P/\sigma_v^2 = -10$ dB, respectively.

The magnitude responses of the two channels are depicted in Fig. 4a. The resulting capacity-MSE functions are illustrated in Fig. 4b. We can see a clear trade-off between the capacity and the MSE in the frequency-selective channel case. In the frequency-flat channel, on the other hand, there is no trade-off between optimal communication and sensing, as one would expect.

In Fig. 4c, we show the allocation of power among sub-carriers for different points on the capacity-MSE curve. Interestingly, in the frequency-selective scenarios we consider, the communication-optimal water-filling solution allocates zero power to sub-carriers with insufficient SNR. As a result, the communication-optimal point on the Pareto boundary is characterized by $\Delta \rightarrow \infty$, as shown in Fig. 4b. For the intermediate trade-off point, indicated by a marker in Fig. 4b, non-zero power is allocated to all sub-carriers to guarantee a finite MSE, while still following the water-filling trend by allocating more power to sub-carriers with higher SNR.

VI. CONCLUSION

We considered the problem of joint communication and sensing in a basic Gaussian MIMO setting, where the sensing task is to estimate a random parameter vector. The random parameter vector, drawn according to a known prior distribution, is assumed to remain fixed throughout the transmission block. We studied the fundamental trade-off between communication and sensing, formulated in terms of the rate of reliable communication (i.e., channel coding rate) against the MSE of parameter estimation. The optimal rate-MSE trade-off in the asymptotic regime of large block-lengths, characterized in terms of a capacity-MSE function, was established and a number of analytical properties were shown. Finally, case studies of exemplary scenarios illustrated this trade-off: one involving beamforming and DoA estimation, the other involving power allocation over OFDM sub-carriers and tapped-delay channel response estimation.

The problem we considered here can be extended in several directions. In the current paper, we focused our attention on non-adaptive (i.e., open-loop) transmission strategies. It is of interest to investigate the potential gains of adaptive (i.e. closed-loop) transmission strategies, where the next channel input is chosen depending on previous sensor observations. Progress along these lines in the case of discrete parameters (i.e., hypothesis testing) has been reported in [28], and it is of interest to extend some of these insights to problems involving continuous parameter

estimation. Another possible direction would be to extend the setting to multiple users, as well as multiple distinct targets, and study trade-off regions involving multiple rates and MSEs. It is also of interest to refine the results and derive characterizations valid in the finite block-length regime, as has been done for channel coding problems [41]–[43]. As a first step in this direction, one can investigate refined asymptotics, such as second-order coding rates and tighter asymptotic MSE characterizations. Our achievability approach, based on Feinstein’s lemma, provides a starting point for enabling such analysis. Finally, relaxing the fixed-parameter assumption—without going as far as the i.i.d. varying regime—is also of interest. Ideally, one would like to investigate models where parameters evolve slowly, as in practical sensing scenarios. Initial attempts to model and study such scenarios in basic settings have been reported in [44]–[46].

APPENDIX A

A. Proof of Lemma 1

Here we show that the function $d(\mathbf{Q}) \triangleq \text{tr}(\mathbb{E}[\mathbf{J}(\boldsymbol{\theta}|\mathbf{Q})^{-1}])$ is monotonically non-increasing and convex in $\mathbf{Q} \in \mathbb{S}_+^M$. To this end, we first note that for any $\boldsymbol{\theta}$, the conditional FIM $\mathbf{J}(\boldsymbol{\theta}|\mathbf{Q})$ is positive semi-definite and linear in $\mathbf{Q}$, which can be verified directly from the definition, i.e., (18). Let $\mathbf{Q}_1$ and $\mathbf{Q}_2$ be two covariance matrices such that $\mathbf{Q}_1 \succeq \mathbf{Q}_2$. We have

$$\mathbf{J}(\boldsymbol{\theta}|\mathbf{Q}_1) = \mathbf{J}(\boldsymbol{\theta}|\mathbf{Q}_1 + \mathbf{Q}_2 - \mathbf{Q}_2) \quad (82)$$

$$= \mathbf{J}(\boldsymbol{\theta}|\mathbf{Q}_2) + \mathbf{J}(\boldsymbol{\theta}|\mathbf{Q}_1 - \mathbf{Q}_2) \quad (83)$$

$$\succeq \mathbf{J}(\boldsymbol{\theta}|\mathbf{Q}_2) \quad (84)$$

where (83) is from linearity, while (84) holds since $\mathbf{Q}_1 - \mathbf{Q}_2 \succeq \mathbf{0}$ and hence $\mathbf{J}(\boldsymbol{\theta}|\mathbf{Q}_1 - \mathbf{Q}_2) \succeq \mathbf{0}$. Therefore, we have shown that $\mathbf{J}(\boldsymbol{\theta}|\mathbf{Q})$ is monotonically non-decreasing in $\mathbf{Q} \in \mathbb{S}_+^M$. Note that (84) directly implies that $\mathbf{J}(\boldsymbol{\theta}|\mathbf{Q}_2)^{-1} \succeq \mathbf{J}(\boldsymbol{\theta}|\mathbf{Q}_1)^{-1}$, which in turn implies

$$d(\mathbf{Q}_2) \geq d(\mathbf{Q}_1). \quad (85)$$

Therefore, we conclude that $d(\mathbf{Q})$ is non-increasing. For convexity, this holds since $\mathbf{J}(\boldsymbol{\theta}|\mathbf{Q})$ is linear in $\mathbf{Q}$, the inverse function $\mathbf{J}(\boldsymbol{\theta}|\mathbf{Q})^{-1}$ is convex, and then the expectation and trace are linear. Therefore, it can be easily verified that the overall function $d(\mathbf{Q})$ is convex.

B. Properties of $C(\Delta)$

We now show that the characterization of $C(\Delta)$ in (21) is continuous, non-decreasing, and concave in Δ . First, it is clear from (21) that the objective is non-decreasing in Δ , since a solution under a smaller MSE constraint Δ is also feasible under a larger constraint. To show concavity, let us consider two MSE constraint values Δ_1 and Δ_2 , and let $\mathbf{Q}_1$ and $\mathbf{Q}_2$ be optimal solutions achieving $C(\Delta_1)$ and $C(\Delta_2)$, respectively, in (21). We have

$$\alpha C(\Delta_1) + (1 - \alpha)C(\Delta_2) \leq \log \det \left(\mathbf{I} + \frac{1}{\sigma_\omega^2} \mathbf{H}(\alpha \mathbf{Q}_1 + (1 - \alpha)\mathbf{Q}_2) \mathbf{H}^H \right) \quad (86)$$

which holds due to the concavity of the Gaussian mutual information function. We observe that

$$d(\alpha \mathbf{Q}_1 + (1 - \alpha)\mathbf{Q}_2) \leq \alpha d(\mathbf{Q}_1) + (1 - \alpha)d(\mathbf{Q}_2) \quad (87)$$

$$\leq \alpha \Delta_1 + (1 - \alpha)\Delta_2 \quad (88)$$

where the first inequality is by convexity of $d(\mathbf{Q})$, while the second inequality follows from the optimality (and hence feasibility) of $\mathbf{Q}_1$ and $\mathbf{Q}_2$ for their corresponding problems. The inequality in (88) implies that $\alpha \mathbf{Q}_1 + (1 - \alpha)\mathbf{Q}_2$ is a feasible solution for (21) with an MSE constraint $\Delta = \alpha \Delta_1 + (1 - \alpha)\Delta_2$. Therefore, (86) is further upper bounded by $C(\alpha \Delta_1 + (1 - \alpha)\Delta_2)$, which proves concavity. Continuity, on interior points of the domain, follows from concavity.

APPENDIX B

PROOF OF LEMMA 4

Starting with the semi-positive definite constraint on the sequence $\tilde{\mathbf{x}}^N \in \mathcal{B}_{N,\delta}(\mathbf{Q})$, i.e.,

$$\mathbf{Q} - \delta \mathbf{I} \preceq \hat{\mathbf{Q}}(\tilde{\mathbf{x}}^N) \preceq \mathbf{Q} + \delta \mathbf{I} \quad (89)$$

the following scalar inequality holds by the definition of positive semi-definiteness

$$\mathbf{a}^H (\mathbf{Q} - \delta \mathbf{I}) \mathbf{a} \leq \mathbf{a}^H \left(\frac{1}{N} \sum_{i=1}^N \tilde{\mathbf{x}}_i \tilde{\mathbf{x}}_i^H \right) \mathbf{a} \leq \mathbf{a}^H (\mathbf{Q} + \delta \mathbf{I}) \mathbf{a}. \quad (90)$$

This is equivalently written as

$$\mathbf{a}^H \mathbf{Q} \mathbf{a} - \delta \|\mathbf{a}\|^2 \leq \frac{1}{N} \sum_{i=1}^N |\mathbf{a}^H \tilde{\mathbf{x}}_i|^2 \leq \mathbf{a}^H \mathbf{Q} \mathbf{a} + \delta \|\mathbf{a}\|^2 \quad (91)$$

for any $\mathbf{a} \in \mathbb{C}^M$ and some $\delta > 0$. Since $\mathbf{Q} \in \mathbb{S}_+^M$ is Hermitian positive semi-definite, there exists some matrix $\mathbf{D} \in \mathbb{C}^{r \times M}$ with $r \triangleq \text{rank}(\mathbf{Q}) \leq M$ such that $\mathbf{Q} = \mathbf{D}^H \mathbf{D}$.

Let $\mathbf{u}_i \sim \mathcal{CN}(\mathbf{0}, \mathbf{I})$ and let $\tilde{\mathbf{x}}_i \triangleq \mathbf{D}^H \mathbf{u}_i$. For $r < M$, i.e., $\mathbf{Q}$ is rank-deficient, the nullspace of $\mathbf{D}$ contains non-zero vectors. When $\mathbf{a}$ is in the nullspace of $\mathbf{D}$, i.e., $\mathbf{D}\mathbf{a} = \mathbf{0}$, (91) becomes

$$-\delta \|\mathbf{a}\|^2 \leq 0 \leq \delta \|\mathbf{a}\|^2 \quad (92)$$

which trivially holds true. When $\mathbf{Q}$ is full-rank, $\mathbf{D}$ is a full-rank square matrix, and the null space consists of only the zero vector. When a non-zero vector $\mathbf{a}$ is not in the nullspace of $\mathbf{D}$, i.e., $\mathbf{D}\mathbf{a} \neq \mathbf{0}$, we can divide (91) by $\mathbf{a}^H \mathbf{Q} \mathbf{a}$, which yields

$$1 - \frac{\delta \|\mathbf{a}\|^2}{\mathbf{a}^H \mathbf{Q} \mathbf{a}} \leq \frac{1}{N} \sum_{i=1}^N \left| \frac{\mathbf{a}^H \tilde{\mathbf{x}}_i}{\sqrt{\mathbf{a}^H \mathbf{Q} \mathbf{a}}} \right|^2 \leq 1 + \frac{\delta \|\mathbf{a}\|^2}{\mathbf{a}^H \mathbf{Q} \mathbf{a}}. \quad (93)$$

Since (91) must be satisfied for any $\mathbf{a} \in \mathbb{C}^M$, a sufficient condition is obtained using the bound

$$\mathbf{a}^H \mathbf{Q} \mathbf{a} \leq \lambda_{\max}(\mathbf{Q}) \|\mathbf{a}\|^2 \quad (94)$$

where $\lambda_{\max}(\mathbf{Q})$ is the largest eigenvalue of $\mathbf{Q}$. Then, we get

$$1 - \frac{\delta}{\lambda_{\max}(\mathbf{Q})} \leq \frac{1}{N} \sum_{i=1}^N \left| \frac{\mathbf{a}^H \tilde{\mathbf{x}}_i}{\sqrt{\mathbf{a}^H \mathbf{Q} \mathbf{a}}} \right|^2 \leq 1 + \frac{\delta}{\lambda_{\max}(\mathbf{Q})}. \quad (95)$$

Defining $\alpha_i \triangleq (\mathbf{a}^H \mathbf{Q} \mathbf{a})^{-\frac{1}{2}} \mathbf{a}^H \tilde{\mathbf{x}}_i$, we observe that $\alpha_i \sim \mathcal{CN}(0, 1)$. Setting $\delta' \triangleq \delta / \lambda_{\max}(\mathbf{Q}) > 0$, we can rewrite (95) as

$$\left| \frac{1}{N} \sum_{i=1}^N |\alpha_i|^2 - 1 \right| \leq \delta' \quad (96)$$

where $|\alpha_i|^2 \sim \text{Exp}(1)$ has an exponential distribution. Let $S_N \triangleq \frac{1}{N} \sum_{i=1}^N |\alpha_i|^2$ denote the empirical mean of N i.i.d. exponential random variables. The mean and the variance of S_N are given by $\mathbb{E}[S_N] = 1$ and $\mathbb{E}[(S_N - 1)^2] = 1/N$, respectively. Finally, for any $\delta > 0$

$$\mathbb{P}[\tilde{\mathbf{x}}^N \notin \mathcal{B}_{N,\delta}(\mathbf{Q})] \leq \mathbb{P}[|S_N - 1| \geq \delta'] \quad (97)$$

$$\leq \frac{1}{\delta'^2 N} \quad (98)$$

using the sufficient condition in (95) as an upper bound in (97) and Chebyshev's inequality in (98). As $N \rightarrow \infty$, we get $\mathbb{P}[\tilde{\mathbf{x}}^N \in \mathcal{B}_{N,\delta}(\mathbf{Q})] \rightarrow 1$, which completes the proof.

APPENDIX C

BCRB-BASED CAPACITY-MSE BOUND

In the appendix, we show how to obtain the capacity-MSE upper bound in (23) using the BCRB lower bound [25]. This derivation in this appendix builds on the converse proof in Section IV-C.

To define the BCRB, we first introduce the Bayesian information matrix associated with estimating $\boldsymbol{\theta}$ from $\mathbf{z}^N$, given that $\mathbf{x}^N = \mathbf{x}^N$ is transmitted. This is written as

$$\mathbf{J}_B(\mathbf{x}^N) \triangleq \mathbb{E}[\mathbf{J}(\boldsymbol{\theta}|\mathbf{x}^N)] + \mathbf{J}_P \quad (99)$$

where $\mathbf{J}_P \triangleq \mathbb{E} \left[[\nabla_{\boldsymbol{\theta}} \ln f_{\boldsymbol{\theta}}(\boldsymbol{\theta})] [\nabla_{\boldsymbol{\theta}} \ln f_{\boldsymbol{\theta}}(\boldsymbol{\theta})]^\top \right]$ is the prior information matrix. The BCRB inequality of van Trees [25, 4.3] states the the MSE is lower bounded as

$$N \operatorname{tr} \left(\boldsymbol{\Sigma} \left(\mathbf{x}^N, \hat{\boldsymbol{\theta}}_N^* \right) \right) \geq \operatorname{tr} \left(\left[\frac{1}{N} \mathbf{J}_B(\mathbf{x}^N) \right]^{-1} \right) \quad (100)$$

$$= \operatorname{tr} \left(\left[\mathbb{E} \left[\mathbf{J}(\boldsymbol{\theta}|\hat{\mathbf{Q}}(\mathbf{x}^N)) \right] + \frac{1}{N} \mathbf{J}_P \right]^{-1} \right). \quad (101)$$

Unlike the lower bound in (32), this bound holds not only asymptotically but for every N . Adopting this in the converse proof in Section IV-C, and continuing from (46), we have

$$\Delta + \epsilon \geq \frac{1}{M_N} \sum_{w \in \mathcal{M}_N} \operatorname{tr} \left(\left[\mathbb{E} \left[\mathbf{J}(\boldsymbol{\theta}|\hat{\mathbf{Q}}(\mathbf{x}^N(w))) \right] + \frac{1}{N} \mathbf{J}_P \right]^{-1} \right) \quad (102)$$

$$\geq \operatorname{tr} \left(\left[\mathbb{E} \left[\mathbf{J}(\boldsymbol{\theta}|\bar{\mathbf{Q}}) \right] + \frac{1}{N} \mathbf{J}_P \right]^{-1} \right) \quad (103)$$

$$\geq \operatorname{tr} \left(\mathbb{E} \left[\mathbf{J}(\boldsymbol{\theta}|\bar{\mathbf{Q}}) \right]^{-1} \right) - \epsilon \quad (104)$$

where (103) holds by convexity and Jensen's inequality, and (104) holds for sufficiently large N by continuity and the fact that $\mathbf{J}_P$ is constant, and hence $\frac{1}{N} \mathbf{J}_P$ vanishes.

It is also possible to adopt the Miller–Chang version of the BCRB [24], as done in [13]. In this formulation, the random codeword $\mathbf{x}^N$ is treated as a nuisance parameter, and the average BCRB with respect to $\mathbf{x}^N$ is considered. Adapting this to our setting amounts to replacing the maximal MSE risk (over codewords) ξ_N with its average counterpart ξ_N^{av} (see Section II-B3). Nevertheless, since $\mathbf{x}^N$ is uniformly distributed on the codebook as the message is uniform, the previous result still applies. In particular, (102) onward remain valid in this case.

REFERENCES

- [1] F. Liu, C. Masouros, A. P. Petropulu, H. Griffiths, and L. Hanzo, “Joint radar and communication design: Applications, state-of-the-art, and the road ahead,” *IEEE Trans. Commun.*, vol. 68, no. 6, pp. 3834–3862, 2020.
- [2] F. Liu, C. Masouros, A. Li, H. Sun, and L. Hanzo, “MU-MIMO communications with MIMO radar: From co-existence to joint transmission,” *IEEE Trans. Wireless Commun.*, vol. 17, no. 4, pp. 2755–2770, 2018.
- [3] F. Liu, L. Zhou, C. Masouros, A. Li, W. Luo, and A. Petropulu, “Toward dual-functional radar-communication systems: Optimal waveform design,” *IEEE Trans. Signal Process.*, vol. 66, no. 16, pp. 4264–4279, 2018.

- [4] X. Liu, T. Huang, N. Shlezinger, Y. Liu, J. Zhou, and Y. C. Eldar, "Joint transmit beamforming for multiuser MIMO communications and MIMO radar," *IEEE Trans. Signal Process.*, vol. 68, pp. 3929–3944, 2020.
- [5] L. Chen, F. Liu, W. Wang, and C. Masouros, "Joint radar-communication transmission: A generalized Pareto optimization framework," *IEEE Trans. Signal Process.*, vol. 69, pp. 2752–2765, 2021.
- [6] T. Wild, V. Braun, and H. Viswanathan, "Joint design of communication and sensing for beyond 5G and 6G systems," *IEEE Access*, vol. 9, pp. 30 845–30 857, 2021.
- [7] C. Sturm and W. Wiesbeck, "Waveform design and signal processing aspects for fusion of wireless communications and radar sensing," *Proc. IEEE*, vol. 99, no. 7, pp. 1236–1259, 2011.
- [8] D. Ma, N. Shlezinger, T. Huang, Y. Liu, and Y. C. Eldar, "Joint radar-communication strategies for autonomous vehicles: Combining two key automotive technologies," *IEEE Signal Process. Mag.*, vol. 37, no. 4, pp. 85–97, 2020.
- [9] F. Liu, Y.-F. Liu, A. Li, C. Masouros, and Y. C. Eldar, "Cramér-Rao bound optimization for joint radar-communication beamforming," *IEEE Trans. Signal Process.*, vol. 70, pp. 240–253, 2022.
- [10] H. Hua, T. X. Han, and J. Xu, "MIMO integrated sensing and communication: CRB-rate tradeoff," *IEEE Trans. Wireless Commun.*, vol. 23, no. 4, pp. 2839–2854, 2024.
- [11] Z. Ren, Y. Peng, X. Song, Y. Fang, L. Qiu, L. Liu, D. W. K. Ng, and J. Xu, "Fundamental CRB-rate tradeoff in multi-antenna ISAC systems with information multicasting and multi-target sensing," *IEEE Trans. Wireless Commun.*, vol. 23, no. 4, pp. 3870–3885, 2024.
- [12] C. Xu and S. Zhang, "MIMO integrated sensing and communication exploiting prior information," *IEEE J. Sel. Areas Commun.*, vol. 42, no. 9, pp. 2306–2321, 2024.
- [13] Y. Xiong, F. Liu, Y. Cui, W. Yuan, T. X. Han, and G. Caire, "On the fundamental tradeoff of integrated sensing and communications under Gaussian channels," *IEEE Trans. Inf. Theory*, vol. 69, no. 9, pp. 5723–5751, 2023.
- [14] H. Joudeh, "Joint communication and target detection with multiple antennas," in *26th Int. ITG Workshop Smart Antennas and 13th Conf. on Syst., Commun. and Coding (WSA & SCC)*, 2023, pp. 1–6.
- [15] E. Telatar, "Capacity of multi-antenna Gaussian channels," *Eur. Trans. Telecomm.*, vol. 10, no. 6, pp. 585–595, 1999.
- [16] G. J. Foschini and M. J. Gans, "On limits of wireless communications in a fading environment when using multiple antennas," *Wireless personal commun.*, vol. 6, no. 3, pp. 311–335, 1998.
- [17] A. Goldsmith, S. Jafar, N. Jindal, and S. Vishwanath, "Capacity limits of MIMO channels," *IEEE J. Sel. Areas Commun.*, vol. 21, no. 5, pp. 684–702, 2003.
- [18] M. Kobayashi, G. Caire, and G. Kramer, "Joint state sensing and communication: Optimal tradeoff for a memoryless case," in *Proc. IEEE Int. Symp. Inf. Theory (ISIT)*, 2018, pp. 111–115.
- [19] M. Kobayashi, H. Hamad, G. Kramer, and G. Caire, "Joint state sensing and communication over memoryless multiple access channels," in *Proc. IEEE Int. Symp. Inf. Theory (ISIT)*, 2019, pp. 270–274.
- [20] M. Ahmadipour and M. Wigger, "An information-theoretic approach to collaborative integrated sensing and communication for two-transmitter systems," *IEEE J. Sel. Areas Inf. Theory*, vol. 4, pp. 112–127, 2023.
- [21] M. Ahmadipour, M. Kobayashi, M. Wigger, and G. Caire, "An information-theoretic approach to joint sensing and communication," *IEEE Trans. Inf. Theory*, vol. 70, no. 2, pp. 1124–1146, 2024.
- [22] O. Günlü, M. R. Bloch, R. F. Schaefer, and A. Yener, "Secure integrated sensing and communication," *IEEE J. Sel. Areas Inf. Theory*, vol. 4, pp. 40–53, 2023.
- [23] H. Joudeh and G. Caire, "Joint communication and state sensing under logarithmic loss," in *Proc. 4th IEEE Int. Symp. Joint Commun. Sensing (JC&S)*, 2024, pp. 1–6.
- [24] R. Miller and C. Chang, "A modified Cramér-Rao bound and its applications (corresp.)," *IEEE Trans. Inf. Theory*, vol. 24, no. 3, pp. 398–400, 1978.

- [25] H. L. Van Trees, K. L. Bell, and Z. Tian, *Detection, Estimation, and Modulation Theory Part I*, 2nd ed. Hoboken, NJ, USA: John Wiley & Sons, 2013.
- [26] H. Joudeh and F. M. J. Willems, "Joint communication and binary state detection," *IEEE J. Sel. Areas Inf. Theory*, vol. 3, no. 1, pp. 113–124, 2022.
- [27] H. Wu and H. Joudeh, "Joint communication and channel discrimination," *Entropy*, vol. 26, no. 12, 2024.
- [28] M.-C. Chang, S.-Y. Wang, T. Erdoğan, and M. R. Bloch, "Rate and detection-error exponent tradeoff for joint communication and sensing of fixed channel states," *IEEE J. Sel. Areas Inf. Theory*, vol. 4, pp. 245–259, 2023.
- [29] D. Seo and S. H. Lim, "Integrated communication and binary state detection under unequal error constraints," *IEEE Trans. Commun.*, vol. 73, no. 8, pp. 5729–5744, 2025.
- [30] D. Seo, "Integrated communication and Bayesian estimation of fixed channel states," *IEEE Commun. Lett.*, vol. 29, no. 3, 2025.
- [31] S. M. Kay, "Fundamentals of statistical signal processing: Estimation theory," 1993.
- [32] T. M. Cover and J. A. Thomas, *Elements of Information Theory*, 2nd ed. Wiley, 2006.
- [33] Y. Wu and S. Verdú, "MMSE dimension," *IEEE Trans. Inf. Theory*, vol. 57, no. 8, pp. 4857–4879, 2011.
- [34] O. Aharon and J. Tabrikian, "Asymptotically tight Bayesian Cramér-Rao bound," *IEEE Trans. Signal Process.*, vol. 72, pp. 3333–3346, 2024.
- [35] H. Wu and H. Joudeh, "On joint communication and channel discrimination," in *Proc. IEEE Int. Symp. Inf. Theory (ISIT)*, 2022, pp. 3321–3326.
- [36] Y. Polyanskiy and Y. Wu, *Information Theory: From Coding to Learning*. Cambridge University Press, 2025.
- [37] A. N. D'Andrea, U. Mengali, and R. Reggiannini, "The modified Cramer-Rao bound and its application to synchronization problems," *IEEE Trans. Commun.*, vol. 42, no. 234, pp. 1391–1399, 2002.
- [38] A. El Gamal and Y.-H. Kim, *Network information theory*. Cambridge University Press, 2011.
- [39] J. Li, L. Xu, P. Stoica, K. W. Forsythe, and D. W. Bliss, "Range compression and waveform optimization for MIMO radar: A Cramér–Rao bound based study," *IEEE Trans. Signal Process.*, vol. 56, no. 1, pp. 218–232, 2008.
- [40] H. L. Van Trees, *Detection, Estimation, and Modulation Theory Part IV: Optimum Array Processing*. New York, USA: John Wiley & Sons, 2002.
- [41] Y. Polyanskiy, H. V. Poor, and S. Verdú, "Channel coding rate in the finite blocklength regime," *IEEE Trans. Inf. Theory*, vol. 56, no. 5, pp. 2307–2359, 2010.
- [42] M. Hayashi, "Information spectrum approach to second-order coding rate in channel coding," *IEEE Trans. Inf. Theory*, vol. 55, no. 11, pp. 4947–4966, Nov 2009.
- [43] W. Yang, G. Durisi, T. Koch, and Y. Polyanskiy, "Quasi-static multiple-antenna fading channels at finite blocklength," *IEEE Trans. Inf. Theory*, vol. 60, no. 7, pp. 4232–4265, 2014.
- [44] S. Li and G. Caire, "On the capacity and state estimation error of "beam-pointing" channels: The binary case," *IEEE Trans. Inf. Theory*, vol. 69, no. 9, pp. 5752–5770, 2023.
- [45] H. Nikbakht, M. Wigger, S. S. Shitz, and H. V. Poor, "A memory-based reinforcement learning approach to integrated sensing and communication," in *Proc. 58th Asilomar Conf. Sig. Syst. Comp.*, 2024, pp. 433–437.
- [46] C. P. Lindstrom and M. R. Bloch, "Rate distortion approach to joint communication and sensing with Markov states: Open loop case," *arXiv:2501.15652*, 2025.