

MMSE Dimension

Yihong Wu

Department of Electrical Engineering
Princeton University
Princeton, NJ 08544, USA
Email: yihongwu@princeton.edu

Sergio Verdú

Department of Electrical Engineering
Princeton University
Princeton, NJ 08544, USA
Email: verd@princeton.edu

Abstract—If N is standard Gaussian, the minimum mean-square error (MMSE) of estimating X based on $\sqrt{\text{snr}}X + N$ vanishes at least as fast as $\frac{1}{\text{snr}}$ as $\text{snr} \rightarrow \infty$. We define the *MMSE dimension* of X as the limit as $\text{snr} \rightarrow \infty$ of the product of snr and the MMSE. For discrete, absolutely continuous or mixed X we show that the MMSE dimension equals Rényi's information dimension. However, for singular X , we show that the product of snr and MMSE oscillates around information dimension periodically in snr (dB). We also show that discrete side information does not reduce MMSE dimension. These results extend considerably beyond Gaussian N under various technical conditions.

I. INTRODUCTION

A. Basic Setup

The minimum mean square error (MMSE) plays a pivotal role in estimation theory and Bayesian statistics. Due to the lack of closed-form expressions for posterior distributions and conditional expectations, exact MMSE formulae are scarce. Asymptotic analysis is more tractable and sheds important insights about how the fundamental estimation-theoretic limits depend on the input and noise statistics. The theme of this paper is the high-SNR scaling law of the MMSE of estimating X based on $\sqrt{\text{snr}}X + N$ when N is independent of X .

The MMSE of estimating X based on Y is given by

$$\text{mmse}(X|Y) = \mathbb{E}[(X - \mathbb{E}[X|Y])^2]. \quad (1)$$

When Y is related to X through an additive-noise channel with gain $\sqrt{\text{snr}}$, i.e.,

$$Y = \sqrt{\text{snr}}X + N \quad (2)$$

where $N \perp X$, we denote

$$\text{mmse}(X, N, \text{snr}) = \text{mmse}(X|\sqrt{\text{snr}}X + N), \quad (3)$$

and, in particular, when the noise is Gaussian (denoted by N_G), we simplify

$$\text{mmse}(X, \text{snr}) = \text{mmse}(X, N_G, \text{snr}). \quad (4)$$

B. Asymptotics of MMSE

The low-SNR asymptotics of $\text{mmse}(X, \text{snr})$ has been studied extensively in [1, Section II.F], where the Taylor expansion of $\text{mmse}(X, \text{snr})$ at $\text{snr} = 0$ has been obtained and the coefficients turn out to depend only on the moments of X . However, the asymptotics in the high-SNR regime have not

received much attention in the literature. The high-SNR behavior depends on the input distribution: For example, for binary X , $\text{mmse}(X, \text{snr})$ decays exponentially, while for standard Gaussian X ,

$$\text{mmse}(X, \text{snr}) = \frac{1}{\text{snr} + 1}. \quad (5)$$

Unlike the low-SNR regime, the high-SNR asymptotics depend on the measure-theoretical structure of the input distribution rather than its moments.

Before defining MMSE dimension, note that

$$0 \leq \text{mmse}(X, N, \text{snr}) \leq \frac{\text{var}N}{\text{snr}}, \quad (6)$$

where the upper bound can be achieved using the affine estimator $f(y) = \frac{y - \mathbb{E}N}{\sqrt{\text{snr}}}$. Assuming N has finite variance, as $\text{snr} \rightarrow \infty$, we have¹:

$$\text{mmse}(X, N, \text{snr}) = O\left(\frac{1}{\text{snr}}\right) \quad (7)$$

Seeking a finer characterization, we are interested in the exact scaling constant in (7). To this end, we define the *upper* and *lower MMSE dimension* of the pair (X, N) as:

$$\overline{\mathcal{D}}(X, N) = \limsup_{\text{snr} \rightarrow \infty} \frac{\text{snr} \cdot \text{mmse}(X, N, \text{snr})}{\text{var}N}, \quad (8)$$

$$\underline{\mathcal{D}}(X, N) = \liminf_{\text{snr} \rightarrow \infty} \frac{\text{snr} \cdot \text{mmse}(X, N, \text{snr})}{\text{var}N}. \quad (9)$$

If $\overline{\mathcal{D}}(X, N) = \underline{\mathcal{D}}(X, N)$, the common value is denoted by $\mathcal{D}(X, N)$, called the *MMSE dimension* of (X, N) . This information measure governs the high-SNR scaling law of MMSE and sharpens (7) to:

$$\text{mmse}(X, N, \text{snr}) = \frac{\mathcal{D}(X, N)\text{var}N}{\text{snr}} + o\left(\frac{1}{\text{snr}}\right). \quad (10)$$

When the side information U (independent of N) is present, replacing $\text{mmse}(X, N, \text{snr})$ by $\text{mmse}(X, N, \text{snr}|U)$, the *conditional MMSE dimension* of (X, N) given U is similarly defined, denoted by $\mathcal{D}(X, N|U)$. When N is Gaussian, we use the notation $\mathcal{D}(X)$ ($\mathcal{D}(X|U)$), called the *(conditional) MMSE dimension* of X . By (6), we have

$$0 \leq \mathcal{D}(X, N|U) \leq \overline{\mathcal{D}}(X, N|U) \leq 1. \quad (11)$$

¹We use the following asymptotic notation: $f(x) = O(g(x))$ if and only if $\limsup \frac{|f(x)|}{|g(x)|} < \infty$. $f(x) = \Omega(g(x))$ if and only if $g(x) = O(f(x))$. $f(x) = \Theta(g(x))$ if and only if $f(x) = O(g(x))$ and $f(x) = \Omega(g(x))$.

Although most of our attention is focused on square integrable noise, $\text{mmse}(X, N, \text{snr})$ is finite even for infinite-variance N , as long as $\text{var}X < \infty$. But (8) and (9) no longer make sense in that case. Hence the scaling law in (7) could fail. Examples can be constructed where $\text{mmse}(X, N, \text{snr})$ decays strictly slower than $\frac{1}{\text{snr}}$, e.g., as $\frac{\log \text{snr}}{\text{snr}}$ or even $\frac{1}{\sqrt{\text{snr}}}$. However, for certain X , $\text{mmse}(X, N, \text{snr})$ can still decay according to $\frac{1}{\text{snr}}$ for some infinite-variance N [2].

C. Asymptotics of Fisher information

In the special case of Gaussian noise it is interesting to draw conclusions on the asymptotic behavior of Fisher's information based on our results. Similarly to the MMSE dimension, we can define the *Fisher dimension* of a random variable X as follows:

$$\overline{\mathcal{J}}(X) = \limsup_{\epsilon \downarrow 0} \epsilon^2 \cdot J(X + \epsilon N_G), \quad (12)$$

$$\underline{\mathcal{J}}(X) = \liminf_{\epsilon \downarrow 0} \epsilon^2 \cdot J(X + \epsilon N_G). \quad (13)$$

In view of the representation of MMSE via the Fisher information of the channel output [3, (1.3.4)], [1, (58)]:

$$\text{snr} \cdot \text{mmse}(X, \text{snr}) = 1 - J(\sqrt{\text{snr}}X + N_G) \quad (14)$$

and $J(aZ) = a^{-2}J(Z)$, we conclude that Fisher dimension and MMSE dimension are complementary of each other:

$$\overline{\mathcal{J}}(X) + \underline{\mathcal{D}}(X) = \underline{\mathcal{J}}(X) + \overline{\mathcal{D}}(X) = 1. \quad (15)$$

D. Connections to asymptotic statistics

The high-SNR asymptotic behavior of $\text{mmse}(X, \text{snr})$ is equivalent to the behavior of the Bayesian risk for the Gaussian location model in the large sample limit, where P_X is the prior distribution and the sample size n plays the role of snr . Let $\{N_i : i \in \mathbb{N}\}$ be a sequence of i.i.d. standard Gaussian random variables independent of X and denote $Y_i = X + N_i$. By the sufficiency of the sample mean $\bar{Y} = \frac{1}{n} \sum_{i=1}^n Y_i$ in Gaussian location models, we have

$$\text{mmse}(X|Y^n) = \text{mmse}(X|\bar{Y}) = \text{mmse}(X, n). \quad (16)$$

Therefore as sample size grows, the Bayesian risk of estimating X vanishes as $O(\frac{1}{n})$ with the scaling constant given by the MMSE dimension of the prior:

$$\mathcal{D}(X) = \lim_{n \rightarrow \infty} \text{mmse}(X|Y^n) n. \quad (17)$$

The asymptotics of $\text{mmse}(X|Y^n)$ has been studied in [4], [5] for absolutely continuous priors and general models where X and Y are not necessarily related by additive Gaussian noise. Those results are compared with ours in Section II.

II. MAIN RESULTS

A. Connections to information dimension

The MMSE dimension is intimately related to the *information dimension* defined by Rényi in [6]:

Definition 1. Let X be a real-valued random variable. Denote for a positive integer m the quantized version of X :

$$\langle X \rangle_m = \frac{\lfloor mX \rfloor}{m}. \quad (18)$$

Define²

$$\underline{d}(X) = \liminf_{m \rightarrow \infty} \frac{H(\langle X \rangle_m)}{\log m} \quad (19)$$

and

$$\overline{d}(X) = \limsup_{m \rightarrow \infty} \frac{H(\langle X \rangle_m)}{\log m}, \quad (20)$$

where $\underline{d}(X)$ and $\overline{d}(X)$ are called *lower* and *upper information dimensions* of X respectively. If $\underline{d}(X) = \overline{d}(X)$, the common value is called the *information dimension* of X , denoted by $d(X)$.

The following theorem reveals a connection between the MMSE dimension and the information dimension of the input:

Theorem 1. If $\mathbb{E}[X^2] < \infty$, then

$$\underline{\mathcal{D}}(X) \leq \underline{d}(X) \leq \overline{d}(X) \leq \overline{\mathcal{D}}(X), \quad (21)$$

Therefore, if $\mathcal{D}(X)$ exists, then $d(X)$ exists, and

$$\mathcal{D}(X) = d(X). \quad (22)$$

Due to space limitations, proofs are referred to [2].

In [7, p.755] it is conjectured that $\underline{\mathcal{J}}(X) = 1 - \overline{d}(X)$, or equivalently, in view of (15), $\overline{\mathcal{D}}(X) = \overline{d}(X)$. This holds for distributions without singular components but not in general. The Cantor distribution gives a counterexample to the general conjecture (see Section II-D). However, whenever X has a discrete-continuous mixed distribution, (22) holds and information dimension governs the high-SNR asymptotics of MMSE.

Next we drop the assumption of $\text{var}X < \infty$ and proceed to give results for various input and noise distributions.

B. Absolutely continuous inputs

Theorem 2. If X has an absolutely continuous distribution with respect to Lebesgue measure, then $\mathcal{D}(X) = 1$, i.e.,

$$\text{mmse}(X, \text{snr}) = \frac{1}{\text{snr}} + o\left(\frac{1}{\text{snr}}\right). \quad (23)$$

In view of (40), Theorem 2 implies that absolutely continuous priors result in procedures whose asymptotic Bayes risk coincides with the minimax level.

Under certain regularity conditions on N (X resp.) we can extend Theorem 2 to show

$$\mathcal{D}(X, N) = 1 \quad (24)$$

²Throughout the paper, natural logarithms are adopted and information units are nats.

for all absolutely continuous X (N resp.). For example, when X has a continuous and bounded density, (24) holds for all square-integrable N even without a density.

A refinement of Theorem 2 entails the second-order expansion for $\text{mmse}(X, \text{snr})$ for absolutely continuous input X . This involves the Fisher information of X . Suppose $J(X) < \infty$, then $J(\sqrt{\text{snr}}X + N_G) \leq J(X)\text{snr}^{-1}$, by the convexity and translation invariance of Fisher information. In view of (14), we have

$$\text{mmse}(X, \text{snr}) = \frac{1}{\text{snr}} + O\left(\frac{1}{\text{snr}^2}\right), \quad (25)$$

which is a slight improvement of (23).

Under stronger conditions the second-order term can be determined exactly. A result in [8] states that: if $J(X) < \infty$ and the density of X satisfies certain regularity conditions [8, (3) – (7)], then

$$J(X + \epsilon N_G) = J(X) + O(\epsilon). \quad (26)$$

Therefore in view of (14), we have

$$\text{mmse}(X, \text{snr}) = \frac{1}{\text{snr}} - \frac{J(X)}{\text{snr}^2} + o\left(\frac{1}{\text{snr}^2}\right). \quad (27)$$

This result can be understood as follows: using (14) and Stam's inequality [9], we have [10, pp. 72–73]

$$\text{mmse}(X, \text{snr}) \geq \frac{1}{J(X) + \text{snr}}, \quad (28)$$

which is also known as the Bayesian Cramér-Rao bound (or the Van Trees inequality). In view of (27), we see that (28) is asymptotically tight for sufficiently regular densities of X .

Instead of using the asymptotic expansion of Fisher information and Stam's inequality, we can show that (27) holds for a much broader class of densities of X and non-Gaussian noise: if X has finite third moment and its density has bounded first two derivatives, then for any finite-variance N ,

$$\text{mmse}(X, N, \text{snr}) = \frac{\text{var}N}{\text{snr}} - \frac{J(X)\text{var}^2N}{\text{snr}^2} + o\left(\frac{1}{\text{snr}^2}\right). \quad (29)$$

This equation reveals a new *operational role* of $J(X)$. The regularity conditions imposed on the input density are much weaker and easier to check than those in [8]; however, (27) is slightly stronger than (29) because the $o(\text{snr}^{-2})$ term in (27) is in fact $O(\text{snr}^{-5/2})$, as a result of (26).

It is interesting to compare (29) to the asymptotic expansion of Bayesian risk in the large sample limit [5]. Under the regularity conditions in [5, Theorem 5.1a], we have

$$\text{mmse}(X|Y_1^n) = \frac{1}{nJ(N)} - \frac{J(X)}{n^2J(N)^2} + o\left(\frac{1}{n^2}\right). \quad (30)$$

On the other hand, by (29) we have

$$\text{mmse}(X|\bar{Y}) = \text{mmse}(X, n) = \frac{\text{var}N}{n} - \frac{J(X)(\text{var}N)^2}{n^2} + o\left(\frac{1}{n^2}\right) \quad (31)$$

When N is Gaussian, (31) agrees with (30), but our proof requires much weaker regularity conditions. When N is not Gaussian, (31) is inferior to (30) on the first term, due to

the Cramér-Rao bound $\text{var}N \geq \frac{1}{J(N)}$. This agrees with the fact that sample mean is asymptotically suboptimal for non-Gaussian noise, and the suboptimality is characterized by the gap in the Cramér-Rao inequality.

To conclude the discussion of absolutely continuous inputs, we give an example where (29) fails:

Example 1. Consider X and N uniformly distributed on $[0, 1]$. It can be shown that $\text{mmse}(X, N, \text{snr}) = \text{var}N \left(\frac{1}{\text{snr}} - \frac{1}{2\text{snr}^{\frac{3}{2}}} \right)$ for $\text{snr} \geq 4$. Note that (29) does not hold because X does not have a differentiable density, hence $J(X) = \infty$ (in the sense of the generalized Fisher information in [11, Definition 4.1]).

C. Conditional MMSE dimension

Next we present results for general mixed distributions, which can be equivalently stated via the conditional MMSE dimension.

Theorem 3 (Conditional MMSE dimension).

$$\overline{\mathcal{D}}(X, N) \geq \overline{\mathcal{D}}(X, N|U), \quad (32)$$

$$\underline{\mathcal{D}}(X, N) \geq \underline{\mathcal{D}}(X, N|U). \quad (33)$$

When N is Gaussian and U is discrete, (32) and (33) hold with equality.

Theorem 3 implies that any *discrete* side information does not change the high-SNR asymptotics of MMSE. Consequently, knowing arbitrarily finitely many digits of X does not reduce its MMSE dimension. Another application of Theorem 3 is to determine the MMSE dimension of inputs with mixed distributions, which are frequently used in statistical models of sparse signals [12], [13], [14]. According to Theorem 3, knowing the support does not decrease the MMSE dimension.

Corollary 1. Let $X = UZ$ where U is independent of Z , taking values in $\{0, 1\}$ with $\mathbb{P}\{U = 1\} = \rho$. Then $\overline{\mathcal{D}}(X) = \rho\overline{\mathcal{D}}(Z)$, $\underline{\mathcal{D}}(X) = \rho\underline{\mathcal{D}}(Z)$.

Having proved results about countable mixtures, the following theorem about MMSE dimension of discrete random variables is just a simple corollary.

Theorem 4. For X discrete (even with countably infinite alphabet), $\mathcal{D}(X) = 0$.

Without recourse to conditional MMSE dimension, Theorem 4 can be extended to $\mathcal{D}(X, N) = 0$ for all discrete X and absolutely continuous N .

Obtained by combining Theorems 2 – 4, the next result gives a complete characterization to the MMSE dimension of discrete-continuous mixtures. Together with Theorem 1, Theorem 5 also provides an MMSE-based proof of Rényi's theorem on information dimension [6] for square integrable inputs.

Theorem 5. If X is distributed according to a mixture of an absolutely continuous and a discrete distribution, then its MMSE dimension equals the weight of the absolutely continuous part.

In [2] we generalize Theorem 5 to non-Gaussian noise.

For square integrable input, the MSE attained by the optimal linear estimator is $\frac{\text{var} N}{\text{snr}} + o\left(\frac{1}{\text{snr}}\right)$. Therefore the linear estimator is dimensionally optimal for estimating absolutely continuous random variables contaminated by additive Gaussian noise, in the sense that it achieves the input MMSE dimension.

To conclude this subsection, we illustrate Theorem 5 by the following examples:

Example 2 (Continuous input). If $X \sim \mathcal{N}(0, 1)$, then $\text{mmse}(X, \text{snr})$ is given in (5) and Theorem 2 holds.

Example 3 (Discrete input). If X is equiprobable on $\{-1, 1\}$, then

$$\text{mmse}(X, \text{snr}) = O\left(\frac{1}{\sqrt{\text{snr}}} e^{-\frac{\text{snr}}{2}}\right) \quad (34)$$

and therefore $\mathcal{D}(X) = 0$ as predicted by Theorem 4. In fact, exponential decay holds for all inputs whose alphabet has no accumulation points. Otherwise, the MMSE can decay polynomially [2].

Example 4 (Mixed input). Let N be uniformly distributed in $[0, 1]$, and let X be distributed according to an equal mixture of a mass at 0 and a uniform distribution on $[0, 1]$. Then $\text{mmse}(X, N, \text{snr}) = \text{var} N \left(\frac{1}{2\text{snr}} + \frac{1}{4\text{snr}^{\frac{3}{2}}} \right) + o\left(\frac{1}{\text{snr}^{\frac{3}{2}}}\right)$, which implies $\mathcal{D}(X, N) = \frac{1}{2}$ as predicted by Theorem 5 for non-Gaussian noise.

D. Singularly continuous inputs

We focus on a special class of inputs with self-similar singular distributions [15, p. 36]: inputs with i.i.d. digits.

Theorem 6. Let $X \in [0, 1]$ a.s., whose M -ary expansion $X = \sum_{j \in \mathbb{N}} (X)_j M^{-j}$ consists of i.i.d. digits $\{(X)_j\}$ with common distribution P . Then for any finite-variance N , there exists a $2 \log M$ -periodic function $\Phi_{X,N} : \mathbb{R} \rightarrow [0, 1]$, such that as $\text{snr} \rightarrow \infty$,

$$\text{mmse}(X, N, \text{snr}) = \frac{\text{var} N}{\text{snr}} \Phi_{X,N}(\log \text{snr}) + o\left(\frac{1}{\text{snr}}\right). \quad (35)$$

The lower and upper MMSE dimension of (X, N) are given by

$$\underline{\mathcal{D}}(X, N) = \limsup_{b \rightarrow \infty} \Phi_{X,N}(b) = \sup_{0 \leq b \leq 2 \log M} \Phi_{X,N}(b), \quad (36)$$

$$\overline{\mathcal{D}}(X, N) = \liminf_{b \rightarrow \infty} \Phi_{X,N}(b) = \inf_{0 \leq b \leq 2 \log M} \Phi_{X,N}(b). \quad (37)$$

Moreover, when $N = N_G$ is Gaussian, the average of Φ_{X,N_G} over one period coincides with the information dimension of X :

$$\frac{1}{2 \log M} \int_0^{2 \log M} \Phi_{X,N_G}(b) db = d(X) = \frac{H(P)}{\log M}. \quad (38)$$

Theorem 6 shows that in the high-SNR regime, the function $\text{snr mmse}(X, N, \text{snr})$ is periodic in snr (dB) with period M^2 (dB). This periodicity arises from the *self-similarity* of the input, and unlike the lower and upper dimension, the period

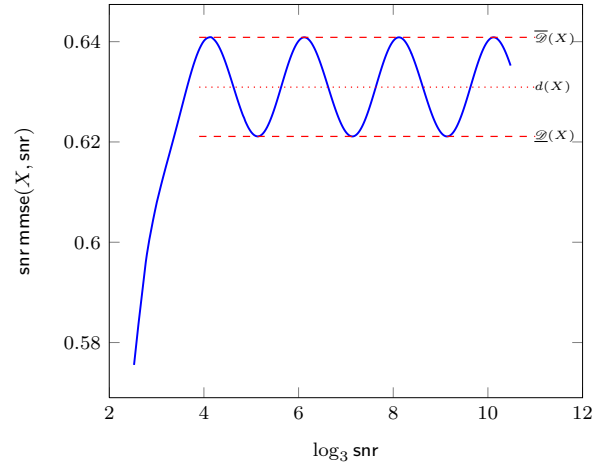


Fig. 1. $\text{snr mmse}(X, \text{snr})$ when X has a ternary Cantor distribution.

depends on the input distribution but not the noise. Although Theorem 6 reveals the oscillatory nature of $\text{mmse}(X, N, \text{snr})$, we do not have a general formula to compute the lower (or upper) MMSE dimension of (X, N) . However, when the noise is Gaussian, Theorem 1 provides a sandwich bound in terms of the information dimension of X , which is reconfirmed by combining (36) – (38).

As an illustrative example, consider Cantor-distributed X whose ternary expansion consists of i.i.d. digits, and for each j , $\mathbb{P}\{(X)_j = 0\} = \mathbb{P}\{(X)_j = 2\} = 1/2$. According to Theorem 6, in the high-SNR regime, $\text{snr mmse}(X, \text{snr})$ oscillates periodically in $\log \text{snr}$ with period $2 \log 3$, as plotted in Fig. 1. The lower and upper MMSE dimensions of the Cantor distribution turn out to be (to six decimals): $\underline{\mathcal{D}}(X) = 0.621102$, $\overline{\mathcal{D}}(X) = 0.640861$. The information dimension $d(X) = \log_3 2 = 0.630930$ is sandwiched between $\underline{\mathcal{D}}(X)$ and $\overline{\mathcal{D}}(X)$, according to Theorem 1. From this and other numerical evidence it is tempting to conjecture that when the noise is Gaussian,

$$d(X) = \frac{\underline{\mathcal{D}}(X) + \overline{\mathcal{D}}(X)}{2}. \quad (39)$$

It should be pointed out that the sandwich bounds in (21) need not hold when N is not Gaussian. For example, for X Cantor distributed and N uniformly distributed in $[0, 1]$, numerical calculation shows that $d(X) = \log_3 2 > \underline{\mathcal{D}}(X, N) = 0.5741$.

III. CONCLUDING REMARKS

Through the high-SNR asymptotics of MMSE in Gaussian channels, we defined a new information measure called MMSE dimension. Although stemming from estimation-theoretic principles, MMSE dimension shares several important features with Rényi's information dimension. By Theorem 3, MMSE dimensions are affine functionals, a fundamental property shared by information dimension [12, Theorem 2]. According to Theorem 1, information dimension is sandwiched between the lower and upper MMSE dimensions. For distributions with no singular components, they coincide to be the weight on the

absolutely continuous part of the distribution. In [12], we have shown that the information dimension plays a pivotal role in almost lossless analog compression, an information theory for noiseless compressed sensing. In fact we have shown that the MMSE dimension serves as the fundamental limit in noisy compressed sensing with stable recovery.

Via the input-noise duality

$$\text{snr} \cdot \text{mmse}(X, N, \text{snr}) = \text{mmse}(N, X, \text{snr}^{-1}) \quad (40)$$

and the following equivalent definition of the MMSE dimension

$$\mathcal{D}(X, N) = \lim_{\epsilon \downarrow 0} \text{mmse}(N, X, \epsilon), \quad (41)$$

studying high-SNR asymptotics provides insights into the behavior of $\text{mmse}(X, N, \text{snr})$ for non-Gaussian noise N . Combining various results from Sections II, we observe that unlike in Gaussian channels where $\text{mmse}(X, \text{snr})$ is decreasing and convex in snr , $\text{mmse}(X, N, \text{snr})$ can behave very irregularly in general. To illustrate this, we consider the case where standard Gaussian input is contaminated by various additive noises. For all N , it is evident that $\text{mmse}(X, N, 0) = \text{var}X = 1$. The behavior of MMSE associated with Gaussian, Bernoulli and Cantor distributed noises is (Fig. 2)

- For standard Gaussian N , $\text{mmse}(X, \text{snr}) = \frac{1}{1+\text{snr}}$ is *continuous* at $\text{snr} = 0$ and decreases monotonically according to $\frac{1}{\text{snr}}$. This monotonicity is due to the MMSE data-processing inequality and the stability of Gaussian distribution.
- For equiprobable Bernoulli N , $\text{mmse}(X, N, \text{snr})$ is *discontinuous* at $\text{snr} = 0$. As $\text{snr} \rightarrow 0$, the MMSE vanishes according to $O\left(\frac{1}{\sqrt{\text{snr}}} e^{-\frac{1}{2\text{snr}}}\right)$, in view of (40) and (34), and since it also vanishes as $\text{snr} \rightarrow \infty$, it is not monotonic with $\text{snr} > 0$.
- For Cantor distributed³ N , $\text{mmse}(X, N, \text{snr})$ is also *discontinuous* at $\text{snr} = 0$. According to Theorem 6, as $\text{snr} \rightarrow 0$, the MMSE oscillates relentlessly between $\overline{\mathcal{D}}(N)$ and $\underline{\mathcal{D}}(N)$ and does not have a limit (See the zoom-in plot in Fig. 3).

Nevertheless, as $\text{snr} \rightarrow \infty$, the MMSE vanishes as $\frac{\text{var}N}{\text{snr}}$ regardless of the noise.

REFERENCES

- [1] D. Guo, S. Shamai, and S. Verdú, "Mutual Information and Minimum Mean-Square Error in Gaussian Channels," *IEEE Transactions on Information Theory*, vol. 51, no. 4, pp. 1261 – 1283, April 2005.
- [2] Y. Wu and S. Verdú, "MMSE Dimension," submitted to *IEEE Transactions on Information Theory*, 2009.
- [3] L. D. Brown, "Admissible estimators, recurrent diffusions, and insoluble boundary value problems," *The Annals of Mathematical Statistics*, vol. 42, no. 3, pp. 855–903, 1971.
- [4] M. V. Burnashev, "Investigation of second order properties of statistical estimators in a scheme of independent observations," *Math. USSR Izvestiya*, vol. 18, no. 3, pp. 439–467, 1982.
- [5] J. K. Ghosh, *Higher Order Asymptotics*. Institute of Mathematical Statistics, 1994.
- [6] A. Rényi, "On the dimension and entropy of probability distributions," *Acta Mathematica Hungarica*, vol. 10, no. 1 – 2, March 1959.

³Since the standard Cantor distribution has variance $\frac{1}{8}$, N has been scaled by $2\sqrt{2}$ to make it unit-variance.

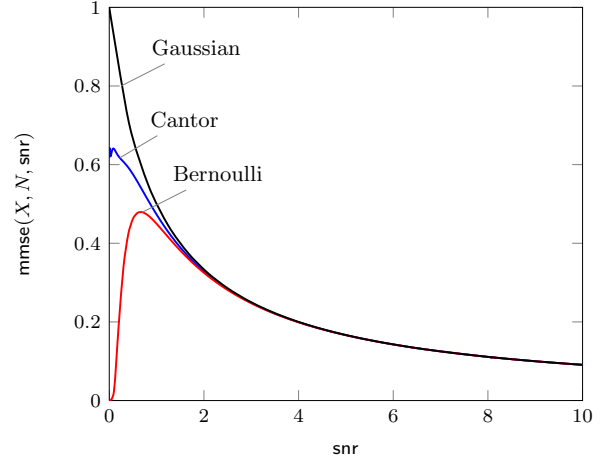


Fig. 2. $\text{mmse}(X, N, \text{snr})$ when X is standard Gaussian and N is standard Gaussian, equiprobable Bernoulli or Cantor distributed.

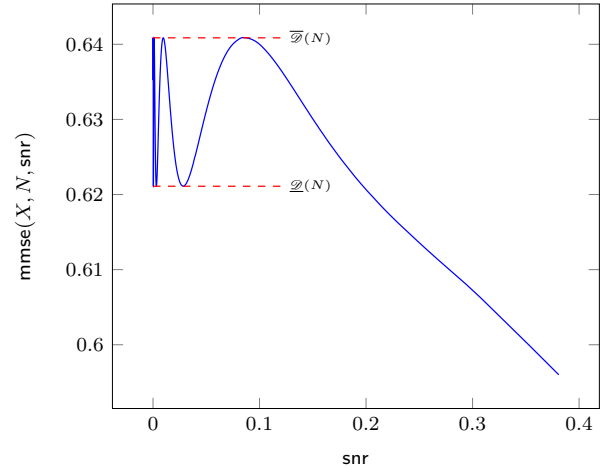


Fig. 3. $\text{mmse}(X, N, \text{snr})$ when X is standard Gaussian and N is Cantor distributed.

- [7] A. Guionnet and D. Shlyakhtenko, "On classical analogues of free entropy dimension," *Journal of Functional Analysis*, vol. 251, no. 2, pp. 738 – 771, 2007.
- [8] V. V. Prelov and E. C. van der Meulen, "Asymptotics of Fisher Information under Weak Perturbation," *Problems of Information Transmission*, vol. 31, no. 1, pp. 14 – 22, 1995.
- [9] A. Stam, "Some inequalities satisfied by the quantities of information of Fisher and Shannon," *Information and Control*, vol. 2, no. 2, pp. 101–112, 1959.
- [10] H. Van Trees, *Detection, Estimation, and Modulation theory, Part I*. New York: Wiley, 1968.
- [11] P. Huber, *Robust Statistics*. Wiley-Interscience, 1981.
- [12] Y. Wu and S. Verdú, "Rényi Information Dimension: Fundamental Limits of Almost Lossless Analog Compression," submitted to *IEEE Transactions on Information Theory*, 2009.
- [13] D. Guo, D. Baron, and S. Shamai, "A single-letter characterization of optimal noisy compressed sensing," in *Proceedings of the Forty-Seventh Annual Allerton Conference on Communication, Control, and Computing*, Monticello, IL, October 2009.
- [14] D. L. Donoho, A. Maleki, and A. Montanari, "Message-passing algorithms for compressed sensing," *Proceedings of the National Academy of Sciences*, vol. 106, no. 45, pp. 18914 – 18919, November 2009.
- [15] K. Falconer, *Techniques in Fractal Geometry*. John Wiley, 1997.