

Estimation of Rate-Distortion Function for Computing with Decoder Side Information

Heasung Kim, Hyeji Kim, and Gustavo de Veciana

Department of Electrical and Computer Engineering

The University of Texas at Austin

Austin, TX, USA

Email: {heasung.kim, hyeji.kim, deveciana}@utexas.edu

Abstract—There has been growing interest in computing rate-distortion functions for real-world data, as they can provide a theoretical benchmark for compression problems. However, a generalized form of rate-distortion that includes side information and coding for computing has been underexplored, despite its relevance in modern compression problems. To address this gap, we propose a new method for estimating the rate-distortion function for computing with side information, using a Lagrangian framework with neural network-parametrized encoding and decoding strategies. This approach enables targeting specific points on the rate-distortion curve through gradient-based optimization. Our methodology is validated in synthetic environments where rate-distortion functions are known, ensuring accuracy in estimation. Additionally, we extend its application to practical, high-dimensional channel state information compression scenarios. We provide rate-distortion estimation results on these scenarios, which in turn enables us to quantify the usefulness of side information in the practical scenarios.¹

I. INTRODUCTION

The rate-distortion function [1], which characterizes the optimal rate-distortion trade-off, serves as a theoretical benchmark for assessing the effectiveness of compression algorithms, as highlighted in recent studies [?], [2]–[4]. However, accurately computing Shannon’s information measures, such as entropy and mutual information which form the basis of rate-distortion function, is notably challenging. This is particularly true in scenarios involving real-world distributions where one must rely solely on samples without any additional knowledge of the distributions, or in cases involving high-dimensional input sources. Closed-form solutions for these measures are generally limited to specific circumstances, e.g., Gaussian sources.

A. Computing rate-distortion functions

An approach to numerically compute the rate-distortion functions and associated information measures for general distributions has been devised based on iterative algorithms in 1972 by Blahut [5] and Arimoto [6]. Known collectively as the Blahut–Arimoto algorithms, they have been adapted to address multiterminal source coding settings [7].

However, these conventional iterative approaches face limitations, especially when applied to high-dimensional or continuous sources [3]. To overcome these challenges, recent

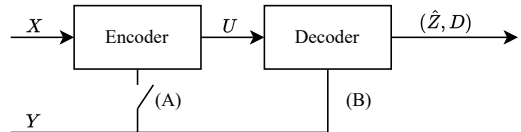


Fig. 1. Coding for Computing with Side Information. We consider a configuration where switch (A) remains open while switch (B) is closed, permitting access to side information at the decoder. The decoder aims to compute a function $Z = g(X, Y)$; we let $\hat{Z}$ and D denote the decoder’s output and distortion, respectively.

studies have explored solutions utilizing neural networks or advanced optimization techniques. The Restricted Boltzmann Machines is integrated with neural networks to estimate rate-distortion functions [2]. In [3], rate-distortion function duality concepts, e.g., [8], is utilized in estimation methods. A sandwich bound for rate-distortion function is introduced in [4] through distribution parameterization with neural networks. In [9], neural estimation methods with generative model framework are used. Notably, the Wasserstein gradient descent algorithm proposed in [?], has demonstrated state-of-the-art performance for rate-distortion estimation, without relying on neural networks.

B. Rate-distortion function for computing with side information

The rate-distortion function concept can be extended to encompass scenarios where side information, correlated with the input source, is available at the decoder, or at both the encoder and decoder, as illustrated in Figure 1. This adaptation is widely recognized as the Wyner-Ziv rate-distortion function [10]. Additionally, the notion of the rate-distortion is further broadened by considering communication systems where the goal is to compute a function of the source. Such applications of the Wyner-Ziv rate-distortion function are commonly referred to as *Coding for Computing* [11].

Such a broader perspective of the rate-distortion function is crucial for evaluating compression algorithms in practical scenarios and understanding side information’s role in various contexts. It also offers a way to quantify the relevance of different types of side information for various sources.

While recent compression techniques increasingly incorporate side information [12]–[14] with specific objectives, research on rate-distortion estimation with side information, especially for continuous or large-dimensional distributions,

¹The code is available at <https://github.com/Heasung-Kim/rate-distortion-side-information>.

remains limited. Prior studies have focused on discrete sources and side information [15], extending the Blahut-Arimoto Algorithm but often struggle with high-dimensional distributions particularly when there is no a priori knowledge of the distributions. Recently, [16] introduced neural network-based methods for estimation of rate-distortion function with side information with discrete codewords, however, prioritizing variational upper bounds of the rate-distortion functions optimization.

Our contributions address these gaps through a neural network-based direct estimation method for the rate-distortion function for computing with side information, along with applicable methodologies:

C. Contributions

We present a generalized framework for estimating the rate-distortion function for computing with side information. We formulate a Lagrangian loss function where the minimization of this function is achieved through specific encoding and decoding schemes that can achieve point(s) on the rate-distortion function. Our algorithm focuses on minimizing the loss by parameterizing the conditional distributions of the codewords for a given source and side information, as well as the decoder. The algorithm is designed to alternatively update these parameters to efficiently minimize the Lagrangian loss.

The effectiveness of our algorithm is validated through numerical evaluation, particularly in cases where the rate-distortion function is known, such as with correlated Gaussian distributions for the source and side information. These simulations demonstrate that our algorithm consistently provides a precise estimation of the rate-distortion function. Extending beyond the synthetic data, we also apply our approach to practical scenarios. This includes assessing the possible gains when side information is available when considering channel state information compression problems and illustrating the practical relevance and applicability of our method to high-dimensional sources.

II. SYSTEM MODEL AND PROPOSED METHOD

Consider the communication system illustrated in Fig. 1, where switch (A) is open and (B) is closed. The primary source and side information pair (X, Y) is assumed to be independently and identically distributed (i.i.d.), following a joint distribution $p_{X,Y}(x, y)$. Here, x and y are realizations of X and Y from the domains $\mathcal{X}$ and $\mathcal{Y}$, respectively. The codeword is represented by U from the domain $\mathcal{U}$ and the output of the decoder is $\hat{Z} \in \mathcal{Z}$, with D denoting a distortion level and d being a distortion measure defined as a mapping as $d : \mathcal{Z} \times \mathcal{X} \rightarrow \mathbb{R}^+$. For readability, we also let $\hat{X}$ denote the output of the decoder when the system aims to reconstruct the original information X . This system implies a Markov chain $Y - X - U$.

In our general framework we assume the objective is to reconstruct a function $Z = g(X, Y)$, where Z is not necessarily identical to X . In this setting, the corresponding rate-distortion function determines the minimum necessary rate to compute

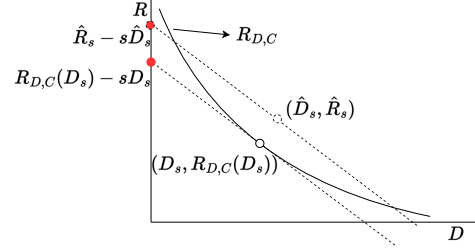


Fig. 2. $R_{D,C}$ is convex with respect to distortion D . For a given slope s , minimizing the y -intercept of a line originating from an achievable point $(\hat{D}_s, \hat{R}_s)$ in the rate-distortion region leads to a new y -intercept, which corresponds to a line that is tangent to the $R_{D,C}$ curve at point(s) with the same slope s .

$g(X, Y)$ within a given distortion threshold D . The rate-distortion function, denoted $R_{D,C}$, is given as follows [17].

Definition 1 (Rate-distortion function for computing with side information).

$$R_{D,C}(D) = \min_{q_{U|X}(u|x), f(u,y): \mathbb{E}[d(Z, \hat{Z})] \leq D} I(X; U|Y) \quad (1)$$

where $q_{U|X}(u|x)$ is a conditional probability distribution of U given X . Z and $\hat{Z}$ are the desired function output and the decoder output, respectively. f is a decoder taking u and side information y as an input pair as $f(u, y) = \hat{z}$.

In this paper, we focus on developing a method to estimate $R_{D,C}(D)$ for a general function, particularly in scenarios where the joint distribution $p_{X,Y}(x, y)$ is unknown and a dataset of N data points $(x_i, y_i)_{i=1}^N$ that are sampled from $p_{X,Y}(x, y)$ is available. This setup is typical in real-world contexts, where the exact distribution underlying a dataset is often not known.

We start with a Lagrangian formulation to address the optimization problem defined in (1) by exploiting convexity and non-increasing property of $R_{D,C}(D)$ with respect to the distortion D . We can formulate an optimization problem for finding the vertical intercept of the tangent with slope $s (\leq 0)$ to the rate-distortion curve as follows.

$$R_{D,C}(D_s) - sD_s = \min_{q_{U|X}, f} \{I(X; U|Y) - s\mathbb{E}[d(Z, \hat{Z})]\}. \quad (2)$$

For a given slope s and a corresponding achievable (distortion, rate) pair, $(\hat{D}_s, \hat{R}_s)$ illustrated in Fig. 2, the y -intercept at this line is $\hat{R}_s - s\hat{D}_s$. This intercept is equivalent to $I(X; U|Y) - s\mathbb{E}[d(Z, \hat{Z})]$, attained by the specific encoding and decoding schemes associated with $q_{U|X}$, f corresponding to $(\hat{D}_s, \hat{R}_s)$. This y -intercept can be minimized through optimization, adjusting the encoding and decoding schemes accordingly.

Due to the convexity of $R_{D,C}$, the lowest achievable value of the vertical intercept corresponds to $R_{D,C}(D_s) - sD_s$ where the distortion D_s and rate $R_{D,C}(D_s)$ is a point lies on the $R_{D,C}$ curve itself. By determining a point on the $R_{D,C}$ curve for each slope s and then varying s , we can estimate the $R_{D,C}$ curve.

Algorithm 1 Estimation of Rate-Distortion Function for Computing with Side Information at Decoder

1: **Input:** Slope s , dataset $\{x_i, y_i\}_{i=1}^N$, initialized sets of parameters $\theta_{\text{po}}, \theta_{\text{pr}}, \theta_{\text{dec}}$, number of iterations T, T' , batch size b
2: **for** $t = 0$ to T **do**
3: Sample minibatch $\mathcal{B} = \{(x_i, y_i)\}_{i=1}^b$ and sample $\{u_i\}_{i=1}^b$ from $\{q_{U|X=x_i}\}_{i=1}^b$
4: Compute $\nabla L_1 = \nabla \frac{1}{b} \sum_{i=1}^b [\log(q_{U|X}(u_i|x_i; \theta_{\text{po}}) - \log q_{U|Y}(u_i|y_i; \theta_{\text{pr}}))] - s[d(g(x_i, y_i), f(u_i, y_i; \theta_{\text{dec}}))]$
5: Update $\theta_{\text{po}} \leftarrow \theta_{\text{po}} - \nabla_{\theta_{\text{po}}} L_1$ and $\theta_{\text{dec}} \leftarrow \theta_{\text{dec}} - \nabla_{\theta_{\text{dec}}} L_1$
6: **for** $t' = 0$ to T' **do**
7: Sample minibatch $\mathcal{B}' = \{(x_i, y_i)\}_{i=1}^b$ and sample $\{u_i\}_{i=1}^b$ from $\{q_{U|X=x_i}\}_{i=1}^b$
8: Compute $\nabla L_2 = \nabla \frac{1}{b} \sum_{i=1}^b [\log(q_{U|X}(u_i|x_i; \theta_{\text{po}}) - \log q_{U|Y}(u_i|y_i; \theta_{\text{pr}}))]$
9: Update $\theta_{\text{pr}} \leftarrow \theta_{\text{pr}} - \nabla_{\theta_{\text{pr}}} L_2$

To facilitate estimation using a given dataset, we reformulate the optimization term as follows.

$$\min_{q_{U|X}, f} \left\{ \mathbb{E}_{X,Y,U} \left[\log \frac{q_{U|X}(U|X)}{q_{U|Y}(U|Y)} \right] - s \mathbb{E}[d(Z, \hat{Z})] \right\}, \quad (3)$$

where $q_{U|Y}(u|y) = \sum_{x \in \mathcal{X}} p_{X|Y}(x|y) q_{U|X}(u|x)$ (when X is a discrete random variable) and $\hat{Z} = f(U, Y)$. This formulation enables the computation of expectation terms using Monte Carlo estimation with data points following the distribution $q_{X,Y,U}(x, y, u)$. Here, we use the notation q to represent a probability distribution influenced by $q_{U|X}$ and f , while p has been used to denote distributions independent of $q_{U|X}$ and f .

To proceed with this approach, we parameterize the key components of the optimization problem using a neural network based model. First, we represent the conditional distribution $q_{U|X}(u|x)$ as $q_{U|X}(u|x; \theta_{\text{po}})$ where θ_{po} denotes a set of parameters for $q_{U|X}$. Similarly, we parameterize the decoding function with a set of parameters θ_{dec} as $f(u, y; \theta_{\text{dec}})$.

It should be noted that the parameterization of $q_{U|X}(u|x; \theta_{\text{po}})$ directly determines the related marginal and joint distributions, such as $q_{U|X,Y}$, $q_{U|Y}$, and $q_{X,Y,U}$ under the fixed $p_{X,Y}$. These distributions, governed by the parameter set θ_{po} , are thus denoted as $q_{U|X,Y;\theta_{\text{po}}}$, $q_{U|Y;\theta_{\text{po}}}$, and $q_{X,Y,U;\theta_{\text{po}}}$.

In summary, we begin with the Lagrangian optimization problem for the rate-distortion function, also known as the supporting hyperplane method. We employ neural networks to parameterize the key components of our loss function. This optimization strategy draws parallels with the conventional Blahut-Arimoto algorithms [5], [6] in terms of formulating the Lagrangian loss, while also drawing inspiration from recent works [4], which has achieved state-of-the-art results in rate-distortion estimation through neural networks.

In the following subsections, we delve into a comprehensive explanation of our proposed algorithm, detailing the steps and techniques involved. We will also discuss the methods used for parameterizing the components in our framework.

A. Algorithm

The proposed method is detailed in Algorithm 1. This algorithm iteratively computes the gradient of the loss function (3) over T training iterations and updates the relevant parameters to minimize the loss.

Line 3. Specifically, in each iteration, a minibatch with size b , $\mathcal{B} = \{(x_i, y_i)\}_{i=1}^b$, is sampled. To estimate the expectation

$\mathbb{E}_{X,Y,U}[\log q_{U|X}(U|X) - \log q_{U|Y}(U|Y)]$, it is necessary to generate data point triples (x_i, y_i, u_i) following the distribution $p_{X,Y}(x, y) q_{U|X}(u|x; \theta_{\text{po}})$. For each sampled pair (x_i, y_i) , a corresponding u_i is drawn from the distribution $q_{U|X}(u|x; \theta_{\text{po}})$. This sampling results in triples (x_i, y_i, u_i) that adhere to the joint distribution $q_{X,Y,U}(x, y, u) = p_{X,Y}(x, y) q_{U|X}(u|x; \theta_{\text{po}})$.

Utilizing these samples, we compute the average gradient of the loss function, which involves the computation of the expected value of $\log \frac{q_{U|X}(U|X)}{q_{U|Y}(U|Y)}$. This corresponds to $\log \frac{q_{U|X}(U|X; \theta_{\text{po}})}{q_{U|Y;\theta_{\text{po}}}(U|Y)}$ based on the parameterization where $q_{U|Y;\theta_{\text{po}}}$ is formulated as

$$\begin{aligned} q_{U|Y;\theta_{\text{po}}}(u|y) &= \sum_{x \in \mathcal{X}} p_{X|Y}(x|y) q_{U|X,Y;\theta_{\text{po}}}(u|x, y) \\ &= \sum_{x \in \mathcal{X}} p_{X|Y}(x|y) q_{U|X}(u|x; \theta_{\text{po}}). \end{aligned} \quad (4)$$

The efficient computation of $q_{U|Y;\theta_{\text{po}}}$ is critical, as it needs to be executed for multiple instances to obtain the average of the log probability. However, this computation of (4) presents a substantial challenge due to the unknown nature of the distribution $p_{X|Y}$, with only sample-based access available. Furthermore, using sampling approaches for the estimation of the sum over $\mathcal{X}$ is non-trivial when domain $\mathcal{X}$ is a high-dimensional space and the data instances are limited. To address this issue, we leverage the following lemma.

Lemma 1. Consider a fixed set of parameters θ_{po} and scenario where the side information Y is available only at the decoder. Then we have

$$\arg \min_{q_{U|Y}} \mathbb{E}_{X,Y,U} \left[\log \frac{q_{U|X}(U|X; \theta_{\text{po}})}{\hat{q}_{U|Y}(U|Y)} \right] = q_{U|Y;\theta_{\text{po}}}. \quad (5)$$

Based on this lemma, we conclude that instead of executing the summation in (4) to derive $q_{U|Y;\theta_{\text{po}}}$ for a given $q_{U|X}(u|x; \theta_{\text{po}})$, we can model the distribution of U given Y as $q_{U|Y}(u|y; \theta_{\text{pr}})$ where θ_{pr} denotes a set of free parameters and then use the parameterized distribution $q_{U|Y}(u|y; \theta_{\text{pr}})$ as an argument for the problem (5). The solution of (5) will lead to $q_{U|Y}(u|y; \theta_{\text{pr}}) = q_{U|Y;\theta_{\text{po}}}(u|y)$ as long as the parametrization of $q_{U|Y}(u|y; \theta_{\text{pr}})$ is expressive enough.

Lines 4-5. By using the parametrized functions $q_{U|X}(u|x; \theta_{\text{po}})$, $q_{U|Y}(u|y; \theta_{\text{pr}})$, and $f(u, y; \theta_{\text{dec}})$, in Line 4,

we compute the gradient of (3). Subsequently, in Line 5, the parameters θ_{po} and θ_{dec} are updated to minimize the loss.

Lines 6-9. At the end of each iteration, we update $q_{U|Y}(u|y; \theta_{\text{pr}})$ by solving (5) based on the newly updated $q_{U|X}(u|x; \theta_{\text{po}})$ to correctly compute the main loss function (3) in the subsequent iteration. Problem (5) can be solved through gradient descent updates of the set of parameters θ_{pr} as described in Lines 7-9 of Algorithm 1. More specifically, for each inner-iteration (occurring T' times), we sample a minibatch and obtain pairs $\{(x_i, y_i, u_i)\}_{i=1}^b$. We then update θ_{pr} to minimize the objective in (5). Practically, we have found that setting $T' = 1$ and reusing the same minibatch $\mathcal{B}$ for $\mathcal{B}'$ not only offers computational efficiency but also provides a tight upper bound on the rate-distortion function relative to theoretical optimality (as detailed in Sec. III).

B. Parameterization

In Algorithm 1, we utilize three distinct parameterized models: $q_{U|X}(u|x; \theta_{\text{po}})$, $q_{U|Y}(u|y; \theta_{\text{pr}})$, and $f(u, y; \theta_{\text{dec}})$. A conventional approach to parameterizing distributions involves assuming a specific distribution form and then parameterizing its moments, such as the mean and variance. Various parameterization setups exist, including Gaussian, uniform distribution-based parameterizations, and more sophisticated forms relevant to modern machine learning research [18]. In our study, we opt for Gaussian distributions for parameterization. For example, in sampling from the distribution $q_{U|X}(u|x; \theta_{\text{po}})$, the random variable U is assumed to follow a Gaussian distribution characterized by mean $\mu(x; \theta_{\text{po}})$ and variance $\Sigma(x; \theta_{\text{po}})$, both of which depend on the given realization x . The functions μ and Σ can be designed in various ways, depending on the specifics of the problem, where they take x as input and output the corresponding mean and variance.

We provide more details on the implementation in Sec. III. The choice of parameterization and the construction of these functions yield a point that represents an upper bound on the rate-distortion curve. This is because the variable spaces for the minimization problem in (3) is constrained by the assumptions inherent in the chosen distribution models. Thus, while these parameterizations facilitate the computational tractability of the problem, they also inherently define the limits of the solution space explored in the optimization process.

III. NUMERICAL EVALUATION

In order to evaluate our algorithm's efficacy, our initial step involves scenarios where the true rate-distortion function for computing with side information is known in closed form. Beyond these environments, we consider estimating the rate-distortion function for practical problem associated with channel state information (CSI) compression [19], incorporating side information.

A. 2-Component White Gaussian Noise

We adapt a scenario from [11, Sec. 21.1], featuring a 2-component White Gaussian Noise (2-WGN(P, ρ)) source, where (X, Y) forms pairs of i.i.d. jointly

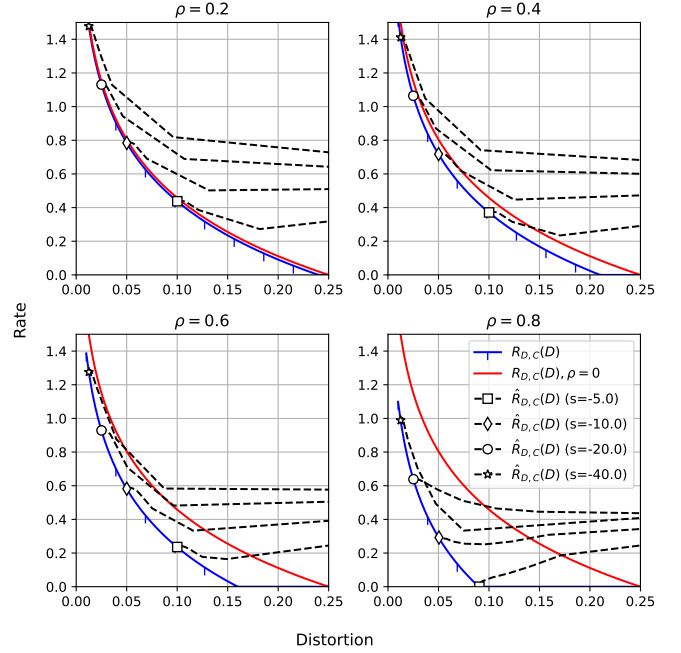


Fig. 3. Compression of Gaussian sources: Rate-distortion functions for computing with side information for various ρ and the estimated points. The solid lines represent the known rate-distortion functions based on (6)

Gaussian random variables. Each pair in the sequence $(X_1, Y_1), (X_2, Y_2), \dots, (X_n, Y_n)$ has zero mean ($\mathbb{E}[X] = \mathbb{E}[Y] = 0$), equal variance ($\mathbb{E}[X^2] = \mathbb{E}[Y^2] = P$), and a correlation coefficient $\rho = \mathbb{E}[XY]/P$. With a squared error distortion measure d and a function $g(X, Y) = (X + Y)/2$, $R_{D,C}$ is given by

$$R_{D,C}(D) = \max \left\{ \frac{1}{2} \log \left(\frac{P(1-\rho^2)}{4D} \right), 0 \right\}. \quad (6)$$

To implement our approach, we employed a multi-layer perceptron (MLP) to model $q_{U|X}(u|x; \theta_{\text{po}})$, $q_{U|Y}(u|y; \theta_{\text{pr}})$, and $f(u, y; \theta_{\text{dec}})$. Specifically, for $q_{U|X}(u|x; \theta_{\text{po}})$ and $q_{U|Y}(u|y; \theta_{\text{pr}})$, we use a single-layer MLP that takes an n -dimensional input and outputs a $2n$ -dimensional vector, half for mean and half for variance, to model an n -dimensional independent multivariate Gaussian distribution. For $f(u, y; \theta_{\text{dec}})$, we used a 2-layer MLP with leaky ReLU activation, which takes (u, y) as an input and outputs an n -dimensional $\hat{z}$.

In Fig. 3, we set $P = 1, n = 100$, and provide simulation results for various ρ values in $\{0.2, 0.4, 0.6, 0.8\}$. Each subplot displays the $R_{D,C}$ curves, alongside four rate-distortion points estimated by our algorithm for different slopes s . We also plot $R_{D,C}$ curves with $\rho = 0$. y -axis has natural units (Nats) and x -axis represents mean squared error distortion. The dashed lines associated with $\hat{R}_{D,C}(D)$ corresponds to the learning trajectory, i.e., the achieved (distortion, rate) points during the training process.

Our algorithm consistently estimates points on $R_{D,C}$ within a small tolerance of less than 5×10^{-3} in average. A decrease in the s value corresponds to points on the left side of the curve, indicating higher rates and lower distortion. As can be seen, a higher correlation ρ results in a reduced rate for a given

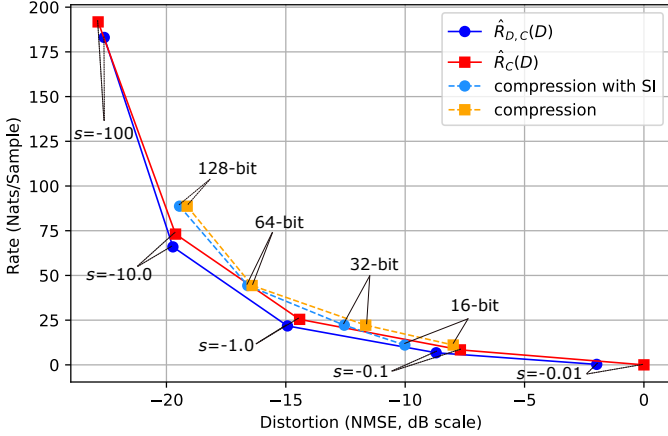


Fig. 4. CSI Compression: Comparison of the estimated rate-distortion function, estimated rate-distortion function with side information, and (distortion, rate) points achieved by the neural compression algorithm.

distortion, and our method effectively estimates points on $R_{D,C}$ regardless of various ρ values.

B. Applications to CSI Compression

This subsection focuses on the compression of Frequency Division Duplex Downlink (DL) Channel State Information (CSI), an area of growing interest in wireless research [19].

1) *Setup*: The objective is to compress the DL CSI, X , at the User Equipment (UE) side. The UE then transmits this compressed information, or codeword U , to the Base Station (BS). The aim is to minimize the Normalized Mean Squared Error (NMSE), defined as $\mathbb{E}[\|X - \hat{X}\|_2^2 / \|X\|_2^2]$, where $\hat{X}$ is the decoder output and $\|\cdot\|_2$ is elementwise square norm.

To enhance compression efficiency, uplink (UL) CSI can be utilized as side information Y . This is based on the observation that UL CSI is typically acquired (available) via pilot transmissions from the UE to BS, and is correlated with DL CSI due to frequency-invariant characteristics [20], [21].

2) *Simulation environment configuration*: A CSI instance X characterized by the setting (n_{tx}, n_{sc}) , with $n_{tx} = 8$ represents the number of transmit (Tx) antennas and $n_{sc} = 667$ denotes the number of subcarriers. UL CSI also has the same parameters. Our numerical evaluation use the CDL-C model specified in [22]. We utilize inception block-based [23] encoding and decoding schemes [24] for the distribution parameterization. In [25], detailed configurations for parameterizing distributions [23], [24], data preprocessing methods, and neural network training methods are provided.

3) *Constructive neural CSI compression algorithms*: For further analysis, we implement CSI compression algorithms having fixed rates (codeword lengths) with the same architecture used in Sec. III-B2. Specifically, replacing the encoder's output with a deterministic function, rather than providing mean and variance for probability distributions, facilitates the design of deterministic codewords. This modification involves adapting the encoder to directly output a codeword. Subsequently, a discretization technique, as outlined in [26], is employed to generate fixed l_{cl} -bit codeword for a given CSI instance.

4) *Results*: In Figure 4, four distinct curves are presented: the estimated rate-distortion curve with side information, $\hat{R}_{D,C}$, the one without side information, $\hat{R}_C$, the rate-distortion curve derived from the constructive compression algorithm with side information (compression with SI), and the curve without side information (compression). The points plotted on $\hat{R}_{D,C}$ and $\hat{R}_C$ denote distinct estimated rate-distortion points, with their positions corresponding to specific s values: -100, -10, -1, -0.1, and -0.01 arranged from left to right. $\hat{R}_C$ is obtained by [4] and using the same neural architectures but which ignores the side information. By adjusting s values, we explore distortion levels from 0dB to approximately -23dB, connecting these points linearly to serve as an upper bound for the estimated rate-distortion curves. The rate-distortion curves from the constructive algorithms are generated by varying the compression bit rates as $l_{cl} \in \{16, 32, 64, 128\}$ and connecting the points.

As expected, introducing UL CSI for DL CSI compression is beneficial as $\hat{R}_{D,C} < \hat{R}_C$, especially at lower CSI feedback rates. For instance, with near 0 nats/sample rate, the BS can still retrieve reasonable information from UL CSI, achieving better than -2dB NMSE. This advantage diminishes with increased feedback resources; for example, at 150 Nats/Sample, the gain is observed to be near zero.

The neural compression algorithm incorporating side information, achieved a rate-distortion curve that establishes an upper bound of $\hat{R}_{D,C}$, and the discrepancy between $\hat{R}_{D,C}$ and the constructive CSI compression algorithm's performance signals room for improvement. For example, in the case of around 80 Nats/sample, we may anticipate a potential improvement about 1dB. Notably, this gap is less pronounced in scenarios with lower rates, allowing one to have a conjecture that the actual performance of the CSI compression algorithms is closer to $\hat{R}_{D,C}$.

IV. DISCUSSION

In this paper, we propose a new algorithm for estimating the generalized rate-distortion function, with a specific emphasis on the rate-distortion function for computing with side information. This approach can offer estimated rates for given distortion levels and also enables the formulation of reliable conjectures about the benefits of side information at varying compression rates. Such a methodology is anticipated to be valuable in practical system design, allowing system designers to effectively measure the potential gains from side information against its processing costs through informed estimations.

V. ACKNOWLEDGEMENT

This work was partly supported by ARO Award W911NF2310062, ONR Award N00014-21-1-2379, National Science Foundation under grant no. 2148224, NSF Award CNS-2008824, funds from OUSD R&E, NIST, industry partners as specified in the Resilient & Intelligent NextG Systems (RINGS) program, and InterDigital, INC. through 6G@UT center within the Wireless Networking and Communications Group (WNCG) at the University of Texas at Austin.

REFERENCES

- [1] T. Berger, "Rate-distortion theory," *Wiley Encyclopedia of Telecommunications*, 2003.
- [2] Q. Li and Y. Chen, "Rate distortion via deep learning," *IEEE Transactions on Communications*, vol. 68, no. 1, pp. 456–465, 2019.
- [3] E. Lei, H. Hassani, and S. S. Bidokhti, "Neural estimation of the rate-distortion function with applications to operational source coding," *IEEE Journal on Selected Areas in Information Theory*, 2023.
- [4] Y. Yang and S. Mandt, "Towards empirical sandwich bounds on the rate-distortion function," *The International Conference on Learning Representations (ICLR)*, 2022.
- [5] R. Blahut, "Computation of channel capacity and rate-distortion functions," *IEEE Trans. Inf. Theory*, vol. 18, no. 4, pp. 460–473, 1972.
- [6] S. Arimoto, "An algorithm for computing the capacity of arbitrary discrete memoryless channels," *IEEE Trans. Inf. Theory*, vol. 18, no. 1, pp. 14–20, 1972.
- [7] Y. Uğur, I. E. Aguerri, and A. Zaidi, "A generalization of blahut-arimoto algorithm to compute rate-distortion regions of multiterminal source coding under logarithmic loss," in *IEEE ITW*. IEEE, 2017, pp. 349–353.
- [8] A. Dembo and I. Kontoyiannis, "Source coding, large deviations, and approximate pattern matching," *IEEE Trans. Inf. Theory*, vol. 48, no. 6, pp. 1590–1615, 2002.
- [9] D. Tsur, B. Huleihel, and H. Permuter, "Rate distortion via constrained estimated mutual information minimization," in *IEEE ISIT*. IEEE, 2023, pp. 695–700.
- [10] A. Wyner and J. Ziv, "The rate-distortion function for source coding with side information at the decoder," *IEEE Trans. Inf. Theory*, vol. 22, no. 1, pp. 1–10, 1976.
- [11] A. El Gamal and Y.-H. Kim, *Network information theory*. Cambridge university press, 2011.
- [12] S. Ayzik and S. Avidan, "Deep image compression using decoder side information," in *Computer Vision—ECCV 2020: 16th European Conference, Glasgow, UK, August 23–28, 2020, Proceedings, Part XVII 16*. Springer, 2020, pp. 699–714.
- [13] N. Mital, E. Özyılkan, A. Garjani, and D. Gündüz, "Neural distributed image compression using common information," in *2022 Data Compression Conference (DCC)*. IEEE, 2022, pp. 182–191.
- [14] Y. Huang, B. Chen, S. Qin, J. Li, Y. Wang, T. Dai, and S.-T. Xia, "Learned distributed image compression with multi-scale patch matching in feature domain," in *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 37, no. 4, 2023, pp. 4322–4329.
- [15] F. Dupuis, W. Yu, and F. M. Willems, "Blahut-arimoto algorithms for computing channel capacity and rate-distortion with side information," in *International Symposium on Information Theory, Proceedings*. IEEE, 2004, p. 179.
- [16] E. Özyılkan, J. Ballé, and E. Erkip, "Learned Wyner-Ziv compressors recover binning," in *IEEE ISIT*. IEEE, 2023, pp. 701–706.
- [17] H. Yamamoto, "Wyner-Ziv theory for a general function of the correlated sources (corresp.)," *IEEE Trans. Inf. Theory*, vol. 28, no. 5, pp. 803–807, 1982.
- [18] J. Ballé, D. Minnen, S. Singh, S. J. Hwang, and N. Johnston, "Variational image compression with a scale hyperprior," *arXiv preprint arXiv:1802.01436*, 2018.
- [19] J. Guo, C.-K. Wen, S. Jin, and G. Y. Li, "Overview of deep learning-based csi feedback in massive mimo systems," *IEEE Transactions on Communications*, vol. 70, no. 12, pp. 8017–8045, 2022.
- [20] D. Vasisht, S. Kumar, H. Rahul, and D. Katabi, "Eliminating channel feedback in next-generation cellular networks," in *Proceedings of the 2016 ACM SIGCOMM Conference*, 2016, pp. 398–411.
- [21] D. Han, J. Park, and N. Lee, "Fdd massive mimo without csi feedback," *IEEE Trans. Wireless Commun.*, 2023.
- [22] G. T. 38.901, "Study on channel model for frequencies from 0.5 to 100 ghz," 2017.
- [23] C. Szegedy, W. Liu, Y. Jia, P. Sermanet, S. Reed, D. Anguelov, D. Erhan, V. Vanhoucke, and A. Rabinovich, "Going deeper with convolutions," in *Proceedings of the IEEE conference on computer vision and pattern recognition*, 2015, pp. 1–9.
- [24] Z. Lu, J. Wang, and J. Song, "Multi-resolution csi feedback with deep learning in massive mimo system," in *IEEE International Conference on Communications*. IEEE, 2020, pp. 1–6.
- [25] H. Kim, "rate-distortion-side-information," <https://github.com/Heasung-Kim/rate-distortion-side-information>, 2024.
- [26] A. Van Den Oord, O. Vinyals *et al.*, "Neural discrete representation learning," *Advances in neural information processing systems*, vol. 30, 2017.