



IRWIN AND JOAN JACOBS
CENTER FOR COMMUNICATION AND INFORMATION TECHNOLOGIES

Simplified Erasure/ List Decoding

Nir Weinberger and Neri Merhav

CCIT Report # 873
December 2014

■ ■ ■ ■ ■ Electronics
■ ■ ■ ■ ■ Computers
■ ■ ■ ■ ■ Communications

DEPARTMENT OF ELECTRICAL ENGINEERING
TECHNION - ISRAEL INSTITUTE OF TECHNOLOGY, HAIFA 32000, ISRAEL



Simplified Erasure/List Decoding

Nir Weinberger and Neri Merhav

Dept. of Electrical Engineering

Technion - Israel Institute of Technology

Technion City, Haifa 3200004, Israel

{nirwein@tx, merhav@ee}.technion.ac.il

Abstract

We consider the problem of erasure/list decoding using certain classes of simplified decoders. Specifically, we assume a class of erasure/list decoders, such that a codeword is in the list if its likelihood is larger than a threshold. This class of decoders both approximates the optimal decoder of Forney, and also includes the following simplified subclasses of decoding rules: The first is a function of the output vector only, but not the codebook (which is most suitable for high rates), and the second is a scaled version of the maximum likelihood decoder (which is most suitable for low rates). We provide single-letter expressions for the exact random coding exponents of any decoder in these classes, operating over a discrete memoryless channel. For each class of decoders, we find the optimal decoder within the class, in the sense that it maximizes the erasure/list exponent, under a given constraint on the error exponent. We establish the optimality of the simplified decoders of the first and second kind for low and high rates, respectively.

Index Terms

Erasure/list decoding, mismatch decoding, random coding, error exponents, decoding complexity.

I. INTRODUCTION

An ordinary decoder must always decide on a single codeword, and consequently, its decision regions form a partition of the space of output vectors. In his seminal paper [1], Forney defined a family of generalized decoders with decision regions which do not necessarily satisfy the above property. An *erasure* decoder has the freedom not to decode, and so, its decision regions are disjoint but not necessarily cover of the space of channel output vectors. A *list* decoder may decide on more than one codeword, thus its decision regions may overlap. An erasure/list decoder can practically be used in cases where an additional mechanism is used to resolve the ambiguity left after a non-decisive decoding instance. For example, if the transmitted codewords at different times are somewhat dependent, then an outer decoder may recover from an erasure by “interpolating” from neighboring codewords, or eliminate all codewords in the list, but one. Such a case is present in concatenated codes where the inner decoder may be an erasure/list decoder. As another example, consider the case where a feedback link may signal the transmitter on an

erasure event. In this case, the transmitter may send repeatedly the codeword until there is no erasure. Indeed, this is the case in rateless codes [2], [3], where at each time instant, the decoder needs to decide whether to continue acquiring additional output symbols, or decide on the transmitted codeword.

When considering list decoders, a distinction should be made between list decoders with a fixed list size (see for example [4], and references therein), and decoders for which the list size is a function of the output vector and thus a random variable. In the latter case, a natural trade-off in designing the decoder is the error probability, i.e., the probability that the correct codeword is not on the list, and the average list size. In [1], Forney has generalized the Neyman-Pearson Lemma and derived the decoder which optimally trades off between error probability and average list-size. This optimal list decoder and the optimal erasure decoder (which optimally trades off between error probability and erasure probability), were found to have the same structure.

In addition, Forney has extended the analysis of Gallager for ordinary decoding [5, Chapter 5], and derived lower bounds on the random coding error exponent and the list size exponent (the normalized logarithm of the expected list size). More recently, Somekh-Baruch and Merhav [6] have found exact expressions for the random coding exponents, inspired by the statistical-mechanics perspective on ensembles of random codes [7]. Over the years, research was aimed towards various extensions. For example, in [8, Chapter 10] [9], [10], [11] universal versions of erasure/list decoders, i.e., decoders that perform well even under channel uncertainty. In [12], the size of the list was gauged by its ρ -th moment, where $\rho = 1$ corresponds to the average list size, as in [1], and an achievable pair of error exponent and list size exponents was found.

In essence, as we shall see, Forney's optimal erasure/list decoder is conceptually simple - a codeword belongs to the output list if its a posteriori probability, conditioned on the output vector, is larger than a *fixed* threshold. However, as is well known, the prospective implementation of such a decoder is notoriously difficult, since it requires the calculation of the sum of likelihoods (c.f. Subsection II-B). These likelihoods typically decrease exponentially as the block length increases, and thus have a large dynamic range. Therefore, an exact computation of this sum is usually not feasible. Moreover, the number of summands is equal to the number of codewords and thus increases exponentially with the block length. In fact, it was proposed already by Forney himself to approximate this sum by its maximal element [1, eq. (11a)].

It is the purpose of this paper to study classes of simplified decoders which avoid the need to confront such numerical difficulties. To this end, we first formulate the optimal erasure/list decoder as a threshold decoder, and propose to replace the optimal threshold with a threshold from a fairly general class. Any threshold from this class may represent some complexity constraint of the decoder. We then consider two subclasses of this class of decoders. For the first subclass, which is well suited for high rates, the threshold function is extremely simple - it depends on the output vector only. For the second subclass, which is well suited for low rates, the threshold is the log-likelihood (plus an additive fixed term), which is the approximation proposed by Forney. Then, for a given discrete memoryless channel (DMC), we derive the exact random coding exponents for the ensemble of fixed composition codebooks, associated with decoders from the above general class and its two subclasses. Next, for

the general class of decoders, and for its two subclasses, we find the optimum decoder, within the given class, which achieves the largest list size exponent for a prescribed error exponent. This enables the establishment of the optimality of the simplified decoders of the first and second kind, for low and high rates, respectively.

The outline of the rest of the paper is as follows. In Section II, we establish notation conventions, provide some background of known results, and define the relevant classes of decoders. In Section III, we derive the exact random coding exponents for all defined decoders, and in Section IV, we find the optimal decoder within each class. In Section V, we discuss the optimality of the various simplified decoders defined, for low and high rates. Finally, in Section VI, we compare the random coding exponents of the optimal decoder, as obtained in [6], with the exponents of the suboptimal decoders, for simple numerical examples. All the proofs are deferred to Appendix A.

II. PROBLEM FORMULATION

A. Notation Conventions

Throughout the paper, random variables will be denoted by capital letters, specific values they may take will be denoted by the corresponding lower case letters, and their alphabets, similarly as other sets, will be denoted by calligraphic letters. Random vectors and their realizations will be denoted, respectively, by capital letters and the corresponding lower case letters, both in the bold face font. Their alphabets will be superscripted by their dimensions. For example, the random vector $\mathbf{X} = (X_1, \dots, X_n)$, (n - positive integer) may take a specific vector value $\mathbf{x} = (x_1, \dots, x_n)$ in $\mathcal{X}^n$, the n -th order Cartesian power of $\mathcal{X}$, which is the alphabet of each component of this vector.

A joint distribution of a pair of random variables (X, Y) on $\mathcal{X} \times \mathcal{Y}$, the Cartesian product alphabet of $\mathcal{X}$ and $\mathcal{Y}$, will be denoted by Q_{XY} and similar forms, e.g. $\tilde{Q}_{XY}$. Usually, we will abbreviate this notation by omitting the subscript XY , and denote, e.g. Q_{XY} by Q . The X -marginal (Y -marginal), induced by Q will be denoted by Q_X (respectively, Q_Y), and the conditional distributions will be denoted by $Q_{Y|X}$ and $Q_{X|Y}$. In accordance with this notation, the joint distribution induced by Q_X and $Q_{Y|X}$ will be denoted by $Q_X \times Q_{Y|X}$, i.e. $Q = Q_X \times Q_{Y|X}$.

For a given vector $\mathbf{x}$, let $\hat{Q}_{\mathbf{x}}$ denote the empirical distribution, that is, the vector $\{\hat{Q}_{\mathbf{x}}(x), x \in \mathcal{X}\}$, where $\hat{Q}_{\mathbf{x}}(x)$ is the relative frequency of the letter x in the vector $\mathbf{x}$. Let $\mathcal{T}_P$ denote the type class associated with P , that is, the set of all sequences $\{\mathbf{x}\}$ for which $\hat{Q}_{\mathbf{x}} = P$. Similarly, for a pair of vectors $(\mathbf{x}, \mathbf{y})$, the empirical joint distribution will be denoted by $\hat{Q}_{\mathbf{xy}}$.

The mutual information of a joint distribution Q will be denoted by $I(Q)$, where Q may also be an empirical joint distribution. The information divergence between Q and P will be denoted by $D(Q||P)$, and the conditional information divergence between the empirical conditional distribution $Q_{Y|X}$ and $P_{Y|X}$, averaged over Q_X , will be denoted by $D(Q_{Y|X}||P_{Y|X}|Q_X)$. Here too, the distributions may be empirical.

The probability of an event $\mathcal{A}$ will be denoted by $\mathbb{P}\{\mathcal{A}\}$, and the expectation operator will be denoted by $\mathbb{E}\{\cdot\}$. Whenever there is room for ambiguity, the underlying probability distribution Q will appear as a subscript, e.g., $\mathbb{E}_Q\{\cdot\}$. The indicator function will be denoted by $\mathbb{I}\{\cdot\}$. Sets will normally be denoted by calligraphic letters. The

probability simplex over an alphabet $\mathcal{X}$ will be denoted by $\mathcal{S}(\mathcal{X})$. The complement of a set $\mathcal{A}$ will be denoted by $\overline{\mathcal{A}}$. Logarithms and exponents will be understood to be taken to the natural base. The notation $[t]_+$ will stand for $\max\{t, 0\}$. For two positive sequences, $\{a_n\}$ and $\{b_n\}$, the notation $a_n \doteq b_n$ will mean asymptotic equivalence in the exponential scale, that is, $\lim_{n \rightarrow \infty} \frac{1}{n} \log(\frac{a_n}{b_n}) = 0$.

B. System Model and Background

Consider a DMC, characterized by a finite input alphabet $\mathcal{X}$, a finite output alphabet $\mathcal{Y}$, and a given matrix of single-letter transition probabilities $\{W(y|x), x \in \mathcal{X}, y \in \mathcal{Y}\}$. Conditioning on the channel input $\mathbf{X} = \mathbf{x}$, the distribution of the output $\mathbf{Y}$ is given by

$$\mathbb{P}(\mathbf{Y} = \mathbf{y} | \mathbf{X} = \mathbf{x}) = \prod_{i=1}^n W(y_i | x_i). \quad (1)$$

We will use the shorthand notations $P(\mathbf{x}, \mathbf{y}) \triangleq \mathbb{P}(\mathbf{X} = \mathbf{x}, \mathbf{Y} = \mathbf{y})$, $P(\mathbf{y} | \mathbf{x}) \triangleq \mathbb{P}(\mathbf{Y} = \mathbf{y} | \mathbf{X} = \mathbf{x})$, and $P(\mathbf{x} | \mathbf{y}) \triangleq \mathbb{P}(\mathbf{X} = \mathbf{x} | \mathbf{Y} = \mathbf{y})$. Let R be the coding rate in nats per channel use, and let $\mathcal{C} = \{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_M\}$, $\mathbf{x}_m \in \mathcal{X}^n$, $m = 1, \dots, M$, $M = e^{nR}$ denote the codebook.

An erasure/list decoder ϕ is a mapping from the space of output vectors $\mathcal{Y}^n$ to the set of subsets of $\{1, \dots, M\}$ (i.e., $\phi : \mathcal{Y}^n \rightarrow \mathcal{P}(\{1, \dots, M\})$, where the latter is the power set of $\{1, \dots, M\}$). Alternatively, an erasure/list decoder ϕ is uniquely defined by a set of $M + 1$ decoding regions $\{\mathcal{R}_m\}_{m=0}^M$ such that $\mathcal{R}_m \subseteq \mathcal{Y}^n$ and $\mathcal{R}_0 = \mathcal{Y}^n \setminus \bigcup_{m=1}^M \mathcal{R}_m$. Given an output vector $\mathbf{y}$, the m th codeword belongs to the list if $\mathbf{y} \in \mathcal{R}_m$, and if $\mathbf{y} \in \mathcal{R}_0$ an erasure is declared.

The average error probability of a list decoder ϕ and a codebook $\mathcal{C}$ is the probability that the actual transmitted codeword does not belong to the list,

$$P_e(\mathcal{C}, \phi) \triangleq \frac{1}{M} \sum_{m=1}^M \sum_{\mathbf{y} \in \mathcal{R}_m} P(\mathbf{y} | \mathbf{x}_m). \quad (2)$$

The average number of erroneous codewords on the list is defined as

$$L(\mathcal{C}, \phi) \triangleq \sum_{\mathbf{y} \in \mathcal{Y}^n} \frac{1}{M} \sum_{m=1}^M P(\mathbf{y} | \mathbf{x}_m) \sum_{l \neq m} \mathbb{I}(\mathbf{y} \in \mathcal{R}_l) \quad (3)$$

$$= \frac{1}{M} \sum_{l=1}^M \sum_{\mathbf{y} \in \mathcal{R}_l} \sum_{m \neq l} P(\mathbf{y} | \mathbf{x}_m). \quad (4)$$

If the decoder never outputs more than one codeword, then it is termed a *pure erasure* decoder. In this case, $P_e(\mathcal{C}, \phi)$ designates the *total error probability* (which is the sum of the *erasure probability* and *undetected error probability*), and $L(\mathcal{C}, \phi)$ designates the *undetected error probability*. In what follows, we shall describe our results with the more general erasure/list terms, but unless otherwise stated, the results are also valid for pure erasure decoders, with the above interpretation.

In [1], Forney has generalized the Neyman-Pearson lemma, and obtained a class of decoders $\Phi^* \triangleq \{\phi_T, T \in \mathbb{R}\}$ which achieves the optimal trade-off between $P_e(\mathcal{C}, \phi)$ and $L(\mathcal{C}, \phi)$. The decoding regions for a decoder $\phi_T \in \Phi^*$ are given by

$$\mathcal{R}_m^* \triangleq \left\{ \mathbf{y} : \frac{P(\mathbf{y}|\mathbf{x}_m)}{\sum_{l \neq m} P(\mathbf{y}|\mathbf{x}_l)} \geq e^{nT} \right\}, \quad 1 \leq m \leq M. \quad (5)$$

The parameter T is called the *threshold*, and it controls the trade-off between $P_e(\mathcal{C}, \phi)$ and $L(\mathcal{C}, \phi)$. As T increases, the list size typically becomes smaller, but in exchange, the error probability increases. To obtain a pure erasure decoder, the threshold T should be chosen non-negative. For $T < 0$, an erasure/list decoder is obtained.

As a figure of merit for a decoder, we will consider its random coding exponents, over the ensemble of fixed composition codes with input distribution P_X . For this ensemble, the $M = e^{nR}$ codewords of $\mathcal{C} = \{\mathbf{x}_1, \dots, \mathbf{x}_M\}$ are selected independently at random under the uniform distribution across the type class $\mathcal{T}_{P_X}$. We denote the averaging over this ensemble of random codebooks by overlines. The random coding error exponent is defined as

$$E_e(R, \phi) \triangleq \liminf_{n \rightarrow \infty} -\frac{1}{n} \log \overline{P_e(\mathcal{C}, \phi)} \quad (6)$$

and the random coding list size (negative) exponent is defined as

$$E_l(R, \phi) \triangleq \liminf_{n \rightarrow \infty} -\frac{1}{n} \log \overline{L(\mathcal{C}, \phi)}. \quad (7)$$

Notice that $E_e(R, \phi) \geq 0$ but $E_l(R, \phi)$ may also be negative, which means that the random coding list size may *increase* exponentially.

In [6], the exact random coding exponents for a decoder $\phi_T^* \in \Phi^*$, which we denote by $E_e^*(R, T)$ and $E_l^*(R, T)$, were found¹. Defining

$$E_a(R, T) \triangleq \min_{\tilde{Q}} \min_{Q: Q_Y = \tilde{Q}_Y, f(Q) + T \geq f(\tilde{Q}), I(Q) \geq R} \left\{ D(\tilde{Q}_{Y|X} \| W | P_X) + I(Q) \right\} - R, \quad (8)$$

and

$$E_b(R, T) \triangleq \min_{\tilde{Q} \in \mathcal{L}} D(\tilde{Q}_{Y|X} \| W | P_X), \quad (9)$$

where

$$\mathcal{L} \triangleq \left\{ \tilde{Q} : f(\tilde{Q}) \leq R + T + \max_{Q: Q_Y = \tilde{Q}_Y, I(Q) \leq R} [f(Q) - I(Q)] \right\}, \quad (10)$$

we have that

$$E_e^*(R, T) = \min\{E_a(R, T), E_b(R, T)\} \quad (11)$$

and [11, Lemma 1]

$$E_l^*(R, T) = E_e^*(R, T) + T. \quad (12)$$

¹See Appendix B for a justification of the applicability of the results in [6] also for the case $T < 0$.

In the rest of the paper, we will consider simplified erasure/list decoding which determine if a codeword belongs to the list, by checking if the likelihood of the codeword is larger than some threshold function. To motivate this approach, we first inspect the optimal decoder (5) in more detail. In its current form, the optimal decoder is essentially infeasible to implement. The reason for this is that the computation of $\sum_{l \neq m} P(\mathbf{y}|\mathbf{x}_l)$ is usually intractable, as it is a sum of exponentially many likelihood terms, where each likelihood term is also exponentially small. This is in sharp contrast to ordinary decoders, based on comparison of single likelihood terms which can be carried out in the logarithmic scale, rendering them numerically feasible. However, the optimal decoder can be simplified without asymptotic loss in exponents. To see this, we need first some notation: for a given joint type Q and output vector $\mathbf{y}$, let $N_m(Q|\mathbf{y})$ be the number of codewords other than $\mathbf{x}_m$ in $\mathcal{C}$, whose joint empirical distribution with $\mathbf{y}$ is Q . Now, define the decision regions

$$\tilde{\mathcal{R}}_m \triangleq \left\{ \mathbf{y} : e^{nf(\hat{Q}_{\mathbf{x}_m \mathbf{y}})} \geq e^{nT} \max_Q N_m(Q|\mathbf{y}) \cdot e^{nf(Q)} \right\}, \quad 1 \leq m \leq M, \quad (13)$$

where $f(Q)$ is the normalized log-likelihood ratio, namely

$$f(Q) \triangleq \sum_{x \in \mathcal{X}, y \in \mathcal{Y}} Q(x, y) \log W(y|x). \quad (14)$$

We have the following observation.

Proposition 1. *The random coding error- and list size exponents of the decoder defined by $\{\tilde{\mathcal{R}}_m\}_{m=1}^M$ are the same as the random coding error- and list size exponents of the optimal decoder (defined by $\{\mathcal{R}_m^*\}_{m=1}^M$).*

We will refer to the right-hand side of (13) as the *threshold* of the decoder. The focus of this paper, is the performance of erasure/list decoders for which the threshold function is not necessarily the optimal threshold dictated by (13). The motivation for such an erasure/list decoder is that the threshold function employed may be simpler to compute. Observe, that the threshold of (13) only depends on the joint type of each competing codeword ($l \neq m$) with the output vector. Therefore, we propose to consider the class of decoders $\Psi \triangleq \{\phi_h, h : \mathcal{S}(\mathcal{X} \times \mathcal{Y}) \rightarrow \mathbb{R}\}$, defined by a continuous function $h(\cdot)$ in some compact domain $\mathcal{G} \subseteq \mathcal{S}(\mathcal{X} \times \mathcal{Y})$, and infinite elsewhere². The decoding regions of $\phi_h \in \Psi$ are given by

$$\mathcal{R}_m = \left\{ \mathbf{y} : e^{nf(\hat{Q}_{\mathbf{x}_m \mathbf{y}})} \geq e^{n \cdot \max_{l \neq m} h(\hat{Q}_{\mathbf{x}_l \mathbf{y}})} \right\}, \quad 1 \leq m \leq M. \quad (15)$$

As discussed above, for this decoding rule, computations may be carried in the logarithmic domain as

$$\mathcal{R}_m = \left\{ \mathbf{y} : f(\hat{Q}_{\mathbf{x}_m \mathbf{y}}) \geq \max_{l \neq m} h(\hat{Q}_{\mathbf{x}_l \mathbf{y}}) \right\}. \quad (16)$$

With a slight abuse of notation, we shall denote the random coding exponents of $\phi_h \in \Psi$ by $E_e(R, h)$ and $E_l(R, h)$. Note that in (15), the threshold function is assumed to be the same for all codebooks in the ensemble (for specific

²The set $\mathcal{G}$ will be a strict subset of $\mathcal{S}(\mathcal{X} \times \mathcal{Y})$ for optimal threshold functions, see Section IV.

rate R). In contrast, for the asymptotically optimal decoder (13), the threshold is tuned to the specific codebook employed. This modification reduces complexity (though only slightly) and may, in general, degrade the exponents. Nonetheless, for types which satisfy $I(Q) \leq R$, it is known that when the m th codeword is sent, $N_m(Q|\mathbf{y})$ concentrates double-exponentially fast around its asymptotic expected value of $e^{n(R-I(Q))}$ [7, Section 6.3]. This hints to the possibility that no loss in exponents is incurred by the restriction of a single threshold function used for all codebooks in the ensemble. Consequently, for a properly chosen $h(\cdot)$, it is possible that ϕ_h achieves the optimal exponents.

Next, we propose two subclasses of Ψ . The first subclass is $\Lambda_1 \triangleq \{\phi_g, g : \mathcal{S}(\mathcal{Y}) \rightarrow \mathbb{R}\}$ where $g(\cdot)$ is a continuous function in some compact domain $\mathcal{G} \subseteq \mathcal{S}(\mathcal{Y})$, and infinite elsewhere. A decoder $\phi_g \in \Lambda_1$ has the following decoding regions

$$\mathcal{R}_m = \left\{ \mathbf{y} : f(\hat{Q}_{\mathbf{x}_m \mathbf{y}}) \geq g(\hat{Q}_{\mathbf{y}}) \right\}, \quad 1 \leq m \leq M. \quad (17)$$

With a slight abuse of notation, we shall denote the random coding exponents of $\phi_g \in \Lambda_1$ by $E_e(R, g)$ and $E_l(R, g)$. Here, the *threshold* function $g(Q_Y)$ does not depend on the codebook, but it may depend on R . In essence, a decoder in Λ_1 approximates $e^{nT} N_m(Q|\mathbf{y}) \cdot e^{nf(Q)}$ by a function that depends on $\mathbf{y}$ only, but not on $\mathcal{C}$. The second subclass is $\Lambda_2 \triangleq \{\phi_T, T \in \mathbb{R}\}$, where T is a threshold parameter. A decoder $\phi_T \in \Lambda_2$ has the following decoding regions,

$$\mathcal{R}_m = \left\{ \mathbf{y} : f(\hat{Q}_{\mathbf{x}_m \mathbf{y}}) \geq T + \max_{l \neq m} f(\hat{Q}_{\mathbf{x}_l \mathbf{y}}) \right\}, \quad 1 \leq m \leq M. \quad (18)$$

With a slight abuse of notation, we shall denote the random coding exponents of $\phi_T \in \Lambda_2$ by $E_e(R, T)$ and $E_l(R, T)$. Observe that for $T < 0$, the list size of this decoder is at least 1, since the codeword with maximum likelihood is always on the list. In essence, this decoder approximates

$$N_m(Q|\mathbf{y}) \cdot e^{nf(Q)} \approx \max_{N_m(Q|\mathbf{y}) \geq 1} N_m(Q|\mathbf{y}) \cdot e^{nf(Q)} \approx \max_{l \neq m} P(\mathbf{y}|\mathbf{x}_l), \quad (19)$$

and so, the threshold, in this case, is the second largest likelihood. This approximation was proposed by Forney [1, eq. (11a)], but was not analyzed rigorously before.

In this paper, we provide *exact* single-letter expressions for the error- and list size exponents for the class Ψ , and obtain, as corollaries, the exponents of the previously defined subclasses Λ_1 and Λ_2 . Then, for each of the classes of decoders, we will derive the *optimal* threshold function, in the sense that for a given requirement on the value of $E_e(R, \phi)$, they provide the largest $E_l(R, \phi)$ within the given class. Finally, we discuss the regimes for which the simplified decoders are close to be asymptotically optimal. For the defined ensemble of random codes, the subclasses Λ_1 and Λ_2 represent two extremes of approximation to the threshold. For low rates, and a typical codebook in the ensemble, the threshold will be dominated by a single codeword,³ which has the maximum likelihood, besides the candidate codeword⁴, and the random coding exponents of $\phi_T \in \Lambda_2$ will be optimal. On the

³At least, a sub-exponential number of codewords.

⁴This could be either the codeword which was actually sent, or an erroneous codeword.

other hand, for high rates, the threshold will tend to concentrate around a deterministic function $g(\hat{Q}_Y)$. Using the function $g(\cdot)$ as a threshold function will be accurate for a typical codebook, and so, in this case too, the random coding exponents of $\phi_g \in \Lambda_1$ will tend to the optimal exponents, for high rates.

Before we continue, we emphasize that our decoders do not assume any specific structure of the codebook (as we have assumed the fixed composition random coding ensemble), and so the simplified decoders do not immediately lead to practical implementation. Nonetheless, as common, the random coding analysis serves as an achievable benchmark for possible exponents.

III. EXPONENTS OF THRESHOLD DECODERS

In this section, we derive the error- and list size exponents for the class Ψ , and then obtain, as special cases, the exponents for the subclasses Λ_1 and Λ_2 . As we have assumed an ensemble of fixed composition codebooks with input distribution P_X , a joint type Q of $(\mathbf{x}_m, \mathbf{y})$ will always have the form $P_X \times Q_{Y|X}$. For brevity, this is not explicitly mentioned henceforth. We begin with our main theorem, which provides the random coding exponents for $\phi_h \in \Psi$. We use the following definitions. For a given $h(\cdot)$, let

$$\mathbf{V}(Q_Y, R) \triangleq \max_{Q': Q'_Y = Q_Y, I(Q') \leq R} h(Q'), \quad (20)$$

and

$$\mathcal{D}(\tilde{Q}) \triangleq \left\{ Q : Q = \tilde{Q}_Y, f(Q) \geq \max \left\{ \mathbf{V}(Q_Y, R), h(\tilde{Q}) \right\} \right\}. \quad (21)$$

Theorem 2. *For a decoder $\phi_h \in \Psi$, the random coding error exponent, with respect to (w.r.t) the ensemble of fixed composition codebooks P_X , is given by*

$$E_e(R, h) = \min_{\tilde{Q}} \min_{Q: Q_Y = \tilde{Q}_Y, h(Q) \geq f(\tilde{Q})} \left\{ D(\tilde{Q}_{Y|X} \| W | P_X) + [I(Q) - R]_+ \right\}, \quad (22)$$

and the list size exponent is given by

$$E_l(R, h) = \min_{\tilde{Q}} \min_{Q \in \mathcal{D}(\tilde{Q})} \left\{ D(\tilde{Q}_{Y|X} \| W | P_X) + I(Q) - R \right\}. \quad (23)$$

Next, we provide the exponents for the subclasses Λ_1 and Λ_2 . These exponents are obtained using the general result of Theorem 2, with appropriate substitutions and manipulations.

Corollary 3. *For a decoder $\phi_g \in \Lambda_1$, the random coding error exponent, w.r.t. the ensemble of fixed composition codebooks P_X , is given by*

$$E_e(R, g) = \min_{Q: g(Q_Y) \geq f(Q)} D(Q_{Y|X} \| W | P_X), \quad (24)$$

and the list size exponent is given by

$$E_l(R, g) = \min_{\tilde{Q}} \min_{Q: Q_Y = \tilde{Q}_Y, g(Q_Y) \leq f(Q)} \left\{ D(\tilde{Q}_{Y|X} \| W | P_X) + I(Q) - R \right\}. \quad (25)$$

Note that, for a given $g(\cdot)$, the error exponent $E_e(R, g)$ does not depend on R , because the decision whether to include a codeword in the list is based only the codeword sent and the output vector received, not on other codewords. Next, for $T < 0$, we define

$$\underline{\mathbf{V}}(Q_Y, R) \triangleq \max_{Q': Q'_Y = Q_Y, I(Q') \leq R} f(Q'), \quad (26)$$

and

$$\underline{\mathcal{D}}(\tilde{Q}, R, T) \triangleq \left\{ Q : Q_Y = \tilde{Q}_Y, f(Q) \geq T + \max \left\{ \underline{\mathbf{V}}(Q_Y, R), f(\tilde{Q}) \right\} \right\}. \quad (27)$$

Corollary 4. For a decoder $\phi_T \in \Lambda_2$, the random coding error exponent, w.r.t. the ensemble of fixed composition codebooks P_X , is given by

$$E_e(R, T) = \min_{\tilde{Q}} \min_{Q: Q_Y = \tilde{Q}_Y, f(Q) + T \geq f(\tilde{Q})} \left\{ D(\tilde{Q}_{Y|X} \| W | P_X) + [I(Q) - R]_+ \right\}, \quad (28)$$

and the list size exponent is given by

$$E_l(R, T) = \min_{\tilde{Q}} \min_{Q \in \underline{\mathcal{D}}(\tilde{Q}, R, T)} \left\{ D(\tilde{Q}_{Y|X} \| W | P_X) + I(Q) - R \right\}. \quad (29)$$

Remark 5. Assume that for some channel $V \neq W$, $f(Q)$ is replaced by

$$\tilde{f}(Q) \triangleq \sum_{x \in \mathcal{X}, y \in \mathcal{Y}} Q(x, y) \log V(y|x) \quad (30)$$

in the optimal exponents (8)–(12), as well as in Theorem 2 and its corollaries 3 and 4. This yields random coding exponents associated with *mismatched decoding*.

IV. OPTIMAL THRESHOLD FUNCTIONS

In the previous section, we have derived the exact exponents for decoders of the classes Ψ , Λ_1 and Λ_2 , where these exponents depend on the actual decoder in the class via the threshold functions $h(\cdot)$, $g(\cdot)$, and the T , respectively. In some scenarios, there is a freedom to choose any decoder within a given class. Clearly, just as for the optimal class Φ^* , a trade-off exists between the error- and list size exponents, which is controlled by the choice of threshold function or parameter. In this section, we assume a given rate R , and a target error exponent E . Under this requirement, we find the threshold function or parameter for a given class of decoders, which yields the maximal list size exponent. In addition, we find expressions for the resulting maximal list size exponent. Obviously, the resulting list size exponent cannot exceed $E_l^*(R, T^*)$, where T^* satisfies $E_e^*(R, T^*) = E$, and the difference between these two exponents is a measure for the sub-optimality of the *class* of decoders. We first define optimal threshold functions.

Definition 6. A threshold function $h^*(Q, R, E)$ is said to be *optimal* for the class Ψ , if the corresponding decoder $\phi_{h^*(Q, R, E)} \in \Psi$ achieves $E_e(R, h) \geq E$ and for any other decoder $\phi_h \in \Psi$ with $E_e(R, h) \geq E$, we have

$$E_l(R, h) \leq E_l(R, h^*).$$

The optimal threshold function $g^*(Q_Y, E)$ for the class Λ_1 , and the optimal parameter $T^*(R, E)$ for the class Λ_2 , are defined analogously. The resulting list size exponent for $h^*(Q, R, E)$ will be denoted by $E_l^*(\Psi, R, E)$, and $E_l^*(\Lambda_1, R, E)$, $E_l^*(\Lambda_2, R, E)$, will denote the analogue exponents for $g^*(Q_Y, E)$ and $T^*(R, E)$. In this section, we will find it more convenient to begin with the simple class Λ_1 , as $h^*(Q, R, E)$ is conveniently represented by $g^*(Q_Y, E)$. We then conclude with $T^*(R, E)$ for the class Λ_2 , which is conveniently represented by $h^*(Q, R, E)$.

Theorem 7. *The optimal threshold function $g^*(Q_Y, E)$ for the class Λ_1 is*

$$g^*(Q_Y, E) = \min_{Q': Q'_Y = Q_Y, D(Q'_Y|X||W|P_X) \leq E} f(Q'), \quad (31)$$

and the resulting list size exponent is

$$E_l^*(\Lambda_1, R, E) = \min_{\tilde{Q}} \min_{Q: Q_Y = \tilde{Q}_Y, g^*(Q_Y, E) \leq f(Q)} \left\{ D(\tilde{Q}_Y|X||W|P_X) + I(Q) - R \right\}. \quad (32)$$

Note that for $E < 0$, the minimization problem in (31) is infeasible and so $g^*(Q_Y, E) = \infty$. While $g^*(Q_Y, E)$ does not depend directly on the rate R , the required error exponent E will usually be chosen as a function of the rate R , and thus $g^*(Q_Y, E)$ will depend indirectly on R . The next lemma states a few simple properties of $g^*(Q_Y, E)$.

Lemma 8. *The optimal threshold function $g^*(Q_Y, E)$ has the following properties:*

- 1) *It is a non-increasing function of E .*
- 2) *If $D(Q_Y|X||W|P_X) \leq E$ then $f(Q) \geq g^*(Q_Y, E)$.*
- 3) *It is a convex function of Q_Y .*

Next, we provide the optimal list size exponent for the class Ψ . We define

$$\mathcal{D}^*(\tilde{Q}, R, E) = \left\{ Q : Q_Y = \tilde{Q}_Y, f(Q) > g^*(Q_Y, E - [I(\tilde{Q}) - R]_+) \right\}. \quad (33)$$

Theorem 9. *The optimal threshold function $h^*(Q, R, E)$ for the class Ψ is*

$$h^*(Q, R, E) = g^*(Q_Y, E - [I(Q) - R]_+) \quad (34)$$

and the resulting list size exponent is

$$E_l^*(\Psi, R, E) = \min_{\tilde{Q}} \min_{Q \in \mathcal{D}^*(\tilde{Q}, R, E)} \left\{ D(\tilde{Q}_Y|X||W|P_X) + I(Q) - R \right\}. \quad (35)$$

It is easily shown, using Lemma 8 (property 2), that in the domain where $h^*(Q, R, E) < \infty$, it is a continuous function of the joint type Q .

Remark 10. Following Remark 5, consider the scenario for which the channel W is not known exactly, but only known to belong to a given class of channels $\mathcal{W}$. For example, suppose we are given a nominal channel V , and it

is known that W is not far from V , where the distance is measured in the $\mathcal{L}_1$ norm. For a moment, let us denote the optimal threshold, as $g_W^*(Q_Y, E)$, i.e. with explicit dependency in the channel W . Then, choosing a threshold function

$$g_W^*(Q_Y, E) \triangleq \min_{W \in \mathcal{W}} g_W^*(Q_Y, E) \quad (36)$$

guarantees that $E_e(R, g_W^*) \geq E$ uniformly over $W \in \mathcal{W}$. The resulting list size exponent

$$E_l(R, g) = \min_{\tilde{Q}} \min_{Q: Q_Y = \tilde{Q}_Y, \min_{W \in \mathcal{W}} g_W^*(Q_Y, E) \leq f(Q)} \left\{ D(\tilde{Q}_{Y|X} \| W | P_X) + I(Q) - R \right\} \quad (37)$$

$$= \min_{W \in \mathcal{W}} \min_{\tilde{Q}} \min_{Q: Q_Y = \tilde{Q}_Y, g_W^*(Q_Y, E) \leq f(Q)} \left\{ D(\tilde{Q}_{Y|X} \| W | P_X) + I(Q) - R \right\} \quad (38)$$

$$= \min_{W \in \mathcal{W}} E_l^*(W, \Lambda_1, R, E) \quad (39)$$

where here, $E_l^*(W, \Lambda_1, R, E)$ denotes the optimal list size exponent of Theorem 7, with explicit dependency in W . Analogue results hold also for the class Ψ . See [9], [10], [11] for related ideas.

We conclude this section with the optimal T for the class Λ_2 .

Theorem 11. *The optimal threshold parameter $T^*(R, E)$ for the class Λ_2 is*

$$T^*(R, E) = \min_Q \{ h^*(Q, R, E) - f(Q) \} \quad (40)$$

and the resulting list size exponent is

$$E_l^*(\Lambda_2, R, E) = \min_{\tilde{Q}} \min_{Q \in \underline{\mathcal{D}}(\tilde{Q}, R, T^*)} \left\{ D(\tilde{Q}_{Y|X} \| W | P_X) + I(Q) - R \right\}. \quad (41)$$

V. OPTIMALITY OF SIMPLIFIED DECODERS FOR LOW AND HIGH RATES

In this section, we discuss the optimality of the two subclasses Λ_1 and Λ_2 in certain regimes of the rate. For Λ_1 , we show that for rates *above* some critical rate, the optimal decoder from the class Λ_1 has the same exponents as the optimal decoder from the class Ψ . In turn, as we have discussed in Subsection II-B, the optimal decoder from the class Ψ has exponents which are close, and sometimes even equal, to the optimal exponents of $\phi_T^* \in \Phi^*$. For Λ_2 , we show that assuming $T > 0$, for rates *below* some critical rate, the exponents of the decoder $\phi_T \in \Lambda_2$ are the same exponents as ϕ_T^* (for the same value of T). From continuity arguments, approximate optimality will be obtained for $T < 0$ with small $|T|$. Moreover, for $T < 0$, or for rates above the critical rate, the exponents of the decoder from Λ_2 will improve if we choose the optimal T^* , according to Theorem 11.

We begin with Λ_1 . Let the optimal solution of $E_l^*(\Lambda_1, R, E)$ be $(\tilde{Q}^*, Q^*)$, and let $\bar{R}_{\text{cr}}(E) \triangleq I(\tilde{Q}^*)$.

Proposition 12. *For $R \geq \bar{R}_{\text{cr}}(E)$*

$$E_l^*(\Lambda_1, R, E) = E_l^*(\Psi, R, E). \quad (42)$$

Namely, a threshold which only depends on the output vector is sufficient to obtain the best exponents of the class

Ψ .

Next, we demonstrate the optimality of the class Λ_2 in case that R is not too large, and $T \geq 0$. Consider a decoder $\phi_T \in \Lambda_2$, and an optimal decoder $\phi_T^* \in \Phi^*$, with the same parameter $T \geq 0$. Let now $(\tilde{Q}^*, Q^*)$ be the optimal solution of

$$\min_{\tilde{Q}} \min_{Q: Q_Y = \tilde{Q}_Y, f(Q) + T \geq f(\tilde{Q})} \left\{ D(\tilde{Q}_{Y|X} \| W | P_X) + I(Q) \right\}, \quad (43)$$

and let $\underline{R}_{\text{cr}}(T) \triangleq I(\tilde{Q}^*)$.

Proposition 13. *For $T \geq 0$ and $R \leq \underline{R}_{\text{cr}}(T)$ both*

$$E_e(R, T) = E_e^*(R, T), \quad (44)$$

and

$$E_l(R, T) = E_l^*(R, T). \quad (45)$$

Namely, $\phi_T \in \Lambda_2$ and the optimal $\phi_T^* \in \Phi^*$ have the same exponents.

VI. NUMERICAL EXAMPLES

We demonstrate the results for binary channels ($|\mathcal{X}| = |\mathcal{Y}| = 2$), assuming a threshold $T = \pm 0.05$. First, for any rate $0 \leq R \leq I(P_X \times W)$, the optimal exponents $E_e^*(R, T)$ and $E_l^*(R, T)$ were computed using (11), and (12). Second, for any rate $0 \leq R \leq I(P_X \times W)$, the target error exponent was set to $E = E_e^*(R, T)$, and the maximal list size exponent was found for the class Ψ , as well as its subclasses Λ_1 and Λ_2 .

For $T = 0.05$, Figure 1 shows the random coding exponents for a binary symmetric channel $W_1(0|1) = W_1(1|0) = 0.01$. In this example, the optimal decoder from Λ_2 has the same exponents as ϕ_T^* for the entire range of rates. It is interesting to note that for the optimal decoder from Λ_1 , the maximal list size is not a monotonic decreasing function of the rate (but naturally always smaller than the optimal list size exponent, for the given error exponent). This is due to the two contradicting effects of the rate on the optimal list size exponent of the class Λ_1 : As the rate increases, the class Λ_1 has improved performance on one hand, but the error exponent requirement is decreasing (as we have assumed a fixed T) on the other hand.

Next, for $T = -0.05$, Figure 2, shows the exponents for a binary asymmetric channel $W_2(0|1) = 0.01$, $W_2(1|0) = 0.4$. For high rates, the optimal decoder from Λ_1 approaches optimal performance, but is rather poor for low rates. For these low, and even intermediate, rates, the optimal decoder from Λ_2 has performance close to optimal. It was also found empirically, that the rate for which the maximum likelihood threshold decoder moves away from optimal performance is larger as the channel is more symmetric. Thus, the performance of the simplified decoders for $T = -0.05$ is even better for the previously defined symmetric channel.

In both examples, it can be observed that for the *entire* range of rates, the optimal decoder from Ψ has essentially the same performance as the optimal decoder from Φ^* . Therefore, as was mentioned in Subsection II-B, using a

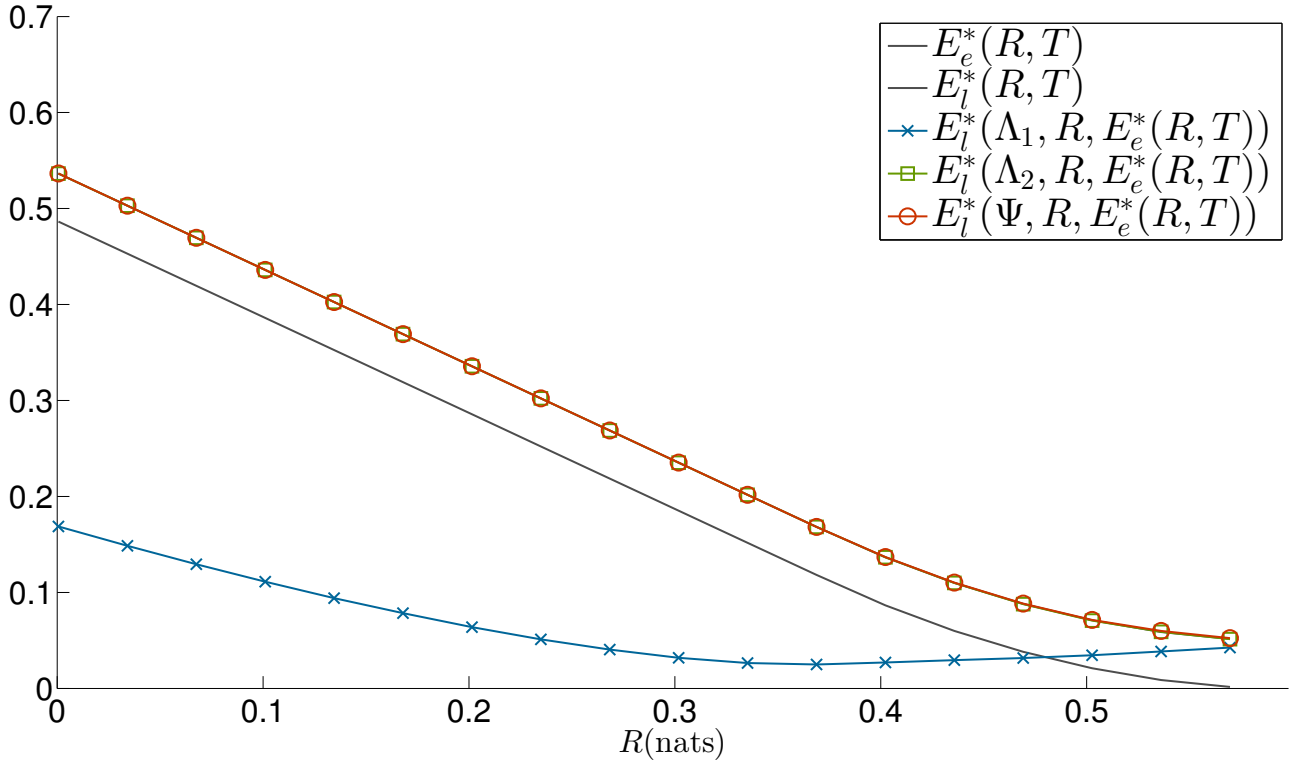


Figure 1. Graphs of random coding exponents for W_1 and $T = 0.05$.

single, properly-optimized, threshold function for all the codebooks in the ensemble may not incur a loss in exponents at all. Finally, we remark that similar behavior was observed for a wide range of T .

APPENDIX A

Proof of Proposition 1: The continuity of the error- and list size exponents in T , evident from (11) (12), implies that if e^{nT} in (5) is replaced by any sequence $\tau_n \doteq e^{nT}$, then the same exponents are achieved. Also, as the number of possible joint types is polynomial in n [8, Chapter 2], we have

$$K_{m,n} \triangleq \frac{\max_Q N_m(Q|\mathbf{y}) \cdot e^{nf(Q)}}{\sum_Q N_m(Q|\mathbf{y}) \cdot e^{nf(Q)}} \doteq 1 \quad (\text{A.1})$$

for all $1 \leq m \leq M$. Thus, $\mathcal{R}_m^*$ and the decoding sets

$$\left\{ \mathbf{y} : P(\mathbf{y}|\mathbf{x}_m) \geq e^{nT} \cdot K_{m,n} \cdot \sum_{m=1}^M P(\mathbf{y}|\mathbf{x}_m) \right\} = \quad (\text{A.2})$$

$$\left\{ \mathbf{y} : e^{nf(\hat{Q}_{\mathbf{x}_m \mathbf{y}})} \geq e^{nT} \cdot K_{m,n} \cdot \sum_Q N_m(Q|\mathbf{y}) \cdot e^{nf(Q)} \right\} = \quad (\text{A.3})$$

$$\left\{ \mathbf{y} : e^{nf(\hat{Q}_{\mathbf{x}_m \mathbf{y}})} \geq e^{nT} \cdot \max_Q N_m(Q|\mathbf{y}) \cdot e^{nf(Q)} \right\} = \tilde{\mathcal{R}}_m \quad (\text{A.4})$$

have the same random coding exponents. ■

Proof of Theorem 2: We assume, without loss of generality, that $\mathbf{x}_1$ was transmitted. For a random codebook,

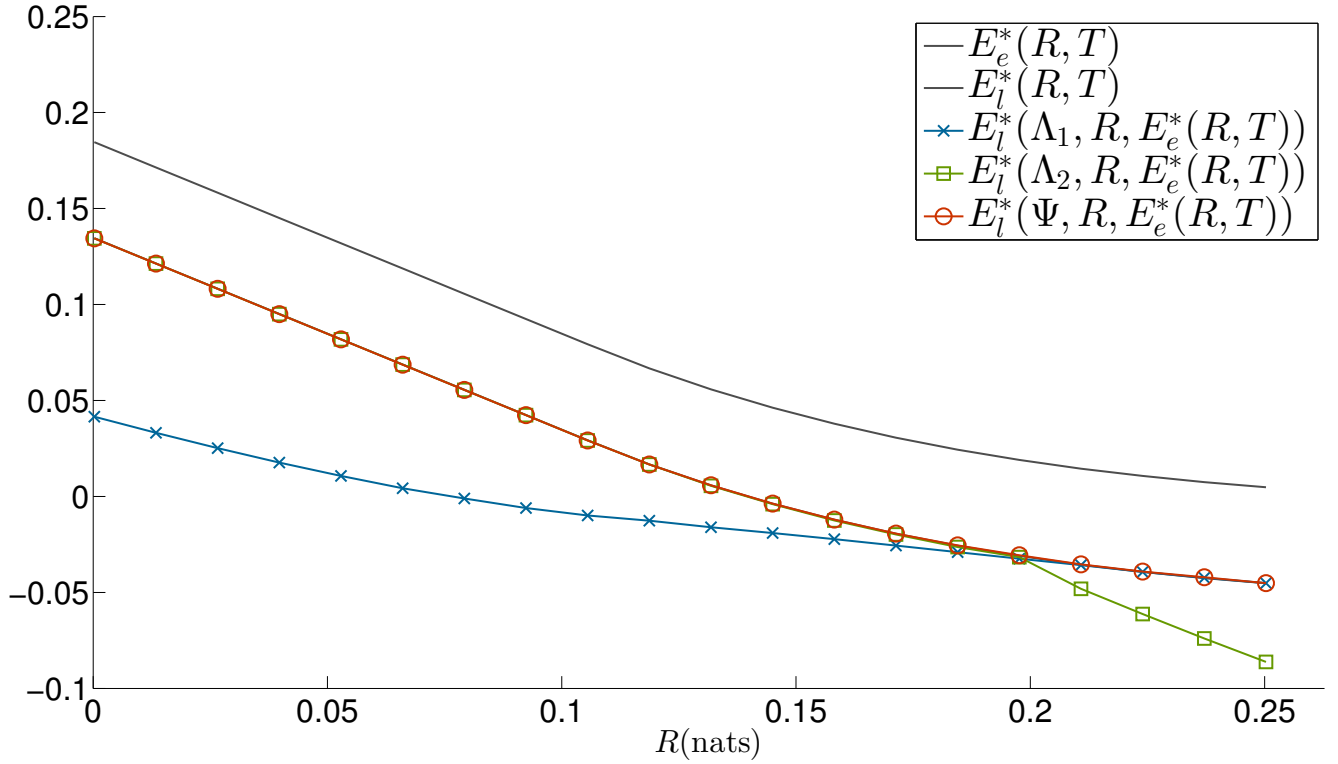


Figure 2. Graphs of random coding exponents for W_2 and $T = -0.05$.

the error probability is

$$\overline{P_e(\mathcal{C}, \phi)} = \sum_{\mathbf{x}_1, \mathbf{y}} P(\mathbf{x}_1, \mathbf{y}) \cdot \mathbb{P}(\text{error} | \mathbf{x}_1, \mathbf{y}) \quad (\text{A.5})$$

$$\stackrel{(a)}{=} \sum_{\mathbf{x}_1, \mathbf{y}} P(\mathbf{x}_1, \mathbf{y}) \cdot \mathbb{P} \left\{ f(\tilde{Q}) < \max_{m \geq 1} h(\hat{Q}_{\mathbf{x}_m \mathbf{y}}) \right\} \quad (\text{A.6})$$

$$= \sum_{\mathbf{x}_1, \mathbf{y}} P(\mathbf{x}_1, \mathbf{y}) \mathbb{P} \bigcup_{m \geq 1} \left\{ f(\tilde{Q}) < h(\hat{Q}_{\mathbf{x}_m \mathbf{y}}) \right\} \quad (\text{A.7})$$

$$\stackrel{(b)}{=} \sum_{\mathbf{x}_1, \mathbf{y}} P(\mathbf{x}_1, \mathbf{y}) \cdot \min \left\{ M \cdot \mathbb{P} \left\{ f(\tilde{Q}) < \exp \left[n \cdot h(\hat{Q}_{\mathbf{x}_2 \mathbf{y}}) \right] \right\}, 1 \right\}, \quad (\text{A.8})$$

where in (a) we have introduced the notation $\tilde{Q} \triangleq \hat{Q}_{\mathbf{x}_1 \mathbf{y}}$, and in (b) we have used the exponential tightness of the union bound (limited by unity) for a union of *exponential* number of events, assuming that they are pairwise independent [3, Lemma A.2]. Now, from the method of types,

$$\mathbb{P} \left\{ f(\tilde{Q}) < \exp \left[n \cdot h(\hat{Q}_{\mathbf{x}_2 \mathbf{y}}) \right] \right\} = \sum_{Q: Q_Y = \tilde{Q}_Y, h(Q) \geq f(\tilde{Q})} e^{-n \cdot I(Q)} \quad (\text{A.9})$$

$$\doteq \exp \left[- \min_{Q: Q_Y = \tilde{Q}_Y, h(Q) \geq f(\tilde{Q})} I(Q) \right] \quad (\text{A.10})$$

and then, it is easy to show that the resulting exponent is as in (22).

Next, let us evaluate the random coding list size exponent

$$\overline{L(\mathcal{C}, \phi)} = \sum_{\mathbf{x}_1, \mathbf{y}} P(\mathbf{x}_1, \mathbf{y}) \cdot \mathbb{E}[L|\mathbf{X}_1 = \mathbf{x}_1, \mathbf{Y} = \mathbf{y}]. \quad (\text{A.11})$$

Given $\mathbf{x}_1, \mathbf{y}$ (with $\hat{Q}_{\mathbf{x}_1, \mathbf{y}} = \tilde{Q}$) we have

$$\mathbb{E}[L|\mathbf{X}_1 = \mathbf{x}_1, \mathbf{Y} = \mathbf{y}] = \mathbb{E}\left[\sum_{m=2}^M \mathbb{I}\left\{P(\mathbf{y}|\mathbf{X}_m) \geq \max\left\{\max_{l>1, l \neq m} \exp\left[n \cdot h(\hat{Q}_{\mathbf{x}_l, \mathbf{y}})\right], \exp\left[n \cdot h(\hat{Q}_{\mathbf{x}_1, \mathbf{y}})\right]\right\}\right\}\right] \quad (\text{A.12})$$

$$= \sum_{m=2}^M \mathbb{P}\left\{P(\mathbf{y}|\mathbf{X}_m) \geq \max\left\{\max_{l>1, l \neq m} \exp\left[n \cdot h(\hat{Q}_{\mathbf{x}_l, \mathbf{y}})\right], \exp\left[n \cdot h(\hat{Q}_{\mathbf{x}_1, \mathbf{y}})\right]\right\}\right\} \quad (\text{A.13})$$

$$\doteq e^{nR} \cdot \mathbb{P}\left\{e^{n \cdot f(\hat{Q}_{\mathbf{x}_2, \mathbf{y}})} \geq \max\left\{\max_{l>2} e^{n \cdot h(\hat{Q}_{\mathbf{x}_l, \mathbf{y}})}, e^{n \cdot h(\hat{Q}_{\mathbf{x}_1, \mathbf{y}})}\right\}\right\} \quad (\text{A.14})$$

$$\stackrel{(a)}{=} e^{nR} \cdot \sum_{Q'} \mathbb{P}(\hat{Q}_{\mathbf{x}_2, \mathbf{y}} = Q') \cdot \mathbb{P}\left\{e^{n \cdot f(Q')} \geq \max\left\{\max_{l>2} e^{n \cdot h(\hat{Q}_{\mathbf{x}_l, \mathbf{y}})}, e^{n \cdot h(\tilde{Q})}\right\}\right\} \quad (\text{A.15})$$

$$= e^{nR} \cdot \sum_{Q': f(Q') \geq h(\tilde{Q})} \mathbb{P}(\hat{Q}_{\mathbf{x}_2, \mathbf{y}} = Q') \cdot \mathbb{P}\left\{e^{n \cdot f(Q')} \geq \max_{l>2} e^{n \cdot h(\hat{Q}_{\mathbf{x}_l, \mathbf{y}})}\right\} \quad (\text{A.16})$$

$$= e^{nR} \cdot \sum_{Q': f(Q') \geq h(\tilde{Q})} \mathbb{P}(\hat{Q}_{\mathbf{x}_2, \mathbf{y}} = Q') \cdot \mathbb{P}\bigcap_{l>2} \left\{e^{n \cdot h(\hat{Q}_{\mathbf{x}_l, \mathbf{y}})} \leq e^{n \cdot f(Q')}\right\} \quad (\text{A.17})$$

where (a) is using the fact that the codewords are drawn independently and with the same distribution, and the law of total probability. Letting $\tilde{Q}_l$ be the joint type of $(\mathbf{X}_l, \mathbf{y})$, for any given $u \in \mathbb{R}$

$$\mathbb{P}\bigcap_{l>2} \left\{e^{n \cdot h(\tilde{Q}_l)} \leq e^{nu}\right\} = \left[\mathbb{P}\left\{e^{n \cdot h(\tilde{Q}_l)} \leq e^{nu}\right\}\right]^{M-2} \quad (\text{A.18})$$

$$= \left[1 - \mathbb{P}\left\{\tilde{Q}_l : h(\tilde{Q}_l) > u\right\}\right]^{M-2} \quad (\text{A.19})$$

$$\doteq \left[1 - e^{-n \cdot \mathbf{U}(\tilde{Q}_Y, u)}\right]^{e^{nR}} \quad (\text{A.20})$$

where

$$\mathbf{U}(\tilde{Q}_Y, u) \triangleq \min_{Q': Q'_Y = \tilde{Q}_Y, h(Q') \geq u} I(Q'). \quad (\text{A.21})$$

Then,

$$\mathbb{P}\bigcap_{l>2} \left\{e^{n \cdot h(\tilde{Q}_l)} \leq e^{nu}\right\} \doteq \begin{cases} 1, & R \leq \mathbf{U}(\tilde{Q}_Y, u) \\ 0, & R > \mathbf{U}(\tilde{Q}_Y, u) \end{cases} \quad (\text{A.22})$$

and so

$$\mathbb{E}[L|\mathbf{X}_1 = \mathbf{x}_1, \mathbf{Y} = \mathbf{y}] \doteq e^{nR} \cdot \sum_{Q': f(Q') \geq h(\tilde{Q})} \mathbb{P}(\hat{Q}_{\mathbf{x}_2, \mathbf{y}} = Q') \cdot \mathbb{I}\left\{\mathbf{U}(\tilde{Q}_Y, f(Q')) \geq R\right\} \quad (\text{A.23})$$

$$\doteq \exp\left[-n \cdot \left(\min_{Q \in \mathcal{D}(\tilde{Q})} I(Q) - R\right)\right] \quad (\text{A.24})$$

where

$$\mathcal{D}(\tilde{Q}) \triangleq \left\{ Q : Q_Y = \tilde{Q}_Y, \mathbf{U}(\tilde{Q}_Y, f(Q)) \geq R, f(Q) \geq h(\tilde{Q}) \right\}. \quad (\text{A.25})$$

To simplify the set $\mathcal{D}(\tilde{Q})$, notice that the condition $\mathbf{U}(\tilde{Q}_Y, f(Q)) \geq R$ is actually

$$\forall Q' : Q'_Y = \tilde{Q}_Y, h(Q') \geq f(Q) \Rightarrow I(Q') \geq R \quad (\text{A.26})$$

and equivalent to

$$\forall Q' : Q'_Y = \tilde{Q}_Y, I(Q') < R \Rightarrow h(Q') < f(Q). \quad (\text{A.27})$$

Thus, from continuity, the set $\mathcal{D}(\tilde{Q})$ can equivalently be written as in (21), where $\mathbf{V}(Q_Y, R)$ is as defined in (20).

After averaging w.r.t. $(\mathbf{x}_1, \mathbf{y})$ as in (A.11), the list size exponent (23) is obtained. $\blacksquare$

Proof of Corollary 3: We use the general expression for the exponents of a decoder $\phi_h \in \Psi$, and obtain the exponents of $\phi_g \in \Lambda_1$ by setting $h(Q) = g(Q_Y)$. For the error exponent, we get

$$E_e(R, g) = \min_{\tilde{Q}} \min_{Q: Q_Y = \tilde{Q}_Y, g(Q_Y) \geq f(\tilde{Q})} \left\{ D(\tilde{Q}_{Y|X} \| W | P_X) + [I(Q) - R]_+ \right\} \quad (\text{A.28})$$

$$= \min_{\tilde{Q}} \min_{Q: Q_Y = \tilde{Q}_Y, g(Q_Y) \geq f(\tilde{Q})} \left\{ D(\tilde{Q}_{Y|X} \| W | P_X) + [I(Q) - R]_+ \right\} \quad (\text{A.29})$$

$$= \min_{\tilde{Q}} \min_{Q: Q_Y = \tilde{Q}_Y, g(\tilde{Q}_Y) \geq f(\tilde{Q})} \left\{ D(\tilde{Q}_{Y|X} \| W | P_X) + [I(Q) - R]_+ \right\} \quad (\text{A.30})$$

$$= \min_{Q'_Y} \min_{Q: Q_Y = Q'_Y} \min_{\tilde{Q}: \tilde{Q}_Y = Q'_Y, g(\tilde{Q}_Y) \geq f(\tilde{Q})} \left\{ D(\tilde{Q}_{Y|X} \| W | P_X) + [I(Q) - R]_+ \right\} \quad (\text{A.31})$$

$$= \min_{Q'_Y} \min_{\tilde{Q}: \tilde{Q}_Y = Q'_Y, g(\tilde{Q}_Y) \geq f(\tilde{Q})} \left\{ D(\tilde{Q}_{Y|X} \| W | P_X) + \min_{Q: Q_Y = Q'_Y} [I(Q) - R]_+ \right\} \quad (\text{A.32})$$

$$= \min_{Q'_Y} \min_{\tilde{Q}: \tilde{Q}_Y = Q'_Y, g(\tilde{Q}_Y) \geq f(\tilde{Q})} D(\tilde{Q}_{Y|X} \| W | P_X) \quad (\text{A.33})$$

which is (24). For the list size exponent, we first obtain

$$\mathcal{D}(\tilde{Q}) = \left\{ Q : Q_Y = \tilde{Q}_Y, f(Q) > \max \{ \mathbf{V}(Q_Y, R), g(Q_Y) \} \right\}, \quad (\text{A.34})$$

where

$$\mathbf{V}(Q_Y, R) = \max_{Q': Q'_Y = Q_Y, I(Q') \leq R} g(Q'_Y) = g(Q'_Y) \quad (\text{A.35})$$

and the last equality is due to the feasibility of the maximization, when setting $Q' = P_X \times Q_Y$. So,

$$\mathcal{D}(\tilde{Q}) = \left\{ Q : Q_Y = \tilde{Q}_Y, f(Q) > g(\tilde{Q}_Y) \right\} \quad (\text{A.36})$$

and (23) implies (25). $\blacksquare$

Proof of Corollary 4: Follows directly by substituting $h(Q) = T + f(Q)$ in (22) and (23). $\blacksquare$

Proof of Theorem 7: From (24), if $g(Q_Y)$ is such that $E_e(R, g) > E$ then equivalently, the following condition

holds:

$$D(Q_{Y|X}||W|P_X) \leq E \Rightarrow g(Q_Y) \leq f(Q). \quad (\text{A.37})$$

Clearly, under this requirement on $E_e(R, g)$, the threshold $g(Q_Y)$ should be chosen as large as possible in order to maximize the list size exponent. Thus, the optimal (maximal) threshold function is given by (31). The resulting list size exponent is immediate from eq. (25). ■

Proof of Lemma 8: The first two properties are straightforward to prove. For convexity in Q_Y , first, since $f(Q)$ is linear in Q , the minimizer Q^* in $g^*(Q_Y, E)$ always achieves the divergence constraint with an equality. Second, let Q_Y^0 and Q_Y^1 be two Y -marginals, and consider

$$Q_Y^\alpha = (1 - \alpha) \cdot Q_Y^0 + \alpha \cdot Q_Y^1. \quad (\text{A.38})$$

Also, let Q_0^* and Q_1^* be the corresponding minimizers in $g^*(Q_Y^0, E)$ and $g^*(Q_Y^1, E)$, respectively. Now, since for any $\alpha \in (0, 1)$ the Y -marginal of

$$Q_\alpha \triangleq (1 - \alpha) \cdot Q_0^* + \alpha \cdot Q_1^* \quad (\text{A.39})$$

is exactly Q_Y^α , and because the divergence is a convex function then

$$D(Q_{\alpha,Y|X}||W|P_X) \leq E, \quad (\text{A.40})$$

we obtain

$$g^*(Q_Y^\alpha, E) \leq f(Q_\alpha) \quad (\text{A.41})$$

$$\stackrel{(a)}{=} (1 - \alpha) \cdot f(Q_0^*) + \alpha \cdot f(Q_1^*). \quad (\text{A.42})$$

$$= (1 - \alpha) \cdot g^*(Q_Y^0, E) + \alpha \cdot g^*(Q_Y^1, E). \quad (\text{A.43})$$

where (a) is due to linearity of $f(Q)$. This proves convexity. ■

Proof of Theorem 9: Assume that $E_e(R, h) \geq E$ for a decoder $\phi_h \in \Psi$. Under this requirement on $E_e(R, h)$ the threshold $h(Q)$ should be chosen as large as possible in order to maximize the list size exponent. Suppose we are given a joint type Q , and notice that the requirement $E_e(R, h) > E$ is equivalent to

$$\forall \tilde{Q} : Q_Y = \tilde{Q}_Y, D(\tilde{Q}_{Y|X}||W|P_X) + [I(Q) - R]_+ \leq E \Rightarrow h^*(Q, R, E) < f(\tilde{Q}) \quad (\text{A.44})$$

or

$$h^*(Q, R, E) \leq \min_{\tilde{Q} : \tilde{Q}_Y = Q_Y, D(\tilde{Q}_{Y|X}||W|P_X) + [I(Q) - R]_+ \leq E} f(\tilde{Q}) \quad (\text{A.45})$$

we obtain (34).

For the optimal threshold function $h^*(Q, R, E)$, the resulting error exponent is E by assumption. Let us find

now the achieved $E_l^*(\Psi, R, E)$ given by

$$E_l^*(\Psi, R, E) = \min_{\tilde{Q}} \min_{Q \in \hat{\mathcal{D}}(\tilde{Q}, R, E)} \left\{ D(\tilde{Q}_{Y|X} \| W | P_X) + I(Q) - R \right\} \quad (\text{A.46})$$

where

$$\hat{\mathcal{D}}(\tilde{Q}, R, E) \triangleq \left\{ Q : Q_Y = \tilde{Q}_Y, f(Q) \geq \max \left\{ \mathbf{V}^*(Q_Y, R, E), h^*(\tilde{Q}, R, E) \right\} \right\}, \quad (\text{A.47})$$

and in this case

$$\mathbf{V}^*(Q_Y, R, E) \triangleq \max_{Q' : Q'_Y = Q_Y, I(Q') \leq R} h^*(Q', R, E) \quad (\text{A.48})$$

$$= g^*(Q_Y, E) \quad (\text{A.49})$$

The set $\hat{\mathcal{D}}(\tilde{Q}, R, E)$ can be simplified to the set $\mathcal{D}^*(\tilde{Q}, R, E)$ in (33), by showing that $\mathbf{V}^*(Q_Y, R, E)$ is never strictly larger than $h^*(\tilde{Q}, R, E)$. This can be verified by separating the outer minimization over $\tilde{Q}$ in (A.46), into two cases, which both satisfy $\mathbf{V}^*(Q_Y, R, E) \leq h^*(\tilde{Q}, R, E)$. For $I(\tilde{Q}) \leq R$

$$\mathbf{V}^*(\tilde{Q}_Y, R, E) = g^*(\tilde{Q}_Y, E) \quad (\text{A.50})$$

$$= h^*(\tilde{Q}, R, E), \quad (\text{A.51})$$

and for $I(\tilde{Q}) > R$,

$$\mathbf{V}^*(\tilde{Q}_Y, R, E) = g^*(\tilde{Q}_Y, E) \quad (\text{A.52})$$

$$\leq g^*(\tilde{Q}_Y, E - I(\tilde{Q}) + R) \quad (\text{A.53})$$

using Lemma 8 (property 1). Thus, we obtain (35). ■

Proof of Theorem 11: Here, the requirement $E_e(R, T) \geq E$ is equivalent to

$$D(\tilde{Q}_{Y|X} \| W | P_X) + [I(Q) - R]_+ \leq E \Rightarrow T < f(\tilde{Q}) - f(Q) \quad (\text{A.54})$$

which leads to (40) using

$$T^*(R, E) = \min_{\tilde{Q}} \min_{Q : \tilde{Q}_Y = Q_Y, D(\tilde{Q}_{Y|X} \| W | P_X) + [I(Q) - R]_+ \leq E} \left\{ f(\tilde{Q}) - f(Q) \right\} \quad (\text{A.55})$$

$$= \min_{\tilde{Q}} \min_{\tilde{Q} : \tilde{Q}_Y = Q_Y, D(\tilde{Q}_{Y|X} \| W | P_X) \leq E - [I(Q) - R]_+} \left\{ f(\tilde{Q}) - f(Q) \right\} \quad (\text{A.56})$$

$$= \min_Q \left\{ g^*(Q_Y, E - [I(Q) - R]_+) - f(Q) \right\} \quad (\text{A.57})$$

$$= \min_Q \left\{ h^*(Q, R, E) - f(Q) \right\}. \quad (\text{A.58})$$

The list size exponent (41) is immediate. ■

Proof of Proposition 12: We have

$$E_l^*(\Psi, R, E) = \min_{\tilde{Q}} \min_{Q \in \mathcal{D}^*(\tilde{Q}, R, E)} \left\{ D(\tilde{Q}_{Y|X} \| W | P_X) + I(Q) - R \right\} \quad (\text{A.59})$$

$$\leq \min_{\tilde{Q}: I(\tilde{Q}) \leq R} \min_{Q: Q_Y = \tilde{Q}_Y, f(Q) > g^*(Q_Y, E)} \left\{ D(\tilde{Q}_{Y|X} \| W | P_X) + I(Q) - R \right\} \quad (\text{A.60})$$

$$\stackrel{(a)}{=} \min_{\tilde{Q}} \min_{Q: Q_Y = \tilde{Q}_Y, f(Q) > g^*(Q_Y, E)} \left\{ D(\tilde{Q}_{Y|X} \| W | P_X) + I(Q) - R \right\} \quad (\text{A.61})$$

$$= E_l^*(\Lambda_1, R, E). \quad (\text{A.62})$$

where (a) is for $R > \bar{R}_{\text{cr}}(E)$. Also, since $\Lambda_1 \subset \Psi$, we have $E_l^*(\Psi, R, E) \geq E_l^*(\Lambda_1, R, E)$ and so equality is obtained. $\blacksquare$

Proof of Proposition 13: For ϕ_T^* and $R \leq \underline{R}_{\text{cr}}(T)$

$$E_e^*(R, T) \leq E_a(R, T) \quad (\text{A.63})$$

$$= \min_{\tilde{Q}} \min_{Q: Q_Y = \tilde{Q}_Y, f(Q) + T \geq f(\tilde{Q}), I(Q) \geq R} D(\tilde{Q}_{Y|X} \| W | P_X) + I(Q) - R \quad (\text{A.64})$$

$$= D(\tilde{Q}_{Y|X}^* \| W | P_X) + I(Q^*) - R, \quad (\text{A.65})$$

and also $E_l^*(R, T) = E_e^*(R, T) + T$.

Now, we consider the optimization problem

$$\min_{\tilde{Q}} \min_{Q: Q_Y = \tilde{Q}_Y} \left\{ D(\tilde{Q}_{Y|X} \| W | P_X) + I(Q) \right\}. \quad (\text{A.66})$$

and show that its solution, which we denote by $(\tilde{Q}_0, Q_0)$, satisfies $\tilde{Q}_0 = Q_0$. To see this, we utilize the identity

$$D(\tilde{Q}_{Y|X} \| W | P_X) + I(Q) = D(Q_{Y|X} \| W | P_X) + I(\tilde{Q}) + f(Q) - f(\tilde{Q}) \quad (\text{A.67})$$

which holds under the assumption $Q_Y = \tilde{Q}_Y$, and can be proved using simple algebraic manipulations. Now, for any given $(Q, \tilde{Q})$ such that $Q_Y = \tilde{Q}_Y$

$$\begin{aligned} D(\tilde{Q}_{Y|X} \| W | P_X) + I(Q) &\stackrel{(a)}{=} \frac{1}{2} \left(D(\tilde{Q}_{Y|X} \| W | P_X) + D(Q_{Y|X} \| W | P_X) + I(Q) + I(\tilde{Q}) \right) \\ &\quad + \frac{1}{2} \left(f(Q) - f(\tilde{Q}) \right) \end{aligned} \quad (\text{A.68})$$

$$\stackrel{(b)}{\geq} D \left(\frac{1}{2} (\tilde{Q}_{Y|X} + Q_{Y|X}) \| W | P_X \right) + I \left(\frac{1}{2} (\tilde{Q} + Q) \right) + \frac{1}{2} \left(f(Q) - f(\tilde{Q}) \right) \quad (\text{A.69})$$

$$\stackrel{(c)}{=} D \left(\frac{1}{2} (\tilde{Q}_{Y|X} + Q_{Y|X}) \| W | P_X \right) + I \left(\frac{1}{2} (\tilde{Q} + Q) \right) + \frac{1}{2} \left(f(Q) + f(\tilde{Q}) \right) - f(\tilde{Q}) \quad (\text{A.70})$$

$$\stackrel{(d)}{\geq} D \left(\frac{1}{2} (\tilde{Q}_{Y|X} + Q_{Y|X}) \| W | P_X \right) + I \left(\frac{1}{2} (\tilde{Q} + Q) \right) + \frac{1}{2} \left(f(Q) + f(\tilde{Q}) \right) \quad (\text{A.71})$$

where (a) is because the right hand side of (A.67) equals the average of both sides of (A.67), (b) is due to

convexity of both the divergence and the mutual information⁵, (c) is due to the linearity of $f(Q)$, and (d) is due to the negativity of $f(Q)$. Equalities are obtained in (b) and (d) by choosing $Q = \tilde{Q}$. Thus,

$$\min_{\tilde{Q}} \min_{Q: Q_Y = \tilde{Q}_Y} \left\{ D(\tilde{Q}_{Y|X} \| W | P_X) + I(Q) \right\} = \min_Q \left\{ D(Q_{Y|X} \| W | P_X) + I(Q) \right\}, \quad (\text{A.72})$$

and so $\tilde{Q}_0 = Q_0$. This implies that for $T \geq 0$, $\tilde{Q}^* = Q^* = \tilde{Q}_0 = Q_0$. Thus, for $R \leq \underline{R}_{\text{cr}}(T)$

$$E_e(R, T) = \min_{\tilde{Q}} \min_{Q: Q_Y = \tilde{Q}_Y, f(Q) + T \geq f(\tilde{Q})} D(\tilde{Q}_{Y|X} \| W | P_X) + [I(Q) - R]_+ \quad (\text{A.73})$$

$$= D(\tilde{Q}_{Y|X}^* \| W | P_X) + I(Q^*) - R \quad (\text{A.74})$$

$$= E_a(R, T) \quad (\text{A.75})$$

$$\geq E_e^*(R, T) \quad (\text{A.76})$$

On the other hand, for the list size exponent,

$$E_l(R, T) = \min_{\tilde{Q}} \min_{Q \in \underline{\mathcal{D}}(\tilde{Q}, R, T)} \left\{ D(\tilde{Q}_{Y|X} \| W | P_X) + I(Q) - R \right\} \quad (\text{A.77})$$

$$\stackrel{(a)}{\geq} \min_{\tilde{Q}} \min_{Q: Q_Y = \tilde{Q}_Y, f(Q) \geq f(\tilde{Q}) + T} \left\{ D(\tilde{Q}_{Y|X} \| W | P_X) + I(Q) - R \right\} \quad (\text{A.78})$$

$$\stackrel{(b)}{\geq} \min_{\tilde{Q}} \min_{Q: Q_Y = \tilde{Q}_Y, f(Q) \geq f(\tilde{Q}) + T} \left\{ D(\tilde{Q}_{Y|X} \| W | P_X) + I(Q) - R + f(\tilde{Q}) + T - f(Q) \right\} \quad (\text{A.79})$$

$$\stackrel{(c)}{=} \min_{\tilde{Q}} \min_{Q: Q_Y = \tilde{Q}_Y, f(Q) \geq f(\tilde{Q}) + T} \left\{ D(Q_{Y|X} \| W | P_X) + I(\tilde{Q}) - R + T \right\} \quad (\text{A.80})$$

$$= \min_{\tilde{Q}} \min_{Q: Q_Y = \tilde{Q}_Y, f(\tilde{Q}) \geq f(Q) + T} \left\{ D(\tilde{Q}_{Y|X} \| W | P_X) + I(Q) - R + T \right\} \quad (\text{A.81})$$

$$\geq \min_{\tilde{Q}} \min_{Q: Q_Y = \tilde{Q}_Y} \left\{ D(\tilde{Q}_{Y|X} \| W | P_X) + I(Q) - R + T \right\} \quad (\text{A.82})$$

$$= D(\tilde{Q}_{Y|X}^* \| W | P_X) + I(Q^*) - R + T \quad (\text{A.83})$$

$$= E_e(R, T) + T \quad (\text{A.84})$$

where inequality (a) is obtained by removing the constraint $f(Q) \geq \underline{V}(Q_Y, R)$ from the set $\underline{\mathcal{D}}(\tilde{Q}, R, T)$, inequality (b) is using the constraint $f(Q) \geq f(\tilde{Q}) + T$. Equality (c) is using the identity (A.67), and (d) is simply a substitution of notation $Q \leftrightarrow \tilde{Q}$. To conclude, we have obtained $E_e(R, T) \geq E_e^*(R, T)$ and $E_l(R, T) = E_e(R, T) + T \geq E_e^*(R, T) + T = E_l^*(R, T)$. Since ϕ_T^* provides the optimal trade-off between the error exponent and list size exponent, we obtain the desired result. ■

APPENDIX B

In [6], random coding exponents were derived for an optimal decoder in the erasure mode, i.e. with $T \geq 0$, and expressions for both $E_e^*(R, T)$ and $E_l^*(R, T)$ were derived independently (in the erasure case, $E_e(R, \phi)$ and

⁵Indeed, when both marginals of Q are constrained, $I(Q) = D(Q \| P_X \times Q_Y)$ and so convexity of the mutual information is implied by the convexity of the divergence.

$E_l(R, \phi)$ represent the probability of erasure, and the probability of undetected error, respectively). The reason for restricting the analysis to the erasure mode is that the fact that the decoding regions overlap in the list mode complicates analysis. Nonetheless, it can be easily verified that this restriction is only needed for the analysis of $E_l^*(R, T)$, and so that analysis of $E_e^*(R, T)$ is valid for this list mode $T < 0$. Moreover, it can be shown that $E_l^*(R, T) = E_e^*(R, T) + T$ must be satisfied in general (which was shown in [6] only for the BSC), see [11, Lemma 1].

In addition, the random ensemble considered in [6] is the i.i.d. ensemble over input distribution P_X , but the analysis may be easily adapted to fixed composition codes with the same input distribution, which may only lead to larger exponents. Basically, divergence terms $D(Q_X || P_X)$, where Q_X is a generic distribution, are omitted from exponential assessment of probabilities, and thus from final expressions.

Therefore, from the above reasoning, the results of [6] are also applicable here, with proper adjustments.

It should also be noted that there is an error at the end of the proof of [6, Theorem 1], where it was claimed that $\min\{E_a(R, T), E_b(R, T)\}$, which may not be true in general. The correct expression is as in (11).

REFERENCES

- [1] G. D. Forney Jr., "Exponential error bounds for erasure, list, and decision feedback schemes," *IEEE Transactions on Information Theory*, vol. 14, no. 2, pp. 206–220, 1968.
- [2] M. V. Burnashev, "Data transmission over a discrete channel with feedback: Random transmission time," *Problems of Information transmission*, pp. 250–265, 1976.
- [3] N. Shulman, "Communication over an unknown channel via common broadcasting," Ph.D. dissertation, Tel Aviv University, 2003, http://www.eng.tau.ac.il/~shulman/papers/Nadav_PhD.pdf.
- [4] N. Merhav, "List decoding - random coding exponents and expurgated exponents," *Information Theory, IEEE Transactions on*, vol. 60, no. 11, pp. 6749–6759, Nov 2014.
- [5] R. G. Gallager, *Information Theory and Reliable Communication*. Wiley, 1968.
- [6] A. Somekh-Baruch and N. Merhav, "Exact random coding exponents for erasure decoding," *IEEE Transactions on Information Theory*, vol. 57, no. 10, pp. 6444–6454, 2011.
- [7] N. Merhav, "Statistical physics and information theory," *Foundations and Trends in Communications and Information Theory*, vol. 6, no. 1-2, pp. 1–212, 2009.
- [8] I. Csiszár and J. Körner, *Information Theory: Coding Theorems for Discrete Memoryless Systems*. Cambridge University Press, 2011.
- [9] P. Moulin, "A Neyman-Pearson approach to universal erasure and list decoding," *IEEE Transactions on Information Theory*, vol. 55, no. 10, pp. 4462–4478, 2009.
- [10] N. Merhav and M. Feder, "Minimax universal decoding with an erasure option," *IEEE Transactions on Information Theory*, vol. 53, no. 5, pp. 1664–1675, 2007.
- [11] W. Huleihel, N. Weinberger, and N. Merhav, "Erasure/list random coding error exponents are not universally achievable," *Submitted to IEEE Transactions on Information Theory*, October 2014, available online: <http://arxiv.org/pdf/1410.7005v1.pdf>.
- [12] I. Telatar, "Exponential bounds for list size moments and error probability," in *Proc. of Information Theory Workshop*, June 1998, pp. 60–.