

Distributional Information Embedding: A Framework for Multi-bit Watermarking

Haiyun He*, Yepeng Liu[†], Ziqiao Wang[‡], Yongyi Mao[§], and Yuheng Bu[†]

*Internet of Things Thrust, Hong Kong University of Science and Technology (Guangzhou), Guangzhou, China

[†]Department of CS, University of California, Santa Barbara, Santa Barbara, CA, USA

[‡]School of Computer Science and Technology, Tongji University, Shanghai, China

[§]School of EECS, University of Ottawa, Ottawa, ON, Canada

Emails: haiyunhe@hkust-gz.edu.cn, {yepengliu, buyuheng}@ucsb.edu, ziqiaowang@tongji.edu.cn, ymao@uottawa.ca

Abstract—This paper introduces a novel problem, **distributional information embedding**, motivated by the practical demands of multi-bit watermarking for large language models (LLMs). Unlike traditional information embedding, which embeds information into a pre-existing host signal, LLM watermarking actively controls the text generation process—adjusting the token distribution—to embed a detectable signal. We develop an information-theoretic framework to analyze this distributional information embedding problem, characterizing the fundamental trade-offs among three critical performance metrics: text quality, detectability, and information rate. In the asymptotic regime, we demonstrate that the maximum achievable rate with vanishing error corresponds to the entropy of the LLM’s output distribution and increases with higher allowable distortion. We also characterize the optimal watermarking scheme to achieve this rate. Extending the analysis to the finite-token case with non-i.i.d. tokens, we identify schemes that maximize detection probability while adhering to constraints on false alarm and distortion.

Index Terms—information embedding, watermarking, large language models, information theory, detection theory

I. INTRODUCTION

The rapid advancement of Large Language Models (LLMs) [1], [2] is revolutionizing numerous fields but also raises concerns about misuse, such as spreading disinformation, creating fake news, and enabling academic dishonesty. The growing prevalence and quality of AI-generated text make it challenging to *distinguish it from human-written content*.

A promising solution is to *actively* embed detectable signals into LLM-generated text, i.e., watermarks, which enable provable detection of AI-generated content. Despite recent advances in watermarking algorithms for LLM [3]–[7], they suffer from significant limitations, for example, many algorithms are heuristically designed where watermark detectability is ensured by introducing noticeable alterations to the generated content that degrade the output quality.

Additionally, most watermarking schemes are “zero-bit” schemes, designed solely to distinguish AI-generated text from human-written content without embedding any additional information. As incorporating meta-information—such as the model’s name, version, and generation time—is increasingly important for forensic analysis of LLM misuse, some multi-bit watermarking algorithms [8]–[10] have been developed

recently. However, these approaches remain heuristic and have a low information embedding rate, with current methods unable to support messages longer than a few bits [11].

Therefore, a *principled theoretical framework* is needed to analyze the fundamental trade-offs among key performance metrics in multi-bit LLM watermarking. These metrics include: (1) **Text quality**: ensuring that the watermarked text generated by LLMs maintains a quality comparable to unwatermarked text; (2) **Detectability**: the probability of missed detection and decoding errors; and (3) **Information rate**: the rate at which information can be embedded and reliably recovered.

Information theory has a long-standing history of guiding the design of digital watermarking, dating back to the early 00s [12]–[16], within the broader framework of the information embedding problem [17]–[22]. As we will demonstrate, watermarking in LLMs introduces a novel form of such a problem, which we term *distributional information embedding*. Unlike traditional information embedding, which focuses on reliably embedding information into a pre-existing host signal while minimizing distortion, LLM watermarking actively controls the generation process—the token distribution—to embed a detectable signal while preserving the original distribution. In other words, traditional information embedding is like writing on dirty paper [23], where the challenge is to convey the message clearly despite the interference from pre-existing marks. In contrast, LLM watermarking resembles generating dirty paper in real time, embedding the message into the very process that creates the marks. This fundamental difference reshapes the problem and introduces novel challenges.

In this paper, we present an information-theoretic analysis of a distributional information embedding problem motivated by multi-bit LLM watermarking. Our goal is to design the watermarking scheme by jointly optimizing the encoder and decoder. The system must distinguish human-written text from AI-generated text while ensuring reliable recovery of the embedded information. All of this must be achieved within a specified distortion constraint to preserve text quality. Our contribution includes:

- **Asymptotic analysis in the independent and identically distributed (i.i.d.) case**: We demonstrate that the maximum information rate with vanishing detection error probability corresponds to the entropy of the LLM’s output distribution

This work was carried out while Haiyun He was a postdoctoral associate at Cornell University.

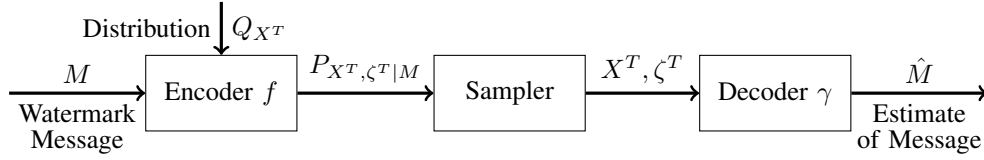


Fig. 1: Illustration of multi-bit watermarking as distributional information embedding with side information.

and increases with higher allowable distortion. Furthermore, we characterize the asymptotically optimal watermarking scheme that achieves this rate.

- **Finite token length analysis in the non-i.i.d. case:** We extend the asymptotic analysis to the practical scenario with a finite token length, aiming to maximize the detection accuracy while satisfying both a worst-case false alarm probability constraint and a distortion constraint. Furthermore, we characterize the minimum achievable detection error and identify the optimal watermarking scheme and decoder for this setting.

II. PROBLEM FORMULATION

a) Distributional Information Embedding with Side Information: Consider a length- T data sequence X^T generated from a joint distribution $Q_{X^T} \in \mathcal{P}(\mathcal{X}^T)$, where $\mathcal{P}(\mathcal{X}^T)$ denotes the probability simplex in $\mathcal{X}^T$. For simplicity, we ignore the potential auto-regressive structure of Q_{X^T} in the current analysis. In the generation process, a message M drawn from $[m] := \{1, \dots, m\}$ needs to be embedded in the data sequence by constructing a dependence structure between X^T and an auxiliary random sequence ζ^T with alphabet $\mathcal{Z}^T$, which serves as side information available to the decoder.

For example, the joint distribution Q_{X^T} can be viewed as the output distribution of an LLM for a length- T token sequence X^T . Most LLM watermarking schemes adopt this distributional information embedding with side information framework. An example of a zero-bit watermarking scheme is provided below.

Example 1 (Existing watermarking schemes as special cases). In the Green-Red List watermarking scheme [4], at each position t , the token vocabulary $\mathcal{X}$ is randomly split into a green list $\mathcal{G}$ and a red list $\mathcal{R}$, with $|\mathcal{G}| = \rho|\mathcal{X}|$. This split is represented by a $|\mathcal{X}|$ -dimensional binary auxiliary variable ζ_t , indexed by $x \in \mathcal{X}$, where $\zeta_t(x) = 1$ means $x \in \mathcal{G}$; otherwise, $x \in \mathcal{R}$. The watermarking scheme is as follows:

- Compute a hash of the previous token X_{t-1} using a hash function $\text{hash} : \mathcal{X} \times \mathbb{R} \rightarrow \mathbb{R}$ and a shared secret key: $\text{hash}(X_{t-1}, \text{key})$.
- Use $\text{hash}(X_{t-1}, \text{key})$ as a seed to uniformly sample the auxiliary variable ζ_t from the set $\{\zeta \in \{0, 1\}^{|\mathcal{X}|} : \|\zeta\|_1 = \rho|\mathcal{X}|\}$ to construct the green list $\mathcal{G}$.
- Sample X_t from the modified token-generating distribution which increases the logit of tokens in $\mathcal{G}$ by $\delta > 0$:

$$P_{X_t | x^{t-1}, \zeta_t}(x) = \frac{Q_{X_t | x^{t-1}}(x) \exp(\delta \cdot \mathbb{I}\{\zeta_t(x) = 1\})}{\sum_{x \in \mathcal{V}} Q_{X_t | x^{t-1}}(x) \exp(\delta \cdot \mathbb{I}\{\zeta_t(x) = 1\})}.$$

More examples of several other watermarking schemes, e.g., Gumbel-Max [3], EXP-Edit [5] and text-adaptive watermark [7], is provided in [24, Appendix A].

In this paper, we focus on studying one usage scenario within this framework: multi-bit watermarking. Below, we formulate the multi-bit watermarking problem during data generation as a distributional information embedding problem with side information, as illustrated in Figure 1.

Definition 1 (Multi-bit Watermarking). A watermarking system is an encoder/decoder pair (f, γ) . The encoder $f : [m] \times \mathcal{P}(\mathcal{X}^T) \rightarrow \mathcal{P}(\mathcal{X}^T \times \mathcal{Z}^T | [m])$ inputs a watermark message M drawn from the index set $[m]$ and the data generation distribution Q_{X^T} , outputting a joint distribution $P_{X^T, \zeta^T | M}$ that creates dependence between the generated data and auxiliary random sequence ζ^T . The decoder receives (X^T, ζ^T) sampled from $P_{X^T, \zeta^T | M}$, and guesses the message M with decoder $\gamma : \mathcal{X}^T \times \mathcal{Z}^T \rightarrow [0 : m]$, i.e., $\hat{M} = \gamma(X^T, \zeta^T)$. If $\hat{M} = 0$, the sequence X^T is decoded as unwatermarked; if $\hat{M} \in [m]$, X^T is decoded as watermarked with message $\hat{M}$. This system defines an (m, T) watermarking scheme with an information rate $R := \log m/T$.

Note that the watermarked sequence is generated from P_{X^T} (induced by the encoder f) instead of the original Q_{X^T} . To measure the distortion level of a watermarking scheme, we use the divergence between these two distributions.

Definition 2 (d -Distorted Watermarking). A watermarking encoder f is d -distorted with respect to the distortion D , if for any $M \in [m]$ and $Q_{X^T} \in \mathcal{P}(\mathcal{X}^T)$, the marginal distribution of the output $P_{X^T, \zeta^T | M}$ satisfies $D(P_{X^T | M}, Q_{X^T}) \leq d$.

Here, D can be any divergence. Common examples of such divergences include total variation, KL divergence, and Wasserstein distance. For $d = 0$, the watermarking scheme is called distortion-free.

Moreover, to ensure the secrecy of the embedded message, we assume that the watermarked sequence X^T should be indistinguishable for any embedded message M , provided the auxiliary sequence is unknown. A distortion-free watermarking scheme satisfies this condition, as it ensures $P_{X^T | M=j} = Q_{X^T}$, for all $j \in [m]$. Additionally, the auxiliary sequence itself should not reveal any information about the message. Otherwise, the message M can be transmitted directly via the dependence between ζ and M , bypassing the need for the generated text.

Assumption 1 (Secrecy of Embedded Message). The encoder

f must ensure that both X^T and ζ^T are statistically independent of the embedded message M .

Under this assumption, the embedded message M cannot be decoded with only X^T or ζ^T and $I(M; X^T, \zeta^T) = I(M; X^T | \zeta^T) = I(M; \zeta^T | X^T)$. To detect if X^T is watermarked, the decoder must exploit the auxiliary sequence ζ^T . This corresponds to decoding with side information.

b) Watermark Detection and Decoding: Under our framework, if the token sequence X^T is unwatermarked, it is independent of ζ^T ; otherwise, (X^T, ζ^T) is jointly distributed according to one of the m distributions $\{P_{X^T, \zeta^T | M=j}\}_{j=1}^m$. Thus, detecting and decoding the watermark message M boils down to the $(m+1)$ -ary hypothesis testing:

- H_0 : X^T is generated by a human, i.e., $(X^T, \zeta^T) \sim \mathbb{P}_0 := P_{X^T, \zeta^T | M=0} = Q_{X^T} \otimes P_{\zeta^T}$;
- $H_j, \forall j \in [m]$: X^T is generated by a watermarked LLM and embedded with message j , $(X^T, \zeta^T) \sim \mathbb{P}_j := P_{X^T, \zeta^T | M=j}$.

Detection performance is measured by the j -th error probability: for any $j \in [0:m]$,

$$\beta_j(\gamma, P_{X^T, \zeta^T | M=j}) := \mathbb{P}_j(\gamma(X^T, \zeta^T) \neq j).$$

Note that for $j \neq 0$, $\beta_j(\gamma, P_{X^T, \zeta^T | M=j}) = \mathbb{P}_j(\gamma(X^T, \zeta^T) = 0) + \mathbb{P}_j(\gamma(X^T, \zeta^T) \in [m] \setminus \{j\})$ is the sum of miss detection error and miss decoding error. For $j = 0$, $\beta_0(\gamma, Q_{X^T} \otimes P_{\zeta^T})$ is the false alarm error. Since human-generated texts can vary widely, we aim to control the *worst-case* false alarm error $\sup_{Q_{X^T}} \beta_0(\gamma, Q_{X^T} \otimes P_{\zeta^T})$ at a given $\alpha \in (0, 1)$.

Our design objective is then three-fold: 1) maximizing the information rate R , 2) ensuring the distortion remains bounded by d , and 3) minimizing $\beta_j(\gamma, P_{X^T, \zeta^T | M=j})$ for all $j \in [m]$ while the *worst-case* false alarm error $\sup_{Q_{X^T}} \beta_0(\gamma, Q_{X^T} \otimes P_{\zeta^T})$ is controlled.

III. ASYMPTOTIC RESULTS WITH IID TOKENS

In this section, we begin with an asymptotic analysis by letting the length of tokens $T \rightarrow \infty$ for the i.i.d. case to build intuition for the optimal design of the watermarking scheme.

Suppose $X_1, \dots, X_T$ are i.i.d. with an identical distribution P_X , and $\zeta_1, \dots, \zeta_T$ are i.i.d. with P_ζ . Under each H_j , $(X_1, \zeta_1), \dots, (X_T, \zeta_T)$ are conditionally i.i.d. with distribution $P_{X, \zeta | M=j}$. Specifically, $P_{X, \zeta | M=0} = Q_X \otimes P_\zeta$. Additionally, we assume a uniform prior distribution of message M on $[m]$.

A. Converse Result

We first analyze the maximum information rate that an (m, T) watermarking scheme can achieve with vanishing decoding error $\Pr(\hat{M} \neq M) = \frac{1}{m} \sum_{j=1}^m \beta_j(\gamma, P_{X^T, \zeta^T | M=j})$.

Lemma 1 (Best Achievable Information Rate). *Given any Q_X, P_X satisfying $D(P_X^T, Q_X^T) \leq d$, and $\{P_{X, \zeta | M=i}\}_{i=0}^m$, if the decoding error $\Pr(\hat{M} \neq M) \rightarrow 0$ as $T \rightarrow \infty$, the information rate of this d -distorted (m, T) watermarking scheme is upper bounded by*

$$R \leq H(P_X) \leq \sup_{P_X: D(P_X^T, Q_X^T) \leq d} H(P_X),$$

where the last bound holds for all d -distorted (m, T) watermarking schemes for the LLM Q_{X^T} .

The proof is provided in [25, Appendix A]. Lemma 1 shows that it is impossible for a distortion-free watermarking to embed more than approximately $2^{T H(Q_X)}$ messages in a length- T i.i.d. token sequence while achieving vanishing j -th error for all $j \in [m]$, regardless of the false alarm probability. As the distortion d increases, a d -distorted watermarking can trade off text quality to achieve a higher rate. Note that when $d = 0$, this maximum information rate coincides with the steganography capacity characterized in classical works [22], [26] when there is no adversary and noise.

B. Achievability Result

Next, we aim to identify the asymptotically optimal watermarking scheme that can achieve vanishing j -th errors and the maximum watermarking rate, while ensuring the false alarm error below α .

To develop intuition for the optimal design, we first present an upper bound for the j -th error exponent under i.i.d. assumptions. Specifically, we extend [27, Lemma 11.8.1] to our $(m+1)$ -ary hypothesis testing setting.

Lemma 2 (Upper Bound for the j -th Error Exponent). *Fix any $j \in [m]$, $\epsilon \in (0, 1/2)$ and any set of distributions $\{\mathbb{P}_i\}_{i \in [0:m] \setminus \{j\}}$. Let $\{\mathcal{B}_{T,i}\}_{i \in [0:m] \setminus \{j\}} \subset \mathcal{X}^T \times \mathcal{Z}^T$ be any collection of m disjoint sets of sequences $((x_i, \zeta_i))_{i=1}^T$ such that $\mathbb{P}_i(\mathcal{B}_{T,i}) \geq 1 - \epsilon$. Let $\mathcal{B}_{T,j} = (\cup_{i \in [0:m] \setminus \{j\}} \mathcal{B}_{T,i})^c$. For any $\mathbb{P}_j$ such that $\max_{i \in [0:m] \setminus \{j\}} D_{\text{KL}}(P_{X, \zeta | M=i} \| P_{X, \zeta | M=j}) < \infty$,*

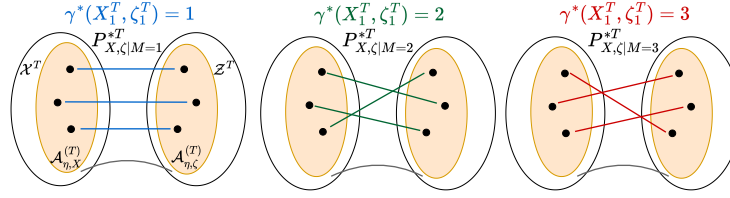
$$-\frac{\log \mathbb{P}_j(\mathcal{B}_{T,j}^c)}{T} \leq E_j^* + \epsilon - \frac{\log(m(1-2\epsilon))}{T},$$

where $E_j^* := \max_{P_X: D(P_X^T, Q_X^T) \leq d} \min_{i \in [0:m] \setminus \{j\}} D_{\text{KL}}(P_{X, \zeta | M=i} \| P_{X, \zeta | M=j})$.

The proof is provided in [25, Appendix B]. Lemma 2 shows that given any watermarking scheme $(\mathbb{P}_0, \dots, \mathbb{P}_m)$, the minimum achievable j -th error probability for all decoders decays exponentially with the rate E_j^* , while other errors are controlled below ϵ . Furthermore, the error exponent depends on the distortion level d (cf. Definition 2), which increases as d increases. If the distortion metric is set as $D_{\text{KL}}(Q_X^T \| P_X^T)$, the rate is further upper bounded by $D_{\text{KL}}(P_\zeta \| P_{\zeta | X, M=j} | Q_X) + d$.

Inspired by Lemma 2, we can design the joint distributions $(P_{X, \zeta | M=i})_{i=0}^m$ by maximizing the error exponent E_j^* . In this way, the j -th error probability decays exponentially to 0 at the fastest rate. One solution is to make the masses of $P_{X, \zeta | M=i}$ and $P_{X, \zeta | M=j}$ concentrated at different locations for $i \neq j$, which leads to $D_{\text{KL}}(P_{X, \zeta | M=i} \| P_{X, \zeta | M=j}) \rightarrow \infty$. This hints that the optimal joint distribution produced by the encoder f should almost deterministically map (X^T, ζ^T) to a message M . Based on this intuition, we construct the asymptotically jointly optimal encoder/decoder pair in the watermarking scheme.

Under any hypothesis H_j and any $P_{X, \zeta | M=j}$, we define the typical sets of sequences $\{(x^T, \zeta^T)\}$.

Fig. 2: Illustration of the asymptotically optimal watermarking scheme when $m = 3$.

Definition 3 (Typical Sets). For arbitrarily small $\eta > 0$, define the typical sets $\mathcal{A}_{\eta,X}^{(T)}$ and $\mathcal{A}_{\eta,\zeta}^{(T)}$ as

$$\mathcal{A}_{\eta,X}^{(T)} := \left\{ x^T \in \mathcal{X}^T : \left| \frac{1}{T} \log \frac{1}{P_X^T(x^T)} - H(P_X) \right| \leq \eta \right\},$$

$$\mathcal{A}_{\eta,\zeta}^{(T)} := \left\{ \zeta^T \in \mathcal{Z}^T : \left| \frac{1}{T} \log \frac{1}{P_\zeta^T(\zeta^T)} - H(P_\zeta) \right| \leq \eta \right\}.$$

The typical sequences in $\mathcal{A}_{\eta,X}^{(T)}$ and $\mathcal{A}_{\eta,\zeta}^{(T)}$ are nearly uniformly distributed and can be mapped with almost deterministic precision. Leveraging the asymptotic equipartition property (AEP), we first present the optimal design when distortion $d = 0$ as follows. Here, we use $\doteq$ to denote equality to the first order in the exponent.

Theorem 3 (Asymptotically Optimal Distortion-Free Watermarking Scheme). Let $P_X^* = Q_X$, $\mathcal{Z} \subset \mathbb{Z}$ and design $P_\zeta^* \in \mathcal{P}(\mathcal{Z})$ such that $H(P_\zeta^*) = H(P_X^*)$. If the number of message bits satisfies $\frac{1}{T}(\log m - \log \alpha) \leq H(P_X^*)$, the asymptotically optimal distortion-free watermarking encoders and decoders exist, given as follows.

Let $\eta = T^{-\frac{1}{4}}$. The class of optimal decoders is given by

$$\Gamma_\eta^* := \left\{ \gamma : \begin{cases} g(x_1^T, \zeta_1^T), & \forall x_1^T \in \mathcal{A}_{\eta,X}^{(T)}, \zeta_1^T \in \mathcal{A}_{\eta,\zeta}^{(T)}, \\ & \text{and } g(x^T, \zeta^T) \leq m; \\ 0, & \text{otherwise,} \end{cases} \right.$$

for some function $g : \mathcal{A}_{\eta,X}^{(T)} \times \mathcal{A}_{\eta,\zeta}^{(T)} \rightarrow [|\mathcal{A}_{\eta,X}^{(T)}|]$ such that $g(x^T, \cdot)$ and $g(\cdot, \zeta^T)$ are bijective, for any fixed x^T and ζ^T .

For any $\gamma^* \in \Gamma_\eta^*$, the corresponding asymptotically optimal encoder f^* outputs $P_{X,\zeta|M}^*$ as follows: for any $i \in [m]$,

- for all $x^T \in \mathcal{A}_{\eta,X}^{(T)}$, $P_{X,\zeta|M}^{*T}(x^T, \zeta^T|i) = \begin{cases} P_X^T(x^T) \doteq e^{-TH(P_\zeta^*)}, & \text{if } \zeta^T \in \mathcal{A}_{\eta,\zeta}^{(T)} \text{ and } \gamma^*(x^T, \zeta^T) = i; \\ 0, & \text{otherwise.} \end{cases}$
 - for all $x^T \notin \mathcal{A}_{\eta,X}^{(T)}$, let $P_{X,\zeta|M}^{*T}(x^T, \zeta^T|i)$ take any non-negative value as long as $\sum_{x^T, \zeta^T} P_{X,\zeta|M}^{*T}(x^T, \zeta^T|i) = 1$.
- Thus, for any $\gamma^* \in \Gamma_\eta^*$ and its corresponding $P_{X,\zeta|M}^*$, as $T \rightarrow \infty$, we have for all $j \in [m]$,

$$\beta_j(\gamma^*, P_{X,\zeta|M=j}^*) \leq \exp(-\Omega(T^{\frac{1}{2}})) \rightarrow 0,$$

$$\text{and } \sup_{Q_X} \beta_0(\gamma^*, Q_X^T \otimes P_\zeta^{*T}) \leq \alpha + \exp(-\Omega(T^{\frac{1}{2}})) \rightarrow \alpha.$$

The proof of Theorem 3 is provided in [25, Appendix C]. The asymptotically optimal decoder deterministically maps a typical

sequence x^T to a typical sequence ζ^T uniquely under different messages M . The corresponding optimal joint distribution output by the encoder f^* assigns probability 1 to such pair of sequences (x^T, ζ^T) , making sure that the detection accuracy is high. Figure 2 illustrates the design using a toy example when $m = 3$.

Remark 1 (Existence of g function and implementations). As $|\mathcal{A}_{\eta,X}^{(T)}| \doteq |\mathcal{A}_{\eta,\zeta}^{(T)}|$, any Latin square operation [28] can serve as a valid g function. Below are two examples.

- **Example 1:** Let $n = |\mathcal{A}_{\eta,X}^{(T)}|$ and index the typical sequences as $\{(x^T)_i\}_{i=1}^n, \{(\zeta^T)_i\}_{i=1}^n$. One can define $g((x^T)_i, (\zeta^T)_{(i+M-2) \bmod n+1}) = M$, for any $i, M \in [n]$, which takes cyclic permutation of $\mathcal{A}_{\eta,\zeta}^{(T)}$ as input (as shown in Figure 2).
- **Example 2:** Consider the typical sets and the message set as finite fields, $\mathcal{A}_{\eta,X}^{(T)} = \mathcal{A}_{\eta,\zeta}^{(T)} = \mathbb{F}_q^T$, for some valid q . One can define $g(x^T, \zeta^T) = (x^T + \zeta^T) \bmod q$.

In general, the optimal design can be implemented by lossless coding schemes where the presence of side information ζ^T ensures that a codeword X^T can be uniquely decoded to one message, e.g., a conditional version of arithmetic coding.

The information rate of this distortion-free (m, T) watermarking scheme is at most

$$R \leq H(Q_X) + \frac{\log \alpha}{T} \xrightarrow{T \rightarrow \infty} H(Q_X),$$

which achieves the maximum rate in Lemma 1 for $d = 0$.

When we allow some distortion $d > 0$ in the watermarking scheme, in Theorem 3, we can change P_X^* to any P_X satisfying $D(P_X^T, Q_X^T) \leq d$, and the one that maximizes the information rate is

$$P_X^* = \arg \max_{P_X : D(P_X^T, Q_X^T) \leq d} H(P_X).$$

When the distortion metric is set as D_{KL} , the optimal P_X^* is the tilting of Q_X as shown in [29, Theorem 1].

Notably, the asymptotic results derived for the i.i.d. case using classical typical set analysis can be extended to the case where X^T, ζ^T are stationary ergodic processes. In this generalization, the entropy $H(P_X^*)$ is replaced by the entropy rate of the stationary ergodic process.

IV. FINITE-LENGTH ANALYSIS

Inspired by the asymptotically optimal design, we are now ready to proceed with our analysis in the finite-length setting under the practical non-i.i.d. scenarios.

We consider the following optimization problem. We aim to minimize the maximum j -th error probability by jointly optimizing the watermarking encoder and decoder, subject to the following constraints: 1) the worst-case false alarm error under control, and 2) the distortion remains bounded:

$$\min_{\gamma, (\mathbb{P}_k)_{k \in [m]}} \max_{j \in [m]} \beta_j(\gamma, \mathbb{P}_j) \quad (\text{P1})$$

$$\text{s.t. } \sup_{Q_{X^T}} \beta_0(\gamma, Q_{X^T} \otimes P_{\zeta^T}) \leq \alpha, \quad D(P_{X^T}, Q_{X^T}) \leq d.$$

Recall that $\mathbb{P}_k := P_{X^T, \zeta^T | M=k}$. Assumption 1 implicitly imposes the constraint that all $(\mathbb{P}_k)_{k \in [m]}$ share the same marginals projected on $\mathcal{X}^T$ and $\mathcal{Z}^T$.

The following theorem characterizes the min-max j -th error probability of this optimization problem—the universal minimax that holds for all watermarking schemes within our framework. The proof is provided in [25, Appendix D].

Theorem 4 (Universal Min-Max j -th Error). *The universal min-max j -th error, denoted by $\beta^*(m, T, \alpha, d)$, from (P1) is*

$$\beta^*(m, T, \alpha, d) = \min_{P_{X^T}: D(P_{X^T}, Q_{X^T}) \leq d} \sum_{x^T} \left(P_{X^T}(x^T) - \frac{\alpha}{m} \right)_+ \quad (1)$$

In the converse proof, we first fix any decoder γ and show a lower bound for the min-max j -th error probability, i.e., (1). The lower bound is independent of γ and thus holds for the optimal value of (P1). In the achievability proof, we present a watermarking scheme that achieves this lower bound, as stated in the next theorem. We observe that $\beta^*(m, T, \alpha, d)$ resembles the universally minimum Type-II error for zero-bit watermarking schemes [24], with α replaced by $\frac{\alpha}{m}$. As the message set size m increases (with α and T fixed), $\beta^*(m, T, \alpha, d)$ increases to 1. Conversely, increasing the false alarm threshold α or the token length T allows for a larger message set m while keeping $\beta^*(m, T, \alpha, d)$ under control. If $m \leq \alpha |\mathcal{X}|^T$, there exist cases where $\beta^*(m, T, \alpha, d) = 0$. Hence, Theorem 4 reveals a fundamental trade-off between detectability, token length and message size. Moreover, as the distortion d increases, the lower bound for β_j^* decreases. This means that a watermarking scheme can trade off text quality for lower detection errors.

The following theorem presents a class of d -distorted (m, T) watermarking schemes for any $m \in \mathcal{M}$ that satisfies Assumption 1 and achieves Theorem 4.

Theorem 5 (Optimal d -Distorted Watermarking Scheme). *Choose $\mathcal{Z}^T \subset \mathbb{Z}^T$ such that $|\mathcal{Z}^T| = |\mathcal{X}^T| + 1$. Randomly pick a redundant sequence $\tilde{\zeta}^T \in \mathcal{Z}^T$. If the message set size $m \leq |\mathcal{X}|^T$, the optimal watermarking encoders and decoders that achieve Theorem 4 exist, given as follows.*

The class of optimal decoders is given by

$$\Gamma_{\tilde{\zeta}^T}^* := \left\{ \gamma \left| \gamma(x^T, \zeta^T) = \begin{cases} h(x^T, \zeta^T), & \text{if } \zeta^T \neq \tilde{\zeta}^T \text{ and} \\ h(x^T, \zeta^T) \leq m, \\ 0, & \text{otherwise,} \end{cases} \right. \right.$$

for some function $h: \mathcal{X}^T \times \mathcal{Z}^T \setminus \{\tilde{\zeta}^T\} \rightarrow [|\mathcal{X}^T|]$ such that $h(x^T, \cdot)$ and $h(\cdot, \zeta_1^T)$ are bijective, given any fixed x^T and ζ_1^T . For any $\gamma^ \in \Gamma_{\tilde{\zeta}^T}^*$, the corresponding optimal encoder f^**

outputs $\mathbb{P}_j^$ as follows: for any $j \in [m]$, $\mathbb{P}_j^*(x^T, \zeta^T) =$*

$$\begin{cases} P_{X^T}^*(x^T) \wedge P_{\zeta^T}^*(\zeta^T), & \text{if } \gamma^*(x^T, \zeta^T) = j; \\ \frac{(P_{X^T}^*(x^T) - P_{\zeta^T}^*(\gamma_j^{*-1}(x^T)))_+ \cdot (P_{\zeta^T}^*(\zeta^T) - P_{X^T}^*(\gamma_j^{*-1}(\zeta^T)))_+}{\beta^*(m, T, \alpha, d)}, & \text{otherwise,} \end{cases} \quad (2)$$

where γ_j^{-1} represents the inverse of γ^* for a fixed $j \in [m]$ ($\gamma_j^{*-1}(\zeta^T) = \emptyset$ in particular),*

$$P_{X^T}^* = \arg \min_{P_{X^T}: D(P_{X^T}, Q_{X^T}) \leq d} \sum_{x^T} \left(P_{X^T}(x^T) - \frac{\alpha}{m} \right)_+, \text{ and}$$

$$P_{\zeta^T}^* = \underbrace{\left(\left(P_{X^T}^*(x^T) \wedge \frac{\alpha}{m} \right)_{x^T \in \mathcal{X}^T} \right)}_{P_{\zeta^T}^*(\zeta^T), \forall \zeta^T \in \mathcal{Z}^T \setminus \{\tilde{\zeta}^T\}} \underbrace{\sum_{x^T \in \mathcal{X}^T} \left(P_{X^T}^*(x^T) - \frac{\alpha}{m} \right)_+}_{P_{\zeta^T}^*(\tilde{\zeta}^T)}.$$

The proof of Theorem 5 is provided in the achievability part in [25, Appendix D]. This optimal watermarking scheme ensures the secrecy of embedded message, i.e., $P_{X^T|M}^* = P_{X^T}^*$ and $P_{\zeta^T|M}^* = P_{\zeta^T}^*$. The construction of the encoder's output distributions $(\mathbb{P}_j^*)_{j \in [m]}$ are adaptive to the original LLM output Q_{X^T} , which can be regarded as an extension of [24, Theorem 2]. It is equivalent to transporting the probability mass from $\mathcal{V}^T$ to $\mathcal{Z}^T$, maximizing $\mathbb{P}_M^*(x^T, \zeta^T)$ for $\gamma(x^T, \zeta^T) = M$, while keeping the worst-case false alarm error below α . Moreover, the introduction of $\tilde{\zeta}^T$ helps to control the worst-case false alarm. If $P_{X^T}^*(x^T) > \frac{\alpha}{m}$ (i.e., low-entropy text), x^T may be mapped to $\tilde{\zeta}^T$ during watermarking, which makes it harder to detect as watermarked. In conclusion, the proposed scheme provides a structured approach to improving the implementation of multi-bit LLM watermarking.

V. DISCUSSION AND FUTURE WORKS

While our theoretical analysis of the distributional information embedding problem does not fully account for all aspects of LLMs (e.g., auto-regressive nature), we believe it provides valuable insights for designing multi-bit watermarking schemes. We rigorously demonstrate that the best achievable rate in the asymptotic regime is determined by the entropy of the text distribution $H(P_X)$, establishing a fundamental limit that serves as a benchmark for evaluating existing multi-bit watermarking schemes.

Moreover, this result implies an inherent connection between the problem of distributional information embedding and lossless compression, where the fundamental limit is also the entropy of the source distribution. Interestingly, [29] proposes a steganography algorithm that exploits this connection by using the decoder of an arithmetic coding scheme¹ as the LLM watermarking encoder to sample from tokens while employing the arithmetic coding encoder as the decoder in our context. This duality between the two problems suggests that new watermarking schemes could be inspired by existing source coding techniques, presenting an intriguing direction for future exploration.

¹The decoder of source coding maps message bit to symbols (tokens), and the encoder maps tokens to message.

REFERENCES

- [1] H. Touvron, T. Lavril, G. Izacard, X. Martinet, M.-A. Lachaux, T. Lacroix, B. Rozière, N. Goyal, E. Hambro, F. Azhar *et al.*, “Llama: Open and efficient foundation language models,” *arXiv preprint arXiv:2302.13971*, 2023.
- [2] A. Q. Jiang, A. Sablayrolles, A. Mensch, C. Bamford, D. S. Chaplot, D. d. I. Casas, F. Bressand, G. Lengyel, G. Lample, L. Saulnier *et al.*, “Mistral 7b,” *arXiv preprint arXiv:2310.06825*, 2023.
- [3] S. Aaronson, “Watermarking of large language models,” <https://simons.berkeley.edu/talks/scott-aaronson-ut-austin-openai-2023-08-17>, 2023, accessed: 2023-08.
- [4] J. Kirchenbauer, J. Geiping, Y. Wen, J. Katz, I. Miers, and T. Goldstein, “A watermark for large language models,” in *International Conference on Machine Learning*. PMLR, 2023, pp. 17 061–17 084.
- [5] R. Kudithipudi, J. Thickstun, T. Hashimoto, and P. Liang, “Robust distortion-free watermarks for language models,” *arXiv preprint arXiv:2307.15593*, 2023.
- [6] X. Zhao, P. Ananth, L. Li, and Y.-X. Wang, “Provable robust watermarking for AI-generated text,” *arXiv preprint arXiv:2306.17439*, 2023.
- [7] Y. Liu and Y. Bu, “Adaptive text watermark for large language models,” in *Forty-first International Conference on Machine Learning*, 2024.
- [8] K. Yoo, W. Ahn, J. Jang, and N. Kwak, “Robust multi-bit natural language watermarking through invariant features,” *arXiv preprint arXiv:2305.01904*, 2023.
- [9] K. Yoo, W. Ahn, and N. Kwak, “Advancing beyond identification: Multi-bit watermark for large language models,” in *Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers)*, 2024, pp. 4031–4055.
- [10] W. Qu, D. Yin, Z. He, W. Zou, T. Tao, J. Jia, and J. Zhang, “Provably robust multi-bit watermarking for AI-generated text via error correction code,” *arXiv preprint arXiv:2401.16820*, 2024.
- [11] X. Zhao, S. Gunn, M. Christ, J. Fairuze, A. Fabrega, N. Carlini, S. Garg, S. Hong, M. Nasr, F. Tramèr *et al.*, “Sok: Watermarking for ai-generated content,” *arXiv preprint arXiv:2411.18479*, 2024.
- [12] P. Moulin and J. A. O’Sullivan, “Information-theoretic analysis of watermarking,” in *2000 IEEE International Conference on Acoustics, Speech, and Signal Processing. Proceedings (Cat. No. 00CH37100)*, vol. 6. IEEE, 2000, pp. 3630–3633.
- [13] N. Merhav, “On random coding error exponents of watermarking systems,” *IEEE Transactions on Information Theory*, vol. 46, no. 2, pp. 420–430, 2000.
- [14] P. Moulin, “The role of information theory in watermarking and its application to image watermarking,” *Signal Processing*, vol. 81, no. 6, pp. 1121–1139, 2001.
- [15] Y. Steinberg and N. Merhav, “Identification in the presence of side information with application to watermarking,” *IEEE Transactions on Information Theory*, vol. 47, no. 4, pp. 1410–1422, 2001.
- [16] A. S. Cohen and A. Lapidot, “The gaussian watermarking game,” *IEEE transactions on Information Theory*, vol. 48, no. 6, pp. 1639–1667, 2002.
- [17] B. Chen and G. W. Wornell, “Quantization index modulation: A class of provably good methods for digital watermarking and information embedding,” *IEEE Transactions on Information theory*, vol. 47, no. 4, pp. 1423–1443, 2001.
- [18] R. J. Barron, B. Chen, and G. W. Wornell, “The duality between information embedding and source coding with side information and some applications,” *IEEE Transactions on Information Theory*, vol. 49, no. 5, pp. 1159–1180, 2003.
- [19] P. Moulin and J. A. O’Sullivan, “Information-theoretic analysis of information hiding,” *IEEE Transactions on information theory*, vol. 49, no. 3, pp. 563–593, 2003.
- [20] J. J. Eggers, R. Bauml, R. Tzschoppe, and B. Girod, “Scalar costa scheme for information embedding,” *IEEE Transactions on signal processing*, vol. 51, no. 4, pp. 1003–1019, 2003.
- [21] Y. Wang and P. Moulin, “Perfectly secure steganography: Capacity, error exponents, and code constructions,” *IEEE Transactions on Information Theory*, vol. 54, no. 6, pp. 2706–2722, 2008.
- [22] J. J. Harmsen and W. A. Pearlman, “Capacity of steganographic channels,” *IEEE Transactions on Information Theory*, vol. 55, no. 4, pp. 1775–1792, 2009.
- [23] M. Costa, “Writing on dirty paper (corresp.),” *IEEE transactions on information theory*, vol. 29, no. 3, pp. 439–441, 1983.
- [24] H. He, Y. Liu, Z. Wang, Y. Mao, and Y. Bu, “Theoretically grounded framework for llm watermarking: A distribution-adaptive approach,” 2025. [Online]. Available: <https://arxiv.org/abs/2410.02890>
- [25] —, “Supplementary material for “Distributional information embedding: A framework for multi-bit watermarking,”,” 2024. [Online]. Available: https://github.com/haiyun-he/PaperAppendices/blob/main/APWDSIT2025_Multibit_WM_Appendix.pdf
- [26] N. Merhav and E. Sabbag, “Optimal watermark embedding and detection strategies under limited detection resources,” *IEEE Transactions on Information Theory*, vol. 54, no. 1, pp. 255–274, 2008.
- [27] T. M. Cover and J. A. Thomas, *Elements of Information Theory (Wiley Series in Telecommunications and Signal Processing)*. USA: Wiley-Interscience, 2006.
- [28] A. D. Keedwell and J. Dénes, *Latin Squares and Their Applications: Latin Squares and Their Applications*. Elsevier, 2015.
- [29] Y.-S. Huang, P. Just, K. Narayanan, and C. Tian, “OD-Stega: LLM-based near-imperceptible steganography via optimized distributions,” 2024. [Online]. Available: <https://arxiv.org/abs/2410.04328>