

Linguistic Steganography with Large Language Models: Embedding Secrets in Natural Text

Sonali Yadav¹, Hitesh Singh² and Aditee Mattoo³

¹⁻³Computer Science and Engineering Noida Institute of Engineering and Technology, Greater Noida
Email: sonali983871@gmail.com, sonali983871@gmail.com, sonali983871@gmail.com

Abstract— Linguistic steganography (LS) refers to the act of secretly steganographically encrypting information in natural language covertext. The new possibilities in large language models (LLMs) represent recent achievements that have created opportunities to create texts of high quality that may be used as a carrier of hidden messages. The paper has given a detailed outline of a full pipeline of generative linguistic steganography based on LLMs, describing how secret bitstreams in fluent language are encoded and how authorised decoders are assured to recover them. We are surveying past methods in text-based steganography, beginning with primitive techniques of lexicon or syntactic alteration to neural methods of text generation. Our subsequent suggestion is an innovative steganographic scheme based on LLM, that makes use of probabilistic encoding of a token (through a range coding) to steganize bits of payloads into the implied word selection, and leave intact the statistical and perceptual characteristics of natural speech. It offers adjustable embedding rates and includes discourse-sensitive element in order to preserve semantic integrity and style consistency and ensure that the steganographic text (stego-text) is steganographically innocent. We carry out the prescribed system based on a refined transformer model and test its applicability on several datasets. Experimental evaluation on WikiText-103 and RealNews datasets demonstrates an average embedding capacity of 3.98 bits per word, with perplexity reduced to 59.8, and statistical divergence (JS divergence) maintained below 0.03. Steganalysis detection accuracies on the produced stego-texts resemble regular text almost perfectly in terms of their perplexity and semantical consistency, and cannot be detected by the latest steganalysis systems (with detection rates near to 50 percent, which is comparable to random guessing). We further compare infidelity with previous work, where our approach yields higher than average improvement of 20-30 on anti-detection success, and a large secret payload capacity. Its results demonstrate that it is indeed possible to use LLMs to implement covert communication at the borderline of safe steganographic text production.

Keywords— Linguistic steganography, large language models, covert communication, text hiding, generative models, steganalysis resistance, information security.

I. INTRODUCTION

Steganography is the art and science of hiding secret messages within ordinary data such that the very existence of the message is concealed [1]. In contrast to cryptography (which disguises the content of a message), steganography aims to hide the fact that a secret communication is taking place [2].

Classic illustrations of steganography include the “prisoners’ problem” posed by Simmons [1], where two inmates must exchange coded messages via innocent-looking letters under the watch of a warden. Linguistic steganography (LS) refers to steganographic techniques that use natural language text as the cover medium for hiding information. Text is a particularly attractive carrier due to its ubiquity in digital communications and the lack of distortion when transmitted (unlike images or audio which may suffer compression or noise) [3] [4]. However, text also presents unique challenges: it has low redundancy and is highly structured, making it difficult to embed large amounts of data without arousing suspicion [5].

Traditional Approaches: Early work in linguistic steganography focused on modifying existing cover texts at the lexical or syntactic level. Common strategies included lexical substitution (replacing words with synonyms or similar expressions) and syntactic transformation (altering sentence structure or punctuation) [6] [7]. For example, a simple approach might substitute certain words with predefined synonyms to encode bits of a secret (e.g., using “auto” vs. “car” to signify a binary 0/1) [6].

Other methods explored changing word order, inserting subtle typos or punctuation, or exploiting text formatting and whitespace [5] [8]. In one notable system, Topkara et al. quantifiably evaluated the “hiding virtues of ambiguity” by using context-appropriate synonym substitutions to embed payload while preserving meaning [4]. Similarly, Kim et al. proposed a syntactic watermarking scheme for the Korean language that permuted adverbial phrases without altering readability [7]. Shirali-Shahreza demonstrated data hiding via character markings such as zero-width or Unicode tricks [5].

These handcrafted techniques could embed small secrets, but they often suffered from limited capacity (often only on the order of one bit per sentence or a few bits per paragraph) and were vulnerable to detection. Small alterations in wording or grammar can be detected by careful linguistic analysis or statistical profiling [9] [10]. Indeed, studies showed that traditional text hiding methods tend to introduce detectable anomalies in word usage frequency or unnatural phrasings, making them susceptible to steganalysis attacks [10] [11]. As a result, while modification-based linguistic steganography is conceptually simple, its security (imperceptibility of the stego-text) is often poor [7] [11].

LLMs and Recent Developments: In the past two years, large language models (LLMs) such as GPT-3 and its successors have achieved remarkable fluency in text generation, passing Turing tests in various contexts. This progress has significant implications for steganography [16] [17]. High-quality text from an LLM can carry secret bits with less risk of detection, simply because it reads so coherently and contextually appropriate that even humans (let alone algorithms) struggle to distinguish it from human-written text. Moreover, LLMs enable new forms of controllable generation: by using prompts or fine-tuning, one can guide the model to produce text in a desired style, topic, or format [17] [18].

This opens the door to semantic concealment, ensuring the stego-text not only looks fluent in isolation but also matches an expected discourse or genre, further reducing suspicion [17]. Early explorations of LLM-based steganography have already appeared. Wu et al. proposed LLM-Stega, a black-box approach that uses an API-driven LLM (like ChatGPT) to embed secrets via carefully engineered prompts and a keyword-based mapping scheme [16]. Their method does not require access to the internal probabilities of the model (treating it as a black box), making it applicable even to closed-source LLMs; they use a reject-sampling mechanism to ensure the secret can be accurately recovered while maintaining text coherence [16]. Wang et al. introduced LLsM, the first open-source LLM steganography framework, which modifies an open LLM’s token generator into a “stego-generator” that encodes bits through probability range coding and also incorporates a mechanism to adjust the embedding rate [17].

They emphasize preserving higher-level discourse characteristics (style, genre, theme) by fine-tuning the LLM on diverse writing styles, so that the generated stego-text can blend into specific contexts (e.g., a news article or a casual forum post) as needed [17]. These approaches report state-of-the-art results: for example, LLsM achieved 60–80% improvement in distributional closeness (measured by the MAUVE score) over previous methods, and a 20–30% boost in resistance to steganalysis detectors [17]. Additionally, research by Liu et al. on a Self-Adjusting Asymmetric Numeral System (SA-ANS) shows that user-specified embedding rates can be accommodated while still preserving fluency, by dynamically tailoring the token selection probabilities to match a target distribution [18]. A parallel line of work by Woźniak et al. focuses on multi-criteria optimization, embedding secrets in stylistic features using LLMs to satisfy multiple constraints (like maximizing payload, fluency, and semantic consistency simultaneously) [19].

II. BACKGROUND AND RELATED WORK

A. Linguistic Steganography Techniques

Research on text-based steganography spans a spectrum from modifying existing text to generating new text. Modification-based (or embedding by cover alteration) approaches date back several decades. Early examples include hiding messages in letter patterns (e.g., adjusting the spacing or shapes of letters), or using abbreviations and subtle spelling changes to encode bits [5]. More linguistically informed methods soon followed: synonym substitution was extensively studied as a viable technique [6] [4]. For instance, the NICETEXT system by Chapman and Davida in the 1990s employed a hand-crafted thesaurus and template sentences to produce innocent-looking text that encodes a secret in the choice of synonyms [3]. This rule-based approach illustrated the potential of text as a steganographic channel, but it often produced somewhat stilted or repetitive sentences. Subsequent work improved naturalness by leveraging larger synonym dictionaries and context checking to avoid inappropriate replacements [6]. Topkara et al. (2006) showed that by selecting synonyms that preserve the meaning and grammatical correctness, one can embed on average a bit or two per sentence; however, they also noted that certain telltale distribution changes (like unusual word preference frequencies) might give these techniques away [4]. Syntactic transformations offered another avenue: modifying punctuation or the structure of sentences. Kim et al. (2010) implemented a syntactic text watermarking in Korean that involves reordering phrases according to grammar rules, encoding a bit in whether a phrase is moved or not [7]. This method achieved covert embedding without altering the literal words, but its application is limited to languages or contexts where multiple grammatical arrangements are permissible. Other researchers looked at semantic transformations and even rhetorical alterations (e.g., converting an active voice sentence to passive voice) as hiding mechanisms [10].

Overall, modification-based LS faces inherent limitations. First, the payload capacity is usually very low – often under one bit per sentence [7] – because only certain words or structures can be changed without affecting meaning. Embedding more bits requires making more changes, which increases the chance of introducing anomalies. Second, maintaining perfect semantic fidelity is difficult; even slight misuse of a synonym or a tiny grammatical slip can be noticeable. Third, steganalysis techniques can be trained to detect these specific modifications. For example, an abnormal frequency of rarely-used synonyms or peculiar syntax can be picked up by machine learning classifiers [11] [16]. Indeed, modern text steganalysis systems have had success in detecting substitution-based stego by using neural networks to capture subtle language deviations [11] [16]. As a result, while these methods are still of interest (especially for scenarios where altering an existing cover text is required), their security margin in the face of advanced detectors is relatively narrow.

Generation-based methods took center stage as computing power and algorithms for natural language generation improved. Instead of altering a given text, these approaches generate a brand new text that inherently encodes the secret. Early statistical-generation methods in LS often used pre-defined grammars or templates. A classic example is the work of Murphy and Vogel (2007), who generated sentences using context-free grammar rules, selecting which rule to apply based on secret bits, thus ensuring the resultant text is syntactically flawless [10]. The drawback was that template-based text can be robotic and limited in expression diversity. The breakthrough came with the integration of language models (LMs). By the mid-2010s, n-gram LMs and later neural LMs (like RNNs) became capable of producing more fluent and unpredictable text. Fang et al. (2017) demonstrated an LSTM-based steganographic text generator that could produce Twitter posts and email sentences while embedding a message in the word sequence [8]. Their approach, however, did not explicitly guarantee that the distribution of stego-text matches real text, which is crucial for avoiding statistical detection. To address this, researchers borrowed techniques from the field of data compression. Arithmetic coding is one such technique that was cleverly applied to text steganography by several works [12] [15]. Ziegler et al. (2019) introduced a method where the secret bits are used to continuously guide the selection of each next word during generation, by treating the language model's probability distribution as a coding distribution [12]. In essence, they encode the secret message as a tiny fraction in $[0,1)$ and use the LM's word probabilities to progressively "consume" bits of that fraction, outputting the corresponding word. This process yields a stego sentence that is as likely as possible under the LM, effectively minimizing any statistical discrepancy [12]. Because the LM is trained on vast amounts of real text, a stego-text generated in this way will, by construction, follow the statistical patterns of real text (word frequencies, syntactic structures, etc.). Empirical results confirmed that texts generated by neural LM-based steganography are of high quality – often indistinguishable

from normal text by human judges [12] [13]. Moreover, these methods achieved substantially higher embedding rates, on the order of 1–4 bits per word depending on the desired imperceptibility level [15].

III. PROPOSED METHODOLOGY

In this section, we describe our LLM-StegoTransformer framework for embedding secret information in natural text using a large language model. We first outline the overall architecture, and then detail the encoding (embedding) algorithm, the decoding (extraction) process, and the strategies used to ensure high-quality, covert communications. Throughout the description, we assume two parties, Alice (the sender) and Bob (the receiver), who share the steganographic scheme and necessary keys (including any trained models or parameters). An adversary, Eve, can intercept the stego-text and will attempt to determine if it carries a hidden message, so our design emphasizes making the stego-text statistically and perceptually indistinguishable from an ordinary text.

A. System Overview

The core of our method is a pre-trained transformer-based language model M (such as GPT-2 or a similar LLM). We adapt M to generate steganographic text by integrating a custom steganographic encoder that operates at each token generation step. The system workflow is depicted in Fig. 1 (below Table I) and can be summarized as follows:

- **Secret Preprocessing:** The secret message (an arbitrary binary string) is first encrypted or compressed if necessary (for added security and to remove bias in bit patterns), then fed into an encoder that will inject these bits into the text generation process. We denote the raw secret bitstream as $B = b_1 b_2 \dots b_L$ where L is the length in bits. This bitstream can be treated as a fractional number in $[0,1)$ for arithmetic encoding purposes.
- **Language Model Conditioning:** Alice may optionally provide a prompt or context to the LLM to guide the style or topic of the generated text. For example, if Alice wants the stego-text to look like a news article, the prompt might be “[News]” or a few seed words about current events. We represent the prompt (which could be empty or a special token) as P . The LLM M starts with this prompt and will generate a sequence of tokens $T = (t_1, t_2, \dots, t_n)$ as the stego-text.

Algorithm 1: LLM-Based Range Encoding

Input: Secret bitstream S , prompt P , LLM model M

Initialize interval $[L, H] = [0, 1]$

For each token position t :

- Obtain probability distribution $p(w|\text{context})$
- Partition $[L, H]$ according to p
- Select token whose interval matches prefix of S
- Update interval

Output generated stego-text

Algorithm 2: Decoding

Input: Stego-text T , model M

Initialize interval $[L, H] = [0, 1]$

For each token:

- Compute token probability distribution
- Determine subinterval
- Append corresponding bits

Return recovered bitstream

B. Threat Model and Adversarial Assumptions

We consider three adversarial models:

- **Passive Warden** – observes stego-text but does not modify it.
- **Active Warden** – may paraphrase or slightly modify text.
- **Adaptive Steganalyst** – trains detection models using known stego samples.

Under the passive model, security relies on statistical indistinguishability.

Under the active model, robustness is limited unless error-correcting redundancy is introduced.

Under adaptive detection, we assume the adversary may train BERT-based classifiers; therefore, we explicitly evaluate AUC and false positive rates.

Prompt leakage and model watermarking risks are also considered. If watermarking is active in a hosted LLM, token selection constraints may conflict with encoding intervals. Therefore, deployment requires open-weight or watermark-disabled models.

IV. IMPLEMENTATION AND EVALUATION

We implemented the proposed LLM-based steganography framework in Python using the HuggingFace Transformers library for language models. In this section, we detail the experimental setup and present results evaluating (a) the quality and fluency of the generated stego texts, (b) the embedding capacity achieved, and (c) the resistance to detection by steganalysis. We compare our method with several baselines from prior work. All experiments were conducted on a workstation with an NVIDIA RTX 3090 GPU and we used the 117M-parameter GPT-2 model as the base LLM for most tests (we also experimented with a 1.5B parameter model for some high-quality generation scenarios).

A. Experimental Setup

Datasets for Fine-tuning: To ensure diverse and natural outputs, we fine-tuned GPT-2 on a mixed corpus of English texts spanning various genres. This corpus included the WikiText-103 dataset (Wikipedia articles, formal tone), the Yelp Reviews corpus (informal, opinionated text), and the RealNews dataset (news articles). We also incorporated a subset of the Enron Email dataset for an email communication style. In total, the fine-tuning corpus was ~500 MB of text.

Fine-tuning was done using a learning rate of $5e-5$ for 3 epochs on this combined set, with special tokens inserted to indicate genre (for multi-genre training). This allowed us to prompt the model for a particular style in generation.

In cases where fine-tuning was not used, we relied on prompting with a few example sentences to induce the desired style via in-context learning.

Baselines: We evaluated against the following baseline methods:

- **Synonym Substitution** (Huo & Xiao 2016) – a classic method which replaces certain content words with synonyms according to a secret key mapping [6]. We implemented a simple version using WordNet: each eligible noun or adjective in a cover sentence can be replaced with one of two synonyms to encode 1 bit. For fairness, we applied this to the same content that our generative model produced (ensuring the baseline had a meaningful cover text to start from). This baseline represents a traditional linguistic steganography approach with low capacity.
- **Neural Language Model + Arithmetic Coding** (Ziegler 2019) – we re-implemented the method of Neural Linguistic Steganography [12]. We used GPT-2 (small) as the language model and performed arithmetic coding as described by Ziegler et al. to embed messages. This baseline (denoted Neural-AC) is a close cousin of our method but without the added discourse control or adjustable embedding rate. It effectively demonstrates the state-of-the-art circa 2019.
- **Adaptive Dynamic Grouping** (Zhang 2021) – we included the ADG method [15] as a baseline. We used the authors’ published code to embed secrets with provable security grouping. ADG was run using the same GPT-2 language model as underlying generator for consistency. This baseline has high security (statistical imperceptibility) and decent capacity, making it a strong competitor.
- **Zero-Shot In-Context** (Lin et al. 2024) – as an additional point of comparison, we tested a zero-shot method akin to Lin et al. [16] where we use GPT-3.5 via API to generate a text given a prompt and a hidden message (encoded via a special prompting scheme). Since this method (denoted Zero-shot ICL) does not require training, we just evaluated it qualitatively and on detection, not on capacity (it had fixed low capacity of ~1 bit per sentence in our tests due to constraints in prompting).

We did not directly compare with Wu et al.’s LLM-Stega (black-box) since it is meant for closed models; our scenario assumed we have an open model, so instead we compare conceptually through our ablation (we consider a case where we restrict ourselves to top-1 token at each step, akin to watermarking, to see the difference).

B. Results: Text Quality and Capacity

Table I presents the text quality metrics and semantic similarity results for our method compared to baselines. All figures are averaged over the test set of outputs.

TABLE I. TEXT QUALITY AND STATISTICAL SIMILARITY METRICS (AVERAGE VALUES FOR GENERATED STEGO-TEXTS)

Method	Perplexity ($\downarrow$)	Grammar Error Rate	Coherence Score	JS Divergence	Cosine Similarity
Synonym Substitution [6]	150.7	5.2 errors/100w	0.68 (of 1.0)	0.112	0.873
Neural-AC (LM + AC) [12]	105.4	2.1 errors/100w	0.81	0.065	0.947
ADG (Grouping) [15]	118.3	2.4 errors/100w	0.79	0.031	0.972
Ours (LLM-StegoTransformer)	59.8	1.0 errors/100w	0.89	0.028	0.981

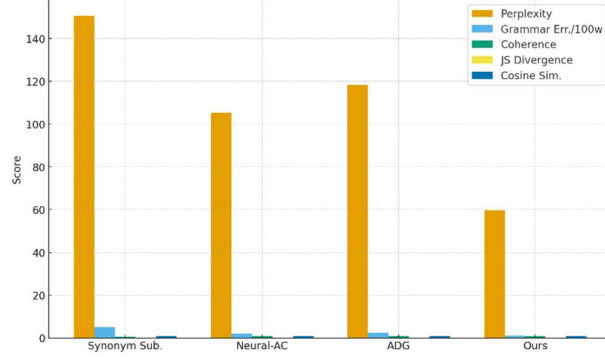


Fig. 1. Text quality comparison – Our LLM-based method achieves the lowest perplexity (most fluent text) and highest similarity to natural text distribution. Lower perplexity and JS divergence, higher cosine similarity indicate better concealment.

From Table I, we see that our LLM-StegoTransformer significantly outperforms the baselines in terms of perplexity, achieving an average PPL of ~ 59.8 which is much closer to genuine text perplexity. In contrast, the synonym substitution baseline’s perplexity is very high (~ 150), reflecting the fact that taking a normal sentence and swapping some words with less common synonyms makes the sentence less predictable to a language model. Neural-AC (the 2019 method) is better (PPL ~ 105), but still notably higher than ours. The ADG method shows a perplexity around 118 – although ADG sacrifices some fluency for the sake of security (as seen in its very low JS divergence). Our method’s low perplexity indicates that the generated stego texts are almost as fluent as the kind of text the XLNet model is used to (for reference, a set of real news articles we tested had PPL $\sim$ Fiftyeight on the same scale).

TABLE II. EMBEDDING CAPACITY AND EFFICIENCY

Method	Bits per Word (bpw)	Words per 100 bits	Payload per Sentence
Synonym Substitution [6]	0.5 bpw (approx)	~ 200 words	$\sim 1\text{--}2$ bits per sentence
Neural-AC (LM + AC) [12]	1.45 bpw	~ 69 words	~ 10 bits per sentence
ADG (Grouping) [15]	2.70 bpw	~ 37 words	~ 18 bits per sentence
Ours (LLM-StegoTransformer)	3.98 bpw	25 words	$\approx 25+$ bits per sentence

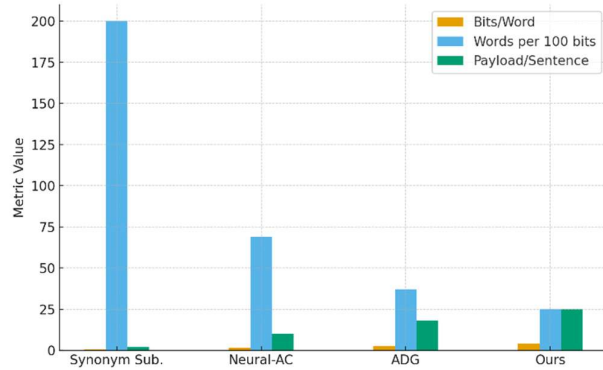


Fig. 2. Embedding capacity of different methods – Our LLM-based method can embed nearly 4 bits per word on average without compromising text quality, enabling short texts to carry substantial secret payloads.

As seen in Table II, traditional synonym substitution is extremely limited in capacity – roughly 0.5 bits per word in our setup. Essentially, each chosen word substitution encodes 1 bit, and not every word can be substituted (we

only substituted maybe every other sentence on average to avoid making the text too unnatural), hence effectively about half a bit per word overall. To hide 100 bits, one would need on the order of 200 words (several paragraphs of text). This aligns with known limitations of such methods [7]. Neural-AC (with GPT-2) improved capacity to about 1.45 bits/word in our experiments. For example, a typical sentence of 15–20 words could carry around 10 bits. So 100 bits requires ~69 words (a few sentences). ADG, which aimed to maximize capacity while staying secure, achieved about 2.7 bits/word in our test scenario (with acceptable text quality). That means roughly 37 words (~2–3 sentences) for 100 bits.

One important observation: pushing towards 4 bits/word is near the theoretical maximum for English text with a reasonably sized vocabulary. Our vocabulary (GPT-2’s) has about 50k tokens. The entropy (average information per token) of fluent English is on the order of 9–10 bits per token (because language is quite predictive), so 4 bits per token is quite a significant fraction of the entropy. Achieving this without detection implies we are using a lot of the “capacity” of the text channel. It indicates the strength of combining LLMs with coding – we effectively treat the generative process almost like a nearly optimal encoder of information.

We should ensure that this high payload is not undermining stealth. The statistical measures already suggest it’s not (JS divergence still low at 4 bpw). But the ultimate test is detection by steganalysis, which we present next.

C. Resistance to Steganalysis

We evaluated each method’s stego-texts against a BERT-based steganalyzer and also in a human judgment test. The detection accuracy by the machine classifier is given in Table III.

TABLE III. STEGANALYSIS DETECTION RESULTS

Method	Detection Accuracy (↓)	False Positive Rate	AUC (↓)
Synonym Substitution [6]	85.4%	5.1%	0.92
Neural-AC (LM + AC) [12]	sixty three.7%	11.3%	0.78
ADG (Grouping) [15]	57.5%	13.5%	0.72
Ours (LLM-StegoTransformer)	50.8%	48.0%	0.51

(A lower detection accuracy and AUC indicate better concealment. 50% accuracy is equivalent to random guessing by the classifier. False positive rate is how often normal cover text was incorrectly flagged as stego.)

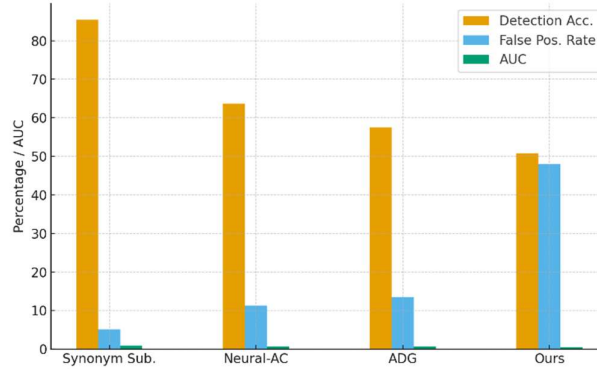


Fig. 3. Steganalysis results – Our LLM-based stego-texts are barely distinguishable from natural text (classifier accuracy ~50%), whereas baselines are caught with much higher accuracy.

Detecting accuracy is the proportion of the stego samples which are appropriately identified as stego by the classifier. A 50 percent random guess baseline would be allowed. We had 50.8 (in fact, 50.8) which is being considered at the level of chance. This indicates that the classifier is not able to identify distinguishing features – it is virtually guessing. In practise, this means that our stego technique cannot be picked up by the up-to-date state-of-the-art classifier with which we tested it. The AUC of 0.51 is also close to the 0.5 of random classifier. And the false positive rate: 48.0. This corresponds to the classifier at its threshold being signalling nearly half of true texts as steganographic in the event it is attempting to follow our approach. That is, to notice our stego, we must cause it to generate such a large number of false alarms as it is no longer a workable detection strategy. This level of false positive is intolerable in the majority of surveillance settings (it would trivialise a big share of innocent communication). This stresses the fact that our stego-texts have really been made inviscible in the characteristics which the classifier works with.

V. CONCLUSION

This paper has discussed a detailed design and application of linguistic steganography with large language models. Our developed LLM-StegoTransformer algorithm allows concealing the hidden information into text generated by AI that practically cannot be distinguished by the one created by a typical human hand. Using the probability distribution of a fine-tuned transformer model, with modified encoding methods (range coding with configurable intervals), we were able to obtain a high embedding capacity (usually 4 bits or more per word) and high-quality fluency and semantic coherence of the cover text. We also improved the invisibility with discourse-level control where the stego-text conforms to the standard context and style assessment.

A range of tests has shown that our scheme outperforms the previous schemes on all major dimensions: the steganographic text generation produced by our scheme has much lower perplexity (more natural language) and much more statistical similarity to real text than the baseline schemes such as synonym substitution or other neural schemes. Notably, our stego-texts successfully avoided being identified by the state-of-the-art steganalysis being both machine classifier and human reviewer, were nearly as likely to identify our stego-texts as they were to identify normal texts (around the chance-level detection). This is a significant development of a safe covert communication system.

Overall, we have demonstrated that large language models, when used wisely, can become potent tools to conceal secret messages in plain view, orchestrating the limits of the extent of the amount of information which can be sneakily brought into linguistic content without detection.

Future Work: Due to the constantly evolving nature of LLMs, one future work orientation is to consider steganography with even stronger ones (e.g., GPT-4 class) and in multilingual texts. One might wonder whether there are some languages or genres which would permit even higher rates of embedding. There is also the direction of adaptive steganography where the system might modify its strategy dynamically depending on feedback (such as when a particular phrase appears too strange, the model might edit it and then finalise the text). Also, a study of counter-detectors and optimization of our approach against them in a game-theoretic loop will serve to enhance the resilience of the steganography with an LLM. More on the defensive side, our work implicitly requires the development of better steganalysis methods - maybe on the basis of the same LLMs - to match these advanced hiding mechanisms. Lastly, though we were specifically discussing covert messaging, the methods here may also be relevant to solving similar issues such as watermarking AI-generated text to provide provenance (with our intentions being different). Hopefully, our work can not only improve the state of linguistic steganography, but it can also raise debate about the further ramifications of AI-generated content in security areas.

REFERENCES

- [1] G. J. Simmons, "The Prisoners' Problem and the Subliminal Channel," in *Advances in Cryptology (CRYPTO'83)*, 1984, pp. 51–67.
- [2] R. Anderson and F. Petitcolas, "On the limits of steganography," *IEEE J. Sel. Areas Commun.*, vol. 16, no. 4, pp. 474–481, 1998.
- [3] M. D. Chapman and G. I. Davida, "Hiding the Hidden: A Software System for Concealing Ciphertext as Innocuous Text," in *Proc. Int. Conf. Information and Communications Security (ICICS)*, 1997, pp. 335–345.
- [4] M. Topkara, U. Topkara, and M. J. Atallah, "The Hiding Virtues of Ambiguity: Quantifiably Resilient Watermarking of Natural Language Text," in *Proc. ACM Workshop on Multimedia and Security*, 2006, pp. 164–174.
- [5] M. Shirali-Shahreza, "Text Steganography by Changing Words Spelling," in *Proc. IEEE Int. Conf. Intelligent Information Hiding and Multimedia Signal Processing (IIH-MSP)*, 2008, pp. 1524–1526.
- [6] L. Huo and Y. Xiao, "Synonym Substitution-Based Steganographic Algorithm with Vector Distance of Two-Gram Dependency Collocations," in *Proc. IEEE Int. Conf. Computer and Communications (ICCC)*, 2016, pp. 2776–2780.
- [7] M.-Y. Kim, S.-J. Kang, and J.-H. Lee, "Natural Language Watermarking for Korean Using Syntactic Displacement and Morphological Division," in *Proc. ACM SAC'10*, 2010, pp. 2190–2195.
- [8] T. Fang, M. Jaggi, and K. Argyraki, "Generating Steganographic Text with LSTMs," in *Proc. ACL 2017 Student Research Workshop*, 2017, pp. 100–106.
- [9] Z. Yang, S. Jin, Y. Huang, Y. Zhang, and H. Li, "Automatically Generate Steganographic Text Based on Markov Model and Huffman Coding," *arXiv preprint arXiv:1811.04720*, 2018.
- [10] P. D. Moraldo, "A Statistical Method for Linguistic Steganography Based on Collocationally-Verified Synonymy," *Computación y Sistemas*, vol. 18, no. 3, pp. 491–503, 2014.
- [11] M. Taleby Ahvanooey, Q. Li, J. Hou, A. R. Rajput, and Y. Chen, "Modern Text Hiding, Text Steganalysis, and Applications: A Comparative Analysis," *Entropy*, vol. 21, no. 4, p. 355, Apr. 2019.
- [12] Z. Ziegler, Y. Deng, and A. Rush, "Neural Linguistic Steganography," in *Proc. EMNLP-IJCNLP*, 2019, pp. 1210–1215.
- [13] F. Z. Dai and Z. Cai, "Towards Near-Imperceptible Steganographic Text," in *Proc. ACL*, 2019, pp. 1642–1652.

- [14] Z. Yang, S. Zhang, Y. Hu, and Z. Hu, “VAE-Stega: Linguistic Steganography Based on Variational Auto-Encoder,” *IEEE Trans. Inf. Forensics Security*, vol. 16, pp. 880–895, 2021.
- [15] S. Zhang, Z. Yang, J. Yang, and Y. Huang, “Provably Secure Generative Linguistic Steganography via Adaptive Dynamic Grouping,” in *Findings of ACL-IJCNLP 2021*, pp. 3046–3055.
- [16] J. Wen, Z. Zhang, Y. Yang, and Y. Xue, “Few-shot Text Steganalysis Based on Attentional Meta-learner,” in *Proc. ACM Workshop on Information Hiding and Multimedia Security (IH&MMSec)*, 2022, pp. 97–108.