

# Generating Fake Cyber Threat Intelligence Using Transformer-Based Models

Priyanka Ranade\*, Aritran Piplai\*, Sudip Mittal†, Anupam Joshi\*, Tim Finin\*,

\*Department of Computer Science & Electrical Engineering, University of Maryland, Baltimore County,

Email: {priyankaranade1, apiplai1, joshi, finin}@umbc.edu

†Department of Computer Science, University of North Carolina, Wilmington,

Email: mittals@uncw.edu

**Abstract**—Cyber-defense systems are being developed to automatically ingest Cyber Threat Intelligence (CTI) that contains semi-structured data and/or text to populate knowledge graphs. A potential risk is that fake CTI can be generated and spread through Open-Source Intelligence (OSINT) communities or on the Web to effect a data poisoning attack on these systems. Adversaries can use fake CTI examples as training input to subvert cyber defense systems, forcing the model to learn incorrect inputs to serve their malicious needs.

In this paper, we automatically generate fake CTI text descriptions using transformers. We show that given an initial prompt sentence, a public language model like GPT-2 with fine-tuning, can generate plausible CTI text with the ability of corrupting cyber-defense systems. We utilize the generated fake CTI text to perform a data poisoning attack on a Cybersecurity Knowledge Graph (CKG) and a cybersecurity corpus. The poisoning attack introduced adverse impacts such as returning incorrect reasoning outputs, representation poisoning, and corruption of other dependent AI-based cyber defense systems. We evaluate with traditional approaches and conduct a human evaluation study with cybersecurity professionals and threat hunters. Based on the study, professional threat hunters were equally likely to consider our fake generated CTI as true.

**Index Terms**—Cybersecurity, Cyber Threat Intelligence, Artificial Intelligence, Data Poisoning Attack

## I. INTRODUCTION

Open-source platforms such as social media, the dark web, security blogs, and news sources play a vital role in providing the cybersecurity community with Cyber Threat Intelligence (CTI). This OSINT based threat intelligence complements sources collected by companies like IBM, Virustotal or Mandiant, by analyzing malware found in the wild, as well as that obtained by the Intelligence community. CTI is information about cybersecurity threats and threat actors that is shared with analysts and systems to help detect and mitigate harmful events. CTI can be shared as text or as semi-structured data with some text fields using formats like Structured Threat Information Expression (STIX) [1] and Malware Information Sharing Platform (MISP) [2]. Recent research has shown how text analysis approaches can be used to transform free text threat information into more structured forms [3]–[10], and even be ingested into defensive systems to enable detection [11].

Although there are many clear benefits to open-source threat intelligence, addressing and handling *misinformation* across

these platforms is a growing concern. The misinformation risk for the security community is the possible dissemination of false CTI by threat actors in an attempt to poison systems that ingest and use the information [12]. In January 2021, Google Threat Analysis Group discovered an ongoing campaign that targets security researchers. Various nation state government-backed threat actors created fake accounts and blog posts with textual cybersecurity information on a variety of exploits in an attempt to divert security researchers from credible CTI sources [13]. There is also additional research that suggests the possibility of future propagation of fake CTI. Maasberg et al. [14] conducted a study of methods in propagating fake cybersecurity news and developed components to categorize it. They did not create fake cyber news, just studied its potential propagation. The widespread generation of fake CTI itself is heavily under-explored, and is a key contribution of this paper.

The widespread propagation of fake CTI primarily impacts cyber analysts who rely on the information to keep up to date with current attack vectors, as well as the cyber defense systems that ingest the information to take correct mitigation steps [11]. Next-generation cyber defense systems are now being developed to automatically ingest and extract data from open source CTI to populate knowledge graphs, that are then used to detect potential attacks or as training data for machine learning systems.

Adversaries can use fake CTI as training input to subvert cyber defense systems. This type of attack is commonly known as a *data poisoning attack* [15]. Many cyber defense systems that rely on this data automatically collect streams of CTI data from common sources. Adversaries can post fake CTI across open sources, infiltrating the training corpus of AI-based cyber defense systems with ease. This fake information will *appear* legitimate to cyber analysts, but will in reality, have false components that contradict the real data. As can be seen from the examples in Table I, convincing fake CTI can be generated that provides incorrect information about the vulnerabilities exploited by an attack, or its consequences. This can cause confusion in analysts on what steps to take to address a threat. In an automated system cyber defense system that is ingesting the CTI, this can also break the reasoning and learning process altogether or force the model to learn incorrect inputs to serve the adversaries' malicious goals. Techniques demonstrated for

open-source CTI can also be applied for covert data, such as proprietary information belonging to a particular company or government entity. In this scenario, potential attack strategies will more than likely be categorized as insider threats, and adversaries will be employees looking to exploit internal systems.

In this paper, we generate realistic fake CTI examples by fine-tuning the public GPT-2 model. Transformer-based methods are state-of-the-art approaches that aid in detecting and generating misinformation on a large scale with minimal human effort [16].

Our generated fake CTI was successfully able to confuse professional threat hunters and led them to label nearly all of the fake CTI as true. We then also use the generated fake CTI examples to demonstrate a data poisoning attack on a Cybersecurity Knowledge Graph (CKG) and a cybersecurity corpus.

Our work makes three main contributions:

- We produce a fine-tuned GPT-2 model that generates fake CTI text (Section III-B),
- We demonstrate a possible poisoning pipeline for infiltrating a CKG (Section IV), and
- We present an evaluation and analysis of the fake and real CTI text (Sections III-C and III-D).

## II. BACKGROUND AND RELATED WORK

In this section, we present a background of transformer architectures and provide related work in the areas of text generation, misinformation, AI-Based cyber systems, knowledge graphs, and adversarial machine learning.

### A. Transformer Models

Encoder-decoder configurations inspired current state-of-the-art language models such as GPT [17] and BERT [18] which utilize the transformer architecture [19]. Similar to Recurrent Neural Network (RNN) based sequence to sequence (Seq2Seq) models, the transformer encoder maps an input sequence into an abstract high dimensional space. The decoder transforms the vector into an output sequence. Unlike its Seq2Seq precursor, the transformer does not utilize any RNN and relies solely on the attention mechanism to generate sequences.

Seq2Seq architectures rely on LSTM cells to process an input sequence one word at a time. In a transformer model, all input words are processed in parallel. Due to this, the transformer introduces the concept of a *positional encoding* in order to capture word ordering information in the  $n$ -dimensional vector of each word. The encoder and decoder components of the transformer also contain a multi-head attention mechanism. This can be shown using the equation below:

$$\underbrace{\text{Attention}(Q, K, V)}_{\text{Queries, Keys, Values}} = \text{softmax} \left( \frac{QK^T}{\sqrt{d_k}} \right) V$$

Where  $Q$  represents queries,  $K$  represents keys, and  $V$  represents values. The complete description of creating these values has been presented by Vaswani et al. [19]. At the start of the encoder, let  $y$  be the initial sentence representation. As it

travels through each layer of the encoder,  $y$  gets updated by different encoder layers. The input  $y$  is utilized to calculate  $Q$ ,  $K$ , and  $V$  in the above equation. Attention is calculated by taking the transpose of the matrix dot product  $QK$  and dividing by the square root of the dimension of the keys  $\sqrt{d_k}$ . Lastly, using the attention weights, we find the weighted sum of values  $V$ . The decoder attention mechanism operates similarly to the encoder, but employs *masked multihead attention*. A linear and softmax layer are also added to produce the output probabilities of each word. In this paper, we focus on the GPT-2 model [20] which exclusively uses decoder blocks.

### B. Transformer based Use-Cases

Generative transformer models have many use-cases such as machine translation [21], question-answering [22] and text summarization [23]. A popular example of a generative transformer model is OpenAI GPT [17]. In recent years, GPT-2 [20] and GPT-3 [24], [25] models have also been developed (At the time of writing this paper, GPT-3 is only accessible by a paywall API, and the model along with its other components are unavailable). GPT models across generations differ from each other in the sizes of data-sets used and number of parameters added. For example, the WebText dataset used to train GPT-2 contains eight million documents.

In this paper, we utilize GPT-2 in our experiments. Unlabeled data is used to *pretrain* an unsupervised GPT model for a generic task. *Fine-tuning* the generic pre-trained models is a common method of extending the architectures for more specific tasks [17]. Lee et al. [26] produced patent claims by fine-tuning the generic pretrained GPT-2 model with U.S. utility patents claims data. Similarly, Feng et al. [27] fine-tuned GPT-2 on a small set of yelp review data-set and used it as a baseline model for various augmentation experiments.

Transformers have been utilized to both detect and generate *misinformation*. Misinformation can be generally categorized as lies, fabricated information, unsupported facts, misunderstandings, and outdated facts and is often used to achieve economic, political, or social gain [28]. Vijjali et al. [29] utilize BERT-based transformers to detect false claims surrounding the COVID-19 pandemic. Similarly, Zellers et al. [30] also use a BERT-based model called Grover, which can detect and generate neural fake news. Their evaluation shows that human beings found machine-generated disinformation more trustworthy than human-written information.

### C. AI-Based Cyber Systems and Knowledge Graphs

Next-generation cyber defense systems use various knowledge representation techniques such as word embeddings and knowledge graphs in order to improve system inference on potential attacks. The use of CTI is an integral component of such systems. Knowledge graphs for cybersecurity have been used before to represent various entities [31]–[33]. Open source CTI has been used to build Cybersecurity Knowledge Graphs (CKG) and other agents to aid cybersecurity analysts working in an organization [3]–[10]. Mittal et al. created Cyber-All-Intel and CyberTwitter [3], [5] which utilizes a variety of

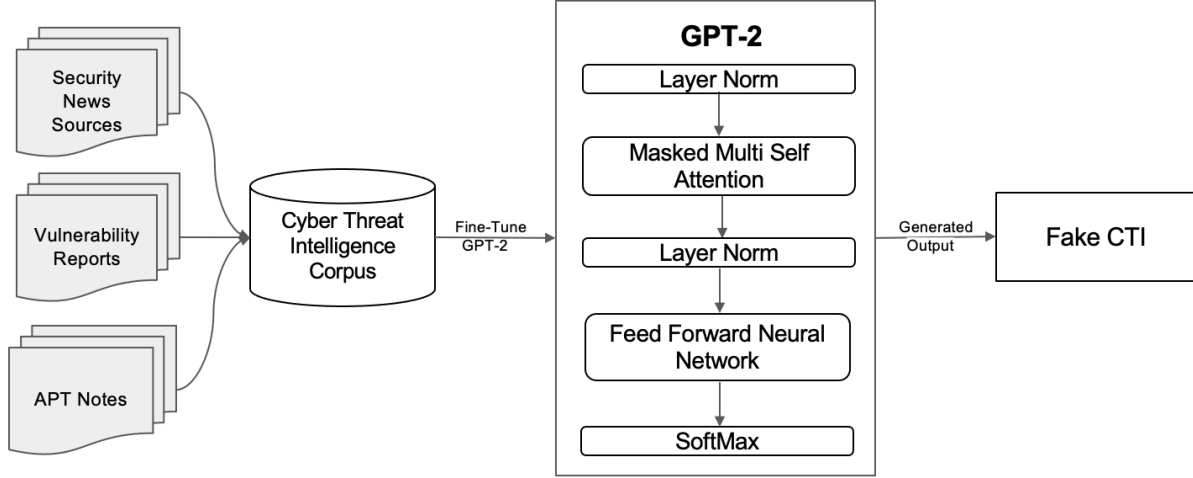


Fig. 1: We collected cybersecurity-related text from several OSINT sources and used it to fine-tune the public GPT-2 model, which generated fake CTI descriptions.

knowledge representations such as a CKG to augment and store CTI.

The use of knowledge graphs for cyber-defense tasks has also been used in malware analysis tasks [34]–[38]. Piplai et al. [32], [39] create a pipeline to extract information from malware after action reports and other unstructured CTI sources and represent that in a CKG. They use this prior knowledge stored in a CKG as input to agents in a reinforcement learning environment [40]. We demonstrate the effects of the poisoning attack, by ingesting fake CTI on CKG using a complete CTI processing pipeline [31], [32].

#### D. Adversarial Machine Learning and Poisoning Attacks

Adversarial machine learning is a technique used to subvert ML systems by providing deceptive inputs to the model. Adversaries use these methods to manipulate AI-based system learning in order to alter protected behavior and serve their own malicious goals [41]. There are several types of adversarial techniques such as evasion, functional extraction, inversion, and poisoning attacks [15]. In this paper, we focus on *data poisoning attack* strategies. Data poisoning attacks are examples of methods that directly compromise the integrity of the learning process of an AI-based system by contaminating the training data-set. These methods rely heavily on the use of synthesized and/or incorrect input data. AI-based cyber defense system can potentially include fake data into their training corpus. The attacker dominates future output by ensuring the system learns fake inputs and performs poorly on actual data.

One example of such an attack is the VirusTotal poisoning attack demonstrated by the McAfee Advanced Threat Research team [42]. This attack compromised several intrusion detection systems that ingest VirusTotal data. The attacker created mutant variants of a ransomware family sample and uploaded

the mutants to the VirusTotal platform. Intrusion detection systems that ingest VirusTotal data classified the mutant files as the particular ransomware family. Similarly, Khurana et al. perform credibility checks on incoming CTI. They develop a reputation score that is used by systems and analysts to evaluate the level of trust for input intelligence data [12]. Duddu et al. survey several methods of using machine learning to model adversary behavior [43].

### III. METHODOLOGY

In this section we describe our fake CTI generation pipeline. Figure 1 presents the overall approach. We begin by creating a cybersecurity corpus in Section III-A. The cybersecurity corpus contains a collection of CTI from a variety of OSINT sources. We then fine-tune the pre-trained GPT-2 model on our cybersecurity corpus (Section III-B). The fine-tuned model allows us to automatically generate large collections of fake CTI samples. We then evaluate our model and describe a poisoning attack against a CTI extraction pipeline.

#### A. Creating a Cybersecurity Corpus

We categorize our CTI collection into three main sources, as shown in Figure 1. We collect security news articles, vulnerability databases, and technical Advanced Persistent Threat (APT) reports. The security news category contains 500 articles from Krebs on Security [44] and CyberWire Daily [45]. The vulnerability reports contain 16,000 Common Vulnerability and Exposures (CVE) records provided by MITRE Corporation and National Vulnerability Database (NVD) from years 2019-2020 [46]. Lastly, we collect 500 technical reports on APTs from the available APTNotes repository [47].

#### B. Fine-Tuning GPT-2 on Cyber Threat Intelligence Data

The original GPT-2 model was trained with the WebText dataset [20]. While the WebText dataset contains some general

| Generation Input                                             | Cyber Threat Intelligence                                                                                                                                                                                                                                                                                                                                                       | True/Fake            |
|--------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----------------------|
| APT41 is a state-sponsored espionage group                   | APT41 is a state-sponsored espionage group which operates against higher education, travel services, and news/media firms which provide some indication that the group also tracks individuals and conducts surveillance.                                                                                                                                                       | True CTI             |
|                                                              | APT41 is a state-sponsored espionage group. The Russian, Moscow-based group is responsible for capturing and stealing computer network communications and intellectual property. APT41 employs social engineering techniques to get around network defenses designed to last for years, and is capable of compromising and stealing dozens of companies                         | Fake CTI (generated) |
| Malicious Domain in SolarWinds Hack Turned into 'Killswitch' | Malicious Domain in SolarWinds Hack Turned into 'Killswitch'. A key malicious domain name used to control potentially thousands of computer systems compromised via the months-long breach at network monitoring software vendor SolarWinds was commandeered by security experts and used as a "killswitch" designed to turn the sprawling cybercrime operation against itself. | True CTI             |
|                                                              | Malicious Domain in SolarWinds Hack Turned into Killswitch Service. This attack took place last year and affected hundreds of organizations. The attackers took advantage of the vulnerability in Win32k framework to attack their targets through crafted code.                                                                                                                | Fake CTI (generated) |
| An issue was discovered in the Quiz and Survey Master plugin | An issue was discovered in the Quiz and Survey Master plugin before 7.0.1 for WordPress. It made it possible for unauthenticated attackers to upload arbitrary files and achieve remote code execution.                                                                                                                                                                         | True CTI             |
|                                                              | An issue was discovered in the Quiz and Survey Master plugin 1.2.1 for WordPress. By setting a CSRF token, an attacker can take advantage of a vulnerability in the survey.php admin panel.                                                                                                                                                                                     | Fake CTI (generated) |

TABLE I: Fake CTI Samples produced by our fine-tuned GPT-2 model.

cybersecurity text, it greatly lacks fine-grained CTI information useful to the security community. To address this problem, we fine-tune the general model with the cybersecurity corpus described above. The diverse CTI sources in our corpus gives the GPT-2 model a variety of examples and the ability to adapt to several aspects of the cybersecurity domain. *Pre-trained* transformer-based language models like GPT-2 are easily adaptable to new domains such as cybersecurity. Instead of training from scratch and initializing with random weights, we initialize the model with pre-trained parameters. We used the publicly released pre-trained GPT-2 model with 117M parameters. The model consists of 12 layers, 786 dimensional states and 12 attention heads.

During our training, we divide the corpus in a 35% train and test split. We set block size as 128, batch size as 64, and learning rate as 0.0001. We utilize the Gaussian Error Linear Unit (GELU) activation function. The GPT-2 architecture shown in Figure 1 consists of normalization layers [48], attention layer, a standard feed forward neural network, and a soft-max layer. The feed forward neural network contains 786\*4 dimensions. We trained the model for twenty three hours (20 epochs) and achieved a perplexity value 35.9. Examples of the generated CTI and more details on our experimentation are given in the next section.

### C. Generating Fake CTI

We use our fine-tuned GPT-2 model to generate fake CTI examples, three of which are shown in Table I. The generation process is initiated with a prompt that is fed as an input to the fine-tuned GPT-2 model (the first column in Table I). The model uses the initial prompt to generate the fake CTI. The generation process is shown in Figure 1. The tokenized prompt is passed through a normalization layer, then through the first block of the attention layer. The block outputs are also passed to a normalization layer and fed to a feed forward neural network, which adds an activation function and dropout. Its output is passed through a softmax layer, which obtains the

positional encoding of the highest probability word inside the vocabulary.

The *first sample* in Table I provides information on APT group APT41. Given the prompt, "APT41 is a state sponsored espionage group", the model was able to form a partially false narrative about APT41. APT41 is a Chinese state-sponsored espionage group, not a Russian group as indicated by the model. Although this is a false fact, the later part of the generated CTI is partially true. Despite some true information, the incorrect nation-state information surrounding APT41 is still present and adds conflicting intelligence if ingested by an AI-based cyber defense system.

In the *second example*, we provide an input prompt from a Krebs on Security article [49]. The model generated fake CTI, which states *kill switch* as an actual service, when in actuality, *kill switch* refers to the method of disconnecting networks from the Internet. In addition, it relates the false service to the *Win32k* framework. This gives the fake CTI enough credibility and seems true to cyber analysts.

Lastly for the *third example*, we provide an input prompt from a 2019 CVE record. The model generated the correct product, but an incorrect associated version and attack type; the true attack was a remote code execution while the generated attack was privilege escalation. While a remote code execution attack can be related to a privilege escalation attack in general, the specific context of using a Cross-Site Request Forgery (CSRF) token to gain access to survey.php is incorrect for this specific product.

### D. Evaluating the generated CTI

We next show that the generated fake CTIs are credible. We use two approaches to show this. First, we evaluate the ability of the fine-tuned model to predict our test data by calculating the perplexity score. Next, we conduct human evaluation studies. The study required a group of cybersecurity professionals and threat hunters to label a collection of generated and actual CTI samples as true or fake. The cybersecurity experience of



the participants range from 2-30 years (in operational settings), with an average experience of 15 years. The idea is to see if professionals in the field can separate real CTI from fake instances generated by our system.

In the context of cybersecurity, human evaluation with potential real-world users of the fake CTI is more indicative than traditional methods such as perplexity scores. The main objective of generating fake CTI is to mislead cyber analysts and bypass intelligence pipelines that they frequently monitor. If the generated CTI does not possess a high range of malformed sentence structure, poor grammar, or incomprehensible text (obvious mistakes indicating the text was produced by a machine), we can assume it has fair potential to appear real to analysts. Perplexity is a common method to determine “uncertainty” in a language model, by assigning probabilities to the test set. Perplexity is measured as the exponentiated average logarithmic loss and ranges from 0-100. The lower the perplexity score, the less uncertainty exists within the model. The base 117M GPT-2 model we fine-tuned has a perplexity score of 24 [26]. We ensure the model is not evaluated on text from the training set by calculating perplexity on a separate test set and achieve a calculated perplexity score of 35.9, showing strong ability of the model to generate plausible text.

In order to evaluate the potential implications of the generated fake CTI in a real world setting, we conduct a study across a group of ten cybersecurity professionals and threat hunters<sup>1</sup>. We provided the participants with an assessment set of both true and fake CTI text samples. Using their own expertise, participants labeled each text sample in the corpus as either true or fake. We created the assessment set by collecting 112 text samples of true CTI drawn from various sources described in Section III-A. We pre-process the text samples by truncating them to the first 500 words and eliminating partial last sentences. We select the first sentence of each sample as an initial prompt to the fine-tuned GPT-2 model and generate a fake CTI example of no more than 500 words. We further divide the 112 samples (56 true CTI and their generated fake counterparts) into two separate annotation sets to ensure true CTI and direct fake counterparts are not part of the same annotation task. Therefore, each annotation task included 28 samples of true text and 28 non-overlapping samples of generated fake data. We randomize the data in each annotation task assigned to the participants.

Participants worked individually, and labeled each of the 56 samples as either true or fake. Participants used their own judgement in labeling each sample, and were prohibited to use external sources like search engines during the assessment. The results of the study are provided in the confusion matrix.

The confusion matrix shows the true positive, false negative, false positive, and true negative rates for 560 CTI samples (including both true and fake data). Of the total 560 samples that were rated, the accuracy (36.8%) was less than chance. The threat hunters predicted 52.5% incorrectly (74 true samples

as false and 220 false statements as true) and 47.5% samples correctly (206 true samples as true and 60 false statements as false). Despite their expertise, the threat hunters were only able to label 60/280 of the generated samples as fake and found the a large majority (78.5%) of the fake samples as true. These results demonstrate the ability of the generated CTI to confuse security experts, and portends trouble if such techniques are widely used.

|             |       | Participant Labels |            |       |
|-------------|-------|--------------------|------------|-------|
|             |       | True               | False      | Total |
| Actual Data | True  | 206 Samples        | 74 Samples | 280   |
|             | False | 220 Samples        | 60 Samples | 280   |
| Total       |       | 426                | 134        |       |

We further investigated the fake samples that were accurately labeled as fake and observed more linguistic errors in the text than in comparison to the fake samples that were labeled as true. Although the majority of the fake CTI contained entities (such as products and attack vectors) that were unrelated to each other, we found if the sentence structure displayed little or no linguistic deficiencies, the data was likely labeled as true. We also noticed sources that lacked substantial context were likely labeled as false.

The generated fake CTI not only has the ability to mislead cybersecurity professionals, but also has the ability to infiltrate cyber defense systems. In the next section, we describe how the generated fake CTI examples can be used to launch a data poisoning attack.

#### IV. DATA POISONING USING FAKE CTI

With the fake CTI examples in Table I we can easily simulate a *data poisoning attack* where the fake CTI is used as training input to subvert knowledge extraction pipelines such as those described by Piplai et al. [32], Mittal et al. [3], [4], Gao et al. [33], [50], and Arnold et al. [10]. Here an attacker can skillfully position fake CTI on multiple OSINT sources like Twitter, Stack Overflow, dark web forums, and blogs.

Many of the systems described above include native crawlers along with cybersecurity concept extractors, entity relationship extractors, and knowledge representation techniques such as word embeddings, tensors, and knowledge graphs. These either use keyword-based methodologies or depend on AI tools to collect and process the CTI. Many of these systems can be easily tricked into including the fake CTI data in a *cybersecurity corpus* along with the true CTI. This is especially possible if the attacker is able to craft the fake CTI in such a way that it “appears very similar” to true CTI. This fake information will then be ingested by a knowledge extraction pipeline utilized to create knowledge representations

<sup>1</sup>Our study protocol was evaluated by UMBC’s IRB and classified as Not Human Subjects Research

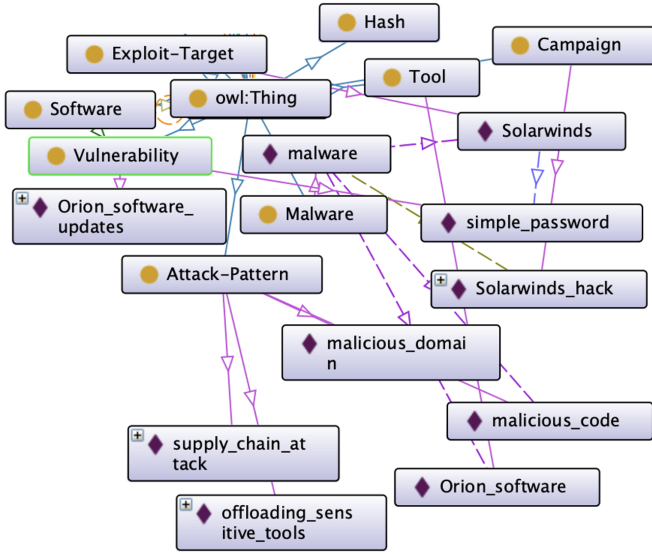


Fig. 2: CKG populated with data from legitimate true CTI sources.

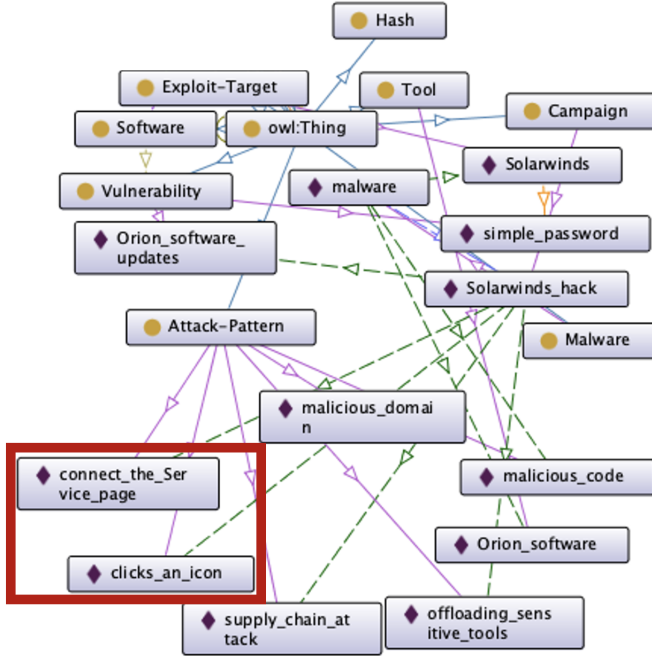


Fig. 3: The poisoned CKG with additional data (red box) extracted from fake CTI.

like, Cybersecurity Knowledge Graphs (CKG). Poisoning a corpus with fake CTI can enable an attacker to contaminate the training data of various AI systems in order to obtain a desired outcome at inference time. With influence over the CTI training data, an attacker can guide the creation of AI models, where an arbitrary input will result in a particular output useful to the attacker.

Next, we describe an attack on a popular knowledge representation technique that involves a CKG [4], [31], [32]. As we already have access to a complete CTI processing

pipeline that outputs a CKG [32], we choose to demonstrate the effects of the poisoning attack on the CKG. Once the fake CTI has been represented in a knowledge representation it can be used to influence other AI systems that depend on these representations. We also discuss the effects of the poisoning attack on the CKG in Section IV-B.

#### A. Processing fake CTI

A CTI ingestion pipeline described in Piplai et al. [32] and similar systems [10], [33], [50] take a CTI source as an input and produces a CKG as an output. The CKG contains cyber entities and their existing relationships. The first stage is a cybersecurity *concept extractor* that takes a CTI and extracts various cyber entities. This is done by using a Named Entity Recognizer (NER) trained on a cybersecurity corpus. The second stage, is a deep-neural network based *relationship extractor* that takes word embeddings of cyber entity pairs as an input and identifies likely relationships. This results in an entity-relationship set that can be asserted into the CKG. As a running example, we use the following *fake CTI* text as input to the extraction pipeline-

*‘Malicious domain in SolarWinds hack turned into killswitch service where the malicious user clicks an icon (i.e., a cross-domain link) to connect the service page to a specific target.’*

When fake CTI is ingested by the pipeline, the cybersecurity concept extractor will output classifications that serve the adversaries’ goals. The concept extractor classifies ‘clicks an icon’, ‘connect the service’ as ‘Attack-Pattern’. It also classifies ‘SolarWinds hack’ as a ‘Campaign’. These entities are extracted from the fake CTI potentially poisoning the CKG.

The relationship extractor while processing the fake CTI above, outputs the following relationships:

- ‘Solarwinds hack’ (Campaign)-uses- ‘clicks an icon’ (Attack-Pattern).
- ‘Solarwinds hack’ (Campaign)- uses - ‘connect the service’ (Attack-Pattern).

The extracted entity relationship set can then be asserted in the CKG. Figures 2 and 3 describe the state of the CKG *before* and *after* asserting knowledge extracted from fake CTI. Figure 2 contains entities and relationships extracted from true CTI samples describing the campaign ‘SolarWinds hack’. We can see entities like ‘Orion Software’, identified as ‘Tool’, and ‘malicious code’ identified as ‘Attack-Pattern’. These entities are used by the malware in the ‘SolarWinds hack’ and are present in the true CTI. We also see ‘simple password’ as a vulnerability. Figure 3 contains additional information extracted from fake CTI generated by our model. These additional entities and relationships have been asserted along with the entity ‘SolarWinds hack’, and are demarcated by the red box. In this figure, we can see additional ‘Attack-Patterns’ like, ‘connect the service page’ and ‘clicks an icon’ being captured in the CKG. These entities have been extracted using the pipeline from the *fake CTI* and are an evidence of

how a poisoned corpus with fake CTI can be ingested and represented in a CKG.

### B. Effects of fake CTI ingestion

The objective of creating a structured knowledge graph from the unstructured CTI text is to aid security professionals in their research. The security professionals can look up past knowledge about cyber incidents, perform reasoning, and retrieve information with the help of queries. However, if generated fake information is ingested by the CKG as part of a data poisoning attack, it can have detrimental impacts such as returning wrong reasoning outputs, bad security alert generation, representation poisoning, model corruption, etc.

For example, if a security professional is interested in knowing which attack campaigns have used ‘click-baits’, they will be misled by the result ‘Solarwinds hack’. As the fake CTI has been ingested and represented in the knowledge representation (See Section IV-A). The following SPARQL [51] query when executed on the CKG,

```
SELECT ?x WHERE {
  ?x a CKG:Campaign;
  CKG:uses CKG:clicks_an_icon.}
```

will result in the following value:

Solarwinds\_hack

If security professionals are interested to know more information about ‘Solarwinds-hack’, they may also receive incorrect information after executing appropriate SPARQL queries.

```
SELECT ?x WHERE {
  ?x a CKG:Attack-Pattern;
  ^CKG:uses CKG:Solarwinds-hack.}
```

This query results in the following values:

malicious\_code, offloading\_sensitive\_tools,  
connect\_the\_service\_page, clicks\_an\_icon

Although we obtained some true results (sourced from true CTI), the presence of fake CTI guided results like, ‘connect the service page’ and ‘clicks an icon’ have the potential to mislead security professionals. Security professionals model cybersecurity attacks and generate network/system detection rules using past available information on the same attacks or similar attacks. They also use these representations to generate alerts for future attacks. For example, a ‘supply chain attack’ exploiting a ‘small password’ vulnerability ‘offloading sensitive tools’ may mean that a new variant of the SolarWinds hack has surfaced. However, if prior knowledge contains fake CTI about the same attack, incorrect alerts can be generated.

Once these knowledge representations are poisoned, additional defense systems can also be adversely impacted by fake cybersecurity information. For example, many of the insights generated by knowledge graphs are useful to other systems like AI-based intrusion detection systems [35], [36], [52], or alert-generators [3], [33], reaching a larger breadth of linked systems and cybersecurity professionals.

## V. CONCLUSION & FUTURE WORK

In this paper, we automatically generated fake CTI text descriptions by fine-tuning the GPT-2 transformer using a

cybersecurity corpus rich in CTI sources. By fine-tuning the GPT-2 transformer with cybersecurity text, we were able to adapt the general model to the cybersecurity domain. Given an initial prompt, the fine-tuned model is able to generate realistic fake CTI text examples. Our evaluation with cybersecurity professionals shows that generated fake CTI could easily mislead cybersecurity experts. We found that cybersecurity professionals and threat hunters labeled the majority of the fake CTI samples as true despite their expertise, showing that they found the fake CTI samples believable.

We use the fake CTI generated by the fine-tuned GPT-2 model to demonstrate a data poisoning attack on a knowledge extraction system that automatically ingests open sourced CTI. We exemplify the impacts of ingesting fake CTI, by comparing the state of the CKG before and after the data poisoning attack. The adverse impacts of these fake CTI sourced assertions include wrong reasoning outputs, representation poisoning, and model corruption.

In ongoing work, we are exploring defences against such data poisoning attacks. One approach is to develop systems that can detect linguistic errors and disfluencies that generative transformers commonly produce, but humans rarely make. A second approach to detecting fake CTI text can use a combination of novelty, consistency, provenance, and trust. CTI sources can be given a score that indicates how much trust the user wants to put in their information.

## ACKNOWLEDGEMENT

This work was supported by a U.S. Department of Defense grant, a gift from IBM research, and National Science Foundation grant #2025685.

We would like to thank various cybersecurity professionals and threat hunters at US defense contractors that took part in our human evaluation study.

## REFERENCES

- [1] Oasis group. Stix 2.0 documentation. <https://oasis-open.github.io/cti-documentation/stix/> May 2013.
- [2] Cynthia Wagner, Alexandre Dulaunoy, Gérard Wagnen, and Andras Iklody. Misp: The design and implementation of a collaborative threat intelligence sharing platform. In *Workshop on Information Sharing and Collaborative Security*, pages 49–56. ACM, 2016.
- [3] Sudip Mittal, Prajit Das, Varish Mulwad, Anupam Joshi, and Tim Finin. Cybertwitter: Using twitter to generate alerts for cybersecurity threats and vulnerabilities. *IEEE/ACM Int. Conf. on Advances in Social Networks Analysis and Mining*, pages 860–867, 2016.
- [4] Sudip Mittal, Anupam Joshi, and Tim Finin. Cyber-all-intel: An AI for security related threat intelligence. *arXiv:1905.02895*, 2019.
- [5] Sudip Mittal, Anupam Joshi, and Tim Finin. Thinking, fast and slow: Combining vector spaces and knowledge graphs. *arXiv:1708.03310*, 2017.
- [6] Lorenzo Neil, Sudip Mittal, and Anupam Joshi. Mining threat intelligence about open-source projects and libraries from code repository issues and bug reports. In *Intelligence and Security Informatics*. IEEE, 2018.
- [7] Priyanka Ranade, Sudip Mittal, Anupam Joshi, and Karuna Joshi. Using deep neural networks to translate multi-lingual threat intelligence. In *International Conference on Intelligence and Security Informatics*, pages 238–243. IEEE, 2018.
- [8] Priyanka Ranade, Sudip Mittal, Anupam Joshi, and Karuna Pande Joshi. Understanding multi-lingual threat intelligence for AI based cyber-defense systems. In *IEEE International Symposium on Technologies for Homeland Security*, 2018.



- [9] Sagar Samtani, Hongyi Zhu, and Hsinchun Chen. Proactively identifying emerging hacker threats from the dark web: A diachronic graph embedding framework (d-gef). *ACM Transactions on Privacy and Security (TOPS)*, 23(4):1–33, 2020.
- [10] Nolan Arnold, Mohammadreza Ebrahimi, Ning Zhang, Ben Lazarine, Mark Patton, Hsinchun Chen, and Sagar Samtani. Dark-net ecosystem cyber-threat intelligence (cti) tool. In *International Conference on Intelligence and Security Informatics*, pages 92–97. IEEE, 2019.
- [11] Sandeep Narayanan, Ashwini Ganesan, Karuna Joshi, Tim Oates, Anupam Joshi, and Tim Finin. Early detection of cybersecurity threats using collaborative cognition. In *2018 IEEE 4th International Conference on Collaboration and Internet Computing (CIC)*, pages 354–363, 2018.
- [12] Nitika Khurana, Sudip Mittal, Aritran Piplai, and Anupam Joshi. Preventing poisoning attacks on AI based threat intelligence systems. In *29th Int. Workshop on Machine Learning for Signal Processing*, pages 1–6. IEEE, 2019.
- [13] Google Threat Analysis Group. New campaign targeting security researchers. <https://blog.google/threat-analysis-group/new-campaign-targeting-security-researchers/>, 2021.
- [14] Michele Maasberg, Emmanuel Ayaburi, Charles Liu, and Yoris Au. Exploring the propagation of fake cyber news: An experimental approach. In *51st Hawaii International Conference on System Sciences*, 2018.
- [15] Yevgeniy Vorobeychik and Murat Kantarcioglu. Adversarial machine learning. *Synthesis Lectures on Artificial Intelligence and Machine Learning*, 12(3):1–169, 2018.
- [16] Aditya Grover and Jure Leskovec. node2vec: Scalable feature learning for networks. In *22nd ACM SIGKDD international conference on Knowledge discovery and data mining*, pages 855–864, 2016.
- [17] Alec Radford, Karthik Narasimhan, Tim Salimans, and Ilya Sutskever. Improving language understanding by generative pre-training. Technical report, OpenAI, 2018.
- [18] Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. Bert: Pre-training of deep bidirectional transformers for language understanding. *arXiv:1810.04805*, 2018.
- [19] Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. Attention is all you need. In *Advances in neural information processing systems*, pages 5998–6008, 2017.
- [20] Alec Radford, Jeffrey Wu, Rewon Child, David Luan, Dario Amodei, and Ilya Sutskever. Language models are unsupervised multitask learners. *OpenAI blog*, 1(8):9, 2019.
- [21] Qiang Wang, Bei Li, Tong Xiao, Jingbo Zhu, Changliang Li, Derek F Wong, and Lidia S Chao. Learning deep transformer models for machine translation. *arXiv:1906.01787*, 2019.
- [22] Taihua Shao, Yupu Guo, Honghui Chen, and Zepeng Hao. Transformer-based neural network for answer selection in question answering. *IEEE Access*, 7:26146–26156, 2019.
- [23] Yang Liu and Mirella Lapata. Text summarization with pretrained encoders. In *Conf. on Empirical Methods in Natural Language Processing and the 9th Int. Joint Conf. on Natural Language Processing*, pages 3721–3731. ACL, 2019.
- [24] Tom B Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, and Amanda Askell. Language models are few-shot learners. *arXiv:2005.14165*, 2020.
- [25] OpenAI. Open AI API. <https://openai.com/blog/openai-api/> 2021.
- [26] Jieh-Sheng Lee and Jieh Hsiang. Patent claim generation by fine-tuning OpenAI GPT-2. *arXiv:1907.02052*, 2019.
- [27] Steven Y Feng, Varun Gangal, Dongyeop Kang, Teruko Mitamura, and Eduard Hovy. Genaug: Data augmentation for finetuning text generators. In *Deep Learning Inside Out (DeeLIO): The First Workshop on Knowledge Extraction and Integration for Deep Learning Architectures*, pages 29–42, 2020.
- [28] Michela Del Vicario, Alessandro Bessi, Fabiana Zollo, Fabio Petroni, Antonio Scala, Guido Caldarelli, H Eugene Stanley, and Walter Quattrociocchi. The spreading of misinformation online. *Proceedings of the National Academy of Sciences*, 113(3):554–559, 2016.
- [29] Rutvik Vijjali, Prathyush Potluri, Siddharth Kumar, and Sundeep Teki. Two stage transformer model for COVID-19 fake news detection and fact checking. *arXiv:2011.13253*, 2020.
- [30] Rowan Zellers, Ari Holtzman, Hannah Rashkin, Yonatan Bisk, Ali Farhadi, Franziska Roesner, and Yejin Choi. Defending against neural fake news. In *Advances in neural information processing systems*, pages 9054–9065, 2019.
- [31] Aditya Pingle, Aritran Piplai, Sudip Mittal, Anupam Joshi, James Holt, and Richard Zak. Relext: Relation extraction using deep learning approaches for cybersecurity knowledge graph improvement. *IEEE/ACM International Conference on Advances in Social Networks Analysis and Mining*, 2019.
- [32] Aritran Piplai, Sudip Mittal, Anupam Joshi, Tim Finin, James Holt, and Richard Zak. Creating cybersecurity knowledge graphs from malware after action reports. *IEEE Access*, 8:211691–211703, 2020.
- [33] Peng Gao, Xiaoyuan Liu, Edward Choi, Bhavna Soman, Chinmaya Mishra, Kate Farris, and Dawn Song. A system for automated open-source threat intelligence gathering and management. *arXiv preprint arXiv:2101.07769*, 2021.
- [34] Jing Liu, Yuan Wang, and Yongjun Wang. The similarity analysis of malicious software. In *Int. Conf. on Data Science in Cyberspace*. IEEE, 2016.
- [35] Younghee Park, Douglas Reeves, Vikram Mulukutla, and Balaji Sundarav. Fast malware classification by automated behavioral graph matching. In *6th Annual Workshop on Cyber Security and Information Intelligence Research*. ACM, 2010.
- [36] Blake Anderson, Daniel Quist, Joshua Neil, Curtis Storlie, and Terran Lane. Graph-based malware detection using dynamic analysis. *Journal in Computer Virology*, 7(1):247–258, 2011.
- [37] Karuna P Joshi, Aditi Gupta, Sudip Mittal, Claudia Pearce, Anupam Joshi, and Tim Finin. Alda: Cognitive assistant for legal document analytics. In *AAAI Fall Symposium*, 2016.
- [38] Maithilee Joshi, Sudip Mittal, Karuna P Joshi, and Tim Finin. Semantically rich, oblivious access control using ABAC for secure cloud storage. In *Int. Conf. on edge computing*, pages 142–149. IEEE, 2017.
- [39] Aritran Piplai, Sudip Mittal, Mahmoud Abdelsalam, Maanak Gupta, Anupam Joshi, and Tim Finin. Knowledge enrichment by fusing representations for malware threat intelligence and behavior. In *International Conference on Intelligence and Security Informatics*. IEEE, 2020.
- [40] Aritran Piplai, Priyanka Ranade, Anantaa Kotal, Sudip Mittal, Sandeep Narayanan, and Anupam Joshi. Using Knowledge Graphs and Reinforcement Learning for Malware Analysis. In *4th International Workshop on Big Data Analytics for Cyber Intelligence and Defense, IEEE International Conference on Big Data*. IEEE, December 2020.
- [41] Anthony D Joseph, Blaine Nelson, Benjamin IP Rubinstein, and JD Tygar. *Adversarial Machine Learning*. Cambridge University Press, 2019.
- [42] MITRE. Virus Total Data Poisoning Case Studies. <http://github.com/mitre/advmthreatmatrix/blob/master/pages/case-studies-page.md#virustotal-poisoning>, 2021.
- [43] Vasisht Duddu. A survey of adversarial machine learning in cyber warfare. *Defence Science Journal*, 68(4), 2018.
- [44] Brian Krebs. Krebs on security. <https://krebsonsecurity.com/> 2021.
- [45] Cyberwire. The CyberWire. <https://thecyberwire.com/> 2021.
- [46] Harold Booth, Doug Rike, and Gregory Witte. The national vulnerability database (nvd): Overview. Technical report, National Institute of Standards and Technology, 2013.
- [47] aptnotes. APTnotes repository. <https://github.com/aptnotes/data> 2021.
- [48] Jimmy Lei Ba, Jamie Ryan Kiros, and Geoffrey E Hinton. Layer normalization. *stat*, 1050:21, 2016.
- [49] Brian Krebs. Malicious Domain in Solarwinds Hack turned into killswitch. <https://krebsonsecurity.com/2020/12/malicious-domain-in-solarwinds-hack-turned-into-killswitch/>, 2021.
- [50] Peng Gao, Fei Shao, Xiaoyuan Liu, Xusheng Xiao, Haoyuan Liu, Zheng Qin, Fengyuan Xu, Prateek Mittal, Sanjeev R Kulkarni, and Dawn Song. A system for efficiently hunting for cyber threats in computer systems using threat intelligence. *arXiv preprint arXiv:2101.06761*, 2021.
- [51] W3. Sparql query language. <https://www.w3.org/TR/rdf-sparql-query/>
- [52] Gulshan Kumar, Krishan Kumar, and Monika Sachdeva. The use of artificial intelligence based techniques for intrusion detection: a review. *Artificial Intelligence Review*, 34(4):369–387, 2010.