

Fake News Detection: Misleading Headlines and Satire

Andrew Wang¹, Phillip Guo², Ethan Wong³

¹andrewxwg@gmail.com

Lower Moreland High School

555 Red Lion Rd, Huntingdon Valley, Pennsylvania, USA 19006

²philliphguo@gmail.com

Montgomery Blair High School

51 University Blvd E, Silver Spring, Maryland, USA 20901

³23ethanw@students.harker.org

The Harker School

500 Saratoga Ave, San Jose, California, USA 95129

Abstract— During recent years, the spread of fake news and incorrect information across the web has been an increasingly severe problem. Since most social media platforms (Twitter included) allow their users to post text and pictures freely, spreading fake news is extremely easy and those who spread fake news are rarely punished. Due to the absence of systems upholding the integrity of social media posts, social media applications treat these posts the same as they would a post with correct information. In addition, misleading headlines can have a dramatic impact on the spread of fake news; headlines can be exaggerated and misleading to a viewer. Experiments have shown that misleading headlines can impact a reader's memory, which further influence's a reader's opinion. Unlike fake news, however, satirical news is used in an entertaining way and is not meant to deceive its viewers. Satirical news is a form of media used to criticize and mock an individual, idea, or a topic. Although satirical news' intention is not to manipulate one's opinion, studies have shown that it indeed can.

Keywords— fake news detection, machine learning, misleading headline, satire, social media

Misleading headlines are prevalent in fake news, especially on social media, because they can trick users into clicking articles and engaging with the fake content. Even outside of fake news, misleading headlines can still cause a variety of problems by altering how users approach an article's subject matter and influencing users who only read the headline without the full article.

Like fake news, satire can be misleading. Instead of deliberately providing users with misinformation, however, Satire news aims to entertain and criticize. Although satire's intention is completely different than fake news' intention, the two can be easily confused. Therefore, satirical detection has been developed to help detect for satire in media.

I. INTRODUCTION

Social media companies prioritize on keeping viewer retention high on their platforms. Social media platforms use algorithms to provide users with desired content that would ensure that they would continue using their platform. While algorithms help generate content that interests a user, fake news posts and media can be easily spread due to the use of the algorithm.

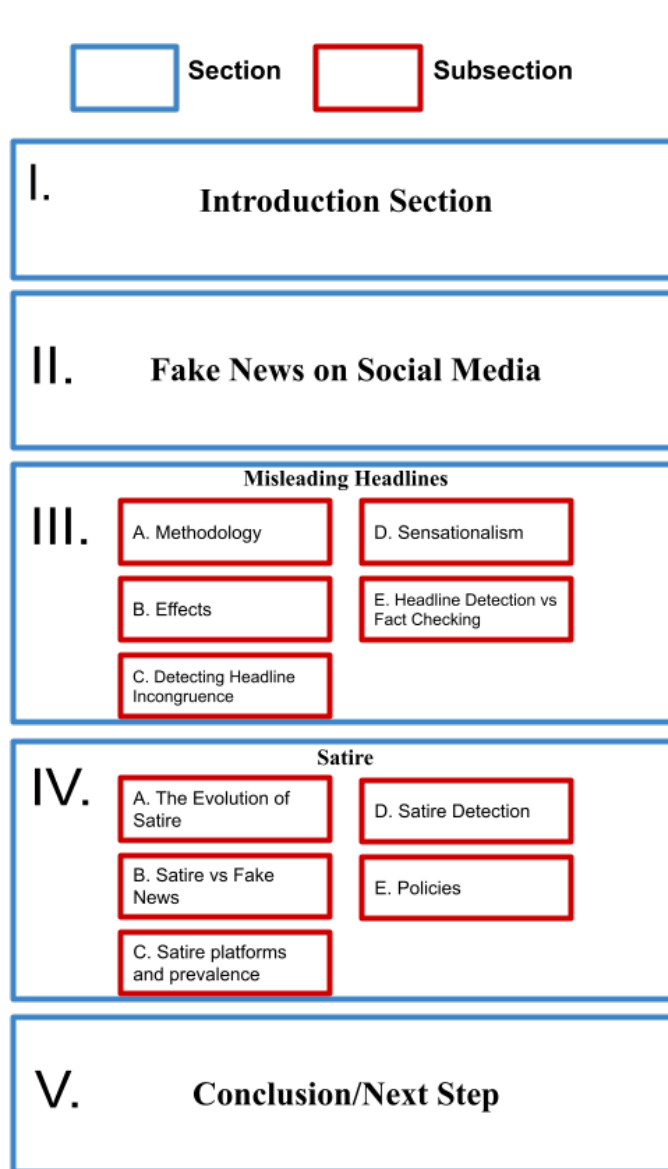


Fig. 1 Site map of the paper

II. FAKE NEWS ON SOCIAL MEDIA

The foundation of how free social media platforms such as TikTok, Facebook, and Twitter profit is by running ads through their platforms. Separate companies who want to boost their interactions frequently use ads put on social media because of the large number of users scrolling through social media every day. In consequence, for social media companies to stay reputable and popular enough for outside businesses to invest in ads on their platform, companies need a method to ensure that users use their platform frequently: the posts that users scroll through must be relevant and interesting. For example, the main way TikTok users interact with the platform is by scrolling

through the “For you” page, a collection of unique user-created videos given to a specific user of the platform that recommends posts that they will enjoy. To filter the millions of different videos on to the “For you” page, TikTok sorts through videos by elements such as the number of interactions, tags, captions, language, country, sound, and even demographic to filter the video into categories that a certain group of users may like. Other platforms like Facebook, Twitter, and Instagram are also similar with their “Explore pages.”

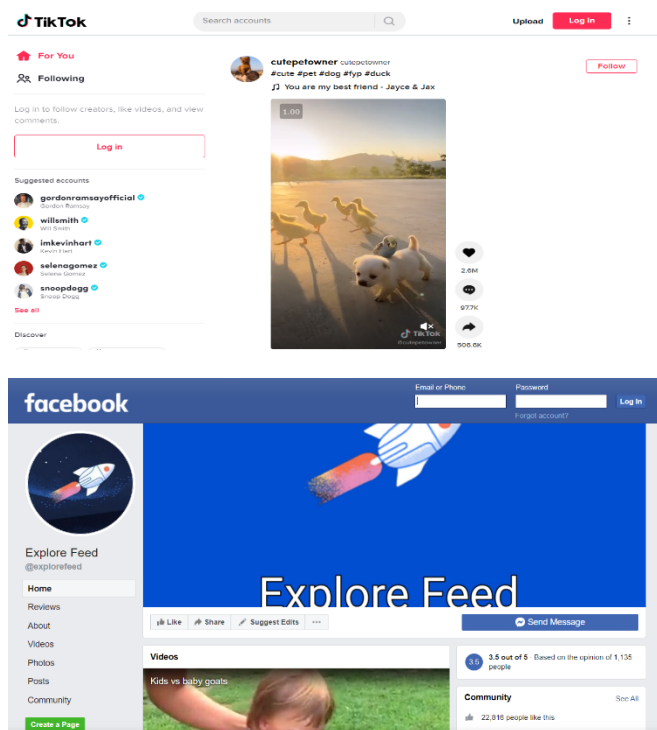


Fig. 2 TikTok's "For You" Page and Facebook's "Explore" Feed [1], [2]

Even regarding topics that do not have a large impact on society, it is easy to see posts on social media that spread misinformation and disinformation that are fueled through an incorrect opinion [3]. Other than real life users putting factually incorrect content in their posts, fake news can also be posted on social media in large amounts by AI bots. Groups on social media companies such as the "Troll Army" on Twitter create bot accounts that are programmed to post disinformation [4]. The programmable aspect of these fake social media accounts allows them to post a large amount of fake news, increasing the probability that a real user will see the post and be affected by it. Other than those two scenarios, there is another minor way for disinformation or misinformation to be posted, which is by a user intentionally saying an incorrect statement for their own personal gain. However, this method (although still dangerous) is not as efficient in spreading fake news as the troll army as they can be programmed to post the same ideas in more numbers. Fortunately, these bots can be countered through fact-checking sources and warnings; Twitter, for example, has taken these measures.

III. MISLEADING HEADLINES

Misleading headlines have been a powerful force driving the spread of fake news, as their usage in trending news articles causes two-fold harm [5]. For one, previous research has demonstrated that a majority of users on different social media

platforms never click on the articles they share, with the number as high as 59% on Twitter [6]. The headline then becomes the only exposure users have to the content of the article, and misleadingly exaggerated or false headlines are then able to reach and influence a large number of people regardless of the actual content.

Second, even for those who do read the articles they share, a biased or slanted headline will affect the way readers perceive the rest of the article's content. Experiments have indicated that misleading headlines can have an impact on a reader's memory, which can influence a reader's opinion and impressions [7]. The researchers found that this was because readers used the initial information from the title to frame the rest of the article, constraining the information they processed to conform with the title, as well as struggling to update memories of misinformation with new corrections.

Currently, to combat the spread of fake news, social media platforms such as Twitter have taken measures to provide corrections and warnings about misleading content and link to fact-checking sources [8]. Previous research has demonstrated that corrections are arguably limited once people have already heard misinformation, with some corrections even backfiring and strengthening the misinformed opinions [9], [10].

Misleading headline detection provides a method of combating the spread of fake news that could be, in many cases, preferable to corrections and warnings. Such a detection system would not actively contradict viewpoints expressed in an article, which is what the majority of resistance to corrections stems from. Instead, it would help users determine whether an article is worth reading and sharing without too much effort on the users' part. This could help inform people and reduce the chance that they share an article with a misleading headline without reading it. This trend has been demonstrated in prior research: in a study, researchers provided a focus group with a headline detection system, finding that 71% of subjects reported a change in how they chose to read articles based on the detection system [11].

A. Methodology

This study made use of Google Scholar and keywords to compile sources.

The following research questions were formed:

1. What are the effects of misleading headlines?
2. How is headline incongruity automatically detected?
3. How is clickbait automatically detected?
4. How is sensationalism detected?
5. How does headline detection compare to other forms of fake-news prevention?
6. How can headlines be automatically generated?

Research questions 1 and 4 were selected to better give readers an overview of the importance of the subject. Research questions 2 and 3 were selected to answer how different forms of misleading headlines are detected, and they were based on (Chesney 2017).

The following table shows the keywords that were used in Google Scholar to find sources for each research question. After finding our initial sources, we read through them and included referenced papers that were also relevant.

TABLE 1
KEYWORDS USED TO FIND SOURCES

Effects	“effects”, “misleading headlines”, “psychological”, “incongruity”, “clickbait”, “social media”, “corrections”
Incongruity	“misleading headlines”, “incongruity”, “model detection”
Clickbait	“clickbait”, “model detection”
Headline Detection Interface	“headline incongruence”, “Clickbait”, “Sensationalism”, “Detection”, “interface”, “extension”
Headline detection vs other forms of fact-checking	“Effectiveness”, “Headline Detection”, “Backfire effect”, “misinformation correction”,
Generation of neutral headlines	“Headline Generation”, “Model”

In cases where there were a limited number of articles that were found via the keywords, Google Scholar’s “Related articles” feature was used.

B. Effects

Headlines serve an important role for all media users. For those who create media, headlines are crucial as a first impression to the content of the article. In addition, headlines allow readers to quickly decide whether an article is worth reading, which is extremely important in today’s oversaturated informational environment.

In recent years, with the advent of social media and the ease of widespread communication, people have been forced to alter the way they consume news. Since there are more news sources and no physical limits on how long articles can be, users have inevitably placed a greater emphasis on headlines when deciding which articles to read or buy [12].

While misleading headlines have been prevalent throughout history, the increased importance of headlines and competition from the burgeoning news industry has invariably led headlines to often be more aggressive, exaggerated, and misleading [12]. This emphasis goes beyond just choosing what to read. For instance, an experimental study found that headlines affect what information readers pay attention to and what inferences they make [13]. Furthermore, experimental studies have also indicated that “misleading headlines affect readers’ memory for news articles or their inferential reasoning, and even their impressions of faces featured in the articles”, due to a variety of cognitive reactions to headlines [7]. Misleading headlines can bias readers to process later information they read in such a way that supports the headline.

The problem of misleading headlines is exacerbated by social media: 55% of US adults get news from social media sometimes or often now, up from 47% from 2018 [14]. On social media, anyone can share news with their own headlines which can accidentally or maliciously misconstrue the article. We have been unable to find studies or quantification for how misleading headlines are on social media, and there is a potential for research in this area.

On social media, users often forgo even reading the articles they see in favor of just looking at headlines: researchers found that 59% of URLs shared on Twitter were never clicked [4]. When misleading headlines are read but their articles are not, readers are undoubtedly misconceived [7].

C. Detecting Headline Incongruence

Incongruent headlines refer to a specific subset of fake news and are “headlines which do not accurately represent the information contained within the article they occur with,” [15]. Subsequent research in incongruence is based on the problem specified by (Chesney 2017) but uses the term “incongruent headlines” more broadly. In (Yoon 2018, Mishra 2020), incongruent headlines describe headlines that range from unrelated to completely contradictory to the article matter, which is larger than the scope of (Chesney 2017) but still separate from other forms of headline misinformation such as clickbait [16], [17].

1) Example

Headline: Air pollution now leading cause of lung cancer

Evidence within article: “We now know that outdoor air pollution is not only a major risk to health in general, but also a leading environmental cause of cancer deaths.” Dr. Kurt Straif, of IARC (Ecker 2014).

This is an example of headline incongruence because of the difference between the perceived meaning of the headline; namely, that air pollution is the number one cause of lung cancer, and the actual body, which only states that air pollution is one of the presumably several causes (Chesney 2017).

2) *Datasets and Preprocessing*

Datasets of existing headline-body pairs of real articles that are labeled with an incongruity rating are crucial when creating a detection model. Existing test data is needed to test the accuracy of a model and how applicable it is to the real world. Large-scale testing helps researchers notice trends in accuracy - a model could work well for certain kinds of articles but not for others, which is hard to diagnose without having many articles of every type. In addition, machine learning (ML) approaches require large datasets for training purposes.

In (Chesney 2017), the authors went through a number of existing headline-body datasets that were labeled with some kind of misinformation rating [15]. They explained that the datasets were inadequate for the incongruent headline problem because their labels did not match their precise definition of incongruence. However, since further research has expanded the scope of the problem, some of the datasets could still be used. Specifically, the 2017 Fake News Challenge provided a dataset of 49972 headline-article pairs that were labeled with one of four stances: the article and headline agreed, disagreed, were unrelated, or the article discussed the headline but did not take a position [18]. The disagree and unrelated stances could be taken as incongruous, as they fit the expanded usage of the term.

Researchers working on the incongruency problem have come up with ways of automatically labeling headline-body data to create new labeled datasets.

In (Yoon 2018), researchers took a corpus of 4 million South Korean and 120 thousand US news articles and made a number of them incongruent by implanting random and likely unrelated content into the body [16]. Then, in (Yoon 2021), the same group of researchers further developed their approach by selecting news stories from trustworthy media outlets to help ensure that the articles were congruent [19]. They generated two datasets: one with the same random implanting method as before and one with implanting from similar news stories. They also checked if two articles were similar by measuring the Euclidean distance of the vector representations of the headlines via fastText, a Python library (see the Approaches subsection).

Because the datasets were generated automatically, they required preprocessing to improve accuracy. (Chesney 2017) notes that it is difficult to train models based on entire headline-body pairs because the body is much longer and more linguistically complex than headlines. This is supported by other research that finds that long word sequences fed into models degrade accuracy [20].

(Chesney 2017) suggests methods to shorten the bodies without compromising too much important information. Key quotes or claims could be extracted from the body using existing NLP approaches and substituted for the entire body, which would greatly reduce the length while making sure that the fundamental ideas were kept.

In (Yoon 2018), researchers split the body of an article into individual paragraphs and used headline-paragraphs as data, which had the additional benefit of increasing the total size of the dataset.

3) *Approaches*

A technique which is used in many NLP approaches is word vectorization - encoding a word as a multidimensional vector of numbers that can be used to quantify relationships with different words. Since word vectors must be able to capture complex semantic connections that exist among all words, ML models have been used to vectorize entire languages (Goldberg 2014).

The fundamental task of automatic incongruence detection is to analyze the relationships between a headline and the article that goes with it [15]. Research in the detection of incongruent headlines uses a variety of semantic and machine learning approaches to solve this task. (Chesney 2017) suggests several approaches based on related literature:

- Researchers could extract key quotes or claims from articles to compare to the headline using existing natural language processing (NLP).
- They could automatically generate an accurate headline (see section 6), then measure how different the potentially incongruent is from the accurate headline.
- They could use argument analysis to determine if the arguments made in potentially incongruent headlines are supported by claims in the body.
- They could use stance detection to check if an article agrees or disagrees with a headline.

In (Mishra 2020), researchers test the headline generation-comparison technique. They use a generative adversarial network to generate a synthetic headline that accurately represents the content of the article [21]. Notably, they use stylized headline generation that results in synthetic headlines being stylistically similar to clickbait and incongruent headlines while still being accurate, in order to more closely compare the headlines.

Once the synthetic headline is generated, they compute a similarity score between the two headlines. They take every word pair of one from the original headline and one from the new headline and calculate the vector difference. They multiply the vector distance by a weight matrix, which is found via ML, and calculate a score matrix that yields the similarity score. They achieve accuracy of .735 and .747 on two datasets [17]. (Mishra 2020) also tries a method that uses cross and dot products along with the vector difference, and they receive slightly better results than with just vector distance.

In the aforementioned Fake News Challenge, teams used a variety of stance-detection methods to detect when articles agreed, disagreed, discussed, or were unrelated to their

headlines [18]. (Lewandowsky 2018) provides explanations for the top three teams (one of which was themselves) and explains their models.

The winning team, Talos Intelligence's SOLAT in the SWEN, made two models and combined their outputs. Its first model, TalosCNN, is a convolutional neural network (CNN) that uses pre-trained word encoding to get word vectors and runs them through a CNN with convolutional, connected, and softmax layers. Its second model, TalosTree, is an XGBoost method that is also a gradient-boosted decision tree. The decision tree considers word count, TF-IDF, sentiment, singular-value decomposition features, and vector embeddings.

The second team, Team Athene, uses a multilayer perceptron model with six hidden layers and one softmax layer. The model considers cosine-similarity of word vectors from the headline and word vectors from the body, unigrams (which are just one word isolated and ignore all context), Dirichlet allocation, and semantic indexing.

The third team, the UCLMR team, also uses a multilayer perceptron model with one hidden layer. The model considers cosine similarity between headline and body vectors, similar to Team Athene and frequency vectors of unigrams of the 5000 most common words.

In (Yoon 2018), they calculate an incongruence score that quantifies the probability that a given headline is incongruent to the body. They also explain four traditional ML models for finding the score: XGBoost from the Fake News Challenge winning team (Lewandowsky 2018), support vector machines with the same features as XGBoost from Talos, a recurrent dual encoder that creates 300-dimensional word vectors and passes them through dual recurrent neural networks, and a convolutional dual encoder that uses convolutional neural networks to vectorize words (see Kim 2014).

They also consider a general data preprocessing method, "Independent Paragraphs", that can be used for each of the traditional ML models. After they split the body into paragraphs, they run their model of choice on each paragraph. The incongruence score for the whole headline-body pair is returned as the maximum incongruence score from all headline-paragraph pairs. This is done to reduce the length-mismatch between the headline and the text, which increases accuracy.

Then, they present their original attentive hierarchical dual encoder model. First, they split body texts into paragraphs and consider headline-paragraph data. They encode the words as vectors using their first word-level RNN. Then, they feed the word vectors from each paragraph into a paragraph-level RNN that is able to learn the relative importance of each paragraph in the whole body text. In addition, the paragraph-level RNN encodes the paragraphs as vectors using the word vectors. The whole body is encoded as a vector using the paragraph vectors, and the model trains weights and biases to translate the body vector into an incongruence score.

They find that paragraph splitting mostly improves the accuracy of each method (with a few exceptions), and that their own methods achieved high accuracies of 0.895 and 0.977.

(Yoon 2021) furthers the approach of their paper from 2018 (Yoon 2018) by incorporating graph-based learning with respect to the paragraphs. Once the headline-paragraph incongruence scores are found, the researchers use graph learning to find the relative importance of each paragraph.

D. Sensationalism

Numerous definitions exist for sensationalism, and the one we settled on for this section is adapted from (Frye 2005) and Oxford Dictionary Online: news that appeals to emotions such as excitement, shock, fear, and astonishment, at the expense of accuracy [22]. Sensationalism has been prevalent in journalism for as long as journalism has existed, and it is as present as ever nowadays because it helps pique reader interest and therefore helps to sell articles and gain readership.

While sensationalism is not necessarily inaccurate, it can still mislead: users may be tricked into believing a certain article is relevant to them because of the sensationally vague or emotionally charged wording. In addition, the barrage of sensationalist headlines present in the media can desensitize readers and cause them to believe that they are well-informed [23].

1) Example

Headline: A sausage a day could lead to cancer: Pancreatic cancer warning over processed meat (Chesney 2017)

This is an example of sensationalism because it draws upon charged terminology (cancer) and draws an exaggerated conclusion.

From following our search methodology, it seems that there is not much dedicated research to the detection of sensationalism in news headlines specifically, so we recommend that there is room for additional research in this area.

One main research paper is (Molek-Kozakowska 2013), which provides a framework for the general detection of sensationalism. In the research, they conduct a focus-group study of sensationalist articles, asking subjects to score how sensationalist an article or headline was and to describe what features of the headline contributed to the score. They found that sensationalist headlines often have one of a few common structures: a narrative structure of climax-complication-resolution-coda, an interrogative structure with a mystery being posed and the truth being revealed in the form of a question, vagueness, negatively charged labels and modifiers, etc. The researcher suggests that these features could be exploited in a detection system.

1) Datasets

Unlike the problem of incongruency, sensationalism detection approaches largely do not need to take the body text of articles into account, as sensationalism is inherently a feature of headlines. Thus, datasets only need to be composed of headlines with sensationalist (positive) or neutral (negative) labels. In addition, since headlines are short, there are fewer parameters and, therefore, less data is required for training.

In (Hoffman 2015), researchers manually applied a framework of balancing five “sensationalist illocutions”: exposing, speculating, generalizing, warning, and extolling, to score the sensationalism (along with other metrics) of news records regarding SARS between 2003 and 2004 [24]. They scored 500 articles with the framework as a training set. They also manually scored 200 other news records and used it as a testing set. The researchers eventually achieved a high level of accuracy using their model, demonstrating the potential effectiveness of a framework of measuring sensationalist features as (Molek-Kozakowska 2013) suggests. However, the manual annotation process is expensive and not generalizable, and, therefore, the framework should be automated for generalized application.

In (Xu 2019), researchers automatically gathered positive data (headlines perceived to be sensationalist) by choosing headlines of articles with many user-submitted comments on Tencent News [25]. For their negative data, they used automatically generated neutral headlines from (See 2017). They also preprocessed the data by converting the headline into Chinese character vectors. Their method’s accuracy may not be ideal, since the number of comments on a news article is not necessarily a marker of its sensationalism.

2) Approaches

(Hoffman 2015) used a logistic regression ML model with their limited size dataset to find the relevance, quality, and sensationalism of news records. Specifically, they used maximum-entropy modeling with constraints given by relationships between the sensationalist data. They found a relatively high accuracy of 73% when scoring sensationalism based on their testing dataset.

In (Xu 2019), the authors covered a method of automatically generating sensational headlines, and in the process developed a method of scoring headlines for their sensationalism. With their positive and negative labeled data, they trained an ML model to score sensationalism. They used a convolutional layer with two filters: a ReLU activation layer and a max-pooling layer. Their resulting model’s accuracy is lower than that of (Hoffman 2015) at only about .65 versus a completely random 0.5. As previously covered, this could be because of the inaccuracy in their method of automatically gathering positive data/sensationalist headlines. Future research could apply the technique of (Xu 2019) with a better dataset, improving the accuracy.

E. Headline Detection vs Fact Checking

Much work in misleading headline detection is done to help prevent the spread of articles with misleading headlines or to help readers avoid being misled by a misleading headline. Despite this, a seemingly small amount of research has been done to actually measure how users react to an automated misleading headline detection system. The only related source we found with our methodology was from (Park 2020): researchers implemented the model from (Yoon 2018) and built it into a web interface that told users how incongruent any given headline was. They ran a focus group of 14 university-age participants, and 10 of them (71%) reported that using the interface affected the articles they chose to read. This is a promising result, but since it was just a focus group with an extremely limited sample size and demographic, it is not enough to make any conclusions.

In comparison, a large amount of research has been done to measure how users respond to reading misinformation and then reading a correction to the misinformation. The topic of misinformation correction is related to headline detection because both are meant to help users determine when an article could potentially mislead them and adjust their reading accordingly. Research into corrections has largely demonstrated that corrections can have limited effectiveness on users.

Once misleading headlines have taken hold of a reader, it can be difficult to correct their beliefs either later on in the article or after the article was published. Several experimental studies have been done to test the effectiveness of corrections on a person who has received inaccurate information [9], [10], [26]. They found that subjects often resist corrections depending on a variety of factors, including how much they have at stake regarding the topic, what they originally believed about the topic, and how strong the correcting evidence and statements seem to be. In addition, they find that corrections often fail to reduce the effect of misinformation, and sometimes, a backfire effect occurs: corrections actually increase subjects’ misconceptions to the contrary.

Even when researchers taught the subjects how to identify fake news headlines, they continued to share fake news in the long term [27].

The prevalence of this backfire effect is debated, but it is accepted that corrections are not reliable ways of mitigating the effects of misleading information [28].

We theorize that headline detection could be more effective at helping change how users consume misleading articles than misinformation correction. This is because while users can sometimes see corrections as contradictions to their own beliefs and therefore resist their effects, headline detection would not criticize the information within the article but rather only point out how the article is trying to mislead them, sidestepping the contradiction problem.

Experimental studies on user reactions to misleading headline detection would serve a big role in helping to determine how effective a system headline detection would be, and it could help convince large social-media organizations or search engines to implement a detection service.

IV. SATIRE

Satire can be described as the use of wit, irony, sarcasm to criticize and expose a form of media [29], [30]. The purpose of satire is not usually malicious or didactic. Instead, satire is used to entertain in a sarcastic manner, to raise awareness of a certain issue, and to challenge certain viewpoints using humor [31].

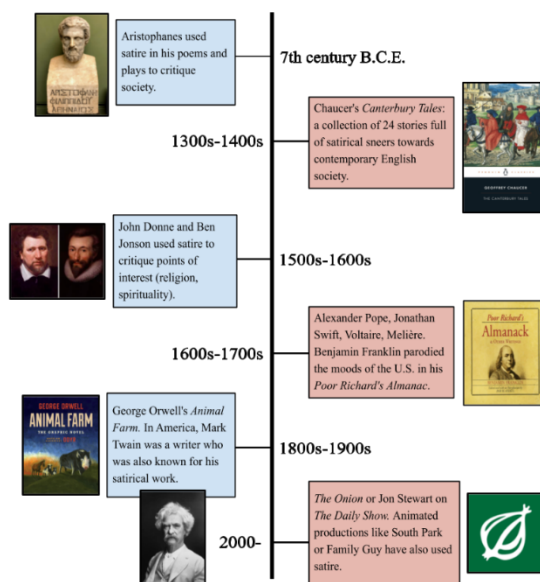
The two types of satire mainly used are Horatian and Juvenalian satire. Horatian satire uses tolerant and witty language to gently ridicule an idea. Horatian satire serves to entertain and poke fun at a certain topic without getting too controversial [32], [33]. For example, the late-night American TV show, Saturday Night Live, uses Horatian satire in its skits by over-exaggerating social issues to mock individuals and events, but in a gentle way without sparking controversy.

On the other hand, Juvenalian satire attacks flaws with contempt. Juvenalian satire is expressed in a much angrier and bitter manner [33]. For example, *1984*, by George Orwell, uses Juvenalian Satire to poke fun at totalitarianism with the government's absurd ideals. For instance, the Ingsoc, the socialist party of the novel, uses slogans such as war is peace, freedom is slavery, and ignorance is strength to control society. These slogans contradict themselves and are paradoxical. In addition, the main character of the story, Winston Smith, does not care about the party's ideals and even rebels against them. Unlike Saturday Night Live, Orwell criticized totalitarian societies, especially communist ones. Orwell's novel was even banned in Jackson County, Florida, for pro-communism [34].

A. The Evolution of Satire

Satire originated from the 7th century B.C.E. through Greek works. Aristophanes, who is also known as the "Father of Comedy," used satire in his poems and plays to critique society at the time. One of his well-known works, the *Babylonians*, was a play that criticized Athens' city officials and ultimately offended Cleon, a powerful figure in Athens [35], [36]. Aristophanes played an influential role in the development of satire. However, after the Greeks and Romans, satire disappeared with the start of the Dark Ages. Satire began to reappear in the 1300s, with works such as Chaucer's *Canterbury Tales*, which criticized English society [37], [38]. In the 1500s during the Renaissance, writers began to use satire to critique and entertain. Poets like John Donne and Ben Jonson used satire to critique various points of interests like religion and spirituality [37]. In America, Benjamin Franklin also used satire in his work, *Poor Richard's Almanac*, where he parodied the moods of America. Franklin's work was influential at the time as

his writing not only made others laugh, but also think about pressing issues [39]. In the 1800s, darker examples of satire began to appear like George Orwell's *Animal Farm*. The novel parodies the events of Russia's Bolshevik Revolution through the use of barn animals who rebel against humans [40]. Today, satire can be found in news articles from websites such as *The Onion* and *The Daily Show*. In addition, cartoons like South Park and Family Guy use satire to entertain and raise awareness about certain social issues. One South Park episode, in particular, focuses on the COVID-19 pandemic. In the episode, the kids of South Park elementary school have broken out of "quarantine." People are shown to be panicking while running around town and stocking up on toilet paper. The episode



clearly exaggerates the events of the pandemic in a humorous way [41].

Fig. 3 Timeline of the evolution of satire

B. Satire vs Fake News

Fake news and satirical news both contain similar elements. Fake news can be identified as news that contains misinformation to deceive or manipulate a viewer's opinion. Likewise, satirical news also contains fake information; however, unlike fake news, satirical news' purpose is to entertain and spread awareness on a specific topic. The use of clickbait is also common in both fake news and satire. Clickbait can include the following: a misleading title, a misleading description or caption, and a luring image.

Fake news and satirical news can also be different. For instance, fake news has the intention to manipulate the public opinion by providing disinformation. Fake news is usually predatory and can sway one's opinion. On the other hand, satire is used in an entertaining way. Satire uses sarcasm, irony, and wit to express ideas and raise awareness. Satire's intent is

usually not malicious and is used for entertainment purposes only.

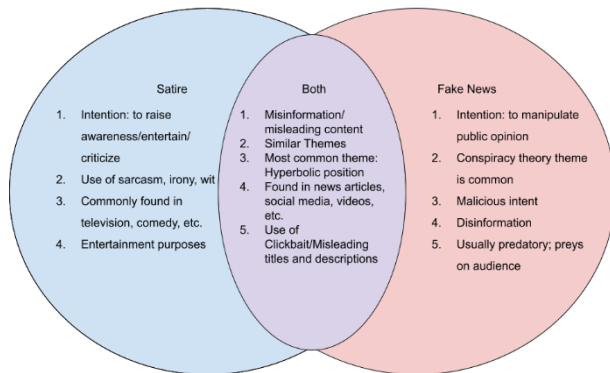


Fig. 4 Venn Diagram of similarities and differences

When examining fake news and satirical news articles, similar themes can be identified. These themes include: a hyperbolic position against or in favor of a group or person, a discrediting of a normal credible source, a sensationalist (an exaggerated story created to attract public attention) crime, racist messaging, paranormal theories, and conspiracy theories. An experiment was performed where 213 articles were analyzed for these specific themes. [42] The paper mentioned two approaches used to analyze the articles: linguistic cues and network analysis. Analyzing linguistic cues was done using bag of words and a Rhetorical Structure theory analytic framework. Network analysis involved using incoming and outgoing links to the articles and relevant topics to create a network. Unfortunately, the authors did not have this information in their dataset [42].

The dataset used contained 283 fake news articles and 203 satirical stories. These articles focused on American politics posted from January 2016 to October 2017 [42]. Overall, the most common themes were conspiracy theories, appearing in 30% of the articles and hyperbolic criticism, appearing in more than two-thirds of the articles analyzed. On the other hand, paranormal themes were the least common themes, appearing in less than 5%. Hyperbolic criticism appeared in satirical articles, albeit slightly. Conspiracy theories were usually more common in fake news articles than satirical ones. In addition, sensationalist crimes were also more common in fake news articles. However, paranormal themes, although rare, were more common in satirical news than in fake news [42].

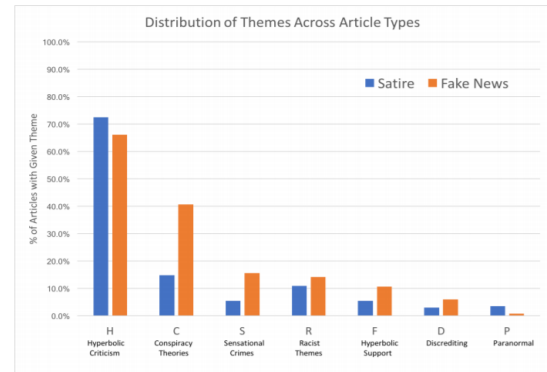


Fig. 5 Distribution of themes across article types [42]

C. Satirical Platforms and Prevalence

Satire can appear in multiple forms of media: News websites, social media, and television. Some popular satirical news websites include *The Onion*, *The World News Daily Report*, *The National Report*, and *The Babylon Bee*. *The Onion*, for example, is a popular satirical news website that can often be misunderstood. *The Onion*'s intention is to provide readers with satirical content for the purpose of entertainment. As of September 2021, *The Onion* averages over 4.27 million daily visits, according to SimilarWeb, a website traffic estimator, ranking the site number 3,056 in the United States [43].

The prevalence of satire may become an issue as satirical news can be difficult to identify. A study was conducted by participants from different age categories, political parties, etc. to determine whether a person could detect satirical news sites and news articles from credible sites [44]. The overall percentage of those who could identify satirical news websites was 11.98% while the overall percentage of those who could identify satirical news articles was 46.5% [44]. Since less than 50% of the participants could not identify satire from legitimate, satire can be difficult to comprehend.

D. Satire Detection

Fortunately, machine learning approaches such as Natural Language Processing can be used to help detect for satire. Recently, machine learning models have been developed to help detect for satire. These models use data from satirical news sites such as *The Onion* and *The Babylon Bee* to train and test their accuracy.

The following table summarizes recent satire detection models in chronological order.

TABLE II
EXISTING SATIRICAL NEWS DETECTION

Title	Model	Results	Year
-------	-------	---------	------

Automatic Satire Detection: Are You Having a Laugh? [45]	Support Vector Machine (SVM), Bag of Words (BoW), Bi-normal separation feature scaling (BNS)	Precision: 0.958 Recall: 0.690 Frequency: 0.798	2009
An Improved Method for Detection of Satire from User-Generated Content [46]	Standard Text Classification Approach: Support Vector Machine (SVM) and Bag of Words (BoW); Bi-Normal Separation Feature Scaling: BNS	BNS: Precision: 0.802 Recall: 0.859 F-Score: 0.829	2015
Fake News or Truth? Using Satirical Cues to Detect Potentially Misleading News. [47]	Natural Language Toolkit (NLTK), WordNet taxonomy, Grammar (gram) feature vector	Base+All Precision: 88 Recall: 82 F-Score: 87 Confidence: 87	2016
Satirical News Detection and Analysis using Attention Mechanism and Linguistic Features [48]	4-level hierarchical network: Support Vector Machine (SVM), Gated Recurrent Unit (GRU), Hierarchical Attention Network (HAN)	4LHNPDP: Precision: 93.51 Accuracy: 98.39 Recall: 79.50 F1: 91.46	2017
A Multi-Modal Method for Satire Detection using Textual and Visual Cues [49]	Multimodal Learning, Image Forgery Detection	Accuracy: 93.8 F1 Score: 92.16 AUC-ROC: 98.03	2020

Since satire can affect people in a negative way, it is important to know when a website is satirical or genuine. Satirical sites and fake news sites usually have similar format and grammar to actual sites to try to throw off the viewer. Satirical news sites also use specific wording in their sentences to create a sarcastic mood for the reader. Linguists who read over satire identify grammatical, stylistic, and structural properties to differentiate writing from satirical and legitimate. In addition, other factors such as headline features, profanity, and slang are identified to determine if a site is legitimate. For instance, legitimate news articles tend to use less profanity and slang. As a result, informality of an article is an important determinant when analyzing a news website [45].

When examining satirical and legitimate news websites, similar features can be found. The following figure displays similar details found in legitimate and satirical news websites. The website on the right, *The Onion*, is the satirical news website while the two on the left (CNN and Fox News) are credible news sources.

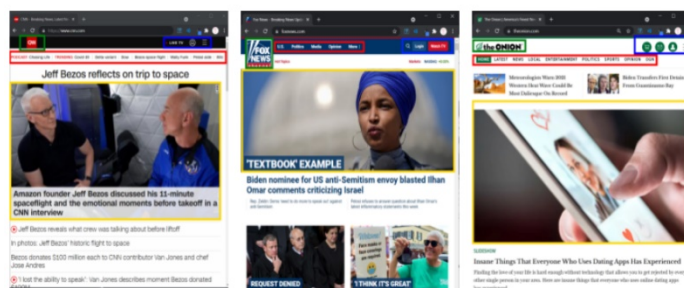


Fig. 6 Satirical and legitimate site comparison [50]–[52]

In the figure, the logos of the websites are located on the top left of the website. This is outlined in green. On the top right, there is a sign in option and dropdown menu icon (three lines). This is outlined in blue. There is also a header for links and topics on each website. This is outlined in red. Finally, all three sites have a headline picture in the center. This is outlined in yellow.

E. Policies

1) United States Policy

In the United States, the first amendment protects Americans' rights to express themselves freely no matter the genuineness in their statement [53]. In past years, the question whether satire is protected by the first amendment. In the *Hustler Magazine, Inc. et al. v. Jerry Falwell* case in 1980, a television minister sued Hustler Magazine for displaying a parody portraying Falwell engaging in indecent acts. Hustler was sued for libel, invasion of privacy, and intentional infliction of emotional distress [53]. While the first amendment protects Americans' right to free speech, satire may come across as defamation. Defamation differs from satire as satire is almost always false and obvious. [54] The court ruled that Hustler's case was not defamation as the work was an obvious parody [54].

Although satirical content is protected by the first amendment, there may be exceptions. For instance, satire can be challenged by copyright. If the satirical content attacks intellectual property that is protected by copyright or a trademark, the content may infringe on copyright laws [54]. Content that is consider "fair use", however, is protected by copyright. Transformative content that adds to the original work can be considered fair use [55].

2) Twitter Policy

Although satire is protected by the first amendment of freedom of speech, social media platforms and news sites have strict policies to control the spread of satire on their platforms. For example, parody accounts on Twitter are required to indicate that their account is not affiliated with a certain subject in their username and bio. Twitter users must meet these

requirements, otherwise, their account may become subjected to termination [56].

If an account does not follow Twitter's requirements, a person may file an impersonation or trademark complaint against a parody account. The account will then be flagged and reviewed by Twitter. If the account violates Twitter's policy on parody accounts, the account may be temporarily suspended. The account may be even permanently suspended if the account violates their policy more than once.

Requirements for parody, newsfeed, commentary, and fan accounts

Here are the requirements for marking your account. All requirements must be met in order to comply with the policy.

- Bio: The bio should clearly indicate that the user is not affiliated with the subject of the account. Non-affiliation can be indicated by incorporating, for example, words such as (but not limited to) "parody," "fake," "fan," or "commentary." Non-affiliation should be stated in a way that can be understood by the intended audience.
- Account name: The account name (note: this is separate from the username, or @handle) should clearly indicate that the user is not affiliated with the subject of the account. Non-affiliation can be indicated by incorporating, for example, words such as (but not limited to) "parody," "fake," "fan," or "commentary." Non-affiliation should be stated in a way that can be understood by the intended audience.

Please note that your account must be fully compliant with the [Twitter Rules](#) and [Terms of Service](#) in addition to meeting these requirements.

Fig. 7 Twitter's policy on parody, newsfeed, commentary, and fan accounts [56]

3) Facebook Policy

Another social media platform for satire, Facebook, also has strict guidelines for satire. On June 20th, 2021, Facebook decided to update their Community Guidelines after an incident regarding the use of satire in a meme post related to The Armenian Genocide [57]. Facebook decided to remove the meme citing its Cruel and Insensitive Community Standard that states that they will remove posts that target victims of serious physical or emotional harm.

Furthermore, another social media platform that is owned by Facebook, Instagram, allows users to flag content as "false information." This content will become excluded from the "explore" and "hashtag" pages to limit discovery [58].

4) Snopes Policy

A popular fact-checking website, Snopes, has a similar feature that labels content as satire. Satirical content will be pre-approved by Snopes' staff. Snopes' Fact Check Ratings list contains ratings such as "True" and "Mostly True" which is rated on content that is mostly credible [58]. On the other hand, categories called "Labeled Satire" and "Originated as Satire" are labeled on content that is satirical. These ratings make it much easier for readers to navigate through Snopes and verify if content is credible or satirical.

	Labeled Satire This rating indicates that a claim is derived from content described by its creator and/or the wider audience as satire. Not all content described by its creator or audience as 'satire' necessarily constitutes satire, and this rating does not make a distinction between 'real' satire and content that may not be effectively recognized or understood as satire despite being labeled as such.
	Originated as Satire This rating refers to content that originally came from a site described as satire, but was later stripped of some of its satirical markings, repackaged, and posted elsewhere. The rating also applies to content not necessarily labeled as satire but that audiences perceived as satirical nonetheless, such as content from The Onion.

Fig. 8 Scopes' labels regarding Satire [59]

V. CONCLUSIONS/NEXT STEP

Fake news can be spread to users through social media. Algorithms are used to generate content for users to enjoy; however, the algorithm can easily spread fake news. Misleading headlines contribute to a variety of problems including the proliferation of fake news. Research has been done on several different methods of automatically determining misleading headlines, but more research should be conducted to determine how effective detection systems might be in preventing their spread.

In addition to fake news, satirical news may also spread rapidly. Satirical news can be easily mixed up with fake news, which may be a problem, as satirical news is used to entertain and criticize in a way that is not supposed to be taken seriously. Fortunately, satire detection can be developed to detect for satire in media. Also, social media platforms have been enforcing new policies to counter satire and parody on their platforms. Still, more research and action are needed to further prevent misinformation from spreading.

REFERENCES

- [1] "Watch trending videos for you," *TikTok*. <https://www.tiktok.com/foryou?lang=en> (accessed Oct. 09, 2021).
- [2] "Explore Feed." <https://www.facebook.com/explorefeed/> (accessed Oct. 09, 2021).
- [3] C. Geeng, S. Yee, and F. Roesner, "Fake News on Facebook and Twitter: Investigating How People (Don't) Investigate," in *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems*, New York, NY, USA: Association for Computing Machinery, 2020, pp. 1–14. Accessed: Aug. 21, 2021. [Online]. Available: <https://doi.org/10.1145/3313831.3376784>
- [4] "The-Kremlins-Troll-Army—Global—The-Atlantic.pdf." Accessed: Aug. 22, 2021. [Online]. Available: <http://rusn340.blogs.wm.edu/files/2015/03/The-Kremlins-Troll-Army-%E2%80%94Global-%E2%80%94The-Atlantic.pdf>
- [5] N. Che Eembi Jamil, I. Ishak, and F. Sidi, "Deception detection approach for data veracity in online digital news: Headlines vs contents," Oct. 2017, vol. 1891, p. 020036. doi: 10.1063/1.5005369.
- [6] M. Gabielkov, A. Ramachandran, A. Chaintreau, and A. Legout, "Social Clicks: What and Who Gets Read on Twitter?," *ACM SIGMETRICS Perform. Eval. Rev.*, vol. 44, no. 1, pp. 179–192, Jun. 2016, doi: 10.1145/2964791.2901462.

- [7] U. K. H. Ecker, S. Lewandowsky, E. P. Chang, and R. Pillai, "The effects of subtle misinformation in news headlines," *J. Exp. Psychol. Appl.*, vol. 20, no. 4, pp. 323–335, Dec. 2014, doi: 10.1037/xap0000028.
- [8] D. Margolin, A. Hannak, and I. Weber, "Political Fact-Checking on Twitter: When Do Corrections Have an Effect?," *Polit. Commun.*, vol. 35, pp. 1–24, Sep. 2017, doi: 10.1080/10584609.2017.1334018.
- [9] B. Nyhan and J. Reifler, "When Corrections Fail: The Persistence of Political Misperceptions," *Polit. Behav.*, vol. 32, no. 2, pp. 303–330, Jun. 2010, doi: 10.1007/s11109-010-9112-2.
- [10] S. Lewandowsky, U. K. H. Ecker, C. M. Seifert, N. Schwarz, and J. Cook, "Misinformation and Its Correction: Continued Influence and Successful Debiasing," *Psychol. Sci. Public Interest*, vol. 13, no. 3, pp. 106–131, Dec. 2012, doi: 10.1177/1529100612451018.
- [11] K. Park, T. Kim, S. Yoon, M. Cha, and K. Jung, "BaitWatcher: A lightweight web interface for the detection of incongruent news headlines," *ArXiv200311459 Cs*, 2020, doi: 10.1007/978-3-030-42699-6.
- [12] J. Reis, F. Benevenuto, P. Vaz de Melo, R. Prates, H. Kwak, and J. An, "Breaking the News: First Impressions Matter on Online News," Jan. 2015.
- [13] J. León, "The effects of headlines and summaries on news comprehension and recall," *Read. Writ.*, vol. 9, pp. 85–106, Jan. 1997, doi: 10.1023/A:1007928221187.
- [14] E. Shearer and E. Grieco, "Americans Are Wary of the Role Social Media Sites Play in Delivering the News," *Pew Research Center's Journalism Project*, Oct. 02, 2019.
<https://www.pewresearch.org/journalism/2019/10/02/americans-are-wary-of-the-role-social-media-sites-play-in-delivering-the-news/> (accessed Aug. 24, 2021).
- [15] S. Chesney, M. Liakata, M. Poesio, and M. Purver, "Incongruent Headlines: Yet Another Way to Mislead Your Readers," in *Proceedings of the 2017 EMNLP Workshop: Natural Language Processing meets Journalism*, Copenhagen, Denmark, Sep. 2017, pp. 56–61. doi: 10.18653/v1/W17-4210.
- [16] "[1811.07066] Detecting Incongruity Between News Headline and Body Text via a Deep Hierarchical Encoder," <https://arxiv.org/abs/1811.07066> (accessed Oct. 09, 2021).
- [17] R. Mishra, P. Yadav, R. Calizzano, and M. Leippold, "MuSeM: Detecting Incongruent News Headlines using Mutual Attentive Semantic Matching," *ArXiv201003617 Cs*, Oct. 2020, Accessed: Oct. 11, 2021. [Online]. Available: <http://arxiv.org/abs/2010.03617>
- [18] A. Hanselowski *et al.*, "A Retrospective Analysis of the Fake News Challenge Stance-Detection Task," in *Proceedings of the 27th International Conference on Computational Linguistics*, Santa Fe, New Mexico, USA, Aug. 2018, pp. 1859–1874. Accessed: Aug. 24, 2021. [Online]. Available: <https://aclanthology.org/C18-1158>
- [19] J. Brownlee, "A Gentle Introduction to the Bag-of-Words Model," *Machine Learning Mastery*, Oct. 08, 2017.
<https://machinelearningmastery.com/gentle-introduction-bag-words-model/> (accessed Feb. 15, 2021).
- [20] "Learning to Detect Incongruence in News Headline and Body Text via a Graph Neural Network | IEEE Journals & Magazine | IEEE Xplore." <https://ieeexplore.ieee.org/document/9363185> (accessed Oct. 09, 2021).
- [21] Y. Goldberg and O. Levy, "word2vec Explained: deriving Mikolov *et al.*'s negative-sampling word-embedding method," *ArXiv14023722 Cs Stat*, Feb. 2014, Accessed: Oct. 11, 2021. [Online]. Available: <http://arxiv.org/abs/1402.3722>
- [22] W. B. Frye, "A qualitative analysis of sensationalism in media," *Grad. Theses Diss. Probl. Rep.*, Dec. 2005, doi: <https://doi.org/10.33915/etd.3218>.
- [23] K. Molek-Kozakowska, "Towards a pragma-linguistic framework for the study of sensationalism in news headlines," *Discourse Commun.*, vol. 7, pp. 173–197, May 2013, doi: 10.1177/1750481312471668.
- [24] S. J. Hoffman and V. Justicz, "Automatically quantifying the scientific quality and sensationalism of news records mentioning pandemics: validating a maximum entropy machine-learning model," *J. Clin. Epidemiol.*, vol. 75, pp. 47–55, Jul. 2016, doi: 10.1016/j.jclinepi.2015.12.010.
- [25] P. Xu, C.-S. Wu, A. Madotto, and P. Fung, "Clickbait? Sensational Headline Generation with Auto-tuned Reinforcement Learning," *ArXiv190903582 Cs*, Sep. 2019, Accessed: Oct. 11, 2021. [Online]. Available: <http://arxiv.org/abs/1909.03582>
- [26] "Undermining the Corrective Effects of Media-Based Political Fact Checking? The Role of Contextual Cues and Naïve Theory - Garrett - 2013 - Journal of Communication - Wiley Online Library." <https://onlinelibrary.wiley.com/doi/abs/10.1111/jcom.12038> (accessed Oct. 11, 2021).
- [27] A. Bor, M. Osmundsen, S. Hebbelstrup Rye Rasmussen, A. Bechmann, and M. Petersen, "Fact-checking videos reduce belief in but not the sharing of 'fake news' on Twitter. 2020. doi: 10.31234/osf.io/a7huq.
- [28] U. K. H. Ecker, S. Lewandowsky, and M. Chadwick, "Can corrections spread misinformation to new audiences? Testing for the elusive familiarity backfire effect," *Cogn. Res. Princ. Implic.*, vol. 5, no. 1, p. 41, Aug. 2020, doi: 10.1186/s41235-020-00241-6.
- [29] "Definition of SATIRE." <https://www.merriam-webster.com/dictionary/satire> (accessed Oct. 09, 2021).
- [30] "What is Satire? || Oregon State Guide to Literary Terms," *College of Liberal Arts*, Oct. 10, 2019. <https://liberalarts.oregonstate.edu/wlf/what-satire> (accessed Oct. 09, 2021).
- [31] "The Purpose and Method of Satire." <https://www.virtualsalt.com/the-purpose-and-method-of-satire/> (accessed Oct. 09, 2021).
- [32] "Horatian satire," *You Do Hoodoo*. <https://wlm3.com/tag/horatian-satire/> (accessed Oct. 09, 2021).
- [33] B. Zimmerman, "The tradition & techniques of satire: the case of Michael Moore," *Humanist Can.*, vol. 37, no. 149, pp. 10–16, Jun. 2004.
- [34] "More Banned Books Week at UC Press," *UC Press Blog*. <https://www.ucpress.edu/blog/52211/more-banned-books-week-at-uc-press/> (accessed Oct. 05, 2021).
- [35] "Aristophanes | Biography, Plays, & Facts," *Encyclopedia Britannica*. <https://www.britannica.com/biography/Aristophanes> (accessed Oct. 05, 2021).
- [36] J. Starkey, "War, Politics, and Aristophanes' Babylonians," *SSRN Electron. J.*, May 2010, doi: 10.2139/ssrn.1607762.
- [37] "0f9672dfd01c4128a1368ee4f34c0b95.pdf." Accessed: Oct. 05, 2021. [Online]. Available: <https://content.schoolinsites.com/api/documents/0f9672dfd01c4128a1368ee4f34c0b95.pdf>
- [38] "Lecture #5--Chaucer Canterbury Tales." <https://www.nku.edu/~rkdrry/202/lecture5.html> (accessed Oct. 05, 2021).
- [39] "Benjamin Franklin . Wit and Wisdom . Franklin Funnies | PBS." https://www.pbs.org/benfranklin/l3_wit_franklin.html (accessed Oct. 05, 2021).
- [40] "Animal Farm | novel by Orwell | Britannica." <https://www.britannica.com/topic/Animal-Farm> (accessed Oct. 05, 2021).
- [41] "South Park Police Equipped To Manage Town in Chaos - SOUTH PARK - YouTube." <https://www.youtube.com/watch?v=uo63QQsm5Dw> (accessed Oct. 05, 2021).
- [42] J. Golbeck *et al.*, "Fake News vs Satire: A Dataset and Analysis," in *Proceedings of the 10th ACM Conference on Web Science*, Amsterdam Netherlands, May 2018, pp. 17–21. doi: 10.1145/3201064.3201100.
- [43] "Theonion.com Traffic Ranking & Marketing Analytics," *Similarweb*. <http://similarweb.com/website/theonion.com/> (accessed Oct. 09, 2021).
- [44] M. Bedard and C. Schoenthaler, "Satire or Fake News: Social Media Consumers' Socio-Demographics Decide," in *Companion of the The Web Conference 2018 on The Web Conference 2018 - WWW '18*, Lyon, France, 2018, pp. 613–619. doi: 10.1145/3184558.3188732.
- [45] C. Burfoot and T. Baldwin, "Automatic Satire Detection: Are You Having a Laugh?," in *Proceedings of the ACL-IJCNLP 2009 Conference*

- Short Papers*, Suntec, Singapore, Aug. 2009, pp. 161–164. Accessed: Oct. 05, 2021. [Online]. Available: <https://aclanthology.org/P09-2041>
- [46] S. Taha, O. Prof, T. Nafis, and S. Khanna, “An Improved Method for Detection of Satire from User-Generated Content.”
 - [47] V. Rubin, N. Conroy, Y. Chen, and S. Cornwell, “Fake News or Truth? Using Satirical Cues to Detect Potentially Misleading News,” in *Proceedings of the Second Workshop on Computational Approaches to Deception Detection*, San Diego, California, Jun. 2016, pp. 7–17. doi: 10.18653/v1/W16-0802.
 - [48] F. Yang, A. Mukherjee, and E. Dragut, “Satirical News Detection and Analysis using Attention Mechanism and Linguistic Features,” *ArXiv170901189 Cs*, Sep. 2017, Accessed: Oct. 05, 2021. [Online]. Available: <http://arxiv.org/abs/1709.01189>
 - [49] L. Li, O. Levi, P. Hosseini, and D. A. Broniatowski, “A Multi-Modal Method for Satire Detection using Textual and Visual Cues,” *ArXiv201006671 Cs*, Oct. 2020, Accessed: Oct. 05, 2021. [Online]. Available: <http://arxiv.org/abs/2010.06671>
 - [50] “CNN - Breaking News, Latest News and Videos.” <https://www.cnn.com/> (accessed Oct. 09, 2021).
 - [51] “Fox News - Breaking News Updates | Latest News Headlines | Photos & News Videos.” <https://www.foxnews.com/> (accessed Oct. 09, 2021).
 - [52] “The Onion | America’s Finest News Source.” <https://www.theonion.com/> (accessed Oct. 09, 2021).
 - [53] “Parody & satire | Freedom Forum Institute.” <https://www.freedomforuminstitute.org/first-amendment-center/topics/freedom-of-speech-2/arts-first-amendment-overview/parody-satire/> (accessed Oct. 11, 2021).
 - [54] J. L. Walker, “Satire.” <https://www.mtsu.edu/first-amendment/article/1015/satire> (accessed Oct. 11, 2021).
 - [55] T. L. Wherry, *The librarian’s guide to intellectual property in the digital age copyrights, patents, and trademarks* / Timothy Lee Wherry. Chicago: American Library Association, 2002.
 - [56] “Parody, newsfeed, commentary, and fan account policy.” <https://help.twitter.com/en/rules-and-policies/parody-account-policy> (accessed Oct. 05, 2021).
 - [57] IANS, “Facebook decides to update its policies on handling satirical content,” *Business Standard India*, Jun. 20, 2021. Accessed: Oct. 11, 2021. [Online]. Available: https://www.business-standard.com/article/international/facebook-decides-to-update-its-policies-on-handling-satirical-content-121062000181_1.html
 - [58] M. Muirhead, “Instagram, Snopes Roll Out New Fact-Checking Features to Address Memes, Satire,” *Karma*, Aug. 19, 2019. <https://karmainmpact.com/instagram-snopes-roll-out-new-fact-checking-features-to-address-memes-satire/> (accessed Oct. 09, 2021).
 - [59] “Fact Check Ratings | Snopes.com.” <https://www.snopes.com/fact-check-ratings/> (accessed Oct. 09, 2021).