

Artificial Intelligence in 2024: A Thematic Analysis of Media Coverage

Luke Warner Williams

Thesis prepared for the faculty of Virginia Polytechnic Institute and State University in partial fulfillment of requirements for the degree of

Master of Arts
in
Communication

Marcus C. Myers, Chair
Jim A. Kuypers
Nneka J. Logan

April 30, 2025
Blacksburg, Virginia

Keywords: framing theory, thematic analysis, artificial intelligence (AI), 2024

Artificial Intelligence in 2024: A Thematic Analysis of Media Coverage

Luke Williams

ACADEMIC ABSTRACT

This thesis investigates how three agenda-setting U.S. newspapers — *The New York Times*, *The Wall Street Journal*, and *The Washington Post* — framed artificial intelligence (AI) during the year 2024. Using an inductive thematic analysis grounded in Braun and Clarke's six-phase procedure, 300 randomly-selected articles (100 per outlet) were analyzed and interpreted as media frames. Eight dominant frames emerged: AI Boom vs. Bubble, Misuse & Misinformation, Ethical & Moral Challenges, Policy & Governance, Societal & Cultural Impact, Work & Automation, Environmental Impact, and Technological Advancements & Future Risks. Across coverage, fear narratives centered on disinformation, job displacement, algorithmic bias, environmental costs, and existential risk. Results suggest that U.S. legacy media have moved beyond early techno-optimism towards a more nuanced discourse that simultaneously fuels investment and adoption, demands regulation and safeguards, and shapes public perception. These findings document how narratives evolve in response to rapid technological change and provide information for scholars, policymakers, and technologists seeking to understand how media discourse may steer AI governance and adoption.

Artificial Intelligence in 2024: A Thematic Analysis of Media Coverage

Luke Williams

GENERAL AUDIENCE ABSTRACT

Artificial intelligence (AI), increasingly capable of creating text, images, audio, and video that is indistinguishable from that produced by humans, dominated news coverage in 2024. Understanding how the media covers AI helps us comprehend public attitudes and policy responses to these powerful technologies. This study analyzed how three major U.S. newspapers—*The New York Times*, *The Wall Street Journal*, and *The Washington Post*—reported on AI throughout the year 2024. By analyzing 300 randomly chosen articles (100 from each paper), the research identified common themes: including debates over whether AI will bring economic prosperity or create an economic bubble, concerns about misuse and misinformation, ethical challenges involving privacy and fairness, impacts on jobs and society, and more. The research also found systematic differences in coverage between newspapers, revealing editorial choices and perspectives that shape narratives. These media narratives impact how the public perceives AI and influences debates around technology adoption. Recognizing these media patterns can help us navigate the current conversation on AI and understand the complex realities and societal impacts of these emerging technologies.

Table of Contents

Academic Abstract	i
General Audience Abstract	ii
Chapter 1: Introduction	1
Chapter 2: Literature Review	3
2.1 Artificial intelligence	3
2.2 U.S. media environment	8
2.3 Framing	21
2.4 Thematic analysis in frame identification	25
2.5 Relevant research	28
Chapter 3: Methodology	37
3.1 Thematic Analysis	37
3.2 Sample	38
3.2.1 Sample Selection: <i>The New York Times</i>	39
3.2.2 Sample Selection: <i>The Wall Street Journal</i>	40
3.2.3 Sample Selection: <i>The Washington Post</i>	41
3.3 Procedure	42
3.4 Trustworthiness	44
3.5 Significance	45
Chapter 4: Results	47
4.1 Research Question 1	47
4.1.1 AI Boom vs. Bubble	48
4.1.2 Misuse and Misinformation	50
4.1.3 Ethical and Moral Challenges	51
4.1.4 Politics and Governance	53
4.1.5 Societal and Cultural Impact	55
4.1.6 Work and Automation	57
4.1.7 Environmental Impact	59
4.1.8 Technological Advancements and Future Risks	61
4.2 Research Question 2	63
4.3 Research Question 3	65
Chapter 5: Discussion	67
5.1 Findings	67
5.2 Limitations	69
5.3 Future research	70
Chapter 6: Conclusion	72
References	75

Chapter 1: Introduction

We stand at a pivotal moment in human history. For the first time, artificial intelligence (AI) systems can produce text, images, audio, and video that are virtually indistinguishable from human-created content. Our online environments are increasingly saturated with AI-generated material, blurring the lines between human and synthetic communication. As AI capabilities advance at breathtaking speed, public understanding of these technologies is filtered through the lens of news coverage. The frames employed by legacy media—particularly how they emphasize risks, benefits, and ethical considerations—shape public discourse, policy debates, and adoption patterns across society.

This study examines how leading U.S. news organizations—*The New York Times*, *The Wall Street Journal*, and *The Washington Post*—framed artificial intelligence through coverage published in 2024. Unlike previous research that often dwells on AI's long-term implications or traces broad historical trends, this thesis provides a timely, focused analysis of media framing during a year marked by significant technological breakthroughs, intense regulatory discussions, and heightened public awareness. Given AI's rapidly evolving nature, an analysis grounded in current media narratives offers unique insights into how public perceptions are being shaped in real time.

Framing theory suggests that media doesn't merely report on AI, but actively constructs meaning through language choices, emphasis patterns, and strategic omissions. Whether AI is portrayed as an existential threat, an economic disruptor, or a revolutionary tool for human progress carries far-reaching consequences. Public perception, policy responses, and investment decisions are all influenced by these media-constructed frames. Through thematic analysis, this

study identifies dominant frames in 2024's AI coverage, illuminating the narratives that are shaping both public understanding and policy discourse.

By analyzing AI framing during this critical year, this research offers a fresh perspective on media influence in an era of technological acceleration. The findings hold relevance not only for communication scholars and media analysts, but also for policymakers, technologists, and journalists navigating the complex, evolving landscape of artificial intelligence—a technology that promises to transform virtually every aspect of human society in the coming decades.

Chapter 2: Literature Review

2.1 Artificial intelligence

Accurately defining artificial intelligence (AI) is a surprisingly difficult task. This is partly because of the inherent difficulty in defining the concept of ‘intelligence’ itself, and partly because of the breathtaking speed AI systems are developing: what was considered novel intelligent behavior by machines a decade ago is hardly noteworthy today. When British mathematician Alan Turing laid the foundation for the modern AI revolution in the mid-20th century, he posed the simple question: “Can machines think?” (Turing, 1950, p. 50). Turing developed what would come to be known as the “Turing Test” for intelligence, still useful today as a benchmark for classifying artificial intelligence: “If a human is interacting with another human and a machine and unable to distinguish the machine from the human, then the machine is said to be intelligent” (Haenlein & Kaplan, 2019, p. 7).

Over the subsequent decades, AI was commonly defined as Turing had: through an act of comparison between a system’s ‘intelligence’ and human intelligence, defined as the “biopsychological potential to process information that can be activated in a cultural setting to solve problems or create products that are of value in a culture” (Gardner, 1999, p. 34). AI pioneers, guided by this definition and conceptualization of the problem, attempted to develop AI systems by replicating human intelligence. This led to some early breakthroughs, notably the ELIZA natural language processing tool for simulating conversation (Weizenbaum, 1966), but progress soon stalled. These systems, classified as *expert systems* in the literature, were built on a fundamental mistake: researchers assumed that human intelligence could be reconstructed as a series of formalized “if-then” statements. While expert systems perform well in tasks with formalized rules of engagement, such as the famous Deep Blue chess program developed by

IBM (M. Campbell et al., 2002), capable of defeating the best human chess players, expert systems are incapable of obtaining, processing, learning from, and then utilizing data to accomplish goals or tasks through flexible adaptation where rules are not formalized (Asemi et al., 2020; Haenlein & Kaplan, 2019). It became clear that these expert systems, however useful to specific applications, were fundamentally incapable of producing ‘true’ adaptable artificial intelligence in the real world.

In 2015, the field was reinvigorated when AlphaGo, a program developed by Google, beat the world champion at the board game Go using a different approach to building artificial intelligence: AlphaGo was built using an artificial neural network that replicates the learning process of neurons in the human brain, known as a “convolutional deep neural network” or more commonly as “deep learning” (Hebb, 2005; Holcomb et al., 2018). This technology laid the early foundation for the AI revolution we are experiencing today. In contrast to the conventional approach to programming used by the developers of Expert Systems, where large problems are deconstructed into smaller defined tasks for a computer to perform, in a neural network, the computer “learns from observational data, figuring out its own solution to the problem at hand” (Nielsen, 2015, p. 3). This is accomplished by processing input data through interconnected layers of artificial neurons that iteratively adjust their parameters based on feedback, allowing the network to refine its internal representations and learn increasingly complex patterns over multiple training cycles (Nielsen, 2015). This imitates the intellectual developmental process in humans, who as children are exposed to vast amounts of “data” (visual, auditory, tactile) and over time develop human cognition and intelligence. Therefore, for the purposes of this study, artificial intelligence (AI) is defined as “a system’s ability to interpret external data correctly, to

learn from such data, and to use those learnings to achieve specific goals and tasks through flexible adaptation” (Kaplan & Haenlein, 2019, p. 5).

Just two years after Google’s AlphaGo victory, the field of AI research was upended yet again, this time in the form of a now-famous paper titled “Attention is All You Need” (Vaswani et al., 2017). In it, the team introduced the groundbreaking Transformer architecture, a novel method for AI development that, while rooted in the traditional convolutional deep neural network approach, eliminated traditional methods like convolution or recurrence and instead relied entirely on self-attention mechanisms (Vaswani et al., 2017). This new methodology launched a revolution in natural language processing, increasing training efficiency, improving accuracy, and perhaps most importantly, giving researchers a clear vision and path towards improving their AI models. Unlike previous models, transformer-based models have been able to “continuously benefit from larger architectures and larger quantities of [training] data” (Bender et al., 2021, p. 611). AI research labs, including the company OpenAI, raced to implement and build on these new findings (Roumeliotis & Tselikas, 2023).

In 2018, OpenAI introduced the first version of what would become the Generative Pre-trained Transformer (GPT) series of models, building on the attention-based innovations introduced by Vaswani et al (2017). This new model was trained on large volumes of text and then fine-tuned with reinforcement learning for specific tasks, enabling the model to understand, predict, and generate remarkably coherent and context-sensitive responses using natural language (Radford et al., 2018). Soon after, GPT-2 and GPT-3 were released, each vastly increasing the number of parameters and the model’s ability to produce human-like text (Brown et al., 2020). These models belong to a broader category now commonly known as Large Language Models (LLMs). These language models (LMs) are systems that are trained on “string

prediction tasks: that is, predicting the likelihood of a token (character, word or string) given either its preceding context or ...its surrounding context” (Bender et al., 2021, p. 611). The difference between language models (LMs) and large language models (LLMs) are simply their size: while GPT-1 was considered large at the time with 117 million parameters, GPT-3 has 175 billion parameters and GPT-4 has ~1.76 trillion parameters (Radford et al., 2018). While the total number of parameters is not the only important factor in model development, larger models tend to exhibit qualitatively richer capabilities than their smaller predecessors.

OpenAI’s 2022 release of ChatGPT, a web-accessible version of the company’s GPT-3 model for public use, marked a pivotal moment in the public's interaction with AI, inciting an extraordinary level of interest and rapid adoption of advanced AI technologies. ChatGPT demonstrated the ability to engage in open-ended conversations, provide creative outputs, and assist with tasks ranging from writing to coding. Unlike earlier AI applications that were confined to specialized domains or designed for enterprise use with a price to match, ChatGPT's capabilities were accessible to the masses, leading to widespread experimentation. ChatGPT set a record for the fastest-growing consumer application in history after reaching an estimated 100 million monthly active users in January, just two months after launch (Hu, 2023). Headlines and social media were soon filled with examples of ChatGPT’s human-like writing ability—from short stories and news summaries to code snippets and creative brainstorming. The free, conversational interface allowed users to experiment with the model’s capabilities in real time, fueling both awe at its responsiveness and concern about its potential drawbacks (Roose, 2022).

While large language models (LLMs) have demonstrated remarkable novel capabilities in text-based tasks—such as generating coherent responses or summarizing large volumes of written information—the underlying transformer architecture has also enabled multi-modal

applications, capable of processing and generating images, audio, and even video. Examples include DALL·E, which can create original images from textual prompts (Ramesh et al., 2021), CLIP, which learns to correlate text and images for tasks like image classification and captioning (Radford et al., 2021), and SORA, which can generate videos of realistic or artistic scenes from text instructions (Liu et al., 2024). By unifying different data streams and modalities within a single model, multi-modal architectures bring AI systems closer to replicating human perception and understanding. Advancements in hardware and the availability of large, labeled datasets have further fueled these developments, enabling richer interactions between text, images, and other sensory information.

Despite the remarkable capabilities expressed by modern AI systems, there remain significant architectural challenges related to interpretability. LLMs are often described as “black boxes” because their internal processes, consisting of billions and trillions of discrete parameters, are difficult to scrutinize, complicating human interpretability of exactly how a particular output is produced (Doshi-Velez & Kim, 2017). As Myers (2025) notes, scaling oversight as the technology scales poses a unique challenge, where “achieving transparency and oversight might not always be technologically feasible” (Myers, 2025). This lack of transparency complicates troubleshooting, risk assessment, and efforts to ensure that model predictions align with intended outcomes. Additionally, large training datasets can inadvertently embed human social and cultural biases into the AI models, which may then manifest in the generated content (Bolukbasi et al., 2016). Multi-modal systems that incorporate images, video, and audio inherit the same vulnerability: if data used for training contains bias, the model may learn and perpetuate stereotypes or discriminatory patterns. Researchers are therefore studying technical strategies (e.g., fine-tuning protocols, data filtering) to mitigate bias, as well as developing interpretability

tools to promote transparency. (Zhao et al., 2018). These ongoing challenges reflect an important dimension of contemporary AI research: how can powerful AI models be understood, trusted, and responsibly deployed?

For the purposes of this study, unless otherwise specified, any future use of the terms "artificial intelligence" or "AI" refer to large language models (LLMs) developed using transformer-based architectures. With acknowledgement to the broader historical context and diverse applications of artificial intelligence as a field, this focused terminology reflects the fact that these modern AI systems are the central subject of investigation in this study and received vast amounts of media coverage in 2024. This narrows the discussion to the most relevant modern systems that have demonstrated unprecedented capabilities in natural language understanding and generation, distinguishes these systems from earlier AI paradigms discussed in the historical overview, and acknowledges that these particular implementations have dominated both public discourse and academic research throughout the year 2024. This specification is particularly important given the rapid evolution of AI technologies, where capabilities and definitions continue to shift alongside technological advancement, and where the term "artificial intelligence" has become increasingly diffuse in both popular and technical usage.

2.2 U.S. Media Environment

The American media landscape has undergone profound transformations as technological development has changed and shaped how the public communicates. As the Canadian communication theorist Marshall McLuhan (2017) famously stated, “the medium is the message,” (McLuhan, 2017, p. 107) arguing that the primary focus of scholars should be the communication medium itself, not simply the messages it carries. A thorough analysis of the

history of the United States media environment reveals how insightful McLuhan truly was: every major technological shift—from newspapers to radio, television to the internet, and now to algorithmically-driven social media and AI—has fundamentally altered not just how information is distributed, but also how it is perceived, processed, and integrated into social consciousness. To fully contextualize how AI was covered by major media outlets in 2024, one must understand how these increasingly rapid and profound technological disruptions have led to our modern moment. Thus, this study will next chronicle the major transformations technology has wrought on the American media environment since the early 1800s.

The invention and increasing proliferation of the telegraph fundamentally transformed communication in the mid-19th century, collapsing geographic distance and enabling near-instantaneous transmission of information across vast territories (Carey, 1983). This technology created the first national news networks and laid the groundwork for wire services such as the Associated Press. The development of the high-speed rotary press simultaneously enabled mass production of newspapers, slashing costs and expanding readership beyond elite audiences (Schudson, 1978). This coincided with a significant rise in literacy rates, as public education expanded throughout the United States and created unprecedented demand for printed material, which newspapers rushed to satisfy. These twin innovations—faster information transmission and cheaper reproduction—established the technological basis for mass media as we understand it today.

The introduction of radio in the 1920s represented the first major technological disruption of print's dominance. Radio's capacity for real-time broadcasting created an entirely new relationship between audience and story—no longer was news bound by the delays of printing and physical distribution. Critically, radio also represented the first "background medium" that

could be consumed while engaged in other activities, fundamentally altering when, where, and how information was consumed (Douglas, 2013). The government, recognizing the power of the new communication medium, passed the Radio Act of 1927 establishing the principle that the electromagnetic spectrum was a public resource requiring government regulation and introducing a regulatory framework that would later extend to television broadcasting (Craig, 2000). Radio inspired the development of new business models based on advertisements, with money increasingly paid not for physical circulation but for access to people's attention. This paved the way for television's rise in the 1950s, a radically transformative technological shift that fundamentally reshaped how Americans experienced current events and emotional engagement with news. Television (TV) creates a fundamentally different experience of current events, privileging emotional impact and visual spectacle over abstract reasoning and detailed analysis. As Neil Postman argued in "Amusing Ourselves to Death" (2005), TV's primary imperative is to entertain rather than inform, a function embedded in its technological DNA. Unlike print media, which could sustain attention through complex ideas and nuanced argumentation, TV demanded visual stimulation and emotional engagement to maintain viewership.

While Postman argues that TV as a medium (with its live sights and sounds) is architecturally bent towards entertainment, the imperative that content be entertaining also came directly from TV's business model: while newspapers generated revenue through both subscriptions and advertising, television relied almost exclusively on advertising dollars, which were allocated based on audience size and demographic composition. This created an inherent tension between journalistic values and financial imperatives. As Les Brown observed in his analysis of television economics, the business of television is not primarily to inform or educate, but to deliver audiences to advertisers (Brown, 1971). This reality fundamentally shaped editorial

decisions about what stories to cover and how to present them, and the consequences for news presentation (and presenters) were profound. TV news directors discovered that certain types of content—particularly stories with strong visual components, human interest angles, and dramatic narratives—attracted and retained viewers more effectively than abstract policy discussions or complex analysis. Coverage of crime, disasters, celebrities, and political conflicts gradually displaced more substantive but visually un compelling topics like economic policy or international diplomacy. The aphorism "if it bleeds, it leads" became standard practice in TV newsrooms precisely because graphic, emotional content drove ratings (Kerbel, 2018).

The introduction of cable news in the 1980s only intensified TV journalism's orientation towards entertainment. Faced with the need to fill 24 hours of programming, networks like CNN pioneered new formats that blurred the line between news and entertainment. Talk shows, panel discussions, and personality-driven commentary expanded to fill airtime at minimal production cost (Zelizer, 1992). The subsequent emergence of explicitly partisan cable channels like Fox News and MSNBC further transformed news into ideological entertainment, catering to audiences seeking affirmation rather than information (Jamieson & Cappella, 2008). As Kuypers (2014) argues, the rise of partisan cable networks represents not an anomaly but a return to the historical norm of American journalism, where editorial bias was a powerful driver of audience loyalty. The true anomaly was the technological medium itself, and its architecturally entertainment-driven approach to news had profound effects on public understanding and civic engagement. As political communication scholars Shanto Iyengar and Donald Kinder (1987) demonstrated through experimental research, television news tends to produce "episodic" rather than "thematic" frames, presenting issues as discrete events rather than ongoing social patterns.

This framing discourages systemic understanding and often leads viewers to attribute responsibility to individuals rather than institutions or policies (Iyengar & Kinder, 1987).

The entertainment imperative also transformed political communication itself. Politicians and their advisors recognized that television rewarded emotional appeals, memorable sound bites, and visual symbolism over detailed policy proposals. The success of figures like Ronald Reagan, who masterfully adapted his communication style to television's requirements, demonstrated how the medium privileged performance over substance (Jamieson, 1988). Political campaigns increasingly hired media consultants and advertising professionals to craft television-friendly messages, further blurring the line between political discourse and entertainment.

The internet's emergence as a mass medium in the 1990s represented perhaps the most fundamental technological disruption to media to date. Unlike previous one-to-many broadcast technologies, the internet introduced a many-to-many communication architecture that fundamentally reconfigured information distribution patterns. This created a "network society" where power flowed through decentralized information channels rather than hierarchical institutions (Castells, 2010). This architectural shift challenged core assumptions about who could publish, how information circulated, and what constituted authoritative knowledge. The internet's near-zero marginal cost of content distribution effectively eliminated scarcity as the economic basis of media. When anyone could publish globally at minimal cost, the traditional justification for large media corporations—their ability to finance expensive printing presses and distribution networks—weakens dramatically (Benkler, 2006). This technological shift also unbundled information from its physical and temporal forms, separating news from newspapers,

music from CDs, and videos from broadcast channels. Unlike broadcast schedules or publishing cycles, online content became perpetually available and asynchronously consumable.

Perhaps most significantly, the internet collapsed the distinction between producers and consumers. In what media scholar Axel Bruns termed "produsage," ordinary users became content creators through blogs, forums, and personal websites (Bruns, 2008). This participatory culture emerged first through rudimentary personal webpages and email lists, expanded through user-generated platforms like Wikipedia (2001), and exploded with the emergence of blogs and comment systems. Services like Blogger (1999) and WordPress (2003) further simplified publishing. By 2004, the rise of what O'Reilly termed "Web 2.0" marked a shift toward platforms designed specifically for user content creation—comments, ratings, uploads, and other participatory features became standard (O'Reilly, 2005). The technological architecture of the internet also fostered profound changes in information consumption patterns. Hyperlinks enabled non-linear reading, allowing users to follow associative pathways rather than editor-determined sequences. The rise of powerful search engines—particularly Google after 1998—transformed information retrieval from a content-centric to a query-centric process. Rather than browsing through editorially curated collections, users could directly query massive indexes to find specific information, bypassing traditional gatekeepers entirely (Halavais, 2009).

For legacy media organizations, these technological disruptions represented an existential threat. The internet undermined their economic foundation through audience fragmentation and advertising disaggregation, and established media organizations no longer monopolized public information flows. Bloggers broke major stories, citizen journalists captured events on mobile phones, and user-generated content platforms facilitated first-person accounts without institutional oversight. The result was paradoxical— simultaneously more democratic and more

chaotic, offering unprecedented access to diverse perspectives while undermining shared frameworks for evaluating information quality (Lemann, 2006). This forced legacy media organizations to fundamentally reinvent themselves: newspapers established digital presences, television networks developed streaming platforms, and journalists acquired multimedia skills. Yet the transition proved difficult, and many organizations struggled to maintain quality journalism while adapting to radically different economic and technological conditions. Between 2008 and 2020, newspaper employment in the United States declined by more than 50%, with a disproportionately severe impact to local journalism (Abernathy, 2020).

Amidst this turmoil the Internet had already wrought, social media platforms emerged in the mid-2000s as the next transformative technology to impact the media landscape. Unlike previous technological innovations that primarily changed content delivery mechanisms, social media reconstructed the entire social architecture of information sharing. These platforms—including Facebook, YouTube, Twitter, Instagram, and later TikTok—created what scholars dubbed "social media logic:" a set of norms, practices, and structures that privilege connectivity, popularity, and programmability over traditional journalistic values (Van Dijck & Poell, 2013). The defining characteristic of social media is its network structure—content spreads not through centralized distribution channels but through social connections and algorithmic amplification. This technological architecture placed ordinary users at the center of information distribution and prioritized "spreadable media"—content designed to circulate through social sharing rather than institutional broadcasting (Jenkins et al., 2013). While early web technologies had already democratized content creation, social media platforms accelerated the shift by instituting easy-to-use interfaces, quick methods to react and comment on posts, and integrated sharing tools that

lowered barriers to participation. The result was a massive scaling of user-generated content across increasingly interconnected networks.

At the heart of social media's power lies algorithmic curation— computational systems that, much like AI, train on the massive volume of user-generated data on social media platforms to “learn” which content should appear in users' feeds to maximize engagement (and profit for the social media companies). These algorithms represented a profound shift from previous media gatekeeping, replacing editorial judgment with opaque computational processes optimized for engagement metrics rather than civic or journalistic values (Gillespie, 2018). Facebook's News Feed algorithm, introduced in 2006 and continually refined thereafter, made individualized determinations about content visibility based on factors including user behavior, social connections, and engagement patterns. This algorithmic gatekeeping created what Eli Pariser termed “filter bubbles”—personalized information environments that reinforce existing beliefs while filtering out contrary perspectives (Pariser, 2011). The economic model underpinning these platforms further transformed media dynamics. Social media companies pioneered “surveillance capitalism”— an economic system that monetizes user data through personalized advertising (Zuboff, 2019). Unlike traditional media, where content and advertising maintained clear separation, social platforms integrated sponsored content directly into attention flows, often blurring the distinction between organic and paid material. This model incentivized maximum user engagement as platforms competed for attention that could be converted into advertising revenue. Research demonstrated that content triggering strong emotional responses—particularly outrage, fear, or tribalism—spread faster and broader on social networks than neutral or balanced information (Vosoughi et al., 2018). All of these architectural decisions had tremendous impact

to the broader media environment, especially as the share of American adults who encountered news through social media feeds increased, reaching a majority (54%) in 2024 (Pew, 2024).

For legacy news organizations, social media introduced an uncomfortable dependency. Platforms like Facebook became critical distribution channels, and many publications started to receive the majority of their referral traffic through social sharing. This placed editorial decisions at least partially under algorithmic control, with news outlets forced to create content specifically engineered to maximize algorithmic amplification. Publications became dependent on platforms they could neither control nor fully understand (Bell et al., 2017). When Facebook adjusted its algorithm in 2018 to prioritize "meaningful social interactions" over news content, many publications experienced dramatic traffic declines, highlighting their vulnerability to platform policies (Ananny, 2018).

Where before journalistic gatekeeping restricted information flow, social media enabled the rapid spread of content regardless of accuracy or quality. This created ideal conditions for misinformation to spread during the 2016 U.S. election, with research finding that fabricated news stories often outperformed legitimate journalism in engagement metrics (Allcott & Gentzkow, 2017). Similarly, organized propaganda efforts— manipulation campaigns using automated accounts, coordinated networks, and targeted advertising— exploited platform vulnerabilities to artificially amplify specific narratives (Howard & Wooley, 2018). Algorithmic amplification also intensified political polarization dynamics, as engagement-maximizing systems learned to promote content that reinforced users' existing beliefs and triggered strong emotional responses. Research by this point had already documented how partisan media fostered "echo chambers" in cable news (Jamieson & Cappella, 2008), but social media algorithms accelerated this process by automatically identifying and reinforcing partisan identity

through its selection of content. Studies demonstrated that social media exposure correlated with increased affective polarization, defined as intensified hostility toward those with opposing partisan beliefs (Lelkes et al., 2017). The result was an increasingly fragmented public sphere where shared factual understanding became increasingly difficult to maintain.

By the late 2010s, these technological dynamics produced what media scholar danah boyd described as the modern “attention economy,” where information competed for finite user attention through increasingly sensational appeals (boyd, 2010). Traditional journalistic gatekeepers were replaced with algorithmic systems that rewarded emotional reactivity, in-group signaling, and conflict narratives. Even high-quality media outlets who had successfully transitioned to the Internet era were forced to compete in the attention economy, optimizing headlines and stories for social sharing rather than pure informational value. For individual information consumers, social media contributed to “attention fragmentation”—the division of cognitive resources across multiple information streams. The average Facebook user scrolled through hundreds of distinct content items daily, with limited cognitive capacity for critical evaluation (Mark et al., 2005). This environment privileged what scientists have called “System 1” thinking—rapid, intuitive judgment—over “System 2” deliberation that requires sustained attention and analytical reasoning (Kahneman, 2011), leading people to increasingly judge content credibility based on social signals rather than independent rational evaluation (Sunstein, 2009).

This complex convergence of social media, algorithmic curation, and increasingly advanced AI systems have conspired to create our modern media environment, where McLuhan's insight that “the medium is the message” takes on novel significance (2017). By 2024, AI technologies were not merely subjects of news coverage but active participants in reshaping the

communication ecosystem itself. As AI systems increasingly mediate our information environment— from content recommendation algorithms to generative models producing text, images, and video— they fundamentally alter how messages are created, distributed, and interpreted. This latest technological shift ensures that AI is destined to become both the subject of media discourse *and* the technological architecture through which that discourse circulates. This recursive relationship between AI as content and AI as medium creates a new unique challenge for media researchers seeking to understand how information about AI is disseminated and perceived by the public.

A final technological media innovation— podcasting— emerged as a dominant medium in the leadup to the 2024 U.S. election cycle, reflecting a broader trend toward more personalized, direct, and conversational forms of news consumption. For the first time in 2024, the majority of American adults 18 years of age and older (53%) listened to podcasts (Buzzsprout, 2024), leading some analysts to dub the 2024 U.S. presidential election as the first “podcast election” (Soto-Vásquez, 2025). Podcasting's architecture—a combination of on-demand audio, subscription-based distribution, and intimate host-listener relationships— represents a fundamental shift in content creation and consumption. Unlike broadcast media's one-to-many model or even social media's many-to-many structure, podcasting creates an “ambient intimacy”—allowing listeners to develop parasocial relationships with hosts through extended, nuanced conversations that traditional sound-bite journalism can’t accommodate (Berry, 2016). The consumption model also enabled listeners to consume substantive political discourse during commutes, exercise, or household tasks, integrating information consumption into daily activities. Podcasting’s relatively low barrier to entry also democratized content creation, allowing political candidates, journalists, and commentators to bypass traditional media

gatekeepers and establish direct communications channels with audiences. The 2024 election saw both major-party candidates appear on podcasts and extended-format interviews, a tacit acknowledgement that podcasting's conversational nature and audience trust offered unique advantages. The rise of political podcasting also coincided with growing distrust in traditional media, with the format's architecture, predisposed towards impressions of transparency and depth, better equipped to evade accusations of biased or superficial coverage in conventional outlets (Soto-Vásquez, 2025).

Despite the numerous technological advances catalogued here that have transformed the American mass media environment, traditional journalistic gatekeepers—including legacy newspapers like *The New York Times*, *The Wall Street Journal*, and *The Washington Post*—continue to exercise significant agenda-setting power through what communication scholars call "intermedia agenda-setting" (McCombs & Guo, 2014) and remain worthy of study, especially when chronicling the spread of information about novel phenomena (including LLMs). Even as direct readership has declined, *The New York Times*, *The Wall Street Journal*, *The Washington Post* and their contemporaries still represent primary sources that shape coverage across the broader media ecosystem. Digital-native outlets, television networks, and even social media discussions frequently respond to reporting first published in these mainstream publications. Studies by Harder et al. (2017) demonstrate that traditional elite newspapers still function as "opinion leaders" for other media, with their framing of issues cascading through the information environment. Katz and Lazarsfeld's "two-step flow" theory argues that information typically flows from media sources to "opinion leaders" who then interpret and disseminate this information to their social networks (2017). Thus, even as direct readership declines, publications like *The New York Times* continue to exert significant influence as their framing of

complex topics like AI proliferates through networks of decision-makers in government, business, and academia. While no single publication can claim to shape public opinion uniformly in today's fragmented media landscape, analyzing how these influential newspapers frame AI provides valuable insight into the dominant narratives available to decision-makers and opinion leaders. These frames establish the initial interpretive frameworks through which AI technologies are understood.

As Postman (2005) described in his extension of McLuhan's thesis, each medium privileges certain types of content and discourse— e.g., newspapers favor analytical depth while social media platforms amplify emotional resonance and controversy. Understanding how AI is framed in legacy media therefore requires attending not just to the explicit content of coverage, but to the architectural affordances of the medium itself. When mainstream media publications like *The New York Times*, *The Wall Street Journal*, and *The Washington Post* report on AI, they are simultaneously describing a phenomenon *and* participating in an information environment increasingly shaped by that phenomenon. As McLuhan argued, "any understanding of social and cultural change is impossible without a knowledge of the way media work as environments" (McLuhan, 2017, p. 26). To understand how AI is being presented to the public—and how those presentations might influence adoption, regulation, and social integration—requires systematically examining the frames through which these technologies are presented and interpreted by journalists.

Media framing theory offers a powerful analytical framework for this investigation. By identifying how journalists select, emphasize, and organize information about AI, we can better understand the "schemata of interpretation" (Goffman, 1974) that guide public understanding of these complex technologies. This analysis focuses on legacy media framing while

acknowledging that many Americans now receive information through other digital channels. The persistence of these publications' agenda-setting role, especially for policy and business elites, makes them particularly relevant for understanding how AI is being contextualized for key decision-makers in our modern and increasingly fragmented “attention economy” (boyd, 2010). Therefore, the following section examines the theoretical foundations of framing analysis and its application to media coverage, establishing the methodological approach that will guide this thematic analysis of AI coverage in three major American newspapers during the year 2024.

2.3 Framing

The evolution of framing theory is commonly divided into three major stages (Ardèvol-Abreu, 2015). The first phase (1974-1990) marked the inception of framing in communication studies, transitioning from psychology into sociological terrain, while framing was still largely seen as a vehicle for understanding media messages. It was the second phase (1990s) that truly crystallized framing's place within media studies as a specialized field. During this period, scholars actively debated whether framing was simply an extension of agenda-setting theory, which tracks the issues the media chooses to emphasize vs. ignore, or if it represented a distinct theory of its own. In the third phase (beginning in the 2000s), framing theory saw major conceptual refinements, transitioning into a unified framework that allowed for more rigorous empirical testing.

Framing theory first emerged from interdisciplinary academic traditions and can be traced to sociology, cognitive psychology, anthropology, linguistics, economics, social movement studies, policy research, and, more recently, communication studies. At its earliest stages, framing theory is rooted in interpretive sociology, which posits that individuals' interpretations of

reality are not created in isolation but are always mediated by social interactions and the shared definitions of situations (Ardèvol-Abreu, 2015). This social construction of reality is fundamentally shaped by intersubjective processes, wherein individuals rely on others' contributions to define situations and guide their actions. These interpretations not only facilitate understanding of one's environment but also shape behavior and influence interpersonal interactions.

The formal use of the term "frame" in the context we understand today can be traced back to Gregory Bateson (1972), a major figure in cognitive psychology, who introduced the term in the 1950s. Bateson perceived frames as cognitive structures that shape one's perception by delineating what is relevant or significant within a particular context. He stated that "any message, which either explicitly or implicitly defines a frame, ipso facto gives the receiver instructions or aids in his attempt to understand the messages included within the frame" (Bateson, 1972, p. 188). He employed the analogy of a picture frame and developed a theory to explain how frames select and exclude information. The picture frame organizes perception by directing attention to what is within its borders while excluding everything outside its scope. Bateson's seminal argument was that frames provide essential contextual understanding to messages, acting as "keys" to interpretation that guide and narrow how reality is understood. This served as one of the primary conceptual backbones for later media studies and framing theory (Tuchman, 1978).

Building upon Bateson's psychological foundation, Erving Goffman (1974) contributed to formalizing framing theory within the context of social interaction. Goffman was instrumental to expanding the original meaning of frame "from the individual to the collective, from the psychological to the sociological realm" (Ardèvol-Abreu, 2015, p. 428), as for Goffman, these

frames allow societies to maintain a shared interpretation of reality and create a shared social discourse. These "schemata of interpretation" become crucial in communication, as they allow individuals to share an understanding of reality. Frames, in this sense, are used collectively and largely shape how society perceives and discusses both everyday events as well as more complex social phenomena. Goffman's work suggested that media outlets, social groups, and institutions heavily rely on these frames to give structure and meaning to external realities, indicating the centrality of framing in communication processes. Freed from its individual and psychological roots, framing was soon applied to the study of media and communication.

In the second phase, Robert Entman (1993) applied and refined the concept of framing to the context of mass communication, positing that framing involves at least four interconnected functions: defining a problem, diagnosing causes, making moral judgments, and suggesting remedies (Entman, 1993). When media outlets choose certain details to emphasize—such as specific angles, language, or imagery—these choices influence audiences' understandings of what the central issues are, who or what is responsible, and how (or whether) problems should be solved. In other words, the process of framing guides the audience to think about a problem in a particular way, even if other interpretations remain possible. Benoit (1995) further introduced a typology for understanding how individuals, organizations, and public figures respond to threats or damage to their reputations. Benoit's focus on strategic messaging, particularly the selection of specific narratives or "accounts" to deflect blame or criticism, parallels the idea of framing. His framework underlines how message producers can highlight or minimize certain aspects of a crisis to shape public perception, ultimately determining how audiences judge the events and actors involved.

Central to these approaches is the notion of *salience*: by making certain information more “noticeable, meaningful, or memorable” (Entman, 1993, p. 53), media frames direct attention and shape how audiences prioritize issues. As Kuypers (2010) emphasizes from the rhetorical perspective, frames operate as persuasive structures that shape public meaning. These frames can be explicit—seen in headlines, images, or repeated keywords—or they can be implicit, embedded in broader narratives and assumptions about the social world. In many instances, the power of frames lies in their subtlety; audiences may not always be aware of how a narrative is guiding them to see an event from a particular angle or how alternative viewpoints have been deemphasized (Kuypers, 2010). Together, Goffman’s foundational framework, Entman’s practical breakdown of media framing, and Benoit’s emphasis on message control underscore how framing is not simply about information but about shaping *interpretation*. In the context of emerging technologies, including artificial intelligence applications like ChatGPT, framing becomes a critical tool for understanding how the public and the press perceive these developments. Whether AI is portrayed as a groundbreaking innovation or as a source of potential societal harm, the frames used have significant implications for public understanding.

While framing theory provides a powerful conceptual framework for understanding how media shape public perception of issues like artificial intelligence, researchers require systematic methodological approaches to identify and analyze these frames in practice. This has led numerous communication researchers to integrate framing theory with other established qualitative methodologies that analyze and systematically detect patterns in textual data. Among these approaches, thematic analysis has emerged as a particularly valuable methodological complement to framing theory, offering both a rigorous procedure to identify frames and the interpretive flexibility needed to capture novel complexities of modern media discourse. The

following section examines how thematic analysis provides the methodological foundation for identifying frames in media coverage, as this study will demonstrate.

2.4 Thematic Analysis in Frame Identification

Thematic analysis has emerged as a foundational qualitative research method for systematically identifying, organizing, and interpreting patterns of meaning (themes) across datasets. This approach is particularly valuable when examining news articles, where meaning is conveyed through language choices, narrative structures, and contextual elements. Braun and Clarke (2006) established one of the most widely adopted frameworks for thematic analysis, describing it as "a method for identifying, analyzing, and reporting patterns (themes) within data" that "organizes and describes [the] data set in [rich] detail" (p. 79). Thematic analysis offers both flexibility and methodological rigor, allowing researchers to adapt their analytical approach to address specific research questions and contexts.

At its core, thematic analysis involves discerning meaningful patterns within data through a process of familiarization, coding, theme development, refinement, and interpretation. Braun and Clarke (2006) identify six phases in this process: familiarization with the data, generating initial codes, searching for themes, reviewing themes, defining and naming themes, and producing the report. This systematic yet flexible approach enables researchers to connect empirical findings with theoretical frameworks. Thematic analysis can be conducted through either an inductive approach, where themes emerge organically from the data without predetermined coding frameworks, or a deductive approach guided by existing theoretical concepts. An inductive approach is particularly valuable for exploratory research into emerging phenomena (such as AI) where predetermined categories might constrain analysis or miss novel

patterns (Fereday & Muir-Cochrane, 2006). This strategy allows researchers to identify themes as they emerge within specific contexts before connecting them to broader theoretical constructs.

The strength of thematic analysis lies in its capacity to identify both semantic content (explicitly stated ideas) and latent content (underlying assumptions or ideologies). When applied to media coverage, this method enables researchers to trace not just what topics receive attention, but how they are characterized and contextualized through specific linguistic and narrative choices. The method acknowledges that meaning emerges not only from what is explicitly stated but also from what is emphasized, deemphasized, or omitted entirely from coverage. While thematic analysis offers flexibility, it maintains methodological rigor through systematic coding processes, reflexive analysis, and transparent documentation of analytical decisions. These properties make it particularly well-suited for examining how complex topics like artificial intelligence are represented and interpreted in news media, revealing the underlying patterns of meaning that structure public discourse around emerging technologies. The marriage of thematic analysis methodology with framing theory represents a powerful analytical approach to assess how news media construct and disseminate narratives about complex phenomena such as artificial intelligence. While framing theory provides the conceptual framework for understanding how media shape public discourse through emphasis and organization of information, thematic analysis offers the methodological tools to systematically identify and analyze these frames through close examination of textual patterns.

The relationship between themes and frames in this integrated approach requires clarification. Themes represent patterns of meaning identified through coding; they are the constituent elements that, when analyzed in relation to each other, reveal broader interpretive frames. As Van Gorp (2010) explains, frames are "more encompassing than a single argument,

metaphor, or theme... The various framing and reasoning devices together form a media package that functions as a whole" (p. 91). Through thematic analysis, researchers identify recurring themes, then interpret how these themes cluster and relate to form coherent frames that organize public understanding of issues. For example, individual themes like "job automation," "efficiency gains," and "productivity metrics" might collectively constitute an "economic impact" frame that interprets AI primarily through its effects on labor markets and business operations. This study will utilize both concepts of "frames" and "themes" in its analysis in accordance with how they have been defined and delineated within this section— with frames operating as broader interpretive sets, and themes situated within these larger frames.

The integration of thematic analysis and framing theory also aligns with Goffman's (1974) foundational conceptualization of frames as "schemata of interpretation" that enable individuals to "locate, perceive, identify, and label" experiences (p. 21). Thematic analysis provides the methodological pathway to identify how these interpretive schemas are constructed in media discourse through recurrent patterns of language, emphasis, and narrative structure. By focusing on both the semantic and latent content of texts, this approach can capture not only explicit framings but also the implicit assumptions and cultural resonances that give frames their persuasive power. The resulting analysis reveals not just what frames dominate media coverage of artificial intelligence, but how these frames are constructed through specific linguistic choices, metaphors, exemplars, and narrative structures—offering a richer understanding of media meaning-making than either approach could provide in isolation.

Before turning to methodology and analysis, this study will assess the current state of relevant research into media framing and thematic analysis of AI coverage. While researchers have previously identified frames in AI media coverage, this study employs an inductive

thematic analysis approach that allows frames to emerge organically from the 2024 dataset rather than imposing predetermined categories. This methodological choice acknowledges that media framing of AI continues to evolve alongside the technology itself— just as AI continually expresses novel capabilities over time, researchers may continually observe novel frames that previous research has not captured. Therefore, while the literature review will discuss frames identified by other researchers, these will not directly inform the coding scheme for the current analysis. This approach provides valuable context while preserving the inductive integrity of the thematic analysis, allowing for the discovery of frames unique to the 2024 media landscape. The following section examines the relevant research into media framing and thematic analysis of AI coverage.

2.5 Relevant Research

Research into how AI is covered in the media has accelerated in recent years, reflecting the technology's growing societal impact. Recent research has explored how media framing can shape public perceptions of AI as either a beneficial innovation or a looming existential threat. This review examines the dominant frames identified across the literature, while acknowledging that the rapidly evolving nature of AI technologies continues to generate new narrative patterns in media coverage— a reality that informs the inductive methodology of this study.

A content analysis of major U.S. newspapers from 2009-2018 found that articles about AI predominantly fell under business and technology topics, often highlighting AI's promise and benefits (Chuan et al., 2019). Notably, the benefits of AI were mentioned more frequently than the risks during this period, though discussions of risks (when they occurred) tended to be more detailed and specific. Common positive frames presented AI as a driver of efficiency, economic

growth, and scientific advancement (e.g., AI breakthroughs curing diseases or boosting productivity). A recent review of 30 studies concluded that mainstream news reporting on AI was largely positive in its evaluations and frequently used an economic or techno-optimist framing (Brause et al., 2023). In other words, the press has historically often portrayed AI as a valuable innovation to be embraced, framing it as an opportunity (for businesses, consumers, or society) rather than primarily as a threat.

Multiple studies point to the strong influence of the technology industry in shaping these media frames. Analyses in the US, UK, and Canada found that AI news was to a significant extent "led by industry sources" and often took corporate claims about AI at face value (Nielsen, 2024). For example, a 2018 study of UK media coverage found that nearly 60% of AI-related news articles were pegged to industry announcements or products, and one-third of all quoted sources were from the tech industry— about double the share of academic sources (Brennen, 2018). A more recent study by these authors found that news sources "construct the expectation of a pseudo-artificial general intelligence: a collective of technologies capable of solving nearly any problem" (Brennan et al., 2022, p. 22). A similar study in Canada noted that technology news, including AI coverage, tended to be "techno-optimistic" and presented more as business news than as critical technology reporting (Dandurand, 2023). This suggests that the framing of AI in the news has often aligned with the narratives promoted by companies and developers: AI is framed as the next big technology breakthrough, and the media echoes this framing found in press releases about new AI-powered products or initiatives. The risk, as some analysts argue, is that journalism can become an unwitting participant in the hype cycle of AI. Rasmus Nielsen observes that coverage frequently "takes [industry] claims about what the technology can and

can't do... at face value", thereby contributing to inflated expectations (Nielsen, 2024) and crowding out more skeptical analysis of AI's limitations or downsides.

Conversely, research has also found a recurring dystopian theme within media coverage. Particularly in recent years, as AI systems have advanced and gained visibility with the public, journalists have also aired concerns: loss of jobs due to automation, threats of surveillance and privacy invasion, algorithmic bias and unfairness, autonomous weapons, and even existential risks from superintelligent AI. These concerns sometimes manifest in dramatic storytelling (e.g. references to unstoppable AI and catastrophic outcomes), often termed the "AI Panic" frame (Weiss-Blatt, 2023). This analysis finds that some outlets tend to fixate on speculative future disasters (such as AI causing human extinction) to generate sensational headlines that draw attention while ignoring present-day ethical issues. For instance, while a news story might run with scientists warning about a hypothetical "superintelligence" scenario, it might downplay current real-world problems like discriminatory bias in facial recognition or the environmental impact of large AI models (Weiss-Blatt, 2023). This tension between future fears and present criticism is a recurring one. Variation by geography and culture has also been observed, with national priorities and cultural outlooks influencing how the public views AI (Wang & Liang, 2024).

A particularly prominent ethical frame concerns algorithmic bias and fairness. Media coverage has highlighted how AI systems can perpetuate or amplify existing social biases, particularly racial and gender discrimination. Logan (2019) argues that corporations should consider what responsibility they have "to human beings and society, especially as it pertains to race" (p. 977), a concern that is echoed in coverage that often calls on AI firms to address racialized harms. Research shows that news stories often emphasize individual cases where

algorithmic systems produced biased outcomes—from discriminatory hiring tools to facial recognition systems that perform poorly on darker skin tones—framing AI as potentially reinforcing structural inequalities rather than providing objective decision-making (Joyce et al., 2021). This coverage tends to emphasize the human consequences of algorithmic bias, often featuring the voices of those negatively affected alongside experts who call for more diversity within development teams and greater robustness of testing protocols.

Another major frame concerns the impact of AI-generated content on artists, creators, and cultural production. As AI systems are now capable of producing convincingly human images, music, text, and even film-like output, stories increasingly explore the tension between AI as a creative tool and as a threat to human creators. One common narrative frames AI-generated art and text as an innovative tool that can augment creativity and productivity. These stories might highlight, for example, how graphic designers use AI image generators to brainstorm ideas, or how filmmakers de-age actors using AI, positioning the technology as empowering or collaborative. However, an opposing frame emphasizes an existential threat or ethical breach: AI-generated content is portrayed as "art-stealing" or "artist-replacing" technology. Reporting on controversies in the art community exemplifies this frame: many news pieces have covered artists' backlash to AI image models trained on millions of online artworks without consent, describing it as a violation of copyright and artistic labor. Viewpoints starkly diverge between AI as a tool for creativity, to an art-destroying, unethical technology trained on millions of copyrighted images used without permission (Epstein et al., 2020). In journalism and publishing, similar debates are framed around AI-written articles or books -- are they a useful aid for writers or a shortcut that produces soulless content and displaces authors? Recent events, such as the 2023 Hollywood writers' and actors' strikes, received substantial media coverage framing AI as a

key issue: studios were seen exploring AI-written scripts or digital replicas of actors, while creative professionals demanded regulations to protect their jobs and likenesses. The dystopian version of this frame warns of AI leading to widespread job loss in creative fields and a devaluation of human artistry. The utopian counter-frame (often put forward by tech advocates) argues that automating drudge work via AI could free artists to focus on higher-level creativity, or that new forms of art will flourish from human-AI collaboration (Hitsuwari et al., 2023). Media stories on AI art auctions, awards won by AI-generated images, or AI music compositions also feed into a legitimacy and cultural value framing: questions are raised about whether AI-generated works can be considered "real art" and who deserves credit. In essence, the framing of AI in creative domains oscillates between augmentation and fears of replacement. The outcome of this framing battle has practical implications for public acceptance of AI in culture and for policy (such as in copyright law reform, which media report is now grappling with whether AI creations can be copyrighted and how to handle training data obtained from existing artworks).

Another emerging but increasingly significant frame in AI media coverage concerns the environmental impact of these technologies. Research indicates that news outlets have begun highlighting the substantial energy consumption and carbon footprint of training and operating large AI models (Crawford, 2021). This frame positions AI development within broader sustainability conversations, emphasizing the material resource demands—electricity, water for cooling data centers, rare minerals for hardware—that underlie seemingly intangible digital technologies. Media stories frequently cite striking statistics about how training a single large language model can emit as much carbon as five cars over their lifetimes or consume water equivalent to that used by a small town for a year (Strubell et al., 2020). Coverage often includes

discussions of potential solutions, including more efficient algorithms, renewable energy for data centers, and calls for transparent reporting of AI's environmental costs.

Underlying many of the above is a broader frame about trust and ethics in the digital public sphere. The emergence of AI-generated content— including deepfakes, synthetic art, and algorithmically produced text— has inspired new framing patterns in media coverage. Media outlets consistently position AI-generated content within a "misinformation crisis" framework, portraying deepfakes as potentially devastating to public trust (Chesney et al., 2019). Coverage emphasizes concrete threats— e.g., AI-generated audio or videos influencing elections or inciting conflict— while citing real-world incidents like deepfake hoaxes in European elections (Meaker, 2023). This framing extends beyond concerns about individual deception to warnings about broader epistemic damage: even unsuccessful deepfakes create an environment where authentic content becomes dismissible as potential fakery, a situation that has been termed the "liar's dividend" (Schiff et al., 2025). The resulting narrative presents AI not merely as a technological innovation but as an existential challenge to shared reality itself. These concerns often have a particular resonance for journalism as a profession, which already faces significant disruption from technological change. Research suggests that news organizations may emphasize the threat of AI-generated content to information ecosystems partly because this framing aligns with journalists' personal concerns about credibility and public trust (Sun et al., 2020). While this does not invalidate the legitimacy of these concerns, it does indicate that media professionals' perspectives likely influence how these technologies are framed for public understanding.

Alongside describing the threat, media coverage frequently discusses what can be done about AI-generated deceptive content, centering policy debates and legislative and regulatory responses. For example, U.S. news outlets have reported on new laws passed in states like

California, Texas, and others that ban the use of deepfakes in elections or require disclosure of AI-manipulated political media, alongside bills introduced in Congress to criminalize certain malicious deepfakes or mandate watermarks on AI-generated videos (Weiner & Norden, 2023) . This frame positions AI as requiring governance: discussions revolve around free speech implications, enforcement challenges, and the balance between curbing misinformation and protecting legitimate uses of AI-generated content. Media also note actions by tech platforms (for instance, policies by Facebook or Twitter regarding deepfake removal or labeling) and the development of detection technologies. The policy framing often features voices of experts from law, policy think tanks, and civil society debating solutions. For instance, articles have cited the Brennan Center's reports on deepfake regulation, which warn that while laws are being updated to address the "evolving threat", care must be taken to avoid unintended consequences for legitimate expression (Weiner & Norden, 2023). The media frame here centers on response: AI-generated content is presented as a problem that lawmakers and institutions are scrambling to address, with an implicit question of whether our legal system can keep pace with the technology.

Media coverage also frequently presents AI development as a geopolitical competition, particularly between the United States and China. This "AI arms race" frame positions technological advancement as a matter of national security and economic dominance rather than purely scientific progress or commercial innovation (Lee, 2018). News stories highlight government investments, policy initiatives, and strategic advantages in AI across different nations, often employing militaristic metaphors and zero-sum thinking. This frame tends to emphasize the potential consequences of "falling behind" in AI development, including economic disadvantage and security vulnerabilities. The geopolitical frame has intensified in

recent years as governments worldwide announce national AI strategies and research initiatives, with media coverage often reflecting national perspectives and priorities in its presentation of these developments.

In conclusion, the body of research and discourse on media framing of AI reveals a complex interplay of narratives. Utopian visions of AI's potential are tempered by dystopian warnings of its pitfalls, and these frames compete in news coverage. Media frames around AI are influenced by powerful actors (industry, government, activists), historical cultural narratives, and emerging real-world events (from elections disrupted by deepfakes to artists protesting AI models). Understanding these frames helps explain how public opinion on AI is being shaped. As AI technologies continue to advance and integrate into society, the frames through which media present them may impact policy outcomes and the trajectory of adoption. Ongoing and future research will need to keep tracking these evolving narratives, especially as new developments in generative AI (e.g. new models of ChatGPT and DALL-E) spark fresh waves of media attention. Future research should aim to critically assess whether media framing is providing a detailed, nuanced understanding of AI's promises and perils—and how such framing could be improved to enhance public discourse on this transformative technology.

While previous studies have identified various frames in AI media coverage, this research employs an inductive thematic analysis approach that allows frames to emerge organically from the 2024 dataset rather than imposing predetermined categories. This methodological choice acknowledges that media framing of AI continues to evolve rapidly alongside technological developments, potentially generating novel frames that previous research has not captured. Nevertheless, awareness of existing literature provides valuable context for interpreting emergent

frames while allowing the analysis to remain open to new patterns specific to the 2024 media landscape.

Chapter 3: Methodology

3.1 Thematic Analysis

This study utilized the qualitative research method of thematic analysis to explore media coverage of AI throughout the year 2024 and identify emergent frames. This analysis sought to understand how specific angles, emphases, and language choices shaped public perception of AI. Given the rapidly evolving nature of AI technologies, a thematic analysis approach as outlined by Braun & Clarke (2006) was selected for its flexible yet systematic approach to capturing the range of narratives that emerged in news discourse during 2024. This choice allowed for inductive code development aligned with Goffman's framing theory (1974) while avoiding the constraints of a predefined coding scheme typical of quantitative content analysis.

In line with the strategy outlined by Braun & Clarke (2006), after gathering a robust data set of news articles, initial readings of the articles focused on identifying repeated keywords, metaphors, and rhetorical devices that indicate patterns of coverage. Once these recurring ideas, themes, and rhetorical features were recognized and coded, they were refined and interpreted as media frames. By comparing results across *The New York Times*, *The Wall Street Journal*, and *The Washington Post*, the analysis aimed to capture insights that are broadly generalizable across the political spectrum while also gauging any recurring differences between media outlets. The thematic analysis was designed to answer the following research questions:

RQ1: How did *The New York Times*, *The Wall Street Journal*, and *The Washington Post* frame AI in 2024?

RQ2: What specific fears about AI were emphasized in *The New York Times*, *The Wall Street Journal*, and *The Washington Post* in 2024?

RQ3: To what extent did the framing of AI in *The New York Times*, *The Wall Street Journal*, and *The Washington Post* encourage or discourage its adoption in 2024?

3.2 Sample

This study focused specifically on legacy media outlets—*The New York Times*, *The Wall Street Journal*, and *The Washington Post*—given their enduring role as inter-media agenda setters (McCombs & Guo, 2014) within public discourse. This selection for the analysis aimed to capture a relatively broad spectrum of mainstream political perspectives in American legacy media. *The New York Times* is often characterized as having a center-left editorial stance, *The Wall Street Journal* tends toward center-right positions particularly on economic issues, and *The Washington Post* generally occupies center to center-left territory with a strong focus on political reporting (Pew, 2022). This diversity of political orientation allows the study to examine how AI framing might vary across different ideological positions in mainstream American journalism. Importantly, each outlet maintains dedicated technology reporting teams with specialized knowledge of AI. Their national reach ensures coverage of AI across diverse sectors and geographic regions, rather than a limited focus on technology hubs like Silicon Valley. This combination of political diversity and national influence makes these three publications particularly valuable for understanding how AI is framed for American audiences across the political spectrum.

Only news articles were considered for inclusion in the dataset (omitting opinion pieces, editorials, podcast transcripts, and other forms of coverage). This deliberate exclusion ensured a clear focus on the journalistic reporting of AI, rather than interpretive commentary or opinion-driven content. By limiting the scope to traditional news articles, the analysis emphasizes

framing as it emerged directly from the newsroom, rather than from the opinion writers or podcast hosts, who typically offer more subjective analysis or commentary on previously reported events. This decision also recognizes that legacy media’s journalistic reporting often serves as the foundational content from which other media forms— including podcasts and editorial opinion pieces— derive their narratives. These articles were accessed using the personal subscription of the researcher, ensuring direct and unmediated engagement with original content sources.

This study employed a systematic random sampling approach to select a total of 300 news articles about artificial intelligence published in 2024 (100 from each publication). After comprehensive searches yielded thousands of potential articles from each publication, the researcher filtered out non-news content (opinion pieces, editorials, arts coverage, etc.) to focus exclusively on factual reporting. Data mining and random sampling techniques were then applied to ensure an unbiased selection of articles for analysis. Each selected article underwent a screening process to confirm that artificial intelligence was indeed its primary focus. The following sections detail the specific procedure followed for each publication— *The New York Times*, *The Wall Street Journal*, and *The Washington Post*— to derive the final dataset of 300 articles, 100 from each media outlet.

3.2.1 Sample Selection: The New York Times

For *The New York Times*, the following procedure was utilized:

1. The researcher visited the *The New York Times* website (nytimes.com).
2. The researcher clicked the menu button in the top left, entered “artificial intelligence” into the search bar, and clicked “Go”.

3. Using the 'Date Range' dropdown menu, the researcher selected the start date of January 1, 2024 and the end date of December 31, 2024. This yielded 2,807 results.
4. Using the 'Type' dropdown menu, the researcher selected "Articles". This filter yielded 2,483 results.
5. The researcher built and utilized a Python web scraping tool to aggregate all 2,483 articles into a single CSV file, capturing the articles' title, subtitle, author, section, URL, and a numerical value from 1 to 2,483 to mark the article in the dataset.
6. Articles in the CSV file that were from the following sections (*Arts, Briefing, Opinion, Podcasts, Style, and The Learning Network*) were omitted from the list. This resulted in 1,382 articles.
7. A Python script was used to randomly select 100 numbers with no duplicate values. The corresponding numbered articles from the dataset were chosen to reach a randomized sample size of 100.
8. Each of the 100 articles was subjected to an initial read and eliminated from the final dataset if the subject of artificial intelligence was not the primary focus of the article and was only tangentially or briefly discussed. Each time this occurred, the researcher chose a new randomly numbered article, subjected it to the same test, and either eliminated or included it in the final dataset.

3.2.2 Sample Selection: The Wall Street Journal

For *The Wall Street Journal*, the following procedure was utilized:

1. The researcher visited the *The Wall Street Journal* website (wsj.com).
2. The researcher clicked the search icon in the top right, entered "artificial intelligence" into the search bar, and clicked "Go".

3. The researcher clicked the ‘Advanced Search’ dropdown menu.
4. Using the ‘Date Range’ menu, the researcher selected the start date of January 1, 2024 and the end date of December 31, 2024.
5. Using the ‘Source’ selection, the researcher selected “WSJ Articles,” then clicked Search. This yielded 3,913 results.
6. The researcher built and utilized a Python web scraping tool to aggregate all 3,913 articles into a single CSV file, capturing the articles’ title, author, section, URL, and a numerical value from 1 to 3,913 to mark the article in the dataset.
7. Articles in the CSV file that were from the following sections (*Opinion, Arts, Lifestyle, Real Estate, Personal Finance, Health, Style, and Sports*) were omitted from the list. This resulted in 2,510 articles.
8. A Python script was used to randomly select 100 numbers with no duplicate values. The corresponding numbered articles from the dataset were chosen to reach a randomized sample size of 100.
9. Each of the 100 articles was subjected to an initial read and eliminated from the final dataset if the subject of artificial intelligence was not the primary focus of the article and was only tangentially or briefly discussed. Each time this occurred, the researcher chose a new randomly numbered article, subjected it to the same test, and either eliminated or included it in the final dataset.

3.2.3 Sample Selection: The Washington Post

For *The Washington Post*, the following procedure was utilized:

1. The researcher utilized Google search (google.com) due to the lack of advanced search features on *The Washington Post* website (washingtonpost.com).

2. The researcher typed in “artificial intelligence site:washingtonpost.com” into the search bar, and clicked “Enter”.
3. The researcher clicked “Tools” and selected “Advanced Search.” Using the ‘Date Range’ menu, the researcher selected the start date of January 1, 2024 and the end date of December 31, 2024, then clicked on the search icon. This yielded 2,160 results.
4. The researcher built and utilized a Python web scraping tool to aggregate all 2,160 articles into a single CSV file, capturing the articles’ title, publication date, URL, and a numerical value from 1 to 2,160 to mark the article in the dataset.
5. A Python script was used to randomly select 100 numbers with no duplicate values. The corresponding numbered articles from the dataset were chosen to reach a randomized sample size of 100.
6. Each of the 100 articles was subjected to an initial read and eliminated from the final dataset if the subject of artificial intelligence was not the primary focus of the article and was only tangentially or briefly discussed. Each time this occurred, the researcher chose a new randomly numbered article, subjected it to the same test, and either eliminated or included it in the final dataset.

3.3 Procedure

In accordance with Braun & Clarke’s (2008) six phases of inductive thematic analysis, the researcher completed the following systematic process of thematic analysis:

1. Familiarization with data: The researcher read each of the 300 articles in the final dataset, taking notes on preliminary impressions of the key themes, arguments, and language

used. This process established familiarity with the material and developed initial impressions that contributed to code generation.

2. Generating initial codes: The researcher thoroughly read all 300 articles again with the purpose of generating codes. Using an inductive approach, codes were identified based on semantic content and latent meanings. Extensive notes were taken during this step, and codes were systematically recorded in a thematic analysis matrix. In total, 27 codes emerged during this process.

3. Searching for themes: After coding, the researcher systematically reviewed all codes to identify patterns and note major recurring themes. This iterative process involved grouping related codes and identifying broader patterns of meaning. This phase often necessitated returning to the original articles to ensure contextual accuracy.

4. Reviewing themes: With the list of themes captured, the researcher reviewed the coded extracts within each theme to ensure accuracy. This involved re-reading notes and returning to the original articles to confirm that the themes accurately captured the data.

5. Defining and naming themes: After review, each theme was clearly defined and reinterpreted as a media frame. Extensive notes were taken to capture information about each frame, its function, and its relationship to the research questions. Definitions were developed to establish the scope and boundaries of each theme.

6. Producing the report: Finally, the analysis was written. The researcher aimed to thoroughly define and analyze each theme within the context of the research questions while integrating characteristic quotes derived from coverage that best illustrated the themes identified.

This analysis was completed using the following tools: the Google Chrome web browser to access media outlets' websites and download articles; Terminal and Python to develop the data

mining and aggregation tools; Microsoft Excel to analyze the aggregated data, randomize the sample set, and develop the thematic analysis matrix; and finally, Microsoft Word to capture notes, document thought processes, and write this study.

3.4 Trustworthiness

To ensure the trustworthiness of this qualitative research, several measures were implemented throughout the study design and analysis process. Following Lincoln and Guba's (1994) established criteria for trustworthiness in qualitative research, this study included measures to address credibility, dependability, confirmability, transferability, and authenticity (Cope, 2013).

To address credibility, the researcher compared coverage across three different publications with distinctive editorial orientations. This approach helped ensure that the identified frames represented broader media narratives rather than isolated perspectives. Additionally, the researcher reviewed a randomized sample of 10% of the articles ($N = 30$) after completing the initial analysis to confirm that the identified frames and themes remained consistent upon a second review.

Dependability was addressed through documentation of the research process. The researcher maintained detailed notes documenting each step of the inductive analysis: familiarization with data, generating initial codes, searching for themes, reviewing themes, defining and naming themes, and producing the report (Braun & Clarke, 2008). This documentation provides transparency about the interpretive process and enables validation of the analysis.

To address confirmability, the researcher practiced “reflexivity” (Lincoln & Guba, 1985) by acknowledging potential biases and preconceptions about AI. The researcher regularly reflected on potential sources of bias (e.g. how a background in communication studies and predisposition towards technology adoption might influence interpretation) to minimize bias in the analysis. The random sampling approach and systematic coding procedures helped minimize selective attention to articles that might confirm preexisting biases.

For transferability, the study provides detailed information about the publications selected, the timeframe examined, and the specific data collection and analysis procedures followed. This thorough description allows readers to translate the applicability of findings to other contexts and media environments.

Finally, to address authenticity, this research aimed to faithfully represent the broad set of divergent viewpoints and perspectives within AI media coverage. The inclusion of publications with varying ideological orientations (center-left, center-right, and center) helped capture a range of positions, and characteristic articles were quoted directly to provide the reader an opportunity to directly observe identified frames. Throughout the analysis, the researcher was attentive to which voices were amplified or marginalized in the coverage and frequently considered how different actors (e.g. tech companies, regulators, the public) were represented in the framing of AI issues.

3.5 Significance

Understanding how major legacy U.S. media outlets framed artificial intelligence (AI) in 2024 provides insight into the narratives shaping public perception and policy decisions during a pivotal year for the technology. This study analyzes the role of media as an “intermedia agenda

setter” (McCombs & Guo, 2014) within a “two-step flow” process (Katz & Lazarsfeld, 2017)—where opinion leaders first receive dominant frames from traditional media and then interpret and disseminate this information across the modern fractured media environment. This study’s inductive thematic approach allows for capturing emergent, nuanced frames that may emerge in response to new developments or advancements in AI. These dominant frames not only shape public perception, they also have the potential to influence the trajectory of AI development and adoption across various sectors.

What makes this 2024 study particularly significant is the unprecedented recursive relationship between AI as subject and AI as medium. For the first time in human history, AI systems are capable of producing convincingly human text, images, audio, and video, making them not merely topics of coverage but active participants in reshaping the communication ecosystem itself. As this technology increasingly mediates our information environment—from existing content recommendation algorithms to novel AI-generated news and entertainment content—AI creates a unique feedback loop where the technology being discussed is simultaneously transforming the very means through which that discussion occurs.

This research also contributes to the theoretical landscape of media framing through its marriage of thematic analysis and framing theory, providing a methodological example that future studies of emergent technologies in media discourse can follow. The insights drawn from this thesis can inform policymakers, technologists, educators, journalists, and citizens.

Chapter 4: Results

4.1 Research Question 1

The first research question in this study asked: How did *The New York Times*, *The Wall Street Journal*, and *The Washington Post* frame AI in 2024? The thematic analysis of news articles from *The New York Times*, *The Wall Street Journal*, and *The Washington Post* yielded eight distinct frames that dominated media coverage of artificial intelligence in 2024. The frames are: AI Boom vs. Bubble; Misuse and Misinformation; Ethical and Moral Challenges; Politics and Governance; Societal and Cultural Impact; Work and Automation; Environmental Impact; and Technological Advancements and Future Risks.

Following the methodological integration of thematic analysis and framing theory outlined in Section 2.4, these identified themes function as frames—coherent interpretive packages that organize meaning and guide audience understanding. As Van Gorp (2010) explains, frames “express culturally shared notions with symbolic significance... [including] stereotypes, values, archetypes, myths, and narratives.” (p. 85). Each frame identified in this analysis contains Entman's (1993) core framing functions: defining problems, diagnosing causes, making moral judgments, and suggesting remedies related to artificial intelligence. The analysis reveals not merely what topics received attention but how these topics were characterized, contextualized, and presented through specific linguistic and narrative choices that construct meaning around AI.

These frames—AI Boom vs. Bubble; Misuse and Misinformation; Ethical and Moral Challenges; Politics and Governance; Societal and Cultural Impact; Work and Automation; Environmental Impact; and Technological Advancements and Future Risks—represent distinct organizing principles through which legacy media outlets structured public discourse about AI in

2024. The analysis of each frame demonstrates how it shapes interpretation through selection and salience (Entman, 1993), highlighting certain aspects of AI while downplaying others, and ultimately guiding readers toward specific understandings of AI's implications for society, economy, and humanity's future. Next, each frame will be defined and analyzed in greater detail.

4.1.1 AI Boom vs. Bubble

A prominent frame that emerged from 2024 media coverage of AI was a “Boom vs. Bubble” dichotomy, concerning whether the frenzy of interest in AI should be seen as a “boom” – a reasonable surge of investment and hype around AI – or as a “bubble” – a speculative period of overinvestment that is destined to crash. Coverage in all three newspapers reflected this tension, balancing both excitement over AI’s economic potential and skepticism about overvaluation and unsustainable spending. This “AI gold rush” (De Vynck, 2024a) driven by Big Tech and venture capital was marked by soaring stock prices and massive research and development (R&D) investments. The WSJ documented this surge of capital investments in several articles and projected that the gold rush would continue, noting that the New York private-equity firm Blackstone “expects demand for around \$2 trillion in generative AI-related investments worldwide” by 2029 (Garcia, 2024). However, the WP noted that tech giants were projected to “...spend around \$60 billion a year... by 2026, but reap only around \$20 billion a year in revenue from AI by that point” (De Vynck, 2024a), and that by mid-2024, financial analysts and even some AI investors started questioning whether the frenzy had outrun the technology’s real value. The WP reported that Wall Street analysts and VCs were “raising concerns about the sustainability of the AI gold rush” (De Vynck, 2024a), arguing the technology might not generate enough revenue to justify the billions invested. Goldman Sachs’s lead tech analyst Jim Covello captured this skepticism, warning in the WP that “overbuilding things the

world doesn't have use for, or is not ready for, typically ends badly" (De Vynck, 2024a). All three outlets highlighted the record-setting valuations (especially of chipmaker Nvidia and AI startups OpenAI and Anthropic) alongside these cautionary voices.

The NYT framed the issue most skeptically, suggesting parallels to past tech manias (like cryptocurrency or the "dot-com" bubble), frequently quoting prominent AI critics (Mickle, 2024), and implicitly asking if AI's rapid rise was driven by genuine utility or investors' fear of missing out. Framing in the WP trended towards the skeptical as well, with one headline pronouncing that "the AI hype bubble is deflating" (De Vynck, 2024b), and another that tempers the evidence of continued success with a cautionary warning: "Nvidia results show AI boom continues despite recent bubble fears" (De Vynck, 2024c). While the WSJ also investigated whether an AI bubble was forming, it did so with more of an optimistic (or at least analytical) tone, primarily emphasizing how difficult it is to accurately predict the outcome of the "boom vs. bubble" debate (Jin et al., 2024) given the complexity of the situation.

Central to this theme were economic fears working in opposing directions— while some worried that the AI financial bubble would burst, leading to massive losses and a collapse in value of over-invested AI companies, others (primarily leading American technology companies) feared underinvesting and missing the window of opportunity. Google CEO Sundar Pichai argued that "the risk of underinvesting is dramatically greater than the risk of overinvesting for [Google]" (De Vynck, 2024a). This echoes the divergence seen in the different coverage when comparing the WSJ and the NYT and WP: whereas companies were encouraged to embrace AI and invest, individuals and the public were warned of the potential risks involved. Since the WSJ caters to a professional audience working in business, it's fitting that their coverage would embrace a more positive framing towards investing in the technology.

4.1.2 Misuse and Misinformation

The "Misuse and Misinformation" frame focuses media coverage through the lens of AI-enabled threats to information ecosystems and trust. This primarily positioned AI with the significant potential to undermine democratic discourse through synthetic media, deepfakes, and automated disinformation campaigns. 2024 was a U.S. election year (and featured major elections globally), and all three newspapers raised alarms about AI's potential for harm to the media environment. The NYT's coverage included this frame most prominently, warning that "...the misinformation and deceptions that A.I. can create could be devastating for democracy..." (Hsu & Metz, 2024) in February of 2024, and releasing a major web interactive piece with the headline "See How Easily A.I. Chatbots Can Be Taught to Spew Disinformation" (White, 2024). The WSJ takes a comparatively cautious stance, simply warning in one headline that "Researchers Warn of Data Poisoning" (Snow, 2024) and using the term "misinformation" only three times across the randomly selected set of 100 articles.

This frame primarily concerns human malicious actors using AI to deceive and the broader "post-truth" dilemma of an AI-saturated media environment. Journalists warned of AI-generated political misinformation, and one particular story— where a subset of New Hampshire residents received AI-generated robocall messages in the voice of President Biden encouraging them not to vote in the primary— was cited several times by the NYT (Metz, 2024; Hsu & Metz, 2024; Frenkel, 2024) and the WP (Verma & Kornfield, 2024; De Vynck, 2024d). Another frequently cited concern was the erosion of trust in images, video, and audio. When the WP wrote about the release of OpenAI's tool Sora, which is capable of generating lifelike videos, the subhead pointed out that "the tool further raises concerns about deepfakes as AI shows up in elections around the world" (De Vynck & Oremus, 2024). Across media coverage, but especially

in the NYT and WP, journalists framed advancements in AI technology within a broader “crisis of trust” – the idea that if any content could be artificially faked, people might cease to believe anything they see or hear. While the WSJ also covered this crisis of trust, it primarily used terms such as “questionable content” or “dubious quality” and focused on impacts to companies and consumers (Mims, 2024) rather than discussing the issue on moral or ethical grounds. The WSJ also highlighted potential solutions or mitigation (like the EU’s labeling requirement or tech firms’ policies) while acknowledging the dystopian possibilities, which brought a moderating pragmatic angle to maintaining information integrity (Mackrael & Schechner, 2024).

4.1.3 Ethical and Moral Challenges

Whereas the previous frame primarily concerned malicious human actors’ use of AI, the “Ethical and Moral Challenges” frame primarily concerns the inherent ethical and moral quandaries within the technology itself— highlighting moral complexities that arise with algorithmic systems. This includes concerns about fairness, accountability, transparency, bias, and the moral implications of AI making decisions or taking actions that affect human lives.

The NYT and WP dominated coverage of bias and fairness in AI outcomes. The NYT, for instance, reported on an AI that was used in Nevada to identify at-risk schoolchildren for extra support and which drastically cut the number of children deemed in need, “leading to tough moral and ethical questions over which children deserve extra assistance” (Closson, 2024), and asked the question “who is left out of the conversation?” when LLMs are primarily trained in English and leave speakers of other languages behind. The WP alternately mentioned “political” bias, “racist and sexist” bias (De Vynck et al., 2024), and even “automation” bias, where people “tend to trust a computer’s decisions” even when it contradicts common sense or training (Hunter, 2024). Each outlet also focused on transparency and consent, pointing to the “black

box” nature of AI and questioning whether people have the right to know when they are interacting with or subject to AI (Roose, 2024b). In one example, the NYT revealed that thousands of physicians were using an AI assistant to draft replies to patient messages, often without the patients’ knowledge. This raises both ethical concerns about informed consent, and whether personal data is going into the training dataset for these systems (Rosenbluth, 2024).

Another significant moral and ethical theme concerned where AI should be permitted to operate unsupervised, especially in the military. The WP reported that internal discussions at OpenAI “expressed discomfort” with the company’s partnership with a weapons manufacturer, Anduril (De Vynck, 2024e). The NYT’s coverage of Meta, which permitted its AI models to be used for U.S. military purposes towards the end of the year, similarly raised moral questions of whether it is ethical for AI to contribute to autonomous weapons or surveillance (Isaac, 2024). The NYT often connected ethical issues to human-interest narratives or case studies to communicate the moral stakes to the audience. This approach grounded AI’s ethical challenges in the present, not as a future theoretical. The WP most frequently framed ethical challenges in terms of accountability and governance, highlighting calls for companies or regulators to impose ethical guidelines (Verma et al., 2024), while the WSJ primarily approached ethical issues from a risk management and policy perspective, communicating ethics in the context of companies and markets.

In addition to concerns about bias, fairness, and transparency, several media accounts in 2024 spotlighted a critical ethical dilemma surrounding copyright and the commodification of creative labor. Several AI companies have been accused of training their models on vast troves of copyrighted material—ranging from journalistic articles to literary works—without obtaining proper consent or providing fair compensation to original creators. *The New York Times* itself

sued OpenAI and its partner, Microsoft, alleging copyright infringement in connection with AI-generated text (Metz, 2024). Critics argue that this practice not only subverts existing copyright laws but also perpetuates a system where tech giants profit immensely from creative outputs that are produced at little or no cost to them, thereby undermining the value of individual authorship and independent media. Such practices raise profound moral questions about the ownership of intellectual labor and the ethical implications of profiting off of content without acknowledging the human work that made AI outputs possible.

4.1.4 Politics and Governance

The "Politics and Governance" frame presented AI development in the context of governments, regulations, geopolitical conflicts, and a global “AI arms race” (Verma & De Vynck, 2024). By 2024, the breakneck speed of AI development had clearly outpaced existing regulations, and media coverage tracked how governments and regulatory bodies around the world – as well as the tech industry itself – responded. This theme includes legislative efforts, international regulations, corporate governance, and the tug-of-war between innovation and regulation.

A major focus was the push for AI regulation at various levels of government. In the United States, much attention was given to the first serious legislative proposals to rein in AI, including California’s AI safety bill (SB 1047) which was described as the nation’s most ambitious attempt to regulate AI (Kang, 2024). The bill would require developers of large AI models to conduct safety tests for “catastrophic harm” (like facilitating cyberattacks or bio-weapons), to implement a “kill switch” to shut down any AI system that ran out of control, and to empower California’s attorney general to sue companies if their AI cause significant harm (death, property damage). All three outlets covered this bill, calling it a “first-of-its-kind” law that could

make California a “standard-bearer” in AI governance. This coincided with a focus on lobbying, with the WSJ reporting that AI startups and tech giants rallied to “kill” the California bill, arguing it “would impose impossibly vague constraints in the name of safety” (Rana, 2024) and framing the opposition to the bill in response to a lack of specifics on how to address compliance issues. The conflict between state initiatives vs. federal action was emphasized: while California forged ahead, “proposals to regulate AI nationally [had] made little progress in Washington” (Rana, 2024).

International coverage focused on the EU’s AI Act and how it restricted the use of “riskier technology” (Kang, 2024). The WSJ explicitly headlined that European lawmakers passed “the world’s most comprehensive legislation yet on artificial intelligence”, detailing how it “sets out sweeping rules” (Mackrael & Schechner, 2024) and imposes new transparency and risk assessment requirements. It noted the law’s global reach (applying to any AI product in the EU market, with penalties up to 7% of global revenue) and quoted EU officials framing it as “a clear path toward a safe and human-centric development of AI.” (Mackrael & Schechner, 2024). The NYT and WP also juxtaposed the EU Act with the inaction of the U.S. government. A core tension within this frame concerned companies’ self-governance vs. external governance—should companies be allowed to set optional safety requirements for themselves? WP coverage described this philosophical rift in the AI community: those who subscribe to “effective altruism” ideals, which “advocates for stricter limits on AI development” (De Vynck & Zakrzewski, 2024) versus those who feel that strict limits will stifle innovation. The NYT highlighted Governor Newsom’s embrace of the latter argument in his veto of the California bill— he argued it “focused too much” on frontier models and potential future risks while overregulating basic uses (Kang, 2024).

This frame also encompasses the global “AI arms race” taking place between companies and nations. All three outlets used this frame, with the WP and WSJ encompassing most of the coverage. Both used the frame to address cooperation and competition between American companies (O’Donovan & Vynck, 2024; Tilley, 2024) as well as competition between American and China (Kao & Huang, 2024; Dou, 2024). That represents a connecting thread within this larger frame: policy and governance attempts primarily characterized by conflict. For instance, in its coverage of the California bill’s passage, WP described a clash that “deepened a rift” in the AI world: on one side, researchers (often aligned with long-term, existential risk concerns and effective altruism) backing strict limits to prevent worst-case scenarios (like AI getting out of control), and on the other side, many tech executives and even some researchers arguing the tech isn’t advanced enough to justify such fears and that heavy regulation would harm innovation and U.S. competitiveness (De Vynck & Zakrzewski, 2024). The WSJ, consistent with its audience, often framed policy in terms of business impact and regulatory philosophy and underscored concerns about overregulation and fragmentation (state vs federal rules). Thus, the media presented a dual set of anxieties driving the governance discussion: fear of AI’s power and fear of losing AI’s promise.

4.1.5 Societal and Cultural Impact

Beyond economics and policy, the coverage often zoomed out to consider AI’s influence on society and culture. This frame captures how AI technologies were portrayed as changing everyday life, social interactions, and cultural norms. It includes both positive stories of AI’s integration into society, as well as anxieties about how AI might alter human behavior, creativity, and social cohesion. One consistent area of focus was AI in daily life and consumer tech, examining how AI assistants and tools became more embedded in routine activities. All three

outlets made mention of Amazon’s plans to overhaul Alexa with more advanced AI to make it “more conversational” (Weise, 2024a). Even such a seemingly benign use case was framed with cultural implications: the WP noted that Alexa and competitors like Google’s Gemini were seemingly programmed to “decline to answer questions about politics” to avoid the appearance of bias or any impact on “a year of consequential global elections” (O’Donovan, 2024).

The media also covered humans using AI chatbots in intimate or social ways, often with a strong negative valence. Both the WP and the NYT (Roose, 2024a) covered the tragic case of a 14-year-old boy dying by suicide “after talking with [an AI] chatbot named after the character Daenerys Targaryen from ‘Game of Thrones’” (Tiku, 2024). Separately, Kevin Roose of the NYT (2024c) reports on the earlier version of an AI that attempted to “break up [his] marriage” and reflected on the worries of “a world where people spend all day talking to chatbots instead of developing human relationships.” A related theme involved AI in creativity and the arts: One WP article described how the AI tool Suno was used to generate songs in the style of a famous band, conveying the mixed feelings this engenders. The musician both saw incredible potential and felt “lingering unease... from a worry that [they would] be screwing the artists [they] love by generating music that sort of sounds like theirs” (Velazco, 2024). This quote encapsulates a fundamental cultural debate mirrored in the coverage— is AI creative output fundamentally different from human output?

A further sub-theme involved social behavior and knowledge: how AI might change the way people learn, make decisions, or interact. One concern was the outsourcing of thinking and the question of whether AI use degrades human intelligence. The NYT cited research and experts worried that users might accept AI outputs uncritically. For instance, an NYT tech columnist observed that users rarely continue their search beyond a quick AI-provided answer, which raises

the risk of misinformation spreading. This suggests a cultural shift towards instant but possibly shallow information consumption. There were also stories about education – schools grappling with students using ChatGPT for homework, or universities integrating AI literacy into curricula, reflecting how cultural norms around intellectual work were evolving.

The NYT tended to frame societal impact stories by exploring how specific groups or domains are adjusting to AI. Its coverage of AI in medicine and education often read as case studies in the social negotiation of technology. While the WP also frequently framed AI's cultural impact in terms of disruption and public reaction, the WSJ approached societal impact with a practical and sometimes upbeat tone, often emphasizing adaptation: workers retooling for the AI era, companies upskilling staff, and everyday tech integrating AI (like new AI features in productivity apps or gadgets). The WSJ did cover controversies (such as the Hollywood writers' strike concerns about AI, which WP also noted), but typically the Journal balanced it with the underlying belief that society will find equilibrium (e.g., highlighting that writers ultimately got protections, indicating a way forward).

4.1.6 Work and Automation

The "Work and Automation" frame positioned AI primarily as an economic force reshaping labor markets, emphasizing tensions between productivity gains and job displacement while questioning how employment and career trajectories will evolve in an AI-integrated economy. The coverage reflected both optimistic visions of increased productivity and dire warnings about job displacement. In one instance, the WP referenced a Goldman Sachs analysis that generative AI could potentially automate "300 million jobs around the world" and significantly boost global gross domestic product (GDP) (Vynck, 2024). This staggering figure was frequently quoted (it made headlines in many outlets) and framed the conversation around

fears that AI might eliminate a large share of work as we know it. WSJ coverage pointed out a surge in demand for AI-related talent even as other tech hiring slowed (Rattner, 2024), highlighting the growing divide between AI-related tech jobs and all other tech jobs. Coverage also focused on how work was being augmented, rather than automated, framing AI as a tool to assist workers rather than replace them. Many articles included examples where AI handled tedious tasks, thus freeing humans for higher-level work. In the field of software development, articles noted AI coding assistants that help programmers be more productive (Cutter, 2024) as well as AI tools that speed up integration, promising to “make setting up and integrating new corporate software systems much faster” (Bousquette, 2024). The WSJ ran several pieces about employees across industries experimenting with GPT-4 or other AI to do parts of their job, often finding increases in efficiency.

There was also significant coverage of organized labor and professional reactions. The Hollywood writers’ and actors’ strikes (settled in late 2023) carried into 2024 coverage, with the WP noting, for example, that Hollywood writers “won protections against being forced to work with AI” in their new contracts— a cultural win that was reported as setting a precedent for other fields (De Vynck, 2024f). Unions and worker groups in other sectors expressed concern with protecting jobs and workers’ rights in the face of AI deployment, with the NYT citing union leaders who pushed for legislation to safeguard workers from unchecked AI (Oreskes, 2024). Another focus was on reskilling and education, with the WSJ prominently covering business schools’ attempts to update curricula (Ellis, 2024) and arguing that employees should take the initiative to upskill in AI (Bindley, 2024). The idea of “AI-proofing” one’s career (by focusing on skills that AI can’t easily replicate, like strategic thinking, creativity, interpersonal skills) was another way the media framed adaptation (Hagerty, 2024).

The NYT often framed coverage critically, questioning whether the AI revolution was truly delivering productivity or just enabling corporations to cut costs. For example, the NYT's profile of Jim Covello (identified as an AI skeptic) warned that "...replacing low wage jobs with tremendously costly technology is basically the polar opposite of [past tech revolutions]" (Weise, 2024b). The NYT tended to pair reporting of companies' enthusiastic adoption of AI with a dose of caution. The WP framed the work and automation debate with a sense of urgency around the human impact. It emphasized real instances of jobs already being affected, noting that "some people who write for a living have already lost their jobs as companies turn to chatbots for advertising or social media copy" (De Vynck, 2024f). The WP also regularly cited polls or studies (like Pew Research) focused on both AI's usage at work and potential to impact jobs. WP articles also mentioned efforts including job training programs for data centers or tech groups aiming to ensure AI doesn't widen inequality. The framing was often one of societal challenge: how do we manage this transition to minimize harm to workers? The WSJ, true to its focus, framed AI in the workplace largely around business strategy and competitiveness. Headlines like "Tech Workers Retool for AI Boom" (Bindley, 2024) or coverage of companies integrating AI into workflows illustrate a forward-looking, adaptive framing. The WSJ thus often leaned into the narrative that AI will change the nature of work (with roles evolving) but tempered fears of displacement with a focus on proactive action workers can take to prepare. It published many practical pieces advising companies on AI adoption in the workplace, encouraging the rapid adoption of AI to stay competitive.

4.1.7 Environmental Impact

The "Environmental Impact" frame contextualized AI development within broader sustainability concerns, highlighting the material resource demands of advanced computing and

raising questions about the ecological footprint of digital infrastructure. Coverage scrutinized the energy consumption, carbon footprint, and resource usage (like water for cooling data centers) associated with AI. In 2024, coverage increasingly asked whether AI's growth is environmentally sustainable, highlighting the energy demands that are projected to increase as AI data centers proliferate. Training large AI models and running them for millions of users requires enormous computational power, which in turn draws a huge amount of electricity. The media reported eye-opening statistics and quotes from industry leaders about this. For example, the WSJ interviewed the CEO of chip-design company Arm, who remarked that AI models are "insatiable in terms of their thirst for electricity" and warned that without efficiency improvements, "by the end of the decade, AI data centers could consume as much as 20% to 25% of U.S. power" (Landers, 2024). The NYT's coverage questioned "whether [companies] can meet the [AI energy] demand while still operating sustainably," (Sisson, 2024) noting the substantial carbon footprint of building and running these facilities. These concerns were often paired with reporting on the search for solutions or mitigations: cleaner energy, more efficient chips, novel cooling techniques, etc. Coverage in all three newspapers included references to sustainable energy, and how some data center operators were planning to build near renewable energy sources (e.g. wind, solar, geothermal, and nuclear). The WP reported that dormant nuclear plants (like at Three Mile Island) were being considered for restart, driven by the "ravenous energy appetites" of AI developers (Halper, 2024). This frame connected AI to broader energy infrastructure decisions. Another common area of focus was the water usage and emissions that contribute to climate change. Large data centers require water for cooling, and multiple articles (particularly in the WP and NYT) mentioned that training a single advanced AI model can consume massive quantities

of water and emit carbon equivalent to many cars. Sustainability experts warned that industry net-zero pledges would be hard to meet with AI's surge.

The NYT and WP tended to frame the environmental impact of AI as a growing challenge that needs foresight, with coverage often focusing on the rapid expansion (lots of new data centers being built) and the question of sustainability. The WP also highlighted OpenAI's calls for government investment in energy infrastructure for AI, acknowledging that even AI's strongest proponents recognize the sustainability issue. The WSJ tended to frame sustainability concerns as an engineering/efficiency problem. Through coverage of Arm's CEO and Nvidia's actions, the WSJ emphasized the need for technological fixes (more efficient chips, better cooling, etc.) to curb AI's energy hunger. It also encouraged investment into sustainable sources of energy, arguing that, since energy costs could become a constraint on AI growth, investors and companies have a vested interest in ensuring there is enough power.

4.1.8 Technological Advancements and Future Risks

The "Technological Advancements and Future Risks" frame deals with how the media framed the frontier of AI technology – the breakthroughs achieved and anticipated – and the future risks those advancements entail. This frame is broad, encompassing everything from new model releases and capabilities in the present to the speculative outlook on artificial general intelligence (AGI), superintelligence, and existential risks in the future. In 2024, coverage oscillated between marveling at AI's progress and warning of its potential future perils. Each outlet covered important AI updates and feature rollouts: including OpenAI's addition of adding voice and image capabilities to ChatGPT, Google's introduction of its Gemini model and AI in search, Meta's release of open-source models, and Elon Musk's AI venture, xAI. The WSJ also ran a forward-looking piece titled "It's the Year 2030. What Will Artificial Intelligence Look

Like?”, gathering experts to share their predictions (Ziegler, 2024). Some predicted gradual progress and doubted near-term “superintelligence” capabilities, while others in that feature envisioned more profound and rapid changes. By presenting these varied predictions, the WSJ framed the future as uncertain but almost certainly transformative.

Competition between AI firms and big tech CEOs also received key coverage. The NYT ran a detailed piece on how a debate between Elon Musk and Google’s Larry Page about AI’s trajectory led to the founding of OpenAI and the current industry boom. The NYT framed it in dramatic terms, noting that “the people who say they are most worried about A.I. are among the most determined to create it” (Metz et al., 2024). A major sub-theme was existential and catastrophic risk – the idea that future AI (if it becomes very powerful, or if misused) could pose extreme dangers to humanity. In 2024, all three papers reported on leading AI scientists and CEOs who were warning about existential risks from AI. The WP noted that in May 2023, dozens of industry leaders released a statement warning “that humanity faced a ‘risk of extinction from AI,” placing this alongside Sam Altman’s congressional testimony that AI could “cause significant harm” (De Vynck et al., 2024). The framing was striking—possibilities that were once relegated to sci-fi (AI wiping out humanity) were now receiving serious discussion in policy circles. The NYT covered this as an ongoing debate in Silicon Valley: “the question of whether artificial intelligence will elevate the world or destroy it” (Metz et al., 2024). Both NYT and WP framed this schism as a philosophical conflict (control vs freedom, caution vs ambition), while the WSJ tended to frame future AI in terms of business and societal outcomes rather than philosophy – e.g., will it boost or harm the economy, or will companies overinvest? The existential discussion was present but not front-and-center in WSJ reporting when compared to NYT and WP.

4.2 Research Question 2

The second research question asked: What specific fears about AI were emphasized in *The New York Times*, *The Wall Street Journal*, and *The Washington Post* in 2024? The analysis reveals several prominent narratives related to fear that appeared consistently across the coverage, though with varying degrees of emphasis between publications.

A dominant fear framed AI with the potential to undermine truth and trust in society through deepfakes, synthetic media, and AI-generated disinformation. This concern was particularly pronounced during the 2024 election cycle, with the NYT and WP emphasizing this threat most strongly. The WP noted that tools like OpenAI's Sora video generator "further raises concerns about deepfakes as AI shows up in elections around the world" (De Vynck & Oremus, 2024). This narrative connected technological capabilities with broader anxieties about democratic discourse and information integrity.

Economic displacement emerged as another significant fear, particularly within the "Work and Automation" frame. Coverage frequently cited a Goldman Sachs analysis suggesting AI could automate "300 million jobs around the world" (De Vynck, 2024a), encapsulating widespread anxiety about labor market disruption. While all three publications acknowledged this concern, the WSJ typically balanced it with more optimistic narratives about adaptation and emerging opportunities, reflecting its business-oriented perspective.

Fears regarding ethical failures and algorithmic bias featured prominently, especially in NYT and WP coverage. These publications emphasized concerns that AI systems would perpetuate or amplify existing societal inequalities through biased decision-making, lack of

transparency, and insufficient accountability mechanisms. Such fears were often illustrated through concrete examples of algorithmic bias in practice rather than abstract speculation.

Environmental sustainability concerns constituted another significant fear narrative, particularly within the "Environmental Impact" frame. The WP's characterization of AI's "ravenous energy appetites" (Halper, 2024) exemplified growing anxiety about the resource demands of AI development and deployment, including energy consumption, water usage for cooling data centers, and subsequent climate impacts.

Although less frequent than other fears, existential risk narratives appeared in the "Technological Advancements and Future Risks" frame. The NYT captured this tension as "the question of whether artificial intelligence will elevate the world or destroy it" (Metz et al., 2024), reflecting industry debates about advanced AI's long-term implications for humanity.

Control and governance failures represented another fear category, appearing across multiple frames but particularly in "Policy and Governance." The WP and NYT emphasized the gap between European regulatory action and American regulatory hesitation, suggesting potential governance failures as AI development outpaces oversight mechanisms. Finally, within the "AI Boom vs. Bubble" frame, financial instability fears emerged, with the NYT and WP particularly warning about parallels to past technological manias and the potential for significant market correction.

These fear narratives were not uniformly distributed across publications. The NYT consistently emphasized ethical, societal, and informational risks; the WSJ focused more on economic and competitive concerns; while the WP often highlighted both immediate practical dangers and longer-term systemic risks. This distribution reflects each publication's broader editorial approach to technology coverage and alignment with their respective audience interests.

4.3 Research Question 3

The third research question asked: To what extent did the framing of AI in *The New York Times*, *The Wall Street Journal*, and *The Washington Post* encourage or discourage its adoption in 2024? The analysis suggests a complex and often contradictory landscape regarding adoption encouragement:

The WSJ consistently presented the most adoption-encouraging frames, emphasizing competitive necessities, productivity gains, and strategic advantages of AI integration. Headlines like "Tech Workers Retool for AI Boom" (Bindley, 2024) and feature articles advising companies on AI implementation reflected an orientation toward practical adoption. The WSJ's business-focused audience likely influenced this approach, as the publication frequently framed AI adoption as essential for maintaining competitive advantage.

The NYT presented the most adoption-discouraging frames, regularly questioning the sustainability of AI investments, highlighting ethical concerns, and featuring critical perspectives on AI's societal impacts. While acknowledging potential benefits, the NYT's framing often emphasized caution, deliberation, and critical evaluation before adoption.

The WP occupied a middle ground, neither consistently encouraging nor discouraging adoption. Instead, it frequently presented balanced perspectives that acknowledged both opportunities and challenges. The WP's coverage emphasized thoughtful, responsible adoption rather than either uncritical embrace or categorical rejection.

Interestingly, adoption encouragement varied not only by publication but also by sector and audience. Business audiences received more adoption-encouraging frames across all three publications, while general public audiences and policymakers received more cautionary,

adoption-discouraging frames. This suggests media publications tailored their framing based on perceived audience needs and contexts.

The framing of AI across these three major publications reveals a media landscape trying to make sense of a rapidly evolving technology with profound implications. The eight identified frames demonstrate how various aspects of AI—economic, ethical, political, societal, environmental, and technological—received attention in 2024 media coverage, with publications emphasizing different dimensions based on their editorial orientations and audience considerations.

Chapter 5: Discussion

5.1 Findings

The thematic analysis of AI media coverage across *The New York Times*, *The Wall Street Journal*, and *The Washington Post* in 2024 reveals several important patterns that contribute to our understanding of how framing shapes public perception of emerging technologies. First, the prevalence of the "AI Boom vs. Bubble" frame aligns with Chuan et al.'s (2019) observation that AI coverage predominantly falls under business and technology categories. However, the 2024 coverage reflects a notable shift from the predominantly optimistic framing identified in previous studies toward a more skeptical evaluation of AI's economic sustainability. While earlier research found that "media narratives have not been dominated by fear-mongering" (Brause et al., 2023), this study identifies a growing emphasis on potential risks and downsides across multiple frames. This evolution likely reflects both the maturation of AI technologies and increased public awareness of their limitations and unintended consequences.

The results also confirm Nielsen's (2024) observation that AI news is significantly influenced by industry sources. The "AI Boom vs. Bubble" frame and "Work and Automation" frame frequently cited corporate announcements and technological advancements, particularly in the WSJ's coverage. However, the 2024 coverage also demonstrates increased journalistic scrutiny of industry claims, especially in the NYT and WP, suggesting a maturation of technology reporting beyond mere corporate boosterism. This represents an important development from earlier findings by Brennan et al. (2018) that 60% of AI articles were pegged to industry announcements, indicating a shift toward more critical, independent analysis.

The prominence of the "Misuse and Misinformation" frame and "Ethical and Moral Challenges" frame indicates a significant focus on AI's potential negative impacts on information

ecosystems and social norms. This aligns with what Weiss-Blatt (2023) termed the "AI Panic" frame but extends beyond speculative future disasters to immediate concerns about electoral integrity and information reliability. The high frequency of these frames suggests that ethical considerations have moved from the periphery to the center of AI discourse in mainstream media.

The findings reveal significant variations in framing between publications, reflecting their distinctive editorial approaches and audience orientations. The WSJ's business-focused perspective, the NYT's critical stance, and the WP's emphasis on societal implications demonstrate how media outlets construct different narratives around the same technological developments. This supports Van Gorp's (2007) assertion that frames are not inherent to issues but are actively constructed through journalistic choices and organizational priorities.

The complex interplay of optimistic and pessimistic frames identified in this study supports the observation by Chuan et al. (2019) that both utopian and dystopian narratives coexist in AI coverage. However, the 2024 coverage reveals a more nuanced landscape than a simple binary. Instead, each publication navigated between multiple dichotomies: innovation versus regulation, economic opportunity versus ethical responsibility, technological progress versus environmental sustainability, automation versus augmentation. These tensions reflect the multifaceted nature of AI's societal implications and the challenges journalists face in capturing this complexity.

Perhaps most significantly, the findings reveal that framing varied not only by publication but also by intended audience and context. While business audiences received more adoption-encouraging frames emphasizing competitive advantage and productivity gains, general audiences encountered more adoption-discouraging frames highlighting risks and ethical

considerations. This contextual variation in framing echoes what Entman (1993) described as the strategic selection and salience of certain aspects of reality to promote particular problem definitions, causal interpretations, moral evaluations, and treatment recommendations.

The identification of the "Environmental Impact" frame as a significant component of AI coverage represents a novel contribution to the literature on AI media framing. This frame highlights growing recognition of AI's material resource demands and ecological footprint— aspects largely absent from previous studies of AI media framing. As technological infrastructure becomes an increasingly visible component of environmental discourse, this frame may become more prominent in future media coverage.

Overall, these findings suggest that media framing of AI in 2024 reflects a maturing discourse that has evolved beyond techno-optimism or simplistic fearmongering toward a more nuanced, multidimensional understanding of AI's implications across various domains of society. The eight identified frames represent distinct interpretive packages through which journalists and audiences make sense of artificial intelligence's ongoing integration into economic, political, and cultural institutions.

5.2 Limitations

While this study provides valuable insights into AI media framing, several limitations must be acknowledged. The explicit focus on three legacy media outlets, while justified by their agenda-setting influence, excludes many other forms of media that increasingly play important roles in shaping public discourse. The fact that podcasting had become a dominant medium by 2024 underscores this. The one-year timeframe (2024) of this study captures a specific moment in AI's developmental trajectory and cannot account for longitudinal shifts in framing that occur

over extended periods of time. The rapid evolution of AI technologies means that media framing can change substantially in short timeframes, which limits the longevity of these findings.

Another limitation was the study's sample size of 300 articles, representing only a fraction of the total AI coverage produced by these outlets during 2024. The randomized selection process mitigated potential sampling bias but cannot guarantee comprehensive representation of all frames present in the broader coverage landscape. Visual elements, headline framing, article placement, and other rhetorical features that contribute to framing were not thoroughly analyzed in this study. Finally, this research examined media frames but did not assess their impacts. The real-world influence of these identified frames remains speculative without further research.

5.3 Future Research

Based on the findings and limitations of this study, several directions for future research are clear. Future researchers could expand this study's scope to include digital-native publications, specialized technology outlets, and alternative media formats (particularly podcasts) to provide a more comprehensive understanding of the modern media ecosystem and how it frames AI. Comparative analyses between legacy and digital-native media could reveal interesting similarities and differences. Longitudinal studies tracking AI framing over extended periods could help identify patterns and determine whether the eight frames identified in this study are long-term features of AI coverage or transient responses to recent developments. Researchers could also conduct experimental studies manipulating exposure to different AI frames in order to establish causal relationships between framing and attitudinal or behavioral outcomes. This research could more confidently answer whether modern media framing is

encouraging or discouraging AI's adoption, and contribute to a more comprehensive, nuanced understanding of how media framing shapes the public's perception and response to AI.

Chapter 6: Conclusion

This study has examined how three influential U.S. news publications—*The New York Times*, *The Wall Street Journal*, and *The Washington Post*—framed artificial intelligence in their 2024 coverage. Through systematic thematic analysis, eight dominant frames were identified: AI Boom vs. Bubble; Misuse and Misinformation; Ethical and Moral Challenges; Policy and Governance; Societal and Cultural Impact; Work and Automation; Environmental Impact; and Technological Advancements and Future Risks. These frames represent distinctive interpretive packages through which journalists and audiences made sense of AI's complex implications across economic, political, ethical, and cultural domains.

The findings reveal a media landscape grappling with the multifaceted nature of artificial intelligence—simultaneously a promising economic frontier, a potential threat to information ecosystems, a challenge to existing governance frameworks, and a transformative force across society. The variations in framing between publications reflect distinctive editorial orientations: the WSJ's business-focused perspective emphasizing strategic adoption, the NYT's critical stance questioning corporate claims and highlighting risks, and the WP's attention to societal implications and human impacts.

This research makes several contributions to the literature on media framing of emerging technologies. It provides a detailed analysis of AI framing at a critical moment in the technology's development and public consciousness. It identifies frames that increasingly dominate coverage, such as the "Environmental Impact" frame, that provide evidence for growing recognition of the technology's material resource demands and other practical considerations as adoption increases. It also demonstrates how media outlets construct different

narratives around AI based on their editorial priorities and audience orientations, confirming the active role of journalism in shaping public discourse about emerging technologies.

The patterns identified in this study suggest a maturing media discourse that has evolved beyond simplistic techno-optimism or fearmongering toward a more nuanced understanding of AI's multifaceted implications. While earlier research found predominantly positive framings of AI with an emphasis on economic benefits, the 2024 coverage reveals increased attention to ethical challenges, governance questions, and potential downsides across multiple domains. This evolution likely reflects both the growing sophistication of AI technologies and increased public awareness of their limitations and unintended consequences.

As artificial intelligence continues to advance and integrate into various domains of society, the frames through which media present these developments will significantly influence public perception, policy responses, and adoption trajectories. By identifying and analyzing these frames, this research contributes to a more reflexive understanding of how public discourse about AI is constructed and how it might evolve to better serve democratic deliberation about this transformative technology.

The news media do not merely report on AI developments but actively participate in constructing their meaning through selection, emphasis, and contextual interpretation. As Entman (1993) noted, to frame is "to select some aspects of a perceived reality and make them more salient in a communicating text, in such a way as to promote a particular problem definition, causal interpretation, moral evaluation, and/or treatment recommendation" (p. 52). The eight frames identified in this study represent distinctive ways of defining AI's problems, diagnosing their causes, making moral judgments about their implications, and suggesting appropriate responses.

By illuminating these framing patterns, this research aims to foster more critical media literacy among audiences and more reflective practices among journalists covering AI. Understanding the frames through which we encounter artificial intelligence is an essential step toward ensuring that public discourse about this technology serves democratic values and the common good rather than merely reflecting industry narratives or fueling unproductive anxieties.

As AI continues to transform our information environments, economic structures, and social interactions, the quality of public deliberation about these changes will depend significantly on the frames available for making sense of them. This study contributes to that deliberation by making visible the interpretive structures through which one of history's most consequential technologies is being presented to the public in this pivotal moment in its development.

References

- Abernathy, P. M. (2020). Will local news survive?. News deserts and ghost newspapers.
- Allcott, H., & Gentzkow, M. (2017). Social media and fake news in the 2016 election. In *Journal of Economic Perspectives* (Vol. 31, pp. 211–236). American Economic Association.
<https://doi.org/10.1257/jep.31.2.211>
- Ananny, M. (2018). The partnership press: Lessons for platform-publisher collaborations as Facebook and news outlets team to fight misinformation.
- Ardèvol-Abreu, A. (2015). *Framing theory in communication research. Origins, development and current situation in Spain*. <https://doi.org/10.4185/RLCS-2015-1053en>
- Asemi, A., Ko, A., & Nowkarizi, M. (2020). Intelligent libraries: a review on expert systems, artificial intelligence, and robot. In *Library Hi Tech* (Vol. 39, Issue 2, pp. 412–434). Emerald Group Holdings Ltd. <https://doi.org/10.1108/LHT-02-2020-0038>
- Bateson, G. (1972). A theory of play and fantasy. *Chandler*.
- Bell, E. J., Owen, T., Brown, P. D., Hauka, C., & Rashidian, N. (2017). The platform press: How Silicon Valley reengineered journalism.
- Benkler, Y., & Nissenbaum, H. (2006). Commons-based peer production and virtue. *Journal of political philosophy*, 14(4).
- Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the dangers of stochastic parrots: Can language models be too big? *FAccT 2021 - Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, 610–623.
<https://doi.org/10.1145/3442188.3445922>
- Benoit, W. L. (1995). Sears' repair of its auto service image: Image restoration discourse in the corporate sector. *Communication Studies*, 46(1–2), 89–105.

- Berry, R. (2016). Podcasting: Considering the evolution of the medium and its association with the word 'radio'. *The Radio Journal—International Studies in Broadcast & Audio Media*, 14(1), 7-22.
- Bindley, K. (2024, May 26). Tech Workers Retool for Artificial-Intelligence Boom. *The Wall Street Journal*. <https://www.wsj.com/tech/ai/ai-skills-tech-workers-job-market-1d58b2dd>
- Bolukbasi, T., Chang, K.-W., Zou, J. Y., Saligrama, V., & Kalai, A. T. (2016). Man is to Computer Programmer as Woman is to Homemaker? Debiasing Word Embeddings. In https://proceedings.neurips.cc/paper_files/paper/2016/hash/a486cd07e4ac3d270571622f4f316ec5-Abstract.html.
- Bousquette, I. (2024, March 11). AI Will Transform One of Corporate Tech's Biggest Cost Areas—Actual Savings TBD. *The Wall Street Journal*. <https://www.wsj.com/articles/ai-will-transform-one-of-corporate-techs-biggest-cost-areasactual-savings-tbd-a6c7306d>
- boyd, d. (2010). Social network sites as networked publics: Affordances, dynamics, and implications. In *A networked self* (pp. 47-66). Routledge.
- Boydston, A. (2013). *Making the News: Politics, the Media & Agenda Setting*.
- Braun, V., & Clarke, V. (2006). Using thematic analysis in psychology. *Qualitative research in psychology*, 3(2), 77-101.
- Brause, S. R., Zeng, J., Schäfer, M. S., & Katzenbach, C. (2023). Media representations of artificial intelligence: surveying the field. *Handbook of Critical Studies of Artificial Intelligence*, 277-288.
- Brennen, J. (2018). An industry-led debate: How UK media cover artificial intelligence. Reuters Institute for the Study of Journalism.

- Brennen, J. S., Howard, P. N., & Nielsen, R. K. (2022). What to expect when you're expecting robots: Futures, expectations, and pseudo-artificial general intelligence in UK news. *Journalism*, 23(1), 22-38.
- Brown, L. (1971). *Television: The Business Behind the Box*.
- Brown, T., Mann, B., Ryder, N., Subbiah, M., Kaplan, J.D., Dhariwal, P., Neelakantan, A., Shyam, P., Sastry, G., Askell, A. and Agarwal, S. (2020). Language models are few-shot learners. *Advances in Neural Information Processing Systems, 2020-December*.
- Bruns, A. (2007, June). Produsage. In Proceedings of the 6th ACM SIGCHI Conference on Creativity & Cognition (pp. 99-106).
- Buzzsprout. (2024, September 17). Podcast statistics and data [September 2024].
<https://www.buzzsprout.com/blog/podcast-statistics>
- Campbell, J. (2001). *Yellow journalism: Puncturing the myths, defining the legacies*.
- Campbell, M., Hoane, A. J., & Hsu, F. H. (2002). Deep Blue. *Artificial Intelligence*, 134, 57–83.
[https://doi.org/10.1016/S0004-3702\(01\)00129-1](https://doi.org/10.1016/S0004-3702(01)00129-1)
- Carey, J. W. (1983). Technology and ideology: The case of the telegraph. *Prospects*, 8, 303-325.
- Castells, M. (2010). Globalisation, networking, urbanisation: Reflections on the spatial dynamics of the information age. *Urban studies*, 47(13), 2737-2745.
- Chesney, B., Citron, D., Wittes, B., Jurecic, Q., Blitz, M., Finney Boylan, J., Bregler, C., Crotoft, R., Fenrich, J., Anne Franks, M., Gleicher, N., Gray, P., Green, Y., Kitchen, K., Hartzog, W., Lin, H., Norton, H., Nossel, S., Schou, A., ... Williams, B. (2019). *Deep Fakes: A Looming Challenge for Privacy, Democracy, and National Security* Silbey for helpful suggestions. *We are grateful to*. <https://doi.org/10.15779/Z38RVOD15J>

- Chuan, C. H., Tsai, W. H. S., & Cho, S. Y. (2019). Framing artificial intelligence in American newspapers. *AIES 2019 - Proceedings of the 2019 AAAI/ACM Conference on AI, Ethics, and Society*, 339–344. <https://doi.org/10.1145/3306618.3314285>
- Closson, T. (2024, October 11). Nevada asked A.I. which students need help. the answer caused an outcry. *The New York Times*. <https://www.nytimes.com/2024/10/11/us/nevada-ai-at-risk-students.html>
- Cook, T. E. (2006). The news media as a political institution: Looking backward and looking forward. In *Political Communication* (Vol. 23, pp. 159–171). <https://doi.org/10.1080/10584600600629711>
- Cope, D. G. (2013). Methods and meanings: Credibility and trustworthiness of qualitative research. Number 1/January 2014, 41(1), 89-91.
- Craig, D. B. (2000). *Fireside politics: Radio and political culture in the United States, 1920-1940*. JHU Press.
- Crawford, K. (2021). *The atlas of AI: Power, politics, and the planetary costs of artificial intelligence*. Yale University Press.
- Cutter, C. (2024, March 10). AI Is Taking On New Work. But Change Will Be Hard—and Expensive. *The Wall Street Journal*. <https://www.wsj.com/tech/ai/office-workers-artificial-intelligence-changes-86d8d4ab>
- Dandurand, G., McKelvey, F., & Roberge, J. (2023). Freezing out: Legacy media's shaping of AI as a cold controversy. *Big Data & Society*, 10(2), 20539517231219242.
- De Vynck, G. (2024a, July 24). Big Tech says AI is booming. Wall Street is starting to see a bubble. *The Washington Post*. <https://www.washingtonpost.com/technology/2024/07/24/ai-bubble-big-tech-stocks-goldman-sachs/>

- De Vynck, G. (2024b, April 25). The AI Hype Bubble is deflating. now comes the hard part. *The Washington Post*. <https://www.washingtonpost.com/technology/2024/04/18/ai-bubble-hype-dying-money/>
- De Vynck, G. (2024c, August 28). Nvidia earnings show Ai Boom is still on, though cracks have formed. *The Washington Post*.
<https://www.washingtonpost.com/technology/2024/08/28/nvidia-earnings-stock-ai-second-quarter/>
- De Vynck, G. (2024d, October 1). OpenAI wants all your apps to talk in its expressive AI voices. *The Washington Post*. <https://www.washingtonpost.com/technology/2024/10/01/openai-realtime-voice-mode-dev-day/>
- De Vynck, G. (2024e, December 4). OpenAI partners with weapons start-up Anduril on military AI. *The Washington Post*. <https://www.washingtonpost.com/technology/2024/12/04/openai-anduril-military-ai/>
- De Vynck, G. (2024f, April 4). Big Tech usually dismisses fears that AI kills jobs. Now it's studying them. *The Washington Post*.
<https://www.washingtonpost.com/technology/2024/04/04/jobs-ai-replace-study-microsoft-google-cisco/>
- De Vynck, G., & Oremus, W. (2024, February 15). OpenAI shows off lifelike videos generated by Sora, its new AI tool. *The Washington Post*.
<https://www.washingtonpost.com/technology/2024/02/15/openai-sora-artificial-intelligence-videos/>
- De Vynck, Garrit, Zakrzewski, C., & Tiku, N. (2024, July 19). California is a battleground for AI bills, as Trump plans to curb regulation. *The Washington Post*.

- <https://www.washingtonpost.com/technology/2024/07/19/biden-trump-ai-regulations-tech-industry/>
- Diakopoulos, N. (2019). *Diakopoulos - Automating the News - chapter 1*.
- Doshi-Velez, F., & Kim, B. (2017). *Towards A Rigorous Science of Interpretable Machine Learning*. <http://arxiv.org/abs/1702.08608>
- Dou, E. (2024, May 13). U.S.-China talks on AI risks set to begin in Geneva. *The Washington Post*. <https://www.washingtonpost.com/technology/2024/05/13/us-china-ai-talks/>
- Douglas, S. J. (2013). *Listening in: Radio and the American imagination*. U of Minnesota Press.
- Ellis, L. (2024, April 3). Business Schools Are Going All In on AI. *The Wall Street Journal*. <https://www.wsj.com/tech/ai/generative-ai-mba-business-school-13199631>
- Entman, R. M. (1993). Framing: Toward clarification of a fractured paradigm. *Journal of Communication*, 43(4), 51–58.
- Epstein, Z., Levine, S., Rand, D. G., & Rahwan, I. (2020). Who gets credit for AI-generated art?. *Iscience*, 23(9).
- Fast, E., & Horvitz, E. (2017). Long-term trends in the public perception of artificial intelligence. *Proceedings of the AAAI Conference on Artificial Intelligence*, 31(1).
- Fereday, J., & Muir-Cochrane, E. (2006). Demonstrating rigor using thematic analysis: A hybrid approach of inductive and deductive coding and theme development. *International journal of qualitative methods*, 5(1), 80-92.
- Floridi, L., & Chiriatti, M. (2020). GPT-3: Its Nature, Scope, Limits, and Consequences. In *Minds and Machines* (Vol. 30, Issue 4, pp. 681–694). Springer Science and Business Media B.V. <https://doi.org/10.1007/s11023-020-09548-1>

- Frenkel, S. (2024, August 21). The year of the A.I. Election that wasn't. *The New York Times*.
<https://www.nytimes.com/2024/08/21/technology/ai-election-campaigns.html>
- Garcia, L. (2024, July 18). Blackstone, Bullish on AI, Charges Into Data Centers. *The Wall Street Journal*. <https://www.wsj.com/articles/ai-bull-blackstone-charges-into-data-centers-5033786a>
- Gestoso, P., & Bakayeva, M. (2024, October 7). What did we learn from “framing deepfakes”? What did we learn from framing deepfakes? <https://weandai.org/what-did-we-learn-from-framing-deepfakes/>
- Gardner, H. (1999). Intelligence reframed: Multiple intelligences for the 21st century. In *Intelligence reframed: Multiple intelligences for the 21st century*. Basic Books.
- Gillespie, T. (2018). Custodians of the Internet: Platforms, content moderation, and the hidden decisions that shape social media. Yale University Press.
- Goffman, E. (1974). Frame analysis: An essay on the organization of experience. *Northeastern UP*.
- Graefe, A. (2016). *Guide to Automated Journalism*.
- Haenlein, M., & Kaplan, A. (2019). A brief history of artificial intelligence: On the past, present, and future of artificial intelligence. *California Management Review*, 61, 5–14.
<https://doi.org/10.1177/0008125619864925>
- Hagerty, J. (2024, November 20). How Students Can AI-Proof Their Careers. *The Wall Street Journal*. <https://www.wsj.com/tech/ai/students-ai-proof-career-tips-ba96fd38>
- Halavais, A. (2017). Search engine society. John Wiley & Sons.
- Halper, E. (2024, July 10). A three mile island nuclear power revival? AI demand for Electricity May Drive IT. *The Washington Post*.

<https://www.washingtonpost.com/business/2024/07/10/three-mile-island-nuclear-artificial-intelligence/>

Harder, R. A., Sevenans, J., & Van Aelst, P. (2017). Intermedia agenda setting in the social media age: How traditional players dominate the news agenda in election times. *The international journal of press/politics*, 22(3), 275-293.

Hasian, M. A. (1998). Freedom of Expression and Propaganda During World War I:

Understanding George Creel and America's Committee on Public Information. *Free Speech Yearbook*, 36(1), 48–60. <https://doi.org/10.1080/08997225.1998.10556225>

Hebb, D. O. (2005). The Organization of Behavior. In *The Organization of Behavior*. Psychology Press. <https://doi.org/10.4324/9781410612403>

Hitsuwari, J., Ueda, Y., Yun, W., & Nomura, M. (2023). Does human–AI collaboration lead to more creative art? Aesthetic evaluation of human-made and AI-generated haiku poetry. *Computers in Human Behavior*, 139, 107502.

Holcomb, S. D., Porter, W. K., Ault, S. V., Mao, G., & Wang, J. (2018). Overview on DeepMind and its AlphaGo Zero AI. *ACM International Conference Proceeding Series*, 67–71. <https://doi.org/10.1145/3206157.3206174>

Howard, P. N., Woolley, S., & Calo, R. (2018). Algorithms, bots, and political communication in the US 2016 election: The challenge of automated political communication for election law and administration. *Journal of information technology & politics*, 15(2), 81-93.

Hsu, T., & Metz, C. (2024, February 16). In Big Election Year, A.I.'s Architects Move Against Its Misuse. *The New York Times*. <https://www.nytimes.com/2024/02/16/technology/ai-elections-defense.html>

- Hu, K. (2023, February 2). ChatGPT sets record for fastest-growing user base - analyst note | reuters. Reuters. <https://www.reuters.com/technology/chatgpt-sets-record-fastest-growing-user-base-analyst-note-2023-02-01/>
- Hunter, T. (2024, March 18). Using AI to spot edible mushrooms could kill you. *The Washington Post*. <https://www.washingtonpost.com/technology/2024/03/18/ai-mushroom-id-accuracy/>
- Isaac, M. (2024, November 5). Meta permits its A.I. models to be used for U.S. military purposes. *The New York Times*. <https://www.nytimes.com/2024/11/04/technology/meta-ai-military.html>
- Iyengar, S., & Kinder, D. R. (2010). News that matters: Television & American opinion. University of Chicago Press.
- Jamieson, K. H. (1988). Eloquence in an electronic age: The transformation of political speechmaking. Oxford University Press.
- Jamieson, K. H., & Cappella, J. N. (2008). Echo chamber: Rush Limbaugh and the conservative media establishment. Oxford University Press.
- Jenkins, H., Ford, S., & Green, J. (2013). Spreadable media: Creating value and meaning in a networked culture. In *Spreadable media*. New York University Press.
- Jin, B., Dotan, T., & Kruppa, M. (2024, August 6). Struggling AI Startups Look for a Bailout From Big Tech. <https://www.wsj.com/tech/ai/struggling-ai-startups-look-for-a-bailout-from-big-tech-3e635927>
- Joyce, K., Smith-Doerr, L., Alegria, S., Bell, S., Cruz, T., Hoffman, S. G., ... & Shestakofsky, B. (2021). Toward a sociology of artificial intelligence: A call for research on inequalities and structural change. *Socius*, 7, 2378023121999581.
- Kahneman, D. (2011). Thinking, fast and slow. macmillan.

- Kang, C. (2024, September 29). California governor vetoes sweeping A.I. legislation. *The New York Times*. <https://www.nytimes.com/2024/09/29/technology/california-ai-bill.html>
- Kao, K., & Huang, R. (2024, August 23). Chips or Not, Chinese AI Pushes Ahead. *The Wall Street Journal*. <https://www.wsj.com/tech/ai/chips-or-not-chinese-ai-pushes-ahead-31034e3d>
- Kaplan, A., & Haenlein, M. (2019). Siri, Siri, in my hand: Who's the fairest in the land? On the interpretations, illustrations, and implications of artificial intelligence. In *Business Horizons* (Vol. 62, pp. 15–25). Elsevier Ltd. <https://doi.org/10.1016/j.bushor.2018.08.004>
- Katz, E., Lazarsfeld, P. F., & Roper, E. (2017). Personal influence: The part played by people in the flow of mass communications. Routledge.
- Kerbel, M. R. (2018). If it bleeds, it leads: An anatomy of television news. Routledge.
- Krystal Hu. (2023). *ChatGPT sets record for fastest-growing user base - analyst note*. <https://www.reuters.com/technology/chatgpt-sets-record-fastest-growing-user-base-analyst-note-2023-02-01/>.
- Kuypers, J. A. (2010). *Framing analysis from a rhetorical perspective*. In P. D'Angelo & J. A. Kuypers (Eds.), *Doing news framing analysis: Empirical and theoretical perspectives* (pp. 301–321). Routledge.
- Kuypers, J. A. (2014). *Partisan journalism: A history of media bias in the United States*. Rowman & Littlefield.
- Landers, P. (2024, April 9). Artificial Intelligence's 'Insatiable' Energy Needs Not Sustainable, Arm CEO Says. *The Wall Street Journal*. <https://www.wsj.com/tech/ai/artificial-intelligences-insatiable-energy-needs-not-sustainable-arm-ceo-says-a11218c9>

- Lazer, D. M. J., Baum, M. A., Benkler, Y., Berinsky, A. J., Greenhill, K. M., Menczer, F., Metzger, M. J., Nyhan, B., Pennycook, G., Rothschild, D., Schudson, M., Sloman, S. A., Sunstein, C. R., Thorson, E. A., Watts, D. J., & Zittrain, J. L. (2018). The science of fake news: Addressing fake news requires a multidisciplinary effort. *Science*, 359, 1094–1096. <https://doi.org/10.1126/science.aao2998>
- Lee, K. F. (2018). *AI superpowers: China, Silicon Valley, and the new world order*. Harper Business.
- Lelkes, Y., & Westwood, S. J. (2017). The limits of partisan prejudice. *The Journal of Politics*, 79(2), 485-501.
- Lemann, N. (2006, July 31). Amateur hour. *The New Yorker*. <https://www.newyorker.com/magazine/2006/08/07/amateur-hour-4>
- Liu, Y., Zhang, K., Li, Y., Yan, Z., Gao, C., Chen, R., Yuan, Z., Huang, Y., Sun, H., Gao, J., He, L., & Sun, L. (2024). *Sora: A Review on Background, Technology, Limitations, and Opportunities of Large Vision Models*.
- Livingstone, S. (2009). On the mediation of everything: ICA presidential address 2008. *Journal of communication*, 59(1), 1-18.
- Logan, N. (2023). Corporate personhood and corporate responsibility to race. In *Encyclopedia of Business and Professional Ethics* (pp. 419-422). Cham: Springer International Publishing.
- Lu, Y. (2019). Artificial intelligence: a survey on evolution, models, applications and future trends. In *Journal of Management Analytics* (Vol. 6, pp. 1–29). Taylor and Francis Ltd. <https://doi.org/10.1080/23270012.2019.1570365>

- Mackrael, K., & Schechner, S. (2024, March 13). European Lawmakers Pass AI Act, World's First Comprehensive AI Law. *The Wall Street Journal*. <https://www.wsj.com/tech/ai/ai-act-passes-european-union-law-regulation-e04ec251>
- Mark, G., Gonzalez, V. M., & Harris, J. (2005, April). No task left behind? Examining the nature of fragmented work. In Proceedings of the SIGCHI conference on Human factors in computing systems (pp. 321-330).
- Marwick, A., & Lewis, R. (2017). *Media Manipulation and Disinformation Online*.
- McLuhan, M. (2017). The medium is the message. In *Communication theory* (pp. 390-402). Routledge.
- McCombs, M. E., & Guo, L. (2014). Agenda-setting influence of the media in the public sphere. *The handbook of media and mass communication theory*, 249-268.
- Meaker, M. (2023). Slovakia's election deepfakes show AI is a danger to democracy. *Wired*. Retrieved March, 15, 2024.
- Metz, C. (2024, March 29). OpenAI unveils A.I. technology that recreates human voices. *The New York Times*. <https://www.nytimes.com/2024/03/29/technology/openai-voice-engine.html>
- Metz, C., Weise, K., Grant, N., & Isaac, M. (2023, December 3). Ego, fear and money: How the A.I. Fuse was lit. *The New York Times*. <https://www.nytimes.com/2023/12/03/technology/ai-openai-musk-page-altman.html>
- Mickle, T. (2024, September 23). Will A.I. be a bust? A wall street skeptic rings the alarm. *The New York Times*. <https://www.nytimes.com/2024/09/23/technology/ai-jim-covello-goldman-sachs.html>

- Mims, C. (2024, February 16). It's the End of the Web as We Know It. *The Wall Street Journal*.
<https://www.wsj.com/tech/ai/its-the-end-of-the-web-as-we-know-it-2ab686a2>
- Myers, C. (2025, Jan.). *Artificial intelligence and public relations: Legal, ethical, and communication trends for 2025*. ITL Thought-Leadership Essay #613. International Public Relations Association.
- Newman, N., Fletcher, R., Robertson, C. T., Ross Arguedas, A., & Nielsen, R. K. (2024). Reuters Institute digital news report 2024. Reuters Institute for the study of Journalism.
- Nguyen, T. T., Hui, P. M., Harper, F. M., Terveen, L., & Konstan, J. A. (2014). Exploring the filter bubble: The effect of using recommender systems on content diversity. *WWW 2014 - Proceedings of the 23rd International Conference on World Wide Web*, 677–686.
<https://doi.org/10.1145/2566486.2568012>
- Nielsen, M. A. (2015). Neural networks and deep learning (Vol. 25, pp. 15-24). San Francisco, CA, USA: Determination press.
- Nielsen, R. (2024, May 20). How news coverage, often uncritical, helps build up the Ai Hype. Reuters Institute for the Study of Journalism.
<https://reutersinstitute.politics.ox.ac.uk/news/how-news-coverage-often-uncritical-helps-build-ai-hype>
- O'Donovan, C. (2024, August 26). Amazon aims to launch delayed AI Alexa subscription in October. *The Washington Post*.
<https://www.washingtonpost.com/technology/2024/08/26/amazon-ai-alexa-launch-subscription-election/>

- O'Donovan, C., & De Vynck, G. (2024, April 11). The future of AI will run on Amazon, company CEO says. *The Washington Post*.
<https://www.washingtonpost.com/technology/2024/04/11/amazon-ai/>
- O'Reilly, T. (2005, October). Web 2.0: compact definition.
- Oreskes, B. (2024, December 20). Hochul weighs legislation limiting A.I. and more than 100 other bills. *The New York Times*. <https://www.nytimes.com/2024/12/20/nyregion/hochul-ai-bills-ny.html>
- Pariser, E. (2011). *The filter bubble: How the new personalized web is changing what we read and how we think*. Penguin.
- Pew Research Center. (2024, September 17). Social Media and News Fact sheet. Pew Research Center. <https://www.pewresearch.org/journalism/fact-sheet/social-media-and-news-fact-sheet/>
- Postman, N. (2005). *Amusing ourselves to death: Public discourse in the age of show business*. Penguin.
- Prior, M. (2013). Media and political polarization. In *Annual Review of Political Science* (Vol. 16, pp. 101–127). <https://doi.org/10.1146/annurev-polisci-100711-135242>
- Radford, A., Kim, J. W., Hallacy, C., Ramesh, A., Goh, G., Agarwal, S., Sastry, G., Askell, A., Mishkin, P., Clark, J., Krueger, G., & Sutskever, I. (2021). *Learning Transferable Visual Models From Natural Language Supervision*.
<https://proceedings.mlr.press/v139/radford21a>.
- Radford, A., Narasimhan, K., Salimans, T., & Sutskever, I. (n.d.). *Improving Language Understanding by Generative Pre-Training*. <https://gluebenchmark.com/leaderboard>

- Radford, A., Narasimhan, K., Salimans, T., & Sutskever, I. (2018). *Improving Language Understanding by Generative Pre-Training*. <https://gluebenchmark.com/leaderboard>
- Ramesh, A., Pavlov, M., Goh, G., Gray, S., Voss, C., Radford, A., Chen, M., & Sutskever, I. (2021). *Zero-Shot Text-to-Image Generation*. <https://proceedings.mlr.press/v139/Ramesh21a.html?ref=journey>.
- Rana, P. (2024, August 7). AI Companies Fight to Stop California Safety Rules. *The Wall Street Journal*. <https://www.wsj.com/tech/ai/ai-regulation-california-bill-29e40745>
- Roose, K. (2022). *The Brilliance and Weirdness of ChatGPT (Published 2022)*. <https://www.nytimes.com/2022/12/05/technology/chatgpt-ai-twitter.html>.
- Roose, K. (2024a, October 23). Can A.I. be blamed for a teen's suicide? *The New York Times*. <https://www.nytimes.com/2024/10/23/technology/characterai-lawsuit-teen-suicide.html>
- Roose, K. (2024b, May 21). A.I.'s black boxes just got a little less mysterious. *The New York Times*. <https://www.nytimes.com/2024/05/21/technology/ai-language-models-anthropic.html>
- Roose, K. (2024c, February 14). The year chatbots were tamed. *The New York Times*. <https://www.nytimes.com/2024/02/14/technology/chatbots-sydney-tamed.html>
- Rosenbluth, T. (2024, September 24). That message from your doctor? it may have been drafted by A.I. *The New York Times*. <https://www.nytimes.com/2024/09/24/health/ai-patient-messages-mychart.html>
- Rattner, N. (2024, March 5). AI Talent Is in Demand as Other Tech Job Listings Decline. *The Wall Street Journal*. <https://www.wsj.com/tech/ai/ai-jobs-demand-tech-layoffs-5b7344c0>
- Roumeliotis, K. I., & Tselikas, N. D. (2023). ChatGPT and Open-AI Models: A Preliminary Review. In *Future Internet* (Vol. 15, Issue 6). MDPI. <https://doi.org/10.3390/fi15060192>

- Schiff, K. J., Schiff, D. S., & Bueno, N. S. (2025). The Liar's Dividend: Can Politicians Claim Misinformation to Evade Accountability? *American Political Science Review*, 119(1), 71–90. doi:10.1017/S0003055423001454
- Schudson, M. (1978). The ideal of conversation in the study of mass media. *Communication research*, 5(3), 320-329.
- Schudson, M. (2001). The objectivity norm in American journalism. *Journalism*, 2, 149–170. <https://doi.org/10.1177/146488490100200201>
- Shao, C., Ciampaglia, G. L., Varol, O., Yang, K. C., Flammini, A., & Menczer, F. (2018). The spread of low-credibility content by social bots. *Nature Communications*, 9. <https://doi.org/10.1038/s41467-018-06930-7>
- Shu, K., Sliva, A., Wang, S., Tang, J., & Liu, H. (2017). Fake News Detection on Social Media. *ACM SIGKDD Explorations Newsletter*, 19, 22–36. <https://doi.org/10.1145/3137597.3137600>
- Sisson, P. (2024, February 29). A.I. Frenzy complicates efforts to keep power-hungry data sites Green. *The New York Times*. <https://www.nytimes.com/2024/02/29/business/artificial-intelligence-data-centers-green-power.html>
- Snow, J. (2024, March 14). As Generative AI Takes Off, Researchers Warn of Data Poisoning. *The Wall Street Journal*. <https://www.wsj.com/tech/ai/as-generative-ai-takes-off-researchers-warn-of-data-poisoning-d394385c>
- Soto-Vásquez, A. D., Fox, K., & Sinnreich, A. The Turn to Podcasts as a Mass Campaign Medium.
- Strauss, A. L. ., & Corbin, J. M. . (1999). *Basics of qualitative research : techniques and procedures for developing grounded theory*. SAGE.

- Stray, J. (2020). Aligning AI Optimization to Community Well-Being. *International Journal of Community Well-Being*, 3(4), 443–463. <https://doi.org/10.1007/s42413-020-00086-3>
- Strubell, E., Ganesh, A., & McCallum, A. (2020, April). Energy and policy considerations for modern deep learning research. In Proceedings of the AAAI conference on artificial intelligence (Vol. 34, No. 09, pp. 13693-13696).
- Sun, S., Zhai, Y., Shen, B., & Chen, Y. (2020). Newspaper coverage of artificial intelligence: A perspective of emerging technologies. *Telematics and Informatics*, 53, 101433.
- Sunstein, C. (2009). *Going to Extremes: How Like Minds Unite and Divide*.
- Tiku, N. (2024, December 6). AI friendships claim to cure loneliness. Some are ending in suicide. *The Washington Post*. <https://www.washingtonpost.com/technology/2024/12/06/ai-companion-chai-research-character-ai/>
- Tilley, A. (2024, June 10). Apple Introduces ‘Apple Intelligence,’ New OpenAI Partnership as AI Takes Center Stage. *The Wall Street Journal*. <https://www.wsj.com/tech/ai/apple-wwdc-2024-ai-release-356c5303>
- Tuchman, G. (1978). The News Net. In *Source: Social Research* (Vol. 45, Issue 2).
- Tufekci, Z. (2017). *ALGORITHMIC HARMS BEYOND FACEBOOK AND GOOGLE: EMERGENT CHALLENGES OF COMPUTATIONAL AGENCY*.
<https://www.facebook.com/akramer/posts/10152987150867796>.
- Turing, A. M. (1950). I.—Computing Machinery and Intelligence. *Mind*, LIX, 433–460.
<https://doi.org/10.1093/mind/LIX.236.433>
- Van Dijck, J., & Poell, T. (2013). Understanding social media logic. *Media and communication*, 1(1), 2-14.

- Van Gorp, B. (2010). Strategies to take subjectivity out of framing analysis. In *Doing news framing analysis* (pp. 100-125). Routledge.
- Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A.N., Kaiser, Ł. and Polosukhin, I. (2017). Attention is all you need. *Advances in neural information processing systems*, 30.
- Velazco, C. (2024, July 5). Creating songs with Ai is a blast, but also uncomfortable. *The Washington Post*. <https://www.washingtonpost.com/technology/2024/07/05/suno-ai-music-iphone-app/>
- Verma, P., & De Vynck, G. (2024, November 13). Trump pledged to gut Biden's AI rules, as OpenAI eyes landmark infusion. *The Washington Post*. <https://www.washingtonpost.com/technology/2024/11/13/openai-nuclear-subsidies-trump-ai-china/>
- Verma, P., & Kornfield, M. (2024, February 26). Democratic operative admits to commissioning Biden AI robocall in New Hampshire. *The Washington Post*. <https://www.washingtonpost.com/technology/2024/02/26/ai-robocall-biden-new-hampshire/>
- Verma, P., Zakrzewski, C., & Tiku, N. (2024, July 23). Senators demand OpenAI detail efforts to make its AI safe. *The Washington Post*. <https://www.washingtonpost.com/technology/2024/07/23/openai-senate-democrats-ai-safe/>
- Vosoughi, S., Roy, D., & Aral, S. (2018). *The spread of true and false news online*. *science*, 359(6380), 1146-1151.
- Wang, S., & Liang, Z. (2024). *What does the public think about artificial intelligence? An investigation of technological frames in different technological context*. *Government Information Quarterly*, 41(2), 101939.

- Weiner, D. I., & Norden, L. (2023). Regulating AI deepfakes and synthetic media in the political arena. Brennan Center for Justice, December, 5.
- Weise, K. (2024a, February 1). Amazon enters chatbot fray with shopping tool. *The New York Times*. <https://www.nytimes.com/2024/02/01/technology/amazon-earnings.html>
- Weise, K. (2024b, August 2). Tech bosses preach patience as they spend and spend on A.I. *The New York Times*. <https://www.nytimes.com/2024/08/02/technology/tech-companies-ai-spending.html>
- Weiss-Blatt, N. (2023). *The AI Doomers' Playbook*. Techdirt.
- White, J. (2024, May 19). See How Easily A.I. Chatbots Can Be Taught to Spew Disinformation. *The New York Times*. <https://www.nytimes.com/interactive/2024/05/19/technology/biased-ai-chatbots.html>
- Weizenbaum, J. (1966). ELIZA—a computer program for the study of natural language communication between man and machine. *Communications of the ACM*, 9(1), 36–45.
- Zelizer, B. (1992). CNN, the Gulf War, and journalistic practice. *Journal of Communication*, 42(1), 66-81.
- Zhang, B., & Dafoe, A. (2019). Artificial intelligence: American attitudes and trends. *Available at SSRN 332874*.
- Zhao, Z., Dua, D., & Singh, S. (2018). Generating natural adversarial examples. *6th International Conference on Learning Representations, ICLR 2018 - Conference Track Proceedings*.
- Ziegler, B. (2024, September 21). It's the Year 2030. What Will Artificial Intelligence Look Like? *The Wall Street Journal*. <https://www.wsj.com/tech/ai/future-of-ai-2030-experts-654fcbfe>

Zuboff, S. (2019, January). Surveillance capitalism and the challenge of collective action. In
New labor forum (Vol. 28, No. 1, pp. 10-29). Sage CA: Los Angeles, CA: sage Publications.