

Rethink Rumor Detection in the Era of LLM: A Review

Anonymous ACL submission

Abstract

The rise of large language models (LLMs) has fundamentally reshaped the technological paradigm of rumor detection, offering transformative opportunities to construct adaptive detection systems while simultaneously ushering in new threats, such as "logically flawless" rumors. This paper focuses on modeling rumor detection in the era of LLMs by unifying existing methods in the field of rumor detection and uncovering their underlying logical mechanisms. From the perspective of complex systems, we innovatively propose a "Cognition-Interaction-Behavior" (CIB) tri-level framework for rumor detection based on collective intelligence and explore the synergistic relationship between LLMs and collective intelligence in rumor governance. Furthermore, we analyze the core challenges in the LLM era and outline future development pathways for social simulation agents. We hope this work lays a theoretical foundation for next-generation rumor detection paradigms and offers valuable insights for advancing the field.

1 Introduction

In the digital era, the widespread adoption of social media and the explosion of user-generated content has enabled rumors to threaten public safety and social trust at unprecedented speeds, scales, and levels of complexity (Shao et al., 2016; Kim and Dennis, 2019). Meanwhile, the rapid advancements in large language models (LLMs) have demonstrated remarkable performance across various fields (Tan et al., 2023; Poldrack et al., 2023), but it has also brought challenges that cannot be ignored. Models like GPT-4 (Achiam et al., 2023) and DeepSeek (Guo et al., 2025), known for their deep semantic understanding and reasoning capabilities, can generate highly credible and logically coherent professional content. However, this ability can also be used to generate "logically perfect rumors" (such as false arguments based on chain reasoning),

which are far more concealed and misleading than traditional generation methods. (Bommasani et al., 2021; Kreps et al., 2022). For example, studies have shown that ChatGPT, when provided with malicious prompts, can not only optimize deceptive text but also proactively enhance their disguise by incorporating additional misleading details (Augenstein et al., 2024). Thus, leveraging the powerful capabilities of LLMs while addressing their inherent limitations has emerged as an urgent challenge in the field of rumor detection.

Existing rumor detection surveys primarily focus on the dissemination mechanisms of rumors on social media (Shu et al., 2017; Del Vicario et al., 2016; Johnson et al., 2020), the psychological mechanisms underlying belief in rumors (Roozenbeek et al., 2020), and effective intervention strategies (Zubiaga et al., 2015; Guess et al., 2020). However, most existing frameworks primarily rely on feature-based or technical classifications, which result in two primary issues: (1) the failure to thoroughly explore the theoretical and logical connections between detection methods and rumor propagation mechanisms, and (2) the inability to effectively reveal the intrinsic relationships between features, especially in the context of research on LLMs in this field (Chen and Shu, 2024).

To bridge this gap, we introduce a three-tiered "Cognition-Interaction-Behavior" (CIB) framework to systematically elucidate the underlying logic of rumor propagation and detection on social networks. The specific contributions of this work include: (1) A new theoretical paradigm for rumor detection. The construction of the CIB framework unifies existing rumor detection methods (as shown in Figure 4) and uncovers the multi-scale coupling mechanisms underlying rumor propagation, including collective knowledge emergence, interactive network evolution, and iterative behavioral patterns. (2) A systematic exploration of LLMs' multifaceted roles in rumor detection and their synergies

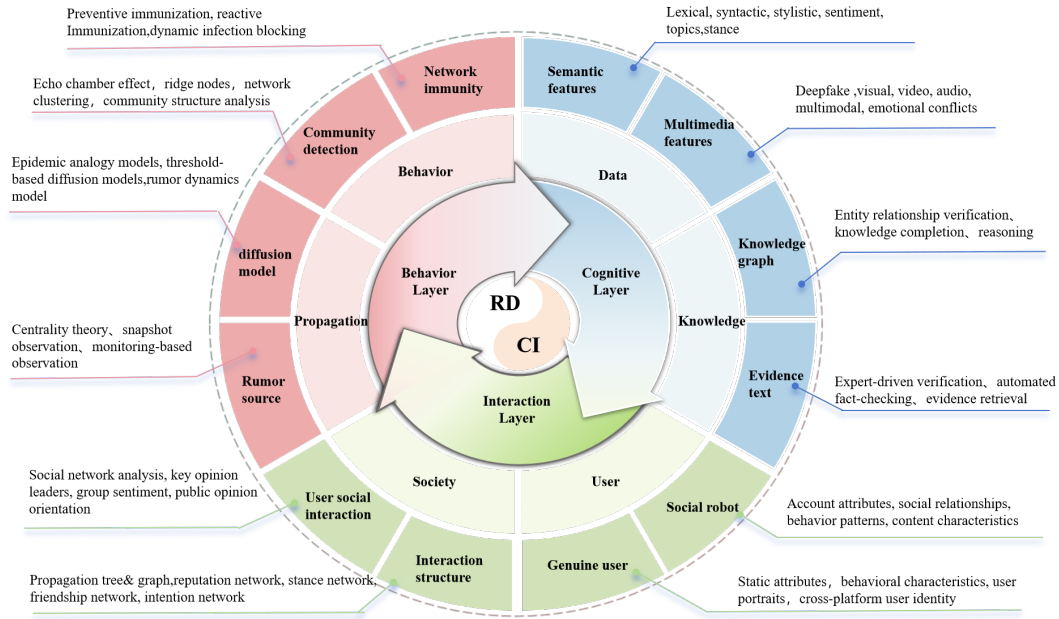


Figure 1: The three-layer architecture operates collaboratively. The cognition layer integrates multi-source evidence to provide informational support for the interaction layer. Through user interactions, the interaction layer facilitates the formation of the behavior layer. The behavior layer, in turn, continuously refines the cognition layer through accumulated experiences and collective cognitive feedback. (RD is rumor detection; CI is collective intelligence, driving information’s dynamic reconstruction and optimization).

with collective intelligence, forming a more comprehensive adaptive governance system. (3) A summary of the core challenges of rumor detection in the LLM era and an outline of future development pathways for socially simulated agents.

2 Collective Intelligence-Based Rumor Detection Framework

In the social media ecosystem, user communities serve dual roles as both disseminators and evaluators, forming the self-organizing foundation of the networked information ecology. Studies have shown that through cross-validation among users and the interplay of opinions, social networks can facilitate collective cognitive correction (Ma et al., 2018). Compared to individual cognition, collective intelligence leverages the integration of diverse knowledge and dynamic interactions, demonstrating superior cognitive capabilities in addressing complex information (Castillo et al., 2011), thereby offering a novel approach to advancing rumor detection (Phan et al., 2023).

From the perspective of complex systems, the emergence of collective intelligence is essentially a self-organizing process driven by the reduction of information entropy. During this process, social media users’ diverse cognition, social connec-

tions, and dynamic behaviors interact, facilitating information flow and collaborative evolution. Rumor diffusion, as a specific form of information dissemination, is often constrained by individuals’ cognitive thresholds (e.g., cognitive abilities, emotional biases) and the topological structure of the social network. At its core, rumor diffusion can be viewed as a staged state of cognitive imbalance: it arises when users, driven by information uncertainty and emotional impetus, engage in social interactions to reduce uncertainty, which in turn drives the continuous evolution of network structures (Allcott and Gentzkow, 2017). This process generates macro-level dissemination behaviors (potentially unintentionally promoting rumor propagation). However, collective intelligence can dynamically correct such imbalanced states in social networks through multi-level knowledge sharing and interaction.

Based on the above theoretical construction, this study proposes a three-order model framework for rumor detection based on collective intelligence, as shown in Figure 1. The cognitive layer facilitates the construction of a collective knowledge system for rumor identification through knowledge sharing and evidence integration among users. It serves as the foundational support for detecting rumors and aggregating multidimensional information. The in-

teraction layer analyzes users' social relationships and interaction behaviors within social networks to capture rumor signals. The behavior layer models the evolution of information dissemination and collective behavior. Finally, the feedback mechanism based on collective intelligence optimizes the dissemination path and reduces the spread of rumors.

2.1 Cognitive Layer

The Cognitive Layer leverages **data-driven analysis** (§2.1.1) to explore the features of multi-source information on social networks deeply, while **knowledge-driven analysis** (§2.1.2) focuses on verifying information sources and effectively integrating multi-perspective evidence. Together, these processes construct collective knowledge as the cognitive foundation for rumor detection.

2.1.1 Data-Driven Analysis

Data-driven analysis focuses on extracting features from multi-source data within social networks, aiming to develop multidimensional approaches for assessing information veracity. Research on linguistic semantics, visual content, and multimodal characteristics provides critical support for more in-depth rumor detection efforts.

In terms of semantic analysis, rumor texts often exhibit multi-dimensionalities. Studies have revealed distinct patterns at various levels by integrating linguistic and psychological theories (Undeutsch, 1967; Zuckerman, 1981). Lexical Level: rumormessages tend to avoid expressing in-depth information (Aich et al., 2022; Horne and Adali, 2017), characterized by lower usage rates of first-person pronouns and higher proportions of adverbs and emotional words. Syntactic Level: rumormessages exhibit a trend toward simplification (Pérez-Rosas et al., 2017), such as reduced lexical diversity and shorter average sentence lengths, abnormal frequency distributions of function words and punctuation, further highlight the "readability" characteristic of rumor texts. Stylistic Level: rumormessages often feature exaggerated headlines, informal language, and frequent insertion of URLs or hashtags to enhance virality and attractiveness (Blass, 1984). These linguistic abnormalities also reveal rumors' strategic use of language to evoke negative emotions (e.g., anger, fear) in public, amplifying their dissemination, particularly in sensitive topics such as politics and health (Vosoughi et al., 2018). Studies have developed various de-

tection methods through linguistic pattern analysis to leverage these characteristics to capture deeper semantic features.

In addition, deepfake technologies have significantly increased the deceptive capabilities and dissemination risks associated with image-based rumors (Vaccari and Chadwick, 2020). Early methods relied on spatial domain (Popescu and Farid, 2004) and frequency domain (Fridrich et al., 2003) analysis to extract pixel-level features for identifying common local manipulations in forged images but suffered from insufficient sensitivity to localized manipulations. Studies (Marra et al., 2019) have shown that statistical properties and spectral responses of GAN-generated content systematically differ from authentic images, providing a technological breakthrough for detection. The introduction of deep learning has further enhanced detection performance. Techniques such as local feature extraction (Bayar and Stamm, 2016; Wang et al., 2020a), temporal modeling (Güera and Delp, 2018), and the use of pre-trained models (Hao et al., 2021; Khan et al., 2022) have been effective in capturing complex visual forgery patterns. For video forgeries, beyond traditional frame-level classification methods (Montserrat et al., 2020), advancements have been made through inter-frame consistency analysis (Amerini et al., 2019), metadata validation (Huh et al., 2018), detection of visual artifacts (Matern et al., 2019), and leveraging biometric signal characteristics (Li et al., 2018). These approaches improve forgery detection performance by revealing inherent flaws in dynamic modeling. To more comprehensively capture the multimodal characteristics of rumor information, multimodal analysis focuses on feature fusion and consistency verification (Wang et al., 2018a; Jin et al., 2017). These methods balance modality-specific features and improve cross-modal detection performance by employing different fusion strategies (Singhal et al., 2019; Qian et al., 2021). Consistency verification is utilized to identify cross-modal information conflicts, including text-image inconsistencies (e.g., emotional conflicts) and audio-visual mismatches (e.g., forged videos) (Agarwal et al., 2020; Chugh et al., 2020). These approaches effectively enhance the performance of multimodal rumor detection.

2.1.2 Knowledge-Driven Analysis

The knowledge-driven analysis leverages external information resources to enrich and validate rumor content, offering critical support for social net-

works' noisy, concise content. Verification methods based on **knowledge graphs** and **evidence texts** are the core research content.

As structured data systems, knowledge graphs provide a networked organization of large-scale entities and relationships, supporting rumor detection through contextual verification and logical reasoning. By matching entities and relationships within the text, knowledge graphs can quickly validate content accuracy and identify potential contradictions (Cui et al., 2020). Several studies (Hu et al., 2021) further integrate semantic analysis and graph reasoning techniques to uncover implicit associations or supporting evidence, thereby improving verification reliability. For rumor content characterized by semantic ambiguity or missing information, knowledge graphs leverage their capabilities in semantic completion and reasoning to explore hidden, deeper-level entity associations. By incorporating contextual information (Dun et al., 2021) and multi-modal data (Wang et al., 2020b), knowledge graphs can fill in critical gaps within implicit or ambiguous statements. Furthermore, they have demonstrated robust adaptability in cross-domain rumor detection tasks (Sun et al., 2022; Zhang et al., 2019).

On the other hand, evidence-based text verification focuses on fact-checking, aiming to validate the factuality of rumor content through authoritative information sources such as news articles, scientific literature, and fact-checking platforms. Traditionally, this task has relied on expert-driven manual verification. Internationally accredited organizations such as the International Fact-Checking Network (IFCN) (Porter and Wood, 2021), the European Fact-Checking Standards Network (EFCN) (Wouters and Opgenhaffen, 2024), and government platforms (e.g., China's Internet Joint Rumor Debunking Platform, Tencent Fact-Check Platform) provide standardized evaluation procedures to assess the authenticity and timeliness of evidence, delivering high-quality validation services to the public (Vlachos and Riedel, 2014).

However, the expert-driven model faces efficiency bottlenecks and struggles to respond rapidly to large-scale, real-time information dissemination demands (Das et al., 2023; Guo et al., 2022). Automated fact-checking has increasingly become a focus of research to address these limitations. Key processes in these systems include multi-source data retrieval, semantic alignment, and logical reasoning. By leveraging deep learning models combined with information retrieval methods

(Hanselowski et al., 2019), semantic relevance between rumors and evidence is extracted from authoritative data sources such as Wikipedia and scientific literature (Schuster et al., 2021; Wadden et al., 2021). Additionally, semantic alignment techniques are employed to assess the extent to which the retrieved evidence supports or refutes the claims embedded in the rumor. For complex, multi-layered claims, deep learning methodologies (Zhong et al., 2019) can generate reliable verification results. Explainability-enhancing techniques are also applied to extract key logic and evidence chains, improving users' understanding of and trust in the verification outcomes (Lu and Li, 2020).

2.2 Interaction layer

The Interaction Layer emphasizes the analysis of **user feature (§2.2.1)** and **social context (§2.2.2)** to provide comprehensive contextual information, including individual behavior patterns, interactions between groups, and the dynamic formation of group consensus. This not only reveals the processes through which information spreads among users but also captures the collaboration and conflicts involved in users' efforts to discern and verify rumors.

2.2.1 User Feature Analysis

User groups within social networks serve as the core driving force behind rumor propagation. They are comprised of genuine users and social bots tasked with content generation and dissemination. Analyzing user characteristics and behavioral patterns can effectively uncover rumor dissemination's underlying mechanisms and risks.

As specialized accounts, social bots often manipulate public opinion through high-frequency content posting and synchronized interactions, significantly accelerating rumor dissemination. Bot detection methods can be broadly categorized into feature-based and network-based approaches. The feature-based analysis identifies non-human attributes, such as anomalous metadata (e.g., default profile pictures, short-lived accounts), high-frequency posting behaviors, and polarized content (e.g., repetitive sentence structures, simple semantics). Network-based detection detects bots by identifying abnormalities in dissemination structures. Social bots often amplify their influence by forming densely interconnected groups or fabricating community structures. In recent years, deep learning techniques, such as Graph Neural

Networks (GNNs), have been employed to model the dependencies within network topologies (Guo et al., 2021). When combined with content features, pre-trained language models have further enhanced the generalization performance of social bot detection in heterogeneous environments (Haider et al., 2023).

Genuine users are the central driving force behind information dissemination (Aral and Walker, 2012). Among them, malicious users deliberately spread rumors for personal or organizational gain (Kahneman, 1979), while ordinary users may inadvertently propagate rumors due to cognitive limitations or emotional triggers. Research has delved into the dissemination patterns of these users by examining their static attributes (e.g., registration time, geographical location, number of friends) and dynamic behavioral characteristics (e.g., posting frequency, emotional traits such as anger and fear, and account credibility scores) (Chu et al., 2012). For example, anomalies such as short-term, high-frequency interactions and highly repeated content releases are often considered potential signals of rumor spreading (Zhao et al., 2014). Additionally, as rumor dissemination increasingly transcends platform boundaries, cross-platform user identity association analysis has emerged as a research focus. For example, by matching user profiles (Iofciu et al., 2011), content geolocation (Riederer et al., 2016), and personalized traits such as posting style (Goga et al., 2013) and interest preferences (Nie et al., 2016), researchers can identify attempts by malicious users to disguise their identities across platforms. Studies have increasingly applied unified analyses of static and dynamic behaviors across broader network ecosystems by integrating deep learning and network analysis techniques (Hamdi et al., 2020; Zhang et al., 2015; Zhou et al., 2015).

2.2.2 Social Context Analysis

Rumor propagation is influenced by individual behaviors and the deeper constraints imposed by social network structures and user interaction patterns. By uncovering the synergies between user interactions and network structures, social context analysis provides critical support for understanding the mechanisms underlying the diffusion of rumors.

The interaction characteristics within social network structures exhibit distinct patterns in rumor propagation. Due to a lack of supervision and information verification mechanisms, Sparse networks

tend to form low-cohesion, flat diffusion structures, which accelerate the spread of false information (Vosoughi et al., 2018). In contrast, dense networks, with their strong connectivity, can partially filter or curb the propagation of rumors. Moreover, the heterogeneity of user roles within a network (e.g., "messenger" bridging multiple communities or "skeptics" constructing local verification networks) dynamically impacts interaction structures (Raponi et al., 2022). For example, user comment chains and resharing behaviors can be modeled as propagation tree structures (Kwon et al., 2013), which are utilized to capture both top-down and bottom-up propagation patterns (Alrubaian et al., 2016; Ma et al., 2015). To more comprehensively represent such complex social contexts, multiple entities such as users, posts, and hashtags can be modeled as propagation graphs (Nguyen et al., 2020; Shu et al., 2020). Studies also introduce graph neural networks (Bian et al., 2020; Min et al., 2022), which aggregate node features to capture complex interaction relationships and diffusion dynamics among users. These significantly enhance the efficacy of modeling rumor propagation patterns and their underlying mechanisms.

User interaction patterns serve as critical clues for uncovering rumor social context. Studies have shown that genuine users, fake news producers, and hybrid users tend to form homogeneous clusters within collaborative networks. The polarization between different clusters reflects regional differences and political and ideological divisions. This phenomenon is particularly pronounced in rumor propagation, where emotion acts as a catalyst. Emotionally charged content, by evoking negative emotions such as anger and fear, triggers group polarization and amplifies the speed and scope of rumor dissemination (Zeng and Zhu, 2019; Pröllochs et al., 2021). According to the "two-step flow" theory in communication studies (Katz, 1957), information flows first from mass media to key opinion leaders (KOLs), who then influence broader audiences. KOLs often act as community bridges (Yang et al., 2018), playing a significant amplifying role in rumor dissemination while also having the potential to contribute effectively to rumor debunking. Modeling KOLs is crucial for understanding their influence mechanisms in rumor propagation (Wei and Meng, 2021). The most common approaches for KOL modeling involve social network analysis techniques, utilizing centrality metrics (Opsahl et al., 2010), statistical models (Amor et al., 2016),

network topology analysis (Zhao et al., 2016), and deep learning methods (Shafiq et al., 2013). These methods aid in designing effective intervention strategies.

2.3 Behavior Layer

Rampant deepfakes and false rumors are often blamed as key culprits in influencing voter behavior. Studies suggest that while changing people’s political opinions is challenging, influencing their actions is comparatively easier (Adam). The spread of rumors on social networks has been extensively studied from both macro and micro perspectives (Xuan et al., 2019). At the macro level, reselevelleverages **propagation pattern modeling (§2.3.1)** to analyze the dynamic processes of rumor propagation in complex networks (Zhu and Huang, 2019). At the micro level, studlevelocus on **behavioral pattern analysis (§2.3.2)** to analyze community characteristics and predict rumor dissemination (Alkhodair et al., 2020).

2.3.1 Propagation Pattern Modeling

Propagation modeling aims to uncover the diffusion patterns of rumors within the framework of a complex sociodynamic system, focusing on the coupled process of information dissemination and collective behaviors. Epidemic analogy models serve as the foundational approach in early studies, where state transition mechanisms (e.g., Susceptible-Infectious) are used to describe the spread of rumors across network topologies (e.g., SI (Kermack and McKendrick, 1927), SIS (Dong and Huang, 2018), SIR (Zhao et al., 2013), and their variants (Wan et al., 2017). Threshold-based diffusion models, such as the Linear Threshold Model (LT) (Chen et al., 2012) and the independent cascade model (IC), shift toward modeling the propagation process from the perspective of audience decision-making. These models are also used to design intervention strategies, such as node blocking or link disruption, to suppress diffusion (Yan et al., 2019). With the development of complex network theory, more studies focus on exploring multidimensional factors that influence rumor propagation in online social networks: temporal dimension (Tripathy et al., 2010), user dimension (Hosni et al., 2018, 2020), network dimension (Wang et al., 2018b), and information dimension (Xiao et al., 2019), providing a more systematic theoretical foundation for studying the diffusion of rumors in complex networks.

Rumor source detection based on propagation models plays a pivotal role in tracing the origins of information dissemination, providing crucial insights for intervention strategies. Early methods, grounded in centrality theory, estimate the importance of the nodes by traversing the global topology of social networks to identify the source nodes (Ali et al., 2020). However, these approaches often suffer from high computational complexity. Snapshot observation methods (Louni and Subbalakshmi, 2018) improve detection efficiency by extracting limited node infection states or propagation paths. These methods effectively perform on homogeneous and heterogeneous propagation models (Cai et al., 2018). Additionally, monitoring-based observation techniques (Qiu et al., 2022) leverage sensor nodes to capture real-time dissemination data, enhancing adaptability to dynamic propagation environments. In scenarios involving multi-source concurrent propagation (Zhu et al., 2022a), research focuses on decoupling infection networks into several independent regions. A divide-and-conquer strategy is then employed to locate the sources iteratively, reducing the computational complexity of detection. With the maturation of deep learning techniques (Wang et al., 2022; Ling et al., 2022) and Graph Neural Networks (GNNs) (Cheng et al., 2024), end-to-end frameworks have emerged as highly effective tools. By integrating propagation paths, temporal dynamics, and node characteristics, these approaches significantly enhance the robustness and accuracy of rumor source detection.

2.3.2 Behavioral Pattern Analysis

In social networks, nodes often cluster into tightly-knit communities, where rumors spread efficiently within communities but rely on bridge nodes or weak ties between communities for cross-community dissemination. Research shows that when bridge nodes are scarce, rumor dissemination remains localized, but when sufficiently abundant, it penetrates communities and reaches a broader audience (Zanette, 2001). To identify these critical nodes, community detection methods (Newman, 2004; Blondel et al., 2008) usually adopt heuristics to find closely related subgroups in the network (Zhang et al., 2018; Yang et al., 2016). Targeted interventions and immunization strategies can then be applied to these key nodes to minimize the spread of rumors.

Network immunization strategies are typically categorized into preventive and counteractive im-

munization: Preventive Immunization focuses on proactively optimizing network structures and key node distributions before harmful information emerges. By analyzing topological properties (e.g., node centrality, spectral attributes) and community structure characteristics, high-risk nodes, and cross-community bridge points can be identified to disseminate accurate information (Petrescu et al., 2021). Counteractive immunization emphasizes real-time intervention during the propagation process. When initial sources or infected nodes are known, dynamic detection and analysis of infected nodes and their neighboring communities are performed. Key high-dissemination nodes are then selected for removal or blocking to efficiently disrupt the propagation chain with minimal cost (Fan et al., 2013). It prioritizes local optimization under resource constraints, such as minimizing the dissemination of malicious information within the network (Tariq et al., 2017), for instance, reducing the probability of some users sharing false content with their network connections.

3 Collective Intelligence-based Rumor Detection in the LLM Era

With the advancement of Natural Language Processing (NLP), rumor detection techniques have evolved from traditional statistical methods (Lim et al., 2017; Rayana and Akoglu, 2015) to deep learning models (Ma et al., 2016; Chen et al., 2019), and more recently to pre-trained language models (PLMs) (Kaliyar et al., 2021; Pelrine et al., 2021), progressively transitioning from static feature extraction to automated dynamic semantic analysis. However, traditional methods often suffer from limited adaptability in scenarios such as early-stage rumor detection and complex environments due to the scarcity of annotated data (He et al., 2021). LLMs further transform the rumor detection landscape with their extensive pre-training and contextual reasoning capabilities, enabling efficient reasoning in resource-poor environments and significantly improving the transparency and interpretability of detection results.

3.1 LLM-enhanced CIB Framework

LLMs play a multifaceted role in existing rumor detection approaches, deeply integrating into various roles ranging from knowledge analysis to adversarial defense, as shown in Figure 2. LLMs have brought revolutionary advancements to traditional

methods, demonstrating significant technical potential and promising application prospects.

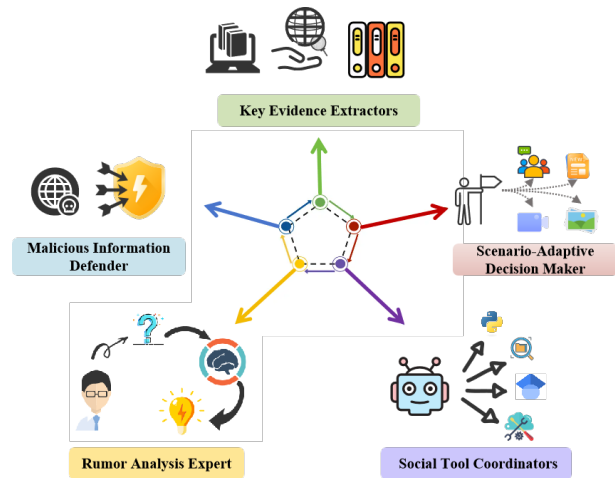


Figure 2: Multiple roles of LLM in rumor detection

At the cognitive layer, LLMs serve as highly efficient **Key Evidence Extractors** through large-scale pretraining. Unlike traditional rumor detection systems that rely on static, structured approaches (such as knowledge graphs) to expand knowledge, LLMs leverage their implicit encoding capabilities to capture deep semantic associations in unstructured information, which is further used as evidence to enhance the generalization ability of traditional language models (Nan et al., 2024; Yang et al., 2023a). Additionally, LLMs operate as **Scenario-Adaptive Decision Makers**, leveraging their zero-shot reasoning abilities to address diversified rumor scenarios efficiently without requiring fine-tuning (Li et al., 2023c; Wu et al., 2023a). Combining Retrieval-Augmented Generation (RAG) with external knowledge bases (Peng et al., 2023; Niu et al., 2024), LLMs can dynamically integrate up-to-date knowledge, resolving limitations in knowledge coverage and timeliness inherent to traditional methods. This integration also effectively reduces the likelihood of hallucination phenomena, thereby enhancing the reliability of detection outputs (Ji et al., 2023; Rawte et al., 2023).

At the interaction layer, LLMs function as **Social Tool Coordinators**, coordinating external tools (e.g., search engines, deepfake detectors) through agents to expand rumor detection capabilities further (Chern et al., 2023; Wan et al., 2024; Li et al., 2024a). Unlike traditional, static social network analysis and modeling methods, LLM-based agents can perceive social environments by combining

short-term (contextual learning) and long-term (external knowledge retrieval) memory. These agents are capable of planning and calling external tools dynamically, improving analytical performance. Furthermore, generative agents (Park et al., 2023) can simulate user interaction behaviors, driving a paradigm shift in rumor detection from static feature modeling to dynamic interaction simulation.

At the behavior layer, LLMs act as **Rumor Analysis Experts** with superior performance in tasks requiring advanced reasoning and cross-domain contextual knowledge. Traditional rumor detection approaches relied heavily on classification tasks, often requiring manually labeled large datasets. In contrast, the emergent abilities of LLMs, such as Chain of Thought (CoT) reasoning, can decompose complex problems into intermediate reasoning steps, thereby significantly improving logical transparency and explainability (Zhang and Gao, 2023a). LLMs also exhibit robust cross-domain transferability (Cao et al., 2023c,b), enabling unified reasoning across multimodal inputs, including text, images, and audio (Yao et al., 2023a). This addresses the limitations of traditional methods in multimodal fusion and shifts detection mechanisms from pattern classification to causal inference (Zhu et al., 2022b; Nan et al., 2021). Additionally, LLMs can act as **Malicious Information Defenders**, showcasing strong robustness in adversarial social network environments. By integrating adversarial training and red-teaming methods (Bhardwaj and Poria, 2023; OpenAI, 2023), LLMs can quickly adapt to evolving forgery techniques, overcoming the lag in model iteration and processing capabilities of traditional approaches (Wu et al., 2024b; Sun et al., 2024). For example, when tackling tasks such as detecting rumor dissemination, stylized language attacks, and deepfake content, this dynamic adaptability further enhances the robustness of rumor detection systems.

3.2 Collective Intelligence-driven LLM Detection

Network users are the primary agents of information dissemination and important participants in rumor detection. The propagation of rumors relies on user attention, trust, and further sharing behaviors. In contrast, user reports and feedback can effectively constrain this diffusion effect, which is critical in rumor detection. This user feedback-based supervisory mechanism helps address the limitations of LLMs in adapting to dynamic sce-

narios by introducing ethical constraints and social consensus at the human cognitive level, thelevelinfusing greater flexibility and human-centricity into the rumor detection process (Kou et al., 2022).

Collective intelligence injects new technical implications into the development of LLM agents. Research demonstrates (Li et al., 2023b) that agents, by simulationing group behaviors modeled in sociology and economics theories, can exhibit emergent corrective mechanisms during collaborative tasks. These emergent behaviors provide theoretical support for rumor detection agents (Zhang et al., 2024b), demonstrating their significant potential in improving the simulation of social media ecosystems and other complex societal environments. For example, agents can produce socially simulated content that is indistinguishable from real-world community behavior (Park et al., 2022), simulate trust-building interactions in social dynamics (Xie et al., 2024), and facilitate harmless discussions that bridge biases and political divides, offering valuable insights into real-world phenomena (Törnberg et al., 2023).

In the field of rumor detection, existing research (Hu et al., 2025) further uses LLM-based multi-agent simulation to explore the trend of rumor propagation and optimize intervention strategies and also uses tools to implement (Li et al., 2024a) real-time evaluation of information credibility based on shallow features such as language style and common sense rules. However, the complexity of the rumor detection task requires a more comprehensive approach that considers multi-level features. We discussed this in the future research route (Appendix B), paving the way for future exploration.

4 Conclusion

Based on the complex system characteristics of collective intelligence, we reconstructed the rumor detection paradigm adapted to the era of LLMs, the "Cognition-Interaction-Behavior" (CIB) framework. We emphasized the important role of LLM in rumor detection and its complementary relationship with collective intelligence. In addition, CIB can use cross-layer dynamic feedback to establish a "Macro-Micro Feedback Loop" to dynamically realize two-way rumor detection and intervention, providing a roadmap for applying LLM-based multi-agent social simulation in rumor detection.

5 Limitations

In the future research section of this paper, we propose a roadmap for collective intelligence-based rumor detection agents under the CIB framework, providing a comprehensive analysis of potential research challenges and corresponding directions. However, further exploration is needed to evaluate the practical application of this framework in large-scale social media environments. Additionally, polarized contexts or anomalous interactions may introduce more significant complexities. To refine and optimize the framework, we will consider enhancing robustness and dynamic adaptability in complex scenarios.

References

- Alberto Acerbi and Joseph M Stubbersfield. 2023. Large language models show human-like content biases in transmission chain experiments. *Proceedings of the National Academy of Sciences*, 120(44):e2313790120.
- Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altenschmidt, Sam Altman, Shyamal Anadkat, et al. 2023. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*.
- David Adam. Misinformation might sway elections-but not in the way that you think. *Nature*.
- Komi Afassinou. 2014. Analysis of the impact of education rate on the rumor spreading mechanism. *Physica A: Statistical Mechanics and Its Applications*, 414:43–52.
- Shruti Agarwal, Hany Farid, Ohad Fried, and Maneesh Agrawala. 2020. Detecting deep-fake videos from phoneme-viseme mismatches. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition workshops*, pages 660–661.
- Muhammad Ahmad, Sarwan Ali, Juvaria Tariq, Imdadullah Khan, Mudassir Shabbir, and Arif Zaman. 2020. Combinatorial trace method for network immunization. *Information Sciences*, 519:215–228.
- Ankit Aich, Souvik Bhattacharya, and Natalie Parde. 2022. Demystifying neural fake news via linguistic feature-based interpretation. In *Proceedings of the 29th International Conference on Computational Linguistics*, pages 6586–6599.
- Abdulrahman I Al-Ghadir, Aqil M Azmi, and Amir Hussain. 2021. A novel approach to stance detection in social media tweets by fusing ranked lists and sentiments. *Information Fusion*, 67:29–40.
- Syed Shafat Ali, Tarique Anwar, and Syed Afzal Murtaza Rizvi. 2020. A revisit to the infection source identification problem under classical graph centrality measures. *Online Social Networks and Media*, 17:100061.
- Sarah A Alkhodair, Steven HH Ding, Benjamin CM Fung, and Junqiang Liu. 2020. Detecting breaking news rumors of emerging topics in social media. *Information Processing & Management*, 57(2):102018.
- Hunt Allcott and Matthew Gentzkow. 2017. Social media and fake news in the 2016 election. *Journal of economic perspectives*, 31(2):211–236.
- Gordon W Allport. 1947. The psychology of rumor. *Henry Holt*.
- Majed Alrubaian, Muhammad Al-Qurishi, Mohammad Mehedi Hassan, and Atif Alamri. 2016. A credibility analysis system for assessing information on twitter. *IEEE Transactions on Dependable and Secure Computing*, 15(4):661–674.
- Irene Amerini, Leonardo Galteri, Roberto Caldelli, and Alberto Del Bimbo. 2019. Deepfake video detection through optical flow based cnn. In *Proceedings of the IEEE/CVF international conference on computer vision workshops*, pages 0–0.
- Benjamin RC Amor, Sabine I Vuik, Ryan Callahan, Ara Darzi, Sophia N Yaliraki, and Mauricio Barahona. 2016. Community detection and role identification in directed networks: understanding the twitter network of the care. data debate. In *Dynamic networks and cyber-security*, pages 111–136. World Scientific.
- Markus Anderljung, Joslyn Barnhart, Anton Korinek, Jade Leung, Cullen O’Keefe, Jess Whittlestone, Shahar Avin, Miles Brundage, Justin Bullock, Duncan Cass-Beggs, et al. 2023. Frontier ai regulation: Managing emerging risks to public safety. *arXiv preprint arXiv:2307.03718*.
- Christopher MJ André, Helene FL Eriksen, Emil J Jakobsen, Luca CB Mingolla, and Nicolai B Thomsen. 2023. Detecting ai authorship: Analyzing descriptive features for ai detection.
- Fares Antaki, Samir Touma, Daniel Milad, Jonathan El-Khoury, and Renaud Duval. 2023. Evaluating the performance of chatgpt in ophthalmology: an analysis of its successes and shortcomings. *Ophthalmology science*, 3(4):100324.
- Dimosthenis Antypas, Jose Camacho-Collados, Alun Preece, and David Rogers. 2021. Covid-19 and misinformation: A large-scale lexical analysis on twitter. In *Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing: Student Research Workshop*, pages 119–126.
- Elena-Simona Apostol, Özgür Coban, and Ciprian-Octavian Truică. 2023. Contain: A community-based algorithm for network immunization. *arXiv preprint arXiv:2303.01934*.

833	Sinan Aral and Dylan Walker. 2012. Identifying influential and susceptible members of social networks. <i>Science</i> , 337(6092):337–341.	889
834		890
835		891
836	Isabelle Augenstein, Timothy Baldwin, Meeyoung Cha, Tanmoy Chakraborty, Giovanni Luca Ciampaglia, David Corney, Renee DiResta, Emilio Ferrara, Scott Hale, Alon Halevy, et al. 2024. Factuality challenges in the era of large language models and opportunities for fact-checking. <i>Nature Machine Intelligence</i> , 6(8):852–863.	892
837		893
838		894
839		895
840		896
841		897
842		898
843	Anton Bakhtin, Sam Gross, Myle Ott, Yuntian Deng, Marc’Aurelio Ranzato, and Arthur Szlam. 2019. Real or fake? learning to discriminate machine from human generated text. <i>arXiv preprint arXiv:1906.03351</i> .	899
844		900
845		901
846		902
847		
848	Eytan Bakshy, Itamar Rosenn, Cameron Marlow, and Lada Adamic. 2012. The role of social networks in information diffusion. In <i>Proceedings of the 21st international conference on World Wide Web</i> , pages 519–528.	903
849		904
850		905
851		906
852		907
853	Ritwik Banerjee, Song Feng, Jun Seok Kang, and Yejin Choi. 2014. Keystroke patterns as prosody in digital writings: A case study with deceptive reviews and essays. In <i>Proceedings of the 2014 conference on empirical methods in natural language processing (EMNLP)</i> , pages 1469–1473.	908
854		909
855		910
856		911
857		912
858		913
859	Yejin Bang, Samuel Cahyawijaya, Nayeon Lee, Wenliang Dai, Dan Su, Bryan Wilie, Holy Lovenia, Ziwei Ji, Tiezheng Yu, Willy Chung, Quyet V. Do, Yan Xu, and Pascale Fung. 2023. A multitask, multilingual, multimodal evaluation of ChatGPT on reasoning, hallucination, and interactivity . In <i>Proceedings of the 13th International Joint Conference on Natural Language Processing and the 3rd Conference of the Asia-Pacific Chapter of the Association for Computational Linguistics (Volume 1: Long Papers)</i> , pages 675–718, Nusa Dua, Bali. Association for Computational Linguistics.	914
860		915
861		916
862		
863		917
864		918
865		919
866		
867		920
868		921
869		922
870		923
871	Jawadul H Bappy, Amit K Roy-Chowdhury, Jason Bunk, Lakshmanan Nataraj, and BS Manjunath. 2017. Exploiting spatial structure for localizing manipulated image regions. In <i>Proceedings of the IEEE international conference on computer vision</i> , pages 4970–4979.	924
872		925
873		926
874		927
875		928
876		929
877	Anthony M Barrett, Dan Hendrycks, Jessica Newman, and Brandie Nonnecke. 2022. Actionable guidance for high-consequence ai risk management: Towards standards addressing ai catastrophic risks. <i>arXiv preprint arXiv:2206.08966</i> .	930
878		
879		931
880		932
881		933
882	Belhassen Bayar and Matthew C Stamm. 2016. A deep learning approach to universal image manipulation detection using a new convolutional layer. In <i>Proceedings of the 4th ACM workshop on information hiding and multimedia security</i> , pages 5–10.	934
883		935
884		936
885		
886		937
887	Yoshua Bengio, Geoffrey Hinton, Andrew Yao, Dawn Song, Pieter Abbeel, Yuval Noah Harari, Ya-Qin	938
888		939
		940
		941
	Zhang, Lan Xue, Shai Shalev-Shwartz, Gillian Hadfield, et al. 2023. Managing ai risks in an era of rapid progress. <i>arXiv preprint arXiv:2310.17688</i> , page 18.	
	Maciej Besta, Nils Blach, Ales Kubicek, Robert Gerstenberger, Michal Podstawski, Lukas Gianinazzi, Joanna Gajda, Tomasz Lehmann, Hubert Niewiadomski, Piotr Nyczyk, et al. 2024. Graph of thoughts: Solving elaborate problems with large language models. In <i>Proceedings of the AAAI Conference on Artificial Intelligence</i> , volume 38, pages 17682–17690.	
	Rishabh Bhardwaj and Soujanya Poria. 2023. Red-teaming large language models using chain of utterances for safety-alignment. <i>arXiv preprint arXiv:2308.09662</i> .	
	Xiuli Bi, Yang Wei, Bin Xiao, and Weisheng Li. 2019. Rru-net: The ringed residual u-net for image splicing forgery detection. In <i>Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition Workshops</i> , pages 0–0.	
	Tian Bian, Xi Xiao, Tingyang Xu, Peilin Zhao, Wenbing Huang, Yu Rong, and Junzhou Huang. 2020. Rumor detection on social media with bi-directional graph convolutional networks. In <i>Proceedings of the AAAI conference on artificial intelligence</i> , volume 34, pages 549–556.	
	Marcel Binz and Eric Schulz. 2023. Using cognitive psychology to understand gpt-3. <i>Proceedings of the National Academy of Sciences</i> , 120(6):e2218523120.	
	Thomas Blass. 1984. Social psychology and personality: Toward a convergence. <i>Journal of Personality and Social Psychology</i> , 47(5):1013.	
	Vincent D Blondel, Jean-Loup Guillaume, Renaud Lambiotte, and Etienne Lefebvre. 2008. Fast unfolding of communities in large networks. <i>Journal of statistical mechanics: theory and experiment</i> , 2008(10):P10008.	
	Rishi Bommasani, Drew A Hudson, Ehsan Adeli, Russ Altman, Simran Arora, Sydney von Arx, Michael S Bernstein, Jeannette Bohg, Antoine Bosselut, Emma Brunskill, et al. 2021. On the opportunities and risks of foundation models. <i>arXiv preprint arXiv:2108.07258</i> .	
	Asgarali Bouyer, Hamid Ahmadi Beni, Bahman Arasteh, Zahra Aghaee, and Reza Ghanbarzadeh. 2023. Fip: A fast overlapping community-based influence maximization algorithm using probability coefficient of global diffusion in social networks. <i>Expert systems with applications</i> , 213:118869.	
	Samuel R Bowman, Jeeyoon Hyun, Ethan Perez, Edwin Chen, Craig Pettit, Scott Heiner, Kamilè Lukošiušė, Amanda Askell, Andy Jones, Anna Chen, et al. 2022. Measuring progress on scalable oversight for large language models. <i>arXiv preprint arXiv:2211.03540</i> .	

942	Arastoo Bozorgi, Saeed Samet, Johan Kwisthout, and	Canyu Chen and Kai Shu. 2023. Can llm-generated	998
943	Todd Wareham. 2017. Community-based influence	misinformation be detected? <i>arXiv preprint</i>	999
944	maximization in social networks under a competitive	<i>arXiv:2309.13788</i> .	1000
945	linear threshold model. <i>Knowledge-Based Systems</i> ,		
946	134:149–158.		
947	Jean-Flavien Bussotti, Luca Ragazzi, Giacomo Frisoni,	Canyu Chen and Kai Shu. 2024. Combating misinfor-	1001
948	Gianluca Moro, and Paolo Papotti. 2024. Unknown	mation in the age of llms: Opportunities and chal-	1002
949	claims: Generation of fact-checking training exam-	enges. <i>AI Magazine</i> , 45(3):354–368.	1003
950	ples from unstructured and structured data. In <i>Pro-</i>		
951	<i>ceedings of the 2024 Conference on Empirical Meth-</i>	Hung-Hsuan Chen, Yan-Bin Ciou, and Shou-De Lin.	1004
952	<i>ods in Natural Language Processing</i> , pages 12105–	2012. Information propagation game: A tool to ac-	1005
953	12122.	quire humanplaying data for multiplayer influence	1006
		maximization on social networks. In <i>Proceedings of</i>	1007
954	Kechao Cai, Hong Xie, and John CS Lui. 2018. Infor-	<i>the 18th ACM SIGKDD international conference on</i>	1008
955	mation spreading forensics via sequential dependent	<i>Knowledge discovery and data mining</i> , pages 1524–	1009
956	snapshots. <i>IEEE/ACM Transactions on Networking</i> ,	1527.	1010
957	26(1):478–491.		
958	Han Cao, Lingwei Wei, Mengyang Chen, Wei Zhou, and	Junyi Chen, Leyuan Liu, and Fan Zhou. 2025. Do not	1011
959	Songlin Hu. 2023a. Are large language models good	wait: Preemptive rumor detection with cooperative	1012
960	fact checkers: A preliminary study. <i>arXiv preprint</i>	llms and accessible social context. <i>Information Pro-</i>	1013
961	<i>arXiv:2311.17355</i> .	<i>cessing & Management</i> , 62(3):103995.	1014
962	Rui Cao, Ming Shan Hee, Adriel Kuek, Wen-Haw		
963	Chong, Roy Ka-Wei Lee, and Jing Jiang. 2023b. Pro-	Lei Chen, Guanying Li, Zhongyu Wei, Yang Yang, Bao-	1015
964	cap: Leveraging a frozen vision-language model for	hua Zhou, Qi Zhang, and Xuanjing Huang. 2022. A	1016
965	hateful meme detection. In <i>Proceedings of the 31st</i>	progressive framework for role-aware rumor resolu-	1017
966	<i>ACM International Conference on Multimedia</i> , pages	<i>tion</i> . In <i>Proceedings of the 29th International Con-</i>	1018
967	5244–5252.	<i>ference on Computational Linguistics</i> , pages 2748–	1019
		2758, Gyeongju, Republic of Korea. International	1020
968	Rui Cao, Roy Ka-Wei Lee, Wen-Haw Chong, and	Committee on Computational Linguistics.	1021
969	Jing Jiang. 2023c. Prompting for multimodal		
970	hateful meme classification. <i>arXiv preprint</i>	Sanxing Chen, Yukun Huang, and Bhuwan Dhingra.	1022
971	<i>arXiv:2302.04156</i> .	2024. Real-time fake news from adversarial feed-	1023
		back. <i>arXiv preprint arXiv:2410.14651</i> .	1024
972	Nicholas Carlini, Milad Nasr, Christopher A Choquette-		
973	Choo, Matthew Jagielski, Irena Gao, Pang Wei W	Weixuan Chen and Daniel McDuff. 2018. Deepphys:	1025
974	Koh, Daphne Ippolito, Florian Tramer, and Ludwig	Video-based physiological measurement using con-	1026
975	Schmidt. 2023. Are aligned neural networks adver-	volutional attention networks. In <i>Proceedings of the</i>	1027
976	sarially aligned? <i>Advances in Neural Information</i>	<i>european conference on computer vision (ECCV)</i> ,	1028
977	<i>Processing Systems</i> , 36:61478–61500.	pages 349–365.	1029
978	Carlos Castillo, Marcelo Mendoza, and Barbara Poblete.		
979	2011. Information credibility on twitter. In <i>Proceed-</i>	Yimin Chen, Niall J Conroy, and Victoria L Rubin. 2015.	1030
980	<i>ings of the 20th international conference on World</i>	Misleading online content: recognizing clickbait as"	1031
981	<i>wide web</i> , pages 675–684.	false news". In <i>Proceedings of the 2015 ACM on</i>	1032
		<i>workshop on multimodal deception detection</i> , pages	1033
982	Chi-Min Chan, Weize Chen, Yusheng Su, Jianxuan Yu,	15–19.	1034
983	Wei Xue, Shanghang Zhang, Jie Fu, and Zhiyuan		
984	Liu. 2023. Chateval: Towards better llm-based eval-	Yixuan Chen, Jie Sui, Liang Hu, and Wei Gong. 2019.	1035
985	uators through multi-agent debate. <i>arXiv preprint</i>	Attention-residual network with cnn for rumor detec-	1036
986	<i>arXiv:2308.07201</i> .	<i>tion</i> . In <i>Proceedings of the 28th ACM international</i>	1037
		<i>conference on information and knowledge manage-</i>	1038
987	Michael Chan, Francis LF Lee, and Hsuan-Ting Chen.	<i>ment</i> , pages 1121–1130.	1039
988	2021. Examining the roles of multi-platform social		
989	media news use, engagement, and connections with	Le Cheng, Peican Zhu, Keke Tang, Chao Gao, and Zhen	1040
990	news organizations and journalists on news literacy:	Wang. 2024. Gin-sd: source detection in graphs with	1041
991	A comparison of seven democracies. <i>Digital Jour-</i>	incomplete nodes via positional encoding and atten-	1042
992	<i>nalism</i> , 9(5):571–588.	tive fusion. In <i>Proceedings of the AAAI Conference</i>	1043
993	Samantha Chan, Pat Pataranutaporn, Aditya Suri,	<i>on Artificial Intelligence</i> , volume 38, pages 55–63.	1044
994	Wazeer Zulfikar, Pattie Maes, and Elizabeth F Loftus.		
995	2024. Conversational ai powered by large language	I Chern, Steffi Chern, Shiqi Chen, Weizhe Yuan, Kehua	1045
996	models amplifies false memories in witness inter-	Feng, Chunting Zhou, Junxian He, Graham Neubig,	1046
997	views. <i>arXiv preprint arXiv:2408.04681</i> .	Pengfei Liu, et al. 2023. Factool: Factuality detec-	1047
		tion in generative ai—a tool augmented framework	1048
		for multi-task and multi-domain scenarios. <i>arXiv</i>	1049
		<i>preprint arXiv:2307.13528</i> .	1050

1051	Zi Chu, Steven Gianvecchio, Haining Wang, and Sushil	Natalia Díaz-Rodríguez, Javier Del Ser, Mark Coeck-	1108
1052	Jajodia. 2012. Detecting automation of twitter ac-	elbergh, Marcos López de Prado, Enrique Herrera-	1109
1053	counts: Are you a human, bot, or cyborg? <i>IEEE</i>	Viedma, and Francisco Herrera. 2023. Connecting	1110
1054	<i>Transactions on dependable and secure computing</i> ,	the dots in trustworthy artificial intelligence: From	1111
1055	9(6):811–824.	ai principles, ethics, and key requirements to respon-	1112
		sible ai systems and regulation. <i>Information Fusion</i> ,	1113
1056	Lynn Chua, Badih Ghazi, Yangsibo Huang, Pritish Ka-	99:101896.	1114
1057	math, Ravi Kumar, Daogao Liu, Pasin Manurangsi,		
1058	Amer Sinha, and Chiyuan Zhang. 2024. Mind the pri-	Peter Sheridan Dodds, Eric M Clark, Suma Desu,	1115
1059	privacy unit! user-level differential privacy for language	Morgan R Frank, Andrew J Reagan, Jake Ryland	1116
1060	model fine-tuning. <i>arXiv preprint arXiv:2406.14322</i> .	Williams, Lewis Mitchell, Kameron Decker Harris,	1117
		Isabel M Kloumann, James P Bagrow, et al. 2015.	1118
1061	Komal Chugh, Parul Gupta, Abhinav Dhall, and Ra-	Human language reveals a universal positivity bias.	1119
1062	manathan Subramanian. 2020. Not made for each	<i>Proceedings of the national academy of sciences</i> ,	1120
1063	other-audio-visual dissonance-based deepfake detec-	112(8):2389–2394.	1121
1064	tion and localization. In <i>Proceedings of the 28th</i>		
1065	<i>ACM international conference on multimedia</i> , pages	Ming Dong, Bolong Zheng, Nguyen Quoc Viet Hung,	1122
1066	439–447.	Han Su, and Guohui Li. 2019. Multiple rumor source	1123
		detection with graph convolutional networks. In <i>Pro-</i>	1124
1067	Thomas H Costello, Gordon Pennycook, and David G	<i>ceedings of the 28th ACM international conference</i>	1125
1068	Rand. 2024. Durably reducing conspiracy	<i>on information and knowledge management</i> , pages	1126
1069	beliefs through dialogues with ai. <i>Science</i> ,	569–578.	1127
1070	385(6714):eadq1814.		
		Suyalatu Dong and Yong-Chang Huang. 2018. Sis ru-	1128
1071	Chaoqun Cui and Caiyan Jia. 2024. Propagation tree	mor spreading model with population dynamics in	1129
1072	is not deep: Adaptive graph contrastive learning	online social networks. In <i>2018 International Confer-</i>	1130
1073	approach for rumor detection. In <i>Proceedings of</i>	<i>ence on Wireless Communications, Signal Processing</i>	1131
1074	<i>the AAAI Conference on Artificial Intelligence</i> , vol-	<i>and Networking (WiSPNET)</i> , pages 1–5. IEEE.	1132
1075	ume 38, pages 73–81.		
		John Dougrez-Lewis, Elena Kochkina, Maria Liakata,	1133
1076	Jian Cui, Kwanwoo Kim, Seung Ho Na, and Seung-	and Yulan He. 2024. Knowledge graphs for real-	1134
1077	won Shin. 2022. Meta-path-based fake news detec-	world rumour verification. In <i>Proceedings of the</i>	1135
1078	tion leveraging multi-level social context information.	<i>2024 Joint International Conference on Computa-</i>	1136
1079	In <i>Proceedings of the 31st ACM international con-</i>	<i>tational Linguistics, Language Resources and Evalua-</i>	1137
1080	<i>ference on information & knowledge management</i> ,	<i>tion (LREC-COLING 2024)</i> , pages 9843–9853.	1138
1081	pages 325–334.		
		Yaqian Dun, Kefei Tu, Chen Chen, Chunyan Hou, and	1139
1082	Limeng Cui, Haeseung Seo, Maryam Tabar, Fenglong	Xiaojie Yuan. 2021. Kan: Knowledge-aware atten-	1140
1083	Ma, Suhang Wang, and Dongwon Lee. 2020. Deter-	tion network for fake news detection. In <i>Proceedings</i>	1141
1084	rent: Knowledge guided graph attention network for	<i>of the AAAI conference on artificial intelligence</i> , vol-	1142
1085	detecting healthcare misinformation. In <i>Proceedings</i>	ume 35, pages 81–89.	1143
1086	<i>of the 26th ACM SIGKDD international conference</i>		
1087	<i>on knowledge discovery & data mining</i> , pages 492–	Yogesh K Dwivedi, Nir Kshetri, Laurie Hughes,	1144
1088	502.	Emma Louise Slade, Anand Jeyaraj, Arpan Kumar	1145
		Kar, Abdullah M Baabdullah, Alex Koohang, Vish-	1146
1089	Anubrata Das, Houjiang Liu, Venelin Kovatchev, and	nupriya Raghavan, Manju Ahuja, et al. 2023. Opin-	1147
1090	Matthew Lease. 2023. The state of human-centered	ion paper: “so what if chatgpt wrote it?” multidisci-	1148
1091	nlp technology for fact-checking. <i>Information pro-</i>	plinary perspectives on opportunities, challenges and	1149
1092	<i>cessing & management</i> , 60(2):103219.	implications of generative conversational ai for re-	1150
		search, practice and policy. <i>International Journal of</i>	1151
1093	Soumita Das, Ravi Kishore Devarapalli, and Anupam	<i>Information Management</i> , 71:102642.	1152
1094	Biswas. 2024. Leveraging cascading information for		
1095	community detection in social networks. <i>Information</i>	Owain Evans, Owen Cotton-Barratt, Lukas Finnve-	1153
1096	<i>Sciences</i> , 674:120696.	den, Adam Bales, Avital Balwit, Peter Wills, Luca	1154
		Righetti, and William Saunders. 2021. Truthful ai:	1155
1097	Luigi De Angelis, Francesco Baglivo, Guglielmo Arzilli,	Developing and governing ai that does not lie. <i>arXiv</i>	1156
1098	Gaetano Pierpaolo Privitera, Paolo Ferragina, Al-	<i>preprint arXiv:2110.06674</i> .	1157
1099	berto Eugenio Tozzi, and Caterina Rizzo. 2023. Chat-		
1100	gpt and the rise of large language models: the new	Kevin Eykholt, Ivan Evtimov, Earlence Fernandes,	1158
1101	ai-driven infodemic threat in public health. <i>Frontiers</i>	Bo Li, Amir Rahmati, Chaowei Xiao, Atul Prakash,	1159
1102	<i>in public health</i> , 11:1166120.	Tadayoshi Kohno, and Dawn Song. 2018. Robust	1160
		physical-world attacks on deep learning visual clas-	1161
1103	Michela Del Vicario, Alessandro Bessi, Fabiana Zollo,	sification. In <i>Proceedings of the IEEE conference</i>	1162
1104	Fabio Petroni, Antonio Scala, Guido Caldarelli, H Eu-	<i>on computer vision and pattern recognition</i> , pages	1163
1105	gene Stanley, and Walter Quattrociocchi. 2016. The	1625–1634.	1164
1106	spreading of misinformation online. <i>Proceedings of</i>		
1107	<i>the national academy of Sciences</i> , 113(3):554–559.		

1165	Yahya H Ezzeldin, Shen Yan, Chaoyang He, Emilio Ferrara, and A Salman Avestimehr. 2023. Fairfed: Enabling group fairness in federated learning. In <i>Proceedings of the AAAI conference on artificial intelligence</i> , volume 37, pages 7494–7502.	1220
1166		1221
1167		1222
1168		1223
1169		1224
1170	Lidan Fan, Zaixin Lu, Weili Wu, Bhavani Thuraisingham, Huan Ma, and Yuanjun Bi. 2013. Least cost rumor blocking in social networks. In <i>2013 IEEE 33rd International Conference on Distributed Computing Systems</i> , pages 540–549. IEEE.	1225
1171		
1172		1226
1173		1227
1174		1228
1175	Mahmoud Fawzi and Walid Magdy. 2024. "pinocchio had a nose, you have a network!": On characterizing fake news spreaders on arabic social media. <i>Proceedings of the ACM on Human-Computer Interaction</i> , 8(CSCW1):1–20.	1229
1176		
1177		1230
1178		1231
1179		1232
1180	Shangbin Feng, Zhaoxuan Tan, Herun Wan, Ningnan Wang, Zilong Chen, Binchi Zhang, Qinghua Zheng, Wenqian Zhang, Zhenyu Lei, Shujie Yang, et al. 2022. Twibot-22: Towards graph-based twitter bot detection. <i>Advances in Neural Information Processing Systems</i> , 35:35254–35269.	1233
1181		1234
1182		1235
1183		1236
1184		1237
1185		
1186	Zhangyin Feng, Xiaocheng Feng, Dezhi Zhao, Maojin Yang, and Bing Qin. 2024. Retrieval-generation synergy augmented large language models. In <i>ICASSP 2024-2024 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)</i> , pages 11661–11665. IEEE.	1238
1187		1239
1188		1240
1189		1241
1190		1242
1191		1243
1192	Santo Fortunato. 2010. Community detection in graphs. <i>Physics reports</i> , 486(3-5):75–174.	1244
1193		1245
1194	Santo Fortunato and Marc Barthelemy. 2007. Resolution limit in community detection. <i>Proceedings of the national academy of sciences</i> , 104(1):36–41.	1246
1195		1247
1196		1248
1197	Santo Fortunato and Darko Hric. 2016. Community detection in networks: A user guide. <i>Physics reports</i> , 659:1–44.	1249
1198		1250
1199		1251
1200	Gabriel Freedman, Adam Dejl, Deniz Gorur, Xiang Yin, Antonio Rago, and Francesca Toni. 2024. Argumentative large language models for explainable and contestable decision-making. <i>arXiv preprint arXiv:2405.02079</i> .	1252
1201		1253
1202		1254
1203		1255
1204		
1205	Jessica Fridrich, David Soukal, Jan Lukas, et al. 2003. Detection of copy-move forgery in digital images. In <i>Proceedings of digital forensic research workshop</i> , volume 3, pages 652–63. Cleveland, OH.	1256
1206		1257
1207		1258
1208		1259
1209	Rinaldo Gagiano, Maria Myung-Hee Kim, Xizhen Jenny Zhang, and Jennifer Biggs. 2021. Robustness analysis of grover for machine-generated news detection. In <i>Proceedings of the 19th Annual Workshop of the Australasian Language Technology Association</i> , pages 119–127.	1260
1210		1261
1211		1262
1212		1263
1213		1264
1214		1265
1215	ŁG Gajewski, Krzysztof Suchecki, and JA Hołyst. 2019. Multiple propagation paths enhance locating the source of diffusion in complex networks. <i>Physica A: Statistical Mechanics and its Applications</i> , 519:34–41.	1266
1216		1267
1217		1268
1218		1269
1219		1270
		1271
	Deep Ganguli, Liane Lovitt, Jackson Kernion, Amanda Askell, Yuntao Bai, Saurav Kadavath, Ben Mann, Ethan Perez, Nicholas Schiefer, Kamal Ndousse, et al. 2022. Red teaming language models to reduce harms: Methods, scaling behaviors, and lessons learned. <i>arXiv preprint arXiv:2209.07858</i> .	1272
		1273
	Maryanne Garry, Way Ming Chan, Jeffrey Foster, and Linda A Henkel. 2024. Large language models (llms) and the institutionalization of misinformation. <i>Trends in cognitive sciences</i> .	1274
		1275
	Jiahui Geng, Fengyu Cai, Yuxia Wang, Heinz Koepl, Preslav Nakov, and Iryna Gurevych. 2024. A survey of confidence estimation and calibration in large language models. In <i>Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers)</i> , pages 6577–6595.	1276
	Moumita Ghosh, Samhita Das, and Pritha Das. 2022. Dynamics and control of delayed rumor propagation through social networks. <i>Journal of Applied Mathematics and Computing</i> , pages 1–30.	
	David Glukhov, Ilia Shumailov, Yarin Gal, Nicolas Papernot, and Vardan Papayan. 2023. Llm censorship: A machine learning challenge or a computer security problem? <i>arXiv preprint arXiv:2307.10719</i> .	
	Oana Goga, Howard Lei, Sree Hari Krishnan Parthasarathi, Gerald Friedland, Robin Sommer, and Renata Teixeira. 2013. Exploiting innocuous activity for correlating users across sites. In <i>Proceedings of the 22nd international conference on World Wide Web</i> , pages 447–458.	
	Jian Guan, Jesse Dodge, David Wadden, Minlie Huang, and Hao Peng. 2023. Language models hallucinate, but may excel at fact verification. <i>arXiv preprint arXiv:2310.14564</i> .	
	David Güera and Edward J Delp. 2018. Deepfake video detection using recurrent neural networks. In <i>2018 15th IEEE international conference on advanced video and signal based surveillance (AVSS)</i> , pages 1–6. IEEE.	
	Andrew M Guess, Michael Lerner, Benjamin Lyons, Jacob M Montgomery, Brendan Nyhan, Jason Reifler, and Neelanjan Sircar. 2020. A digital media literacy intervention increases discernment between mainstream and false news in the united states and india. <i>Proceedings of the National Academy of Sciences</i> , 117(27):15536–15545.	
	Bin Guo, Yasan Ding, Lina Yao, Yunji Liang, and Zhiwen Yu. 2020. The future of false information detection on social media: New perspectives and trends. <i>ACM Computing Surveys (CSUR)</i> , 53(4):1–36.	
	Daya Guo, Dejian Yang, Haowei Zhang, Junxiao Song, Ruoyu Zhang, Runxin Xu, Qihao Zhu, Shirong Ma, Peiyi Wang, Xiao Bi, et al. 2025. Deepseek-r1: Incentivizing reasoning capability in llms via reinforcement learning. <i>arXiv preprint arXiv:2501.12948</i> .	

1277	Qinglang Guo, Haiyong Xie, Yangyang Li, Wen Ma, and Chao Zhang. 2021. Social bots detection via fusing bert and graph convolutional networks. <i>Symmetry</i> , 14(1):30.	1330
1278		1331
1279		1332
1280		1333
1281	Zhijiang Guo, Michael Schlichtkrull, and Andreas Vlachos. 2022. A survey on automated fact-checking. <i>Transactions of the Association for Computational Linguistics</i> , 10:178–206.	1334
1282		
1283		
1284		
1285	Philipp Hacker. 2024. Sustainable ai regulation. <i>Common Market Law Review</i> , 61(2).	
1286		
1287	Philipp Hacker, Andreas Engel, and Marco Mauer. 2023. Regulating chatgpt and other large generative ai models. In <i>Proceedings of the 2023 ACM Conference on Fairness, Accountability, and Transparency</i> , pages 1112–1123.	1335
1288		1336
1289		1337
1290		1338
1291		1339
1292	Thilo Hagendorff. 2024. Deception abilities emerged in large language models. <i>Proceedings of the National Academy of Sciences</i> , 121(24):e2317967121.	
1293		
1294		
1295	Samar Haider, Luca Luceri, Ashok Deb, Adam Badawy, Nanyun Peng, and Emilio Ferrara. 2023. Detecting social media manipulation in low-resource languages. In <i>Companion Proceedings of the ACM Web Conference 2023</i> , pages 1358–1364.	1340
1296		1341
1297		1342
1298		1343
1299		1344
1300	Patrick Haller, Ansar Aynedinov, and Alan Akbik. 2023. Opiniongpt: Modelling explicit biases in instruction-tuned llms. <i>arXiv preprint arXiv:2309.03876</i> .	1345
1301		1346
1302		1347
1303	Tarek Hamdi, Hamda Slimi, Ibrahim Bounhas, and Yahya Slimani. 2020. A hybrid approach for fake news detection in twitter based on user features and graph embedding. In <i>Distributed Computing and Internet Technology: 16th International Conference, ICDIT 2020, Bhubaneswar, India, January 9–12, 2020, Proceedings 16</i> , pages 266–280. Springer.	1348
1304		1349
1305		
1306		
1307		
1308		
1309		
1310	Ahmed Abdeen Hamed. 2023. Improving detection of chatgpt-generated fake science using real publication text: Introducing xfakebibs a supervised learning network algorithm.	1350
1311		1351
1312		1352
1313		1353
1314	Kateřina Haniková, David Chudán, Vojtěch Svátek, Peter Vajdečka, Raphaël Troncy, Filip Vencovský, and Jana Syrovátková. 2024. Towards fact-check summarization leveraging on argumentation elements tied to entity graphs. In <i>Companion Proceedings of the ACM on Web Conference 2024</i> , pages 1473–1481.	1354
1315		1355
1316		
1317		
1318		
1319		
1320	Hans WA Hanley and Zakir Durumeric. 2024. Machine-made media: Monitoring the mobilization of machine-generated articles on misinformation and mainstream news websites. In <i>Proceedings of the International AAAI Conference on Web and Social Media</i> , volume 18, pages 542–556.	1356
1321		1357
1322		1358
1323		1359
1324		1360
1325		
1326	Andreas Hanselowski, Christian Stab, Claudia Schulz, Zile Li, and Iryna Gurevych. 2019. A richly annotated corpus for different tasks in automated fact-checking. <i>arXiv preprint arXiv:1911.01214</i> .	1361
1327		1362
1328		1363
1329		1364
		1365
		1366
		1367
		1368
		1369
		1370
		1371
		1372
		1373
		1374
		1375
		1376
		1377
		1378
		1379
		1380
		1381
		1382
		1383
		1384
		1385

1386	<i>Computational Linguistics and the 11th International</i>	Kathleen H Jamieson. 2008. <i>Echo chamber: Rush Lim-</i>	1440
1387	<i>Joint Conference on Natural Language Processing</i>	<i>baugh and the conservative media establishment.</i> Ox-	1441
1388	(Volume 1: Long Papers), pages 754–763.	ford University Press.	1442
1389	Nan Hu, Zirui Wu, Yuxuan Lai, Xiao Liu, and Yansong	Danish Javed, NZ Jhanjhi, Navid Ali Khan, Sayan Ku-	1443
1390	Feng. 2022. Dual-channel evidence fusion for fact	mar Ray, Alanoud Al Mazroa, Farzeen Ashfaq, and	1444
1391	verification over texts and tables. In <i>Proceedings of</i>	Shampa Rani Das. 2024. Towards the future of	1445
1392	<i>the 2022 Conference of the North American Chap-</i>	bot detection: A comprehensive taxonomical review	1446
1393	<i>ter of the Association for Computational Linguistics:</i>	and challenges on twitter/x. <i>Computer Networks</i> ,	1447
1394	<i>Human Language Technologies</i> , pages 5232–5242.	254:110808.	1448
1395	Tianrui Hu, Dimitrios Liakopoulos, Xiwen Wei, Radu	Ziwei Ji, Nayeon Lee, Rita Frieske, Tiezheng Yu, Dan	1449
1396	Marculescu, and Neeraja J Yadwadkar. 2025. Simu-	Su, Yan Xu, Etsuko Ishii, Ye Jin Bang, Andrea	1450
1397	lating rumor spreading in social networks using llm	Madotto, and Pascale Fung. 2023. Survey of halluci-	1451
1398	agents. <i>arXiv preprint arXiv:2502.01450</i> .	nation in natural language generation. <i>ACM Comput-</i>	1452
1399	Di Huang, Jinbao Song, and Xingyu Zhang. 2025a.	<i>ing Surveys</i> , 55(12):1–38.	1453
1400	Semi-supervised social bot detection with relational	Zhengbao Jiang, Frank F Xu, Luyu Gao, Zhiqing	1454
1401	graph attention transformers and characteristics of	Sun, Qian Liu, Jane Dwivedi-Yu, Yiming Yang,	1455
1402	the social environment. <i>Information Fusion</i> , page	Jamie Callan, and Graham Neubig. 2023. Ac-	1456
1403	102956.	tive retrieval augmented generation. <i>arXiv preprint</i>	1457
1404	Linan Huang and Quanyan Zhu. 2023. An introduction	<i>arXiv:2305.06983</i> .	1458
1405	of system-scientific approaches to cognitive security.	Zhiwei Jin, Juan Cao, Han Guo, Yongdong Zhang, and	1459
1406	<i>arXiv preprint arXiv:2301.05920</i> .	Jiebo Luo. 2017. Multimodal fusion with recurrent	1460
1407	Yue Huang and Lichao Sun. 2023. Harnessing the	neural networks for rumor detection on microblogs.	1461
1408	power of chatgpt in fake news: An in-depth explo-	In <i>Proceedings of the 25th ACM international con-</i>	1462
1409	ration in generation, detection and explanation. <i>arXiv</i>	<i>ference on Multimedia</i> , pages 795–816.	1463
1410	<i>preprint arXiv:2310.05046</i> .	Zhiwei Jin, Juan Cao, Yongdong Zhang, and Jiebo Luo.	1464
1411	Yue Huang and Lichao Sun. 2024. Fakegpt: fake news	2016. News verification by exploiting conflicting	1465
1412	generation, explanation and detection of large lan-	social viewpoints in microblogs. In <i>Proceedings of</i>	1466
1413	guage models. <i>arxiv.org</i> .	<i>the AAAI conference on artificial intelligence</i> , vol-	1467
1414	Yuheng Huang, Jiayang Song, Zhijie Wang, Shengming	ume 30.	1468
1415	Zhao, Huaming Chen, Felix Juefei-Xu, and Lei Ma.	Matthew Jiwa, Patrick S Cooper, Trevor TJ Chong, and	1469
1416	2025b. Look before you leap: An exploratory study	Stefan Bode. 2023. Hedonism as a motive for infor-	1470
1417	of uncertainty analysis for large language models.	mation search: biased information-seeking leads to	1471
1418	<i>IEEE Transactions on Software Engineering</i> .	biased beliefs. <i>Scientific Reports</i> , 13(1):2086.	1472
1419	Minyoung Huh, Andrew Liu, Andrew Owens, and	Anna Jobin, Marcello Ienca, and Effy Vayena. 2019.	1473
1420	Alexei A Efros. 2018. Fighting fake news: Image	The global landscape of ai ethics guidelines. <i>Nature</i>	1474
1421	splice detection via learned self-consistency. In <i>Pro-</i>	<i>machine intelligence</i> , 1(9):389–399.	1475
1422	<i>ceedings of the European conference on computer</i>	Neil F Johnson, Nicolas Velásquez, Nicholas John-	1476
1423	<i>vision (ECCV)</i> , pages 101–117.	son Restrepo, Rhys Leahy, Nicholas Gabriel, Sara	1477
1424	Hongwen Hui, Chengcheng Zhou, Xing Lü, and Jiarong	El Oud, Minzhang Zheng, Pedro Manrique, Stefan	1478
1425	Li. 2020. Spread mechanism and control strategy of	Wuchty, and Yonatan Lupu. 2020. The online com-	1479
1426	social network rumors under the influence of covid-	petition between pro-and anti-vaccination views. <i>Nat-</i>	1480
1427	19. <i>Nonlinear Dynamics</i> , 101:1933–1949.	<i>ture</i> , 582(7811):230–233.	1481
1428	Tereza Iofciu, Peter Fankhauser, Fabian Abel, and Ker-	Daniel Kahneman. 1979. Prospect theory: An analysis	1482
1429	stin Bischoff. 2011. Identifying users across social	of decisions under risk. <i>Econometrica</i> , 47:278.	1483
1430	tagging systems. In <i>Proceedings of the International</i>	Rohit Kumar Kaliyar, Anurag Goswami, and Pratik	1484
1431	<i>AAAI Conference on Web and Social Media</i> , vol-	Narang. 2021. Fakebert: Fake news detection in so-	1485
1432	ume 5, pages 522–525.	cial media with a bert-based deep learning approach.	1486
1433	Daphne Ippolito, Daniel Duckworth, Chris Callison-	<i>Multimedia tools and applications</i> , 80(8):11765–	1487
1434	Burch, and Douglas Eck. 2020. Automatic detec-	11788.	1488
1435	tion of generated text is easiest when humans are	Mohammadamin Kanaani. 2024. Triple-r: Automatic	1489
1436	fooled . In <i>Proceedings of the 58th Annual Meeting of</i>	reasoning for fact verification using language mod-	1490
1437	<i>the Association for Computational Linguistics</i> , pages	els. In <i>Proceedings of the 2024 Joint International</i>	1491
1438	1808–1822, Online. Association for Computational	<i>Conference on Computational Linguistics, Language</i>	1492
1439	Linguistics.	<i>Resources and Evaluation (LREC-COLING 2024)</i> ,	1493
		pages 16831–16840.	1494

1602	Miaoran Li, Baolin Peng, Michel Galley, Jianfeng Gao, and Zhu Zhang. 2023c. Self-checker: Plug-and-play modules for fact-checking with large language models. <i>arXiv preprint arXiv:2305.14623</i> .	1657	<i>and Development in Information Retrieval</i> , pages 3001–3004.	1658
1603				
1604				
1605				
1606	Sha Li, Ruining Zhao, Manling Li, Heng Ji, Chris Callison-Burch, and Jiawei Han. 2023d. Open-domain hierarchical event schema induction by incremental prompting and verification. <i>arXiv preprint arXiv:2307.01972</i> .	1659	Hui Liu, Wenya Wang, Haoru Li, and Haoliang Li. 2024b. Teller: A trustworthy framework for explainable, generalizable and controllable fake news detection. <i>arXiv preprint arXiv:2402.07776</i> .	1660
1607		1661		
1608		1662		
1609				
1610		1663	Xuannan Liu, Peipei Li, Huaibo Huang, Zekun Li, Xing Cui, Jiahao Liang, Lixiong Qin, Weihong Deng, and Zhaofeng He. 2024c. Fka-owl: Advancing multimodal fake news detection through knowledge-augmented lvlms. In <i>Proceedings of the 32nd ACM International Conference on Multimedia</i> , pages 10154–10163.	1664
1611	Weimin Li, Xiaokang Zhou, Chao Yang, Yuting Fan, Zhao Wang, and Yanxia Liu. 2022. Multi-objective optimization algorithm based on characteristics fusion of dynamic social networks for community discovery. <i>Information Fusion</i> , 79:110–123.	1665		
1612		1666		
1613		1667		
1614		1668		
1615		1669		
1616	Xinyi Li, Yongfeng Zhang, and Edward C Malthouse. 2024a. Large language model agent for fake news detection. <i>arXiv preprint arXiv:2405.01593</i> .	1670	Yi Liu, Gelei Deng, Yuekang Li, Kailong Wang, Zihao Wang, Xiaofeng Wang, Tianwei Zhang, Yepang Liu, Haoyu Wang, Yan Zheng, et al. 2023. Prompt injection attack against llm-integrated applications. <i>arXiv preprint arXiv:2306.05499</i> .	1671
1617		1672		
1618		1673		
1619	Yuezun Li, Ming-Ching Chang, and Siwei Lyu. 2018. In ictu oculi: Exposing ai created fake videos by detecting eye blinking. In <i>2018 IEEE International workshop on information forensics and security (WIFS)</i> , pages 1–7. Ieee.	1674		
1620		1675	Alireza Louni and KP Subbalakshmi. 2018. Who spread that rumor: Finding the source of information in large online social networks with probabilistically varying internode relationship strengths. <i>IEEE transactions on computational social systems</i> , 5(2):335–343.	1676
1621		1677		
1622		1678		
1623		1679		
1624	Yupeng Li, Haorui He, Jin Bai, and Dacheng Wen. 2024b. Mcfend: a multi-source benchmark dataset for chinese fake news detection. In <i>Proceedings of the ACM on Web Conference 2024</i> , pages 4018–4027.	1680	Yi-Ju Lu and Cheng-Te Li. 2020. Gcan: Graph-aware co-attention networks for explainable fake news detection on social media. <i>arXiv preprint arXiv:2004.11648</i> .	1681
1625		1682		
1626		1683		
1627				
1628	Zhuohang Li, Jiaxin Zhang, Chao Yan, Kamalika Das, Sricharan Kumar, Murat Kantarcioglu, and Bradley A Malin. 2024c. Do you know what you are talking about? characterizing query-knowledge relevance for reliable retrieval augmented generation. <i>arXiv preprint arXiv:2410.08320</i> .	1684	Qing Lyu, Shreya Havaldar, Adam Stein, Li Zhang, Delip Rao, Eric Wong, Marianna Apidianaki, and Chris Callison-Burch. 2023. Faithful chain-of-thought reasoning . In <i>Proceedings of the 13th International Joint Conference on Natural Language Processing and the 3rd Conference of the Asia-Pacific Chapter of the Association for Computational Linguistics (Volume 1: Long Papers)</i> , pages 305–329, Nusa Dua, Bali. Association for Computational Linguistics.	1685
1629		1686		
1630		1687		
1631		1688		
1632		1689		
1633		1690		
1634	Gang Liang, Wenbo He, Chun Xu, Liangyin Chen, and Jinquan Zeng. 2015. Rumor identification in microblogging systems based on users’ behavior. <i>IEEE Transactions on Computational Social Systems</i> , 2(3):99–108.	1691		
1635		1692		
1636		1693		
1637				
1638				
1639	Wee Yong Lim, Mong Li Lee, and Wynne Hsu. 2017. Ifact: An interactive framework to assess claims from tweets. In <i>Proceedings of the 2017 ACM on Conference on Information and Knowledge Management</i> , pages 787–796.	1694	Jing Ma, Wei Gao, Shafiq Joty, and Kam-Fai Wong. 2019. Sentence-level evidence embedding for claim verification with hierarchical attention networks. Association for Computational Linguistics.	1695
1640		1696		
1641		1697		
1642				
1643				
1644	Stephanie Lin, Jacob Hilton, and Owain Evans. 2022. Teaching models to express their uncertainty in words. <i>arXiv preprint arXiv:2205.14334</i> .	1698	Jing Ma, Wei Gao, Prasenjit Mitra, Sejeong Kwon, Bernard J Jansen, Kam-Fai Wong, and Meeyoung Cha. 2016. Detecting rumors from microblogs with recurrent neural networks.	1699
1645		1700		
1646		1701		
1647	Chen Ling, Junji Jiang, Junxiang Wang, and Zhao Liang. 2022. Source localization of graph diffusion via variational autoencoders for graph inverse problems. In <i>Proceedings of the 28th ACM SIGKDD conference on knowledge discovery and data mining</i> , pages 1010–1020.	1702	Jing Ma, Wei Gao, Zhongyu Wei, Yueming Lu, and Kam-Fai Wong. 2015. Detect rumors using time series of social context information on microblogging websites. In <i>Proceedings of the 24th ACM international on conference on information and knowledge management</i> , pages 1751–1754.	1703
1648		1704		
1649		1705		
1650		1706		
1651		1707		
1652				
1653	Aiwei Liu, Qiang Sheng, and Xuming Hu. 2024a. Preventing and detecting misinformation generated by large language models. In <i>Proceedings of the 47th International ACM SIGIR Conference on Research</i>	1708	Jing Ma, Wei Gao, and Kam-Fai Wong. 2018. Rumor detection on twitter with tree-structured recursive neural networks. Association for Computational Linguistics.	1709
1654		1710		
1655		1711		
1656				

1712	Zihan Ma, Minnan Luo, Hao Guo, Zhi Zeng, Yiran Hao, and Xiang Zhao. 2024. Event-radar: Event-driven multi-view learning for multimodal fake news detection. In <i>Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)</i> , pages 5809–5821.	Fengqing Zhu, et al. 2020. Deepfakes detection with automatic face weighting. In <i>Proceedings of the IEEE/CVF conference on computer vision and pattern recognition workshops</i> , pages 668–669.	1768
1713			1769
1714			1770
1715			1771
1716			
1717			
1718	Peihua Mai, Ran Yan, Zhe Huang, Youjia Yang, and Yan Pang. 2023. Split-and-denoise: Protect large language model inference with local differential privacy. <i>arXiv preprint arXiv:2310.09130</i> .	Xuemei Mou, Wei Xu, Yangfu Zhu, Qian Li, and Yunpeng Xiao. 2022. A social topic diffusion model based on rumor, anti-rumor, and motivation-rumor. <i>IEEE Transactions on Computational Social Systems</i> , 10(5):2644–2659.	1772
1719			1773
1720			1774
1721			1775
1722			1776
1723	Shrikant Malviya and Stamos Katsigiannis. 2024. Evidence retrieval for fact verification using multi-stage reranking . In <i>Findings of the Association for Computational Linguistics: EMNLP 2024</i> , pages 7295–7308, Miami, Florida, USA. Association for Computational Linguistics.	Qiong Nan, Juan Cao, Yongchun Zhu, Yanyan Wang, and Jintao Li. 2021. Mdfend: Multi-domain fake news detection. In <i>Proceedings of the 30th ACM International Conference on Information & Knowledge Management</i> , pages 3343–3347.	1777
1724			1778
1725			1779
1726			1780
1727			1781
1728	Francesco Marra, Diego Gragnaniello, Luisa Verdoliva, and Giovanni Poggi. 2019. Do gans leave artificial fingerprints? In <i>2019 IEEE conference on multimedia information processing and retrieval (MIPR)</i> , pages 506–511. IEEE.	Qiong Nan, Qiang Sheng, Juan Cao, Beizhe Hu, Danding Wang, and Jintao Li. 2024. Let silence speak: Enhancing fake news detection with generated comments from large language models. In <i>Proceedings of the 33rd ACM International Conference on Information and Knowledge Management</i> , pages 1732–1742.	1782
1729			1783
1730			1784
1731			1785
1732			1786
1733			1787
1734	David Martín-Gutiérrez, Gustavo Hernández-Peñaloza, Alberto Belmonte Hernández, Alicia Lozano-Diez, and Federico Álvarez. 2021. A deep learning approach for robust detection of bots in twitter using transformers. <i>IEEE Access</i> , 9:54591–54601.	Mark EJ Newman. 2004. Fast algorithm for detecting community structure in networks. <i>Physical Review E—Statistical, Nonlinear, and Soft Matter Physics</i> , 69(6):066133.	1788
1735			1789
1736			1790
1737			1791
1738	Falko Matern, Christian Riess, and Marc Stamminger. 2019. Exploiting visual artifacts to expose deepfakes and face manipulations. In <i>2019 IEEE Winter Applications of Computer Vision Workshops (WACVW)</i> , pages 83–92. IEEE.	Lynnette Hui Xian Ng and Kathleen M. Carley. 2022. Online coordination: Methods and comparative case studies of coordinated groups across four events in the united states . In <i>Proceedings of the 14th ACM Web Science Conference 2022, WebSci '22</i> , page 12–21, New York, NY, USA. Association for Computing Machinery.	1792
1739			1793
1740			1794
1741			1795
1742			1796
1743	SC Matz, JD Teeny, Sumer S Vaid, H Peters, GM Harari, and M Cerf. 2024. The potential of generative ai for personalized persuasion at scale. <i>Scientific Reports</i> , 14(1):4692.	Van-Hoang Nguyen, Kazunari Sugiyama, Preslav Nakov, and Min-Yen Kan. 2020. Fang: Leveraging social context for fake news detection using graph representation. In <i>Proceedings of the 29th ACM international conference on information & knowledge management</i> , pages 1165–1174.	1797
1744			1798
1745			1799
1746			1800
1747	Mohit Mayank, Shakshi Sharma, and Rajesh Sharma. 2022. Deap-faked: Knowledge graph based approach for fake news detection. In <i>2022 IEEE/ACM International Conference on Advances in Social Networks Analysis and Mining (ASONAM)</i> , pages 47–51. IEEE.	Raymond S Nickerson. 1998. Confirmation bias: A ubiquitous phenomenon in many guises. <i>Review of general psychology</i> , 2(2):175–220.	1801
1748			1802
1749			1803
1750			1804
1751			1805
1752	Nikhil Mehta, María Leonor Pacheco, and Dan Goldwasser. 2022. Tackling fake news detection by continually improving social context representations using graph neural networks. In <i>Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)</i> , pages 1363–1380.	Yuanping Nie, Yan Jia, Shudong Li, Xiang Zhu, Aiping Li, and Bin Zhou. 2016. Identifying users across social networks based on dynamic core interests. <i>Neurocomputing</i> , 210:107–115.	1806
1753			1807
1754			1808
1755			1809
1756			1810
1757			1811
1758			1812
1759	Erxue Min, Yu Rong, Yatao Bian, Tingyang Xu, Peilin Zhao, Junzhou Huang, and Sophia Ananiadou. 2022. Divide-and-conquer: Post-user interaction network for fake news detection on social media. In <i>Proceedings of the ACM web conference 2022</i> , pages 1148–1158.	Cheng Niu, Yang Guan, Yuanhao Wu, Juno Zhu, Juntong Song, Randy Zhong, Kaihua Zhu, Siliang Xu, Shizhe Diao, and Tong Zhang. 2024. Veract scan: Retrieval-augmented fake news detection with justifiable reasoning. <i>arXiv preprint arXiv:2406.10289</i> .	1813
1760			1814
1761			1815
1762			1816
1763			1817
1764			1818
1765	Daniel Mas Montserrat, Hanxiang Hao, Sri K Yarlaga, Sriram Baireddy, Ruiting Shao, János Horváth, Emily Bartusiak, Justin Yang, David Guera,	Cheng Niu, Yuanhao Wu, Juno Zhu, Siliang Xu, Kashun Shum, Randy Zhong, Juntong Song, and Tong Zhang. 2023. Ragtruth: A hallucination corpus for developing trustworthy retrieval-augmented language models. <i>arXiv preprint arXiv:2401.00396</i> .	1819
1766			1820
1767			1821
		R OpenAI. 2023. Gpt-4 technical report. arxiv 2303.08774. <i>View in Article</i> , 2(5).	1822
			1823

1824	Tore Opsahl, Filip Agneessens, and John Skvoretz.	Ethan Perez, Saffron Huang, Francis Song, Trevor Cai,	1880
1825	2010. Node centrality in weighted networks: Gener-	Roman Ring, John Aslanides, Amelia Glaese, Nat	1881
1826	alizing degree and shortest paths. <i>Social networks</i> ,	McAleese, and Geoffrey Irving. 2022. Red team-	1882
1827	32(3):245–251.	ing language models with language models. <i>arXiv</i>	1883
		<i>preprint arXiv:2202.03286</i> .	1884
1828	Liangming Pan, Xiaobao Wu, Xinyuan Lu, Anh Tuan	Verónica Pérez-Rosas, Bennett Kleinberg, Alexandra	1885
1829	Luu, William Yang Wang, Min-Yen Kan, and Preslav	Lefevre, and Rada Mihalcea. 2017. Automatic detec-	1886
1830	Nakov. 2023a. Fact-checking complex claims with	tion of fake news. <i>arXiv preprint arXiv:1708.07104</i> .	1887
1831	program-guided reasoning . In <i>Proceedings of the</i>		
1832	<i>61st Annual Meeting of the Association for Computa-</i>	Heinrich Peters and Sandra Matz. 2024. Large language	1888
1833	<i>tional Linguistics (Volume 1: Long Papers)</i> , pages	models can infer psychological dispositions of social	1889
1834	6981–7004, Toronto, Canada. Association for Com-	media users. <i>PNAS Nexus</i> , page pgae231.	1890
1835	putational Linguistics.		
1836	Yikang Pan, Liangming Pan, Wenhui Chen, Preslav	Alexandru Petrescu, Ciprian-Octavian Truică, Elena-	1891
1837	Nakov, Min-Yen Kan, and William Yang Wang. 2023b.	Simona Apostol, and Panagiotis Karras. 2021.	1892
1838	On the risk of misinformation pollution with large	Sparse shield: Social network immunization vs.	1893
1839	language models . In <i>Findings of the Association</i>	harmful speech. In <i>Proceedings of the 30th ACM</i>	1894
1840	<i>for Computational Linguistics: EMNLP 2023</i> , pages	<i>International Conference on Information & Knowl-</i>	1895
1841	1389–1403, Singapore. Association for Computa-	<i>edge Management</i> , pages 1426–1436.	1896
1842	tional Linguistics.		
1843	Yikang Pan, Liangming Pan, Wenhui Chen, Preslav	Huyen Trang Phan, Ngoc Thanh Nguyen, and Dosam	1897
1844	Nakov, Min-Yen Kan, and William Yang Wang.	Hwang. 2023. Fake news detection: A survey of	1898
1845	2023c. On the risk of misinformation pollu-	graph neural network methods. <i>Applied Soft Comput-</i>	1899
1846	tion with large language models. <i>arXiv preprint</i>	<i>ing</i> , 139:110235.	1900
1847	<i>arXiv:2305.13661</i> .		
1848	Joon Sung Park, Joseph O’Brien, Carrie Jun Cai, Mered-	Moritz Pilarski, Kirill Olegovich Solovov, and Nico-	1901
1849	ith Ringel Morris, Percy Liang, and Michael S Bern-	las Pröllochs. 2024. Community notes vs. snoping:	1902
1850	stein. 2023. Generative agents: Interactive simulacra	How the crowd selects fact-checking targets on so-	1903
1851	of human behavior. In <i>Proceedings of the 36th an-</i>	cial media. In <i>Proceedings of the International AAAI</i>	1904
1852	<i>annual acm symposium on user interface software and</i>	<i>Conference on Web and Social Media</i> , volume 18,	1905
1853	<i>technology</i> , pages 1–22.	pages 1262–1275.	1906
1854	Joon Sung Park, Lindsay Popowski, Carrie Cai, Mered-	Sanjaikanth E Vadakkethil Somanathan Pillai. 2024. En-	1907
1855	ith Ringel Morris, Percy Liang, and Michael S Bern-	hancing misinformation detection through semantic	1908
1856	stein. 2022. Social simulacra: Creating populated	analysis and knowledge graphs. In <i>2024 4th Interna-</i>	1909
1857	prototypes for social computing systems. In <i>Proceed-</i>	<i>tional Conference on Data Engineering and Commu-</i>	1910
1858	<i>ings of the 35th Annual ACM Symposium on User</i>	<i>nication Systems (ICDECS)</i> , pages 1–5. IEEE.	1911
1859	<i>Interface Software and Technology</i> , pages 1–18.		
1860	Peter S Park, Simon Goldstein, Aidan O’Gara, Michael	Clara Pizzuti. 2017. Evolutionary computation for com-	1912
1861	Chen, and Dan Hendrycks. 2024. Ai deception: A	munity detection in networks: A review. <i>IEEE Trans-</i>	1913
1862	survey of examples, risks, and potential solutions.	<i>actions on Evolutionary Computation</i> , 22(3):464–	1914
1863	<i>Patterns</i> , 5(5).	483.	1915
1864	E Parliament. 2023. Artificial intelligence act: deal on	Russell A Poldrack, Thomas Lu, and Gašper Beguš.	1916
1865	comprehensive rules for trustworthy ai. <i>Pressemit-</i>	2023. Ai-assisted coding: Experiments with gpt-4.	1917
1866	<i>teilung vom</i> , 9.	<i>arXiv preprint arXiv:2304.13187</i> .	1918
1867	Kellin Pelrine, Jacob Danovitch, and Reihaneh Rabbany.	Pascal Pons and Matthieu Latapy. 2005. Computing	1919
1868	2021. The surprising performance of simple base-	communities in large networks using random walks.	1920
1869	lines for misinformation detection. In <i>Proceedings</i>	In <i>Computer and Information Sciences-ISCIS 2005:</i>	1921
1870	<i>of the Web Conference 2021</i> , pages 3432–3441.	<i>20th International Symposium, Istanbul, Turkey, Oc-</i>	1922
		<i>ttober 26-28, 2005. Proceedings 20</i> , pages 284–293.	1923
1871	Baolin Peng, Michel Galley, Pengcheng He, Hao Cheng,	Springer.	1924
1872	Yujia Xie, Yu Hu, Qiuyuan Huang, Lars Liden, Zhou	Alin C Popescu and Hany Farid. 2004. Exposing digital	1925
1873	Yu, Weizhu Chen, et al. 2023. Check your facts and	forgeries by detecting duplicated image regions.	1926
1874	try again: Improving large language models with		
1875	external knowledge and automated feedback. <i>arXiv</i>	Ethan Porter and Thomas J Wood. 2021. The global ef-	1927
1876	<i>preprint arXiv:2302.12813</i> .	fectiveness of fact-checking: Evidence from simulta-	1928
1877	Matjaž Perc, Mahmut Ozer, and Janja Hojnik. 2019.	neous experiments in argentina, nigeria, south africa,	1929
1878	Social and juristic challenges of artificial intelligence.	and the united kingdom. <i>Proceedings of the National</i>	1930
1879	<i>Palgrave Communications</i> , 5(1).	<i>Academy of Sciences</i> , 118(37):e2104235118.	1931

1932	Martin Potthast, Johannes Kiesel, Kevin Reinartz, Janek Bevendorff, and Benno Stein. 2017. A stylometric inquiry into hyperpartisan and fake news. <i>arXiv preprint arXiv:1702.05638</i> .	1985
1933		1986
1934		1987
1935		1988
		1989
1936	Ofir Press, Muru Zhang, Sewon Min, Ludwig Schmidt, Noah A Smith, and Mike Lewis. 2022. Measuring and narrowing the compositionality gap in language models. <i>arXiv preprint arXiv:2210.03350</i> .	1990
1937		1991
1938		1992
1939		1993
1940	Nicolas Pröllochs, Dominik Bär, and Stefan Feuerriegel. 2021. Emotions in online rumor diffusion. <i>EPJ Data Science</i> , 10(1):51.	1994
1941		1995
1942		1996
1943	Peng Qi, Zehong Yan, Wynne Hsu, and Mong Li Lee. 2024. Sniffer: Multimodal large language model for explainable out-of-context misinformation detection . <i>Preprint</i> , arXiv:2403.03170.	1997
1944		
1945		1998
1946		1999
		2000
1947	Xiangyu Qi, Yi Zeng, Tinghao Xie, Pin-Yu Chen, Ruoxi Jia, Prateek Mittal, and Peter Henderson. 2023. Fine-tuning aligned language models compromises safety, even when users do not intend to! <i>Preprint</i> , arXiv:2310.03693.	2001
1948		2002
1949		
1950		2003
1951		2004
		2005
1952	Shengsheng Qian, Jinguang Wang, Jun Hu, Quan Fang, and Changsheng Xu. 2021. Hierarchical multi-modal contextual attention network for fake news detection. In <i>Proceedings of the 44th international ACM SIGIR conference on research and development in information retrieval</i> , pages 153–162.	2006
1953		2007
1954		
1955		2008
1956		2009
1957		2010
		2011
1958	Liqing Qiu, Shiqi Sai, and Moji Wei. 2022. Bpsl: a new rumor source location algorithm based on the time-stamp back propagation in social networks. <i>Applied Intelligence</i> , pages 1–13.	2012
1959		2013
1960		
1961		2014
		2015
1962	Ori Ram, Yoav Levine, Itay Dalmedigos, Dor Muhlgay, Amnon Shashua, Kevin Leyton-Brown, and Yoav Shoham. 2023. In-context retrieval-augmented language models. <i>Transactions of the Association for Computational Linguistics</i> , 11:1316–1331.	2016
1963		2017
1964		
1965		2018
1966		2019
		2020
1967	Yuan Rao and Jiangqun Ni. 2016. A deep learning approach to detection of splicing and copy-move forgeries in images. In <i>2016 IEEE international workshop on information forensics and security (WIFS)</i> , pages 1–6. IEEE.	2021
1968		2022
1969		2023
1970		2024
1971		
1972	Simone Raponi, Zeinab Khalifa, Gabriele Oligeri, and Roberto Di Pietro. 2022. Fake news propagation: A review of epidemic models, datasets, and insights. <i>ACM Transactions on the Web (TWEB)</i> , 16(3):1–34.	2025
1973		2026
1974		2027
1975		2028
		2029
1976	Hannah Rashkin, Eunsol Choi, Jin Yea Jang, Svitlana Volkova, and Yejin Choi. 2017. Truth of varying shades: Analyzing language in fake news and political fact-checking. In <i>Proceedings of the 2017 conference on empirical methods in natural language processing</i> , pages 2931–2937.	2030
1977		2031
1978		2032
1979		2033
1980		
1981		2034
1982	Vipula Rawte, Amit Sheth, and Amitava Das. 2023. A survey of hallucination in large foundation models. <i>arXiv preprint arXiv:2309.05922</i> .	2035
1983		2036
1984		2037
	Shebuti Rayana and Leman Akoglu. 2015. Collective opinion spam detection: Bridging review networks and metadata. In <i>Proceedings of the 21th acm sigkdd international conference on knowledge discovery and data mining</i> , pages 985–994.	
	Veronica Red, Eric D Kelsic, Peter J Mucha, and Mason A Porter. 2011. Comparing community structure to characteristics in online collegiate social networks. <i>SIAM review</i> , 53(3):526–543.	
	Zhicheng Ren, Zhiping Xiao, and Yizhou Sun. 2024. Do we trust what they say or what they do? a multi-modal user embedding provides personalized explanations. <i>arXiv preprint arXiv:2409.02965</i> .	
	Christopher Riederer, Yunsung Kim, Augustin Chain-treau, Nitish Korula, and Silvio Lattanzi. 2016. Linking users across domains with location data: Theory and validation. In <i>Proceedings of the 25th international conference on world wide web</i> , pages 707–719.	
	Hamid Roghani, Asgarali Bouyer, and Esmaeil Nourani. 2021. Pldls: A novel parallel label diffusion and label selection-based community detection algorithm based on spark in social networks. <i>Expert Systems with Applications</i> , 183:115377.	
	Jon Roozenbeek, Claudia R Schneider, Sarah Dryhurst, John Kerr, Alexandra LJ Freeman, Gabriel Recchia, Anne Marthe Van Der Bles, and Sander Van Der Linden. 2020. Susceptibility to misinformation about covid-19 around the world. <i>Royal Society open science</i> , 7(10):201199.	
	Martin Rosvall and Carl T Bergstrom. 2008. Maps of random walks on complex networks reveal community structure. <i>Proceedings of the national academy of sciences</i> , 105(4):1118–1123.	
	Shailendra Sahu and T Sobha Rani. 2022. A neighbour-similarity based community discovery algorithm. <i>Expert Systems with Applications</i> , 206:117822.	
	Chiman Salavati, Alireza Abdollahpouri, and Zhaleh Manbari. 2019. Ranking nodes in complex networks based on local structure and improving closeness centrality. <i>Neurocomputing</i> , 336:36–45.	
	Amine Sallah, Said Agoujl, Mudasir Ahmad Wani, Mohamed Hammad, Ahmed A Abd El-Latif, Yassine Maleh, et al. 2024. Fine-tuned understanding: Enhancing social bot detection with transformer-based classification. <i>IEEE Access</i> .	
	Dietram A Scheufele and Nicole M Krause. 2019. Science audiences, misinformation, and fake news. <i>Proceedings of the National Academy of Sciences</i> , 116(16):7662–7669.	
	Kaylyn Jackson Schiff, Daniel S Schiff, and Natália S Bueno. 2023. The liar’s dividend: Can politicians claim misinformation to evade accountability? <i>American Political Science Review</i> , pages 1–20.	

2038	Tal Schuster, Adam Fisch, and Regina Barzilay. 2021.	framework for fake news detection. In <i>2019 IEEE</i>	2093
2039	Get your vitamin c! robust fact verification with con-	<i>fifth international conference on multimedia big data</i>	2094
2040	trastive evidence. <i>arXiv preprint arXiv:2103.08541</i> .	(<i>BigMM</i>), pages 39–47. IEEE.	2095
2041	M Zubair Shafiq, Muhammad U Ilyas, Alex X Liu,	Xing Su, Shan Xue, Fanzhen Liu, Jia Wu, Jian Yang,	2096
2042	and Hayder Radha. 2013. Identifying leaders and	Chuan Zhou, Wenbin Hu, Cecile Paris, Surya Nepal,	2097
2043	followers in online social networks. <i>IEEE Journal on</i>	Di Jin, et al. 2022. A comprehensive survey on com-	2098
2044	<i>Selected Areas in Communications</i> , 31(9):618–628.	munity detection with deep learning. <i>IEEE Transac-</i>	2099
2045	Chintan Shah, Nima Dehmamy, Nicola Perra, Mat-	<i>tions on Neural Networks and Learning Systems</i> .	2100
2046	teo Chinazzi, Albert-László Barabási, Alessandro	Mengzhu Sun, Xi Zhang, Jiaqi Zheng, and Guixiang	2101
2047	Vespignani, and Rose Yu. 2020. Finding patient zero:	Ma. 2022. Ddgc: Dual dynamic graph convolu-	2102
2048	Learning contagion source with graph neural net-	tional networks for rumor detection on social media.	2103
2049	works. <i>arXiv preprint arXiv:2006.11913</i> .	In <i>Proceedings of the AAAI conference on artificial</i>	2104
2050	Chengcheng Shao, Giovanni Luca Ciampaglia, Alessan-	<i>intelligence</i> , volume 36, pages 4611–4619.	2105
2051	dro Flammini, and Filippo Menczer. 2016. Hoaxy: A	Yanshen Sun, Jianfeng He, Limeng Cui, Shuo Lei, and	2106
2052	platform for tracking online misinformation. In <i>Pro-</i>	Chang-Tien Lu. 2024. Exploring the deceptive power	2107
2053	<i>ceedings of the 25th international conference com-</i>	of llm-generated fake news: A study of real-world de-	2108
2054	<i>panion on world wide web</i> , pages 745–750.	tection challenges. <i>arXiv preprint arXiv:2403.18249</i> .	2109
2055	Chengcheng Shao, Giovanni Luca Ciampaglia, Onur	Tetsuro Takahashi and Nobuyuki Igata. 2012. Rumor	2110
2056	Varol, Kai-Cheng Yang, Alessandro Flammini, and	detection on twitter. In <i>The 6th International Con-</i>	2111
2057	Filippo Menczer. 2018. The spread of low-credibility	<i>ference on Soft Computing and Intelligent Systems,</i>	2112
2058	content by social bots. <i>Nature communications</i> ,	<i>and The 13th International Symposium on Advanced</i>	2113
2059	9(1):1–9.	<i>Intelligence Systems</i> , pages 452–457. IEEE.	2114
2060	Junming Shao, Zhichao Han, Qinli Yang, and Tao Zhou.	Shulong Tan, Ziyu Guan, Deng Cai, Xuzhen Qin, Jiajun	2115
2061	2015. Community detection based on distance dy-	Bu, and Chun Chen. 2014. Mapping users across	2116
2062	namics . In <i>Proceedings of the 21th ACM SIGKDD</i>	networks by manifold alignment on hypergraph. In	2117
2063	<i>International Conference on Knowledge Discovery</i>	<i>Proceedings of the AAAI Conference on Artificial</i>	2118
2064	<i>and Data Mining</i> , KDD ’15, page 1075–1084, New	<i>Intelligence</i> , volume 28.	2119
2065	York, NY, USA. Association for Computing Machin-	Yiming Tan, Dehai Min, Yu Li, Wenbo Li, Nan Hu,	2120
2066	ery.	Yongrui Chen, and Guilin Qi. 2023. Can chatgpt	2121
2067	Kai Shu, Limeng Cui, Suhang Wang, Dongwon Lee,	replace traditional kbqa models? an in-depth analysis	2122
2068	and Huan Liu. 2019. defend: Explainable fake news	of the question answering performance of the gpt llm	2123
2069	detection. In <i>Proceedings of the 25th ACM SIGKDD</i>	family. In <i>International Semantic Web Conference</i> ,	2124
2070	<i>international conference on knowledge discovery &</i>	pages 348–367. Springer.	2125
2071	<i>data mining</i> , pages 395–405.	Zhen Tan, Chengshuai Zhao, Raha Moraffah, Yifan Li,	2126
2072	Kai Shu, Yichuan Li, Kaize Ding, and Huan Liu. 2021.	Song Wang, Jundong Li, Tianlong Chen, and Huan	2127
2073	Fact-enhanced synthetic news generation. In <i>Pro-</i>	Liu. 2024. Blue pizza and eat rocks - exploiting vul-	2128
2074	<i>ceedings of the AAAI Conference on Artificial Intelli-</i>	nerabilities in retrieval-augmented generative models .	2129
2075	<i>gence</i> , volume 35, pages 13825–13833.	In <i>Proceedings of the 2024 Conference on Empiri-</i>	2130
2076	Kai Shu, Deepak Mahudeswaran, Suhang Wang, and	<i>cal Methods in Natural Language Processing</i> , pages	2131
2077	Huan Liu. 2020. Hierarchical propagation networks	1610–1626, Miami, Florida, USA. Association for	2132
2078	for fake news detection: Investigation and exploita-	Computational Linguistics.	2133
2079	tion. In <i>Proceedings of the international AAAI con-</i>	Juvaria Tariq, Muhammad Ahmad, Imdadullah Khan,	2134
2080	<i>ference on web and social media</i> , volume 14, pages	and Mudassir Shabbir. 2017. Scalable approximation	2135
2081	626–637.	algorithm for network immunization. <i>arXiv preprint</i>	2136
2082	Kai Shu, Amy Sliva, Suhang Wang, Jiliang Tang, and	<i>arXiv:1711.00784</i> .	2137
2083	Huan Liu. 2017. Fake news detection on social me-	Kassym-Jomart Tokayev. 2023. Ethical implications	2138
2084	dia: A data mining perspective. <i>ACM SIGKDD ex-</i>	of large language models a multidimensional explo-	2139
2085	<i>plorations newsletter</i> , 19(1):22–36.	ration of societal, economic, and technical concerns.	2140
2086	Ronit Singhal, Pransh Patwa, Parth Patwa, Aman	<i>International Journal of Social Analytics</i> , 8(9):17–	2141
2087	Chadha, and Amitava Das. 2024. Evidence-backed	33.	2142
2088	fact checking using rag and few-shot in-context learn-	Guangmo Tong, Weili Wu, and Ding-Zhu Du. 2018.	2143
2089	ing with llms. <i>arXiv preprint arXiv:2408.12060</i> .	Distributed rumor blocking with multiple positive	2144
2090	Shivangi Singhal, Rajiv Ratn Shah, Tanmoy	cascades. <i>IEEE Transactions on Computational So-</i>	2145
2091	Chakraborty, Ponnuram Kumaraguru, and	<i>cial Systems</i> , 5(2):468–480.	2146
2092	Shin’ichi Satoh. 2019. Spotfake: A multi-modal		

2147	Petter Törnberg, Diliara Valeeva, Justus Uitermark, and Christopher Bail. 2023. Simulating social media using large language models to evaluate alternative news feed algorithms. <i>arXiv preprint arXiv:2310.05984</i> .	2203
2148		2204
2149		2205
2150		2206
2151		
2152	Rudra M Tripathy, Amitabha Bagchi, and Sameep Mehta. 2010. A study of rumor control strategies on social networks. In <i>Proceedings of the 19th ACM international conference on Information and knowledge management</i> , pages 1817–1820.	2207
2153		2208
2154		2209
2155		2210
2156		2211
2157	U Undeutsch. 1967. Beurteilung der glaubhaftigkeit von aussagen [veracity assessment of statements]. undeutsch. <i>Handbuch der Psychologie</i> , 11:26–181.	
2158		
2159		
2160	Cristian Vaccari and Andrew Chadwick. 2020. Deep-fakes and disinformation: Exploring the impact of synthetic political video on deception, uncertainty, and trust in news. <i>Social media+ society</i> , 6(1):2056305120903408.	2212
2161		2213
2162		2214
2163		2215
2164		2216
2165	Karthik Valmeekam, Matthew Marquez, Alberto Olmo, Sarath Sreedharan, and Subbarao Kambhampati. 2023. Planbench: An extensible benchmark for evaluating large language models on planning and reasoning about change. <i>Advances in Neural Information Processing Systems</i> , 36:38975–38987.	2217
2166		2218
2167		2219
2168		2220
2169		2221
2170		2222
2171	Neeraj Varshney, Wenlin Yao, Hongming Zhang, Jian-shu Chen, and Dong Yu. 2023. A stitch in time saves nine: Detecting and mitigating hallucinations of llms by validating low-confidence generation. <i>arXiv preprint arXiv:2307.03987</i> .	2223
2172		2224
2173		2225
2174		2226
2175		2227
2176		2228
2177		2229
2178		
2179	Nikhita Vedula and Srinivasan Parthasarathy. 2021. Face-keg: Fact checking explained using knowledge graphs. In <i>Proceedings of the 14th ACM International Conference on Web Search and Data Mining</i> , pages 526–534.	2230
2180		2231
2181		2232
2182	Andreas Vlachos and Sebastian Riedel. 2014. Fact checking: Task definition and dataset construction. In <i>Proceedings of the ACL 2014 workshop on language technologies and computational social science</i> , pages 18–22.	2233
2183		
2184		
2185		
2186	Soroush Vosoughi, Deb Roy, and Sinan Aral. 2018. The spread of true and false news online. <i>science</i> , 359(6380):1146–1151.	2234
2187		2235
2188		2236
2189	Ivan Vykopal, Matúš Pikuliak, Ivan Srba, Robert Moro, Dominik Macko, and Maria Bielikova. 2023. Disinformation capabilities of large language models. <i>arXiv preprint arXiv:2311.08838</i> .	2237
2190		
2191		
2192		
2193	David Wadden, Kyle Lo, Lucy Lu Wang, Arman Cohan, Iz Beltagy, and Hannaneh Hajishirzi. 2021. Multivers: Improving scientific claim verification with weak supervision and full-document context. <i>arXiv preprint arXiv:2112.01640</i> .	2238
2194		2239
2195		2240
2196		2241
2197		2242
2198		
2199	Nathan Walter and Riva Tukachinsky. 2020. A meta-analytic examination of the continued influence of misinformation in the face of correction: How powerful is it, why does it happen, and how to stop it? <i>Communication research</i> , 47(2):155–177.	2243
2200		2244
2201		2245
2202		2246
		2247
		2248
	Chen Wan, Tao Li, and Zhicheng Sun. 2017. Global stability of a seir rumor spreading model with demographics on scale-free networks. <i>Advances in Difference Equations</i> , 2017(1):253.	2249
		2250
		2251
		2252
	Herun Wan, Shangbin Feng, Zhaoxuan Tan, Heng Wang, Yulia Tsvetkov, and Minnan Luo. 2024. Dell: Generating reactions and explanations for llm-based misinformation detection. <i>arXiv preprint arXiv:2402.10426</i> .	2253
		2254
		2255
		2256
		2257
	Biao Wang, Ge Chen, Luoyi Fu, Li Song, and Xinbing Wang. 2017. Drimux: Dynamic rumor influence minimization with user experience in social networks. <i>IEEE Transactions on Knowledge and Data Engineering</i> , 29(10):2168–2181.	
	Bo Wang, Jing Ma, Hongzhan Lin, Zhiwei Yang, Ruichao Yang, Yuan Tian, and Yi Chang. 2024a. Explainable fake news detection with large language model via defense among competing wisdom. In <i>Proceedings of the ACM on Web Conference 2024</i> , pages 2452–2463.	
	Bo Wang, Jing Ma, Hongzhan Lin, Zhiwei Yang, Ruichao Yang, Yuan Tian, and Yi Chang. 2024b. Explainable fake news detection with large language model via defense among competing wisdom. In <i>Proceedings of the ACM Web Conference 2024</i> , WWW ’24, page 2452–2463, New York, NY, USA. Association for Computing Machinery.	
	Chenxi Wang, Zongfang Liu, Dequan Yang, and Xiuying Chen. 2024c. Decoding echo chambers: Llm-powered simulations revealing polarization in social networks. <i>arXiv preprint arXiv:2409.19338</i> .	
	Junxiang Wang, Junji Jiang, and Liang Zhao. 2022. An invertible graph diffusion neural network for source localization. In <i>Proceedings of the ACM Web Conference 2022</i> , pages 1058–1069.	
	Sheng-Yu Wang, Oliver Wang, Richard Zhang, Andrew Owens, and Alexei A Efros. 2020a. Cnn-generated images are surprisingly easy to spot... for now. In <i>Proceedings of the IEEE/CVF conference on computer vision and pattern recognition</i> , pages 8695–8704.	
	Yaqing Wang, Fenglong Ma, Zhiwei Jin, Ye Yuan, Guangxu Xun, Kishlay Jha, Lu Su, and Jing Gao. 2018a. Eann: Event adversarial neural networks for multi-modal fake news detection. In <i>Proceedings of the 24th acm sigkdd international conference on knowledge discovery & data mining</i> , pages 849–857.	
	Yi Wang, Jinde Cao, Xiaodi Li, and Ahmed Alsaedi. 2018b. Edge-based epidemic dynamics with multiple routes of transmission on random networks. <i>Non-linear Dynamics</i> , 91:403–420.	
	Youze Wang, Shengsheng Qian, Jun Hu, Quan Fang, and Changsheng Xu. 2020b. Fake news detection via knowledge-driven multimodal graph convolutional networks. In <i>Proceedings of the 2020 international conference on multimedia retrieval</i> , pages 540–547.	

2258	Zecong Wang, Jiayi Cheng, Chen Cui, and Chenhao Yu.	Yang Xia, Haijun Jiang, Shuzhen Yu, and Zhiyong Yu.	2310
2259	2023. Implementing bert and fine-tuned roberta to	2024. The dynamic analysis of the rumor spreading	2311
2260	detect ai generated news by chatgpt. <i>arXiv preprint</i>	and behavior diffusion model with higher-order inter-	2312
2261	<i>arXiv:2306.07401</i> .	actions. <i>Communications in Nonlinear Science and</i>	2313
		<i>Numerical Simulation</i> , 138:108186.	2314
2262	Alicia Wanless and Michael Berk. 2020. The audience	Bin Xiao, Yang Wei, Xiuli Bi, Weisheng Li, and Jian-	2315
2263	is the amplifier: Participatory propaganda. <i>The SAGE</i>	feng Ma. 2020. Image splicing forgery detection	2316
2264	<i>handbook of propaganda</i> , pages 85–104.	combining coarse to refined convolutional neural net-	2317
		work and adaptive clustering. <i>Information Sciences</i> ,	2318
2265	Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten	511:172–191.	2319
2266	Bosma, Fei Xia, Ed Chi, Quoc V Le, Denny Zhou,	Yunpeng Xiao, Diqiang Chen, Shihong Wei, Qian Li,	2320
2267	et al. 2022. Chain-of-thought prompting elicits rea-	Haohan Wang, and Ming Xu. 2019. Rumor propaga-	2321
2268	soning in large language models. <i>Advances in neural</i>	dynamic model based on evolutionary game and	2322
2269	<i>information processing systems</i> , 35:24824–24837.	anti-rumor. <i>Nonlinear Dynamics</i> , 95:523–539.	2323
2270	Jianliang Wei and Fei Meng. 2021. How opinion dis-	Yunpeng Xiao, Jinsong Yang, Wanjing Zhao, Qian Li,	2324
2271	ortion appears in super-influencer dominated so-	and Yucai Pang. 2024. Cross-domain social rumor-	2325
2272	cial network. <i>Future Generation Computer Systems</i> ,	propagation model based on transfer learning. <i>IEEE</i>	2326
2273	115:542–552.	<i>Transactions on Neural Networks and Learning Sys-</i>	2327
		<i>tems</i> .	2328
2274	Teresa Weikmann and Sophie Lecheler. 2023. Visual	Chengxing Xie, Canyu Chen, Feiran Jia, Ziyu Ye,	2329
2275	disinformation in a digital age: A literature synthe-	Kai Shu, Adel Bibi, Ziniu Hu, Philip Torr, Bernard	2330
2276	sis and research agenda. <i>New Media & Society</i> ,	Ghanem, and Guohao Li. 2024. Can large language	2331
2277	25(12):3696–3713.	model agents simulate human trust behaviors? <i>arXiv</i>	2332
		<i>preprint arXiv:2402.04559</i> .	2333
2278	Ferre Wouters and Michaël Opgenhaffen. 2024. Re-	Jierui Xie, Stephen Kelley, and Boleslaw K Szymanski.	2334
2279	gional facts matter: A comparative perspective of	2013. Overlapping community detection in networks:	2335
2280	sub-state fact-checking initiatives in europe. <i>Media</i>	The state-of-the-art and comparative study. <i>Acm com-</i>	2336
2281	<i>and Communication</i> , 12.	<i>puting surveys (csur)</i> , 45(4):1–35.	2337
2282	Feijie Wu, Zitao Li, Yaliang Li, Bolin Ding, and Jing	Yueqi Xie, Jingwei Yi, Jiawei Shao, Justin Curl,	2338
2283	Gao. 2024a. Fedbiot: Llm local fine-tuning in feder-	Lingjuan Lyu, Qifeng Chen, Xing Xie, and Fangzhao	2339
2284	ated learning without full model. In <i>Proceedings of</i>	Wu. 2023. Defending chatgpt against jailbreak at-	2340
2285	<i>the 30th ACM SIGKDD Conference on Knowledge</i>	tack via self-reminders. <i>Nature Machine Intelligence</i> ,	2341
2286	<i>Discovery and Data Mining</i> , pages 3345–3355.	5(12):1486–1496.	2342
2287	Guangyang Wu, Weijie Wu, Xiaohong Liu, Kele Xu,	Miao Xiong, Zhiyuan Hu, Xinyang Lu, Yifei Li, Jie	2343
2288	Tianjiao Wan, and Wenyi Wang. 2023a. Cheap-fake	Fu, Junxian He, and Bryan Hooi. 2023. Can llms	2344
2289	detection with llm using prompt engineering. In <i>2023</i>	express their uncertainty? an empirical evaluation	2345
2290	<i>IEEE International Conference on Multimedia and</i>	of confidence elicitation in llms. <i>arXiv preprint</i>	2346
2291	<i>Expo Workshops (ICMEW)</i> , pages 105–109. IEEE.	<i>arXiv:2306.13063</i> .	2347
2292	Jiaying Wu, Jiafeng Guo, and Bryan Hooi. 2024b. Fake	Junhao Xu, Longdi Xian, Zening Liu, Mingliang Chen,	2348
2293	news in sheep’s clothing: Robust fake news detection	Qiuyang Yin, and Fenghua Song. 2024a. The future	2349
2294	against llm-empowered style attacks. In <i>Proceedings</i>	of combating rumors? retrieval, discrimination, and	2350
2295	<i>of the 30th ACM SIGKDD conference on knowledge</i>	generation. <i>arXiv preprint arXiv:2403.20204</i> .	2351
2296	<i>discovery and data mining</i> , pages 3367–3378.	Rongwu Xu, Brian Lin, Shujian Yang, Tianqi Zhang,	2352
2297	Jun Wu, Xuesong Ye, and Chengjie Mou. 2023b. Bot-	Weiyan Shi, Tianwei Zhang, Zhixuan Fang, Wei Xu,	2353
2298	shape: A novel social bots detection approach via be-	and Han Qiu. 2024b. The earth is flat because...: Investigating LLMs’ belief towards misinformation via persuasive conversation . In <i>Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)</i> , pages 16259–16303, Bangkok, Thailand. Association for Computational Linguistics.	2354
2299	havioral patterns. <i>arXiv preprint arXiv:2303.10214</i> .		2355
2300	Ke Wu, Song Yang, and Kenny Q Zhu. 2015. False		2356
2301	rumors detection on sina weibo by propagation struc-		2357
2302	tures. In <i>2015 IEEE 31st international conference on</i>		2358
2303	<i>data engineering</i> , pages 651–662. IEEE.		2359
2304	Siyuan Wu, Yue Huang, Chujie Gao, Dongping Chen,	Rongwu Xu, Brian S Lin, Shujian Yang, Tianqi Zhang,	2360
2305	Qihui Zhang, Yao Wan, Tianyi Zhou, Xiangliang	Weiyan Shi, Tianwei Zhang, Zhixuan Fang, Wei	2361
2306	Zhang, Jianfeng Gao, Chaowei Xiao, et al. 2024c.	Xu, and Han Qiu. 2023. The earth is flat be-	2362
2307	Unigen: A unified framework for textual dataset gen-	cause...: Investigating llms’ belief towards misinfor-	2363
2308	eration using large language models. <i>arXiv preprint</i>	mation via persuasive conversation. <i>arXiv preprint</i>	2364
2309	<i>arXiv:2406.18966</i> .	<i>arXiv:2312.09085</i> .	2365
			2366

2367	Ziwei Xu, Sanjay Jain, and Mohan Kankanhalli.	Liang Yang, Xiaochun Cao, Dongxiao He, Chuan Wang,	2421
2368	2024c. Hallucination is inevitable: An innate lim-	Xiao Wang, and Weixiong Zhang. 2016. Modularity	2422
2369	itation of large language models. <i>arXiv preprint</i>	based community detection with deep learning. In	2423
2370	<i>arXiv:2401.11817</i> .	<i>IJCAI</i> , volume 16, pages 2252–2258.	2424
2371	Qi Xuan, Xincheng Shu, Zhongyuan Ruan, Jinbao	Ruichao Yang, Jing Ma, Wei Gao, and Hongzhan Lin.	2425
2372	Wang, Chenbo Fu, and Guanrong Chen. 2019. A	2025. Llm-enhanced multiple instance learning for	2426
2373	self-learning information diffusion model for smart	joint rumor and stance detection with social context	2427
2374	social networks. <i>IEEE Transactions on Network Sci-</i>	information. <i>ACM Transactions on Intelligent Sys-</i>	2428
2375	<i>ence and Engineering</i> , 7(3):1466–1480.	<i>tems and Technology</i> .	2429
2376	Ruidong Yan, Yi Li, Weili Wu, Deying Li, and Yongcai	Tianbao Yang, Rong Jin, Yun Chi, and Shenghuo Zhu.	2430
2377	Wang. 2019. Rumor blocking through online link	2009. Combining link and content for community	2431
2378	deletion on social networks. <i>ACM Transactions on</i>	detection: a discriminative approach. In <i>Proceedings</i>	2432
2379	<i>Knowledge Discovery from Data (TKDD)</i> , 13(2):1–	<i>of the 15th ACM SIGKDD international conference</i>	2433
2380	26.	<i>on Knowledge discovery and data mining</i> , pages 927–	2434
2381	Chang Yang, Xia Yu, JiaYi Wu, BoZhen Zhang, and	936.	2435
2382	HaiBo Yang. 2024a. Graph-aware multi-feature in-	Xianjun Yang, Xiao Wang, Qi Zhang, Linda Petzold,	2436
2383	teracting network for explainable rumor detection on	William Yang Wang, Xun Zhao, and Dahua Lin.	2437
2384	social network. <i>Expert Systems with Applications</i> ,	2023c. Shadow alignment: The ease of subvert-	2438
2385	249:123687.	ing safely-aligned language models. <i>arXiv preprint</i>	2439
2386	Chang Yang, Peng Zhang, Hui Gao, and Jing Zhang.	<i>arXiv:2310.02949</i> .	2440
2387	2024b. Deciphering rumors: A multi-task learning	Zhengyuan Yang, Linjie Li, Kevin Lin, Jianfeng	2441
2388	approach with intent-aware hierarchical contrastive	Wang, Chung-Ching Lin, Zicheng Liu, and Lijuan	2442
2389	learning. In <i>Proceedings of the 2024 Conference on</i>	Wang. 2023d. The dawn of lmms: Preliminary	2443
2390	<i>Empirical Methods in Natural Language Processing</i> ,	explorations with gpt-4v (ision). <i>arXiv preprint</i>	2444
2391	pages 4471–4483.	<i>arXiv:2309.17421</i> , 9(1):1.	2445
2392	Chang Yang, Peng Zhang, Wenbo Qiao, Hui Gao, and	Barry Menglong Yao, Aditya Shah, Lichao Sun, Jin-	2446
2393	Jiaming Zhao. 2023a. Rumor detection on social	Hee Cho, and Lifu Huang. 2023a. End-to-end multi-	2447
2394	media with crowd intelligence and chatgpt-assisted	modal fact-checking and explanation generation: A	2448
2395	networks. In <i>Proceedings of the 2023 Conference on</i>	challenging dataset and models. In <i>Proceedings of</i>	2449
2396	<i>Empirical Methods in Natural Language Processing</i> ,	<i>the 46th International ACM SIGIR Conference on</i>	2450
2397	pages 5705–5717.	<i>Research and Development in Information Retrieval</i> ,	2451
2398	Fan Yang, Yang Liu, Xiaohui Yu, and Min Yang. 2012.	pages 2733–2743.	2452
2399	<i>Automatic detection of rumor on sina weibo</i> . In <i>Pro-</i>	Shunyu Yao, Dian Yu, Jeffrey Zhao, Izhak Shafran,	2453
2400	<i>ceedings of the ACM SIGKDD Workshop on Mining</i>	Tom Griffiths, Yuan Cao, and Karthik Narasimhan.	2454
2401	<i>Data Semantics</i> , MDS '12, New York, NY, USA.	2023b. Tree of thoughts: Deliberate problem solving	2455
2402	Association for Computing Machinery.	with large language models. <i>Advances in neural</i>	2456
2403	Fan Yang, Shiva K Pentiyala, Sina Mohseni, Mengnan	<i>information processing systems</i> , 36:11809–11822.	2457
2404	Du, Hao Yuan, Rhema Linder, Eric D Ragan, Shui-	Xiaopeng Yao, Yue Gu, Chonglin Gu, and Hejiao	2458
2405	wang Ji, and Xia Hu. 2019. Xfake: Explainable fake	Huang. 2022. Fast controlling of rumors with limited	2459
2406	news detector with visualizations. In <i>The world wide</i>	cost in social networks. <i>Computer Communications</i> ,	2460
2407	<i>web conference</i> , pages 3600–3604.	182:41–51.	2461
2408	Jaewon Yang and Jure Leskovec. 2013. Overlapping	Yining Ye, Xin Cong, Shizuo Tian, Jiannan Cao, Hao	2462
2409	community detection at scale: a nonnegative matrix	Wang, Yujia Qin, Yaxi Lu, Heyang Yu, Huadong	2463
2410	factorization approach. In <i>Proceedings of the sixth</i>	Wang, Yankai Lin, et al. 2023. Proagent: From	2464
2411	<i>ACM international conference on Web search and</i>	robotic process automation to agentic process au-	2465
2412	<i>data mining</i> , pages 587–596.	tomation. <i>arXiv preprint arXiv:2311.10751</i> .	2466
2413	Kaijun Yang, Qing Bao, and Hongjun Qiu. 2023b.	Shuzhen Yu, Zhiyong Yu, Haijun Jiang, and Shuai Yang.	2467
2414	Identifying multiple propagation sources with motif-	2021. The dynamics and control of 2i2sr rumor	2468
2415	based graph convolutional networks for social net-	spreading models in multilingual online social net-	2469
2416	works. <i>IEEE Access</i> , 11:61630–61645.	works. <i>Information Sciences</i> , 581:18–41.	2470
2417	Li Yang, Yafeng Qiao, Zhihong Liu, Jianfeng Ma, and	Damián H Zanette. 2001. Critical behavior of propa-	2471
2418	Xinghua Li. 2018. Identifying opinion leader nodes	gation on small-world networks. <i>Physical Review E</i> ,	2472
2419	in online social networks with a new closeness evalu-	64(5):050901.	2473
2420	ation algorithm. <i>Soft Computing</i> , 22:453–464.	Runxi Zeng and Di Zhu. 2019. A model and simulation	2474
		of the emotional contagion of netizens in the process	2475
		of rumor refutation. <i>Scientific reports</i> , 9(1):14164.	2476

2477	Huaiwen Zhang, Quan Fang, Shengsheng Qian, and	Jian Zhao, Nan Cao, Zhen Wen, Yale Song, Yu-Ru Lin,	2534
2478	Changsheng Xu. 2019. Multi-modal knowledge-	and Christopher Collins. 2014. # fluxflow: Visual	2535
2479	aware event memory network for social media rumor	analysis of anomalous information spreading on so-	2536
2480	detection. In <i>Proceedings of the 27th ACM interna-</i>	cial media. <i>IEEE transactions on visualization and</i>	2537
2481	<i>tional conference on multimedia</i> , pages 1942–1951.	<i>computer graphics</i> , 20(12):1773–1782.	2538
2482	Litian Zhang, Xiaoming Zhang, Ziyi Zhou, Feiran	Laijun Zhao, Hongxin Cui, Xiaoyan Qiu, Xiaoli Wang,	2539
2483	Huang, and Chaozhuo Li. 2024a. Reinforced adap-	and Jiajia Wang. 2013. Sir rumor spreading model in	2540
2484	tive knowledge learning for multimodal fake news	the new media age. <i>Physica A: Statistical Mechanics</i>	2541
2485	detection. In <i>Proceedings of the AAAI Conference</i>	<i>and its Applications</i> , 392(4):995–1003.	2542
2486	<i>on Artificial Intelligence</i> , volume 38, pages 16777–		
2487	16785.		
2488	Litian Zhang, Xiaoming Zhang, Ziyi Zhou, Xi Zhang,	Ruochen Zhao, Hailin Chen, Weishi Wang, Fangkai	2543
2489	S Yu Philip, and Chaozhuo Li. 2025. Knowledge-	Jiao, Xuan Long Do, Chengwei Qin, Bosheng	2544
2490	aware multimodal pre-training for fake news detec-	Ding, Xiaobao Guo, Minzhi Li, Xingxuan Li, et al.	2545
2491	tion. <i>Information Fusion</i> , 114:102715.	2023a. Retrieving multimodal information for aug-	2546
		mented generation: A survey. <i>arXiv preprint</i>	2547
2492	Mingqing Zhang, Haisong Gong, Qiang Liu, Shu Wu,	<i>arXiv:2303.10868</i> .	2548
2493	and Liang Wang. 2024b. Breaking event rumor detec-		
2494	tion via stance-separated multi-agent debate. <i>arXiv</i>	Wayne Xin Zhao, Kun Zhou, Junyi Li, Tianyi Tang,	2549
2495	<i>preprint arXiv:2412.04859</i> .	Xiaolei Wang, Yupeng Hou, Yingqian Min, Beichen	2550
		Zhang, Junjie Zhang, Zican Dong, et al. 2023b. A	2551
2496	Qi Zhang, Yuan Li, Jialing Zou, Jianming Zhu, Dingn-	survey of large language models. <i>arXiv preprint</i>	2552
2497	ing Liu, and Jianbin Jiao. 2024c. Unifying multi-	<i>arXiv:2303.18223</i> .	2553
2498	modal interactions for rumor diffusion prediction		
2499	with global hypergraph modeling. <i>Knowledge-Based</i>	Yiran Zhao, Jinghan Zhang, I Chern, Siyang Gao,	2554
2500	<i>Systems</i> , 301:112246.	Pengfei Liu, Junxian He, et al. 2023c. Felm: Bench-	2555
		marking factuality evaluation of large language mod-	2556
2501	Si Zhang and Hanghang Tong. 2016. Final: Fast at-	els. <i>Advances in Neural Information Processing Sys-</i>	2557
2502	tributed network alignment. In <i>Proceedings of the</i>	<i>tems</i> , 36:44502–44523.	2558
2503	<i>22nd ACM SIGKDD international conference on</i>		
2504	<i>knowledge discovery and data mining</i> , pages 1345–	Yuxin Zhao, Shenghong Li, and Feng Jin. 2016. Identi-	2559
2505	1354.	fication of influential nodes in social networks with	2560
		community structure based on label propagation.	2561
2506	Xuan Zhang and Wei Gao. 2023a. Towards llm-based	<i>Neurocomputing</i> , 210:34–44.	2562
2507	fact verification on news claims with a hierarchi-		
2508	cal step-by-step prompting method. <i>arXiv preprint</i>	Wanjun Zhong, Jingjing Xu, Duyu Tang, Zenan Xu,	2563
2509	<i>arXiv:2310.00305</i> .	Nan Duan, Ming Zhou, Jiahai Wang, and Jian Yin.	2564
		2019. Reasoning over semantic-level graph for fact	2565
2510	Xuan Zhang and Wei Gao. 2023b. Towards LLM-based	checking. <i>arXiv preprint arXiv:1909.03745</i> .	2566
2511	fact verification on news claims with a hierarchical		
2512	step-by-step prompting method. In <i>Proceedings of</i>	Jiawei Zhou, Yixuan Zhang, Qianni Luo, Andrea G	2567
2513	<i>the 13th International Joint Conference on Natural</i>	Parker, and Munmun De Choudhury. 2023. Synthetic	2568
2514	<i>Language Processing and the 3rd Conference of the</i>	lies: Understanding ai-generated misinformation and	2569
2515	<i>Asia-Pacific Chapter of the Association for Computa-</i>	evaluating algorithmic and human solutions. In <i>Pro-</i>	2570
2516	<i>tational Linguistics (Volume 1: Long Papers)</i> , pages	<i>ceedings of the 2023 CHI Conference on Human</i>	2571
2517	996–1011, Nusa Dua, Bali. Association for Computa-	<i>Factors in Computing Systems</i> , pages 1–20.	2572
2518	tational Linguistics.		
2519	Yao Zhang and B Aditya Prakash. 2015. Data-aware	Jie Zhou, Xu Han, Cheng Yang, Zhiyuan Liu,	2573
2520	vaccine allocation over large networks. <i>ACM Trans-</i>	Lifeng Wang, Changcheng Li, and Maosong Sun.	2574
2521	<i>actions on Knowledge Discovery from Data (TKDD)</i> ,	2019a. Gear: Graph-based evidence aggregating	2575
2522	10(2):1–32.	and reasoning for fact verification. <i>arXiv preprint</i>	2576
		<i>arXiv:1908.01843</i> .	2577
2523	Yuan Zhang, Tianshu Lyu, and Yan Zhang. 2018. Co-	Peng Zhou, Xintong Han, Vlad I Morariu, and Larry S	2578
2524	sine: Community-preserving social network embed-	Davis. 2018. Learning rich features for image ma-	2579
2525	ding from information diffusion cascades. In <i>Pro-</i>	manipulation detection. In <i>Proceedings of the IEEE</i>	2580
2526	<i>ceedings of the AAAI conference on artificial intelli-</i>	<i>conference on computer vision and pattern recogni-</i>	2581
2527	<i>gence</i> , volume 32.	<i>tion</i> , pages 1053–1061.	2582
2528	Yutao Zhang, Jie Tang, Zhilin Yang, Jian Pei, and	Xiaoping Zhou, Xun Liang, Haiyan Zhang, and Yuefeng	2583
2529	Philip S Yu. 2015. Cosnet: Connecting heteroge-	Ma. 2015. Cross-platform identification of anony-	2584
2530	neous social networks with local and global consis-	ymous identical users in multiple social media net-	2585
2531	tency. In <i>Proceedings of the 21th ACM SIGKDD in-</i>	works. <i>IEEE transactions on knowledge and data</i>	2586
2532	<i>ternational conference on knowledge discovery and</i>	<i>engineering</i> , 28(2):411–424.	2587
2533	<i>data mining</i> , pages 1485–1494.		

Xinyi Zhou and Reza Zafarani. 2018. Fake news: A survey of research, detection methods, and opportunities. *arXiv preprint arXiv:1812.00315*, 2.

Zhixuan Zhou, Huankang Guan, Meghana Moorthy Bhat, and Justin Hsu. 2019b. Fake news detection via nlp is vulnerable to adversarial attacks. *arXiv preprint arXiv:1901.09657*.

Biru Zhu, Xingyao Zhang, Ming Gu, and Yangdong Deng. 2021. Knowledge enhanced fact checking and verification. *IEEE/ACM Transactions on Audio, Speech, and Language Processing*, 29:3132–3143.

Lingxuan Zhu, Weiming Mou, and Peng Luo. 2024. Potential of large language models as tools against medical disinformation. *JAMA Internal Medicine*, 184(4):450–450.

Linhe Zhu and Xiaoyuan Huang. 2019. Sis model of rumor spreading in social network with time delay and nonlinear functions. *Communications in Theoretical Physics*, 72(1):015002.

Peican Zhu, Le Cheng, Chao Gao, Zhen Wang, and Xue-long Li. 2022a. Locating multi-sources in social networks with a low infection rate. *IEEE Transactions on Network Science and Engineering*, 9(3):1853–1865.

Yongchun Zhu, Qiang Sheng, Juan Cao, Shuokai Li, Danding Wang, and Fuzhen Zhuang. 2022b. Generalizing to the future: Mitigating entity bias in fake news detection. In *Proceedings of the 45th International ACM SIGIR Conference on Research and Development in Information Retrieval*, pages 2120–2125.

Linlin Zong, Jiahui Zhou, Wenmin Lin, Xinyue Liu, Xianchao Zhang, and Bo Xu. 2024. Unveiling opinion evolution via prompting and diffusion for short video fake news detection. In *Findings of the Association for Computational Linguistics ACL 2024*, pages 10817–10826.

Arkaitz Zubiaga, Maria Liakata, Rob Procter, Kalina Bontcheva, and Peter Tolmie. 2015. Towards detecting rumours in social media. In *Workshops at the Twenty-Ninth AAAI conference on artificial intelligence*.

M Zuckerman. 1981. Verbal and nonverbal communication of deception. *Advances in experimental social psychology/Academic Press*.

A Rumor Definition and Related Tasks

The core characteristic of rumors lies in their "unverified ambiguity and uncertainty," which makes them highly prone to misinterpretation or misuse during the dissemination process. Unlike debunked false information (misinformation) (Scheufele and Krause, 2019; Kumar and Geethakumari, 2014), deliberately fabricated falsehoods (disinformation)

(Guo et al., 2020), or fake news that adopts the form of journalistic reporting to deliberately mislead the public (Shu et al., 2017, 2019) (Detecting fake news with NLP)], the uniqueness of rumors lies in the dynamic evolution of their verification status. Currently, rumor detection in a broad sense largely focuses on the verification of rumor veracity, emphasizing the description of the potential risks posed by false rumors to societal trust (Takahashi and Igata, 2012; Wu et al., 2015; Liang et al., 2015). From a narrower perspective, studies on rumors also consider their dissemination characteristics and societal impacts (Allport, 1947; Zubiaga et al., 2015). This provides theoretical support for uncovering the deeper logic underpinning rumor propagation while laying the foundational framework for research in rumor detection.

B Future Research Directions

B.1 LLM-based Multi-agent Social Simulation in Rumor Detection

While existing studies have shed light on the potential of rumor detection agents, there remains a lack of in-depth research into the complexities of the task. Critical gaps persist in addressing the key aspects of handling the complexity of rumor propagation, which can be summarized into the following three areas:

Lack of a deep understanding of the mechanisms of rumor propagation. Current research often relies on shallow features to classify information but fails to account for rumor dissemination’s intricate cognitive behavioral patterns and socio-dynamic characteristics. For instance, individuals tend to process information in ways that align with their pre-existing beliefs (cognitive consistency theory) while exhibiting significant non-objective and selective cognitive biases (Nickerson, 1998). Emotional drivers (such as fear and anger) (Chan et al., 2021), and social pressures (Blass, 1984) amplify the effects of rumor propagation. Moreover, some users spread false information not purely based on judgments of its truthfulness but rather due to hedonistic motives (Jiwa et al., 2023) or a sense of group identity (Lewandowsky, 2022; Wanless and Berk, 2020). Social networks’ distributed and decentralized nature further reduces individuals’ capacity to discern false information. It increases the likelihood of social cascades, forming extreme attitudes (Jamieson, 2008). This aspect has also received attention, e.g., simulating social phenomena

like echo chambers (Wang et al., 2024c).

Over-simplified modeling of rumor propagation and intervention. Current models are often restricted to static propagation roles (e.g., spreaders, bystanders, fact-checkers) and fail to comprehensively capture the dynamic changes in individual behaviors during the dissemination process and the dynamic role shifts in group interactions. In contrast, the information propagation dynamics extensively studied in information epidemiology offer valuable inspiration, such as integrating complex social variables (educational levels, forgetting mechanisms) to model dynamic processes more accurately. At the same time, intervention measures also lack flexibility and dynamic optimization regarding accurate verification and influencer blocking. This limits their ability to adapt to the gradual and complex evolution of rumor propagation, which often requires intervention strategies that are adaptable and sensitive to contextual changes.

Insufficient global modeling of the dynamic attributes of rumors. Existing studies are often confined to a specific dimension of rumors, such as content, propagation, or interaction, without systematically integrating these interdependent elements within a logical framework. Furthermore, the combined effect of multimodal information on rumor dissemination and the critical role of social bots as key propagation drivers have not been sufficiently considered. Additionally, active user behaviors, such as retrieval and learning of controversial information, remain underexplored. Although individual differences exist, such behaviors (such as actively verifying the authenticity of information and subsequently choosing to share, report, or ignore it) significantly impact information perception and dissemination.

B.1.1 Multi-Agent Systems Under the CIB Framework

To tackle existing challenges, we propose a research roadmap for multi-agent systems within the CIB framework, as illustrated in Figure 3. This approach utilizes cross-layer dynamic feedback to establish a **Macro-Micro Feedback Loop**, enabling the modeling of macro-level information dissemination based on micro-level individual cognition. It facilitates dynamic, bidirectional interventions between macro and micro levels, fostering evolution driven by deeper cognitive insights.

At the cognitive layer, agents can utilize LLMs and multimodal analysis tools to achieve a deep se-

mantic understanding of text, images, videos, and other content associated with rumors, also performing real-time monitoring and dynamic analysis of content flow on social media platforms. By incorporating psychological models, agents can dynamically assess the authenticity of information and precisely quantify user cognitive biases (e.g., selective processing or emotion-driven behaviors). For instance, agents can infer malicious intents behind false information through context-aware reasoning and identify the potential motivations of target users for spreading such information. This enables agents to provide information support for subsequent collaborative actions.

At the interaction layer, multi-agent systems leverage sophisticated communication protocols to collaborate efficiently, simulating the diversity of user behaviors in social networks, such as information sharing, commenting, and reporting. Moreover, they can also simulate external factors such as social bots, constructing dynamic environments that better reflect real-world propagation patterns. By comprehensively modeling collective knowledge and social interactions, these dynamic simulations help address critical questions, such as: How can we generate more targeted intervention strategies with greater accuracy? How can we improve collaborative efficiency when users participate in rumor reporting or debunking efforts?

At the behavior layer, multi-agent systems can capture the nonlinear propagation paths of collective behavior and dynamically generate personalized debunking content and intervention strategies through cognitive modeling. For example, By analyzing a target audience’s cognitive and behavioral characteristics, agents can produce debunking content that is more persuasive and tailored to the audience. In response to the dynamic evolution of rumor propagation, agents can adaptively adjust intervention models, enabling more efficient and precise countermeasures.

Furthermore, in addition to modeling information dissemination of cognition to behavior, the framework enables a dynamic rumor detection and intervention mechanism in reverse. At the behavior layer, agents detect anomalous communities as initial targets. Subsequently, agents at the interaction layer analyze the suspicious users and interaction structures within these communities. Building on this, agents in the cognition layer perform collective knowledge analysis and social context reasoning on suspicious conversational threads to identify

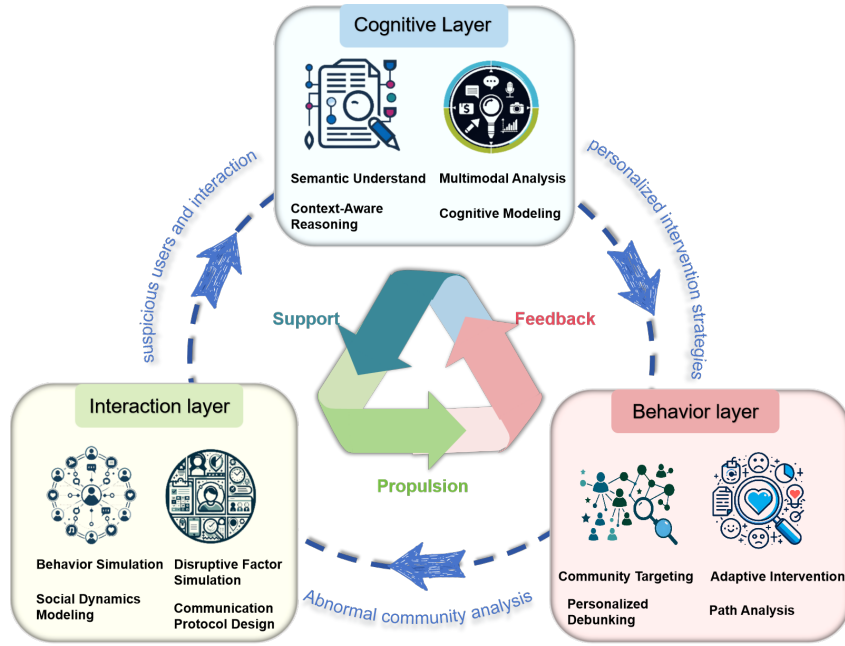


Figure 3: The CIB framework establishes a Macro-Micro Feedback Loop that integrates cross-layer dynamic feedback to bridge macro-level information dissemination with micro-level individual cognition. This enables bidirectional interventions, wherein macro-level propagation dynamics inform micro-level behavior modeling, while individual insights refine system-wide strategies. Beyond modeling the flow from cognition to behavior, the framework supports a reverse feedback mechanism for dynamic rumor detection and intervention, driving continuous adaptation and iterative improvement in complex, evolving misinformation environments.

critical evidence. Finally, the behavior layer intervenes promptly, generating personalized debunking content tailored to the cognitive characteristics of the target audience as part of belief-based interventions. This feedback mechanism allows rumor detection and intervention models to continuously improve and optimize themselves, enhancing their adaptability to the dynamically evolving nature of rumor propagation.

B.2 Cross-disciplinary Collaborative Optimization

LLMs also face significant challenges across three dimensions: content credibility, cognitive alignment, and technology adaptability (Liu et al., 2024a; De Angelis et al., 2023). First, the intertwining of hallucinated content generated by LLMs with malicious misinformation substantially complicates assessing content reliability (Pan et al., 2023c; Chen and Shu, 2024; Shu et al., 2021). Technologies such as deepfakes mislead the public and stigmatize genuine content, further eroding transparency in public discourse (e.g., real but negative content dismissed as synthetic and subsequently discredited (Schiff et al., 2023)). Second, limitations in LLMs’ semantic alignment and ability to

perform complex reasoning may reinforce users’ preexisting cognitive biases (Xu et al., 2023; Garry et al., 2024; Hosseini et al., 2023). In complex scenes, the robustness and dynamic adaptability of the model (Carlini et al., 2023; Haller et al., 2023), as well as early benchmarks (Zhou et al., 2023; Chen and Shu, 2023), show weak performance. (He et al., 2023; Li et al., 2024b). To address these challenges, future research must promote cross-disciplinary collaborative optimization:

Cognition layer: Enhancing content credibility and knowledge attribution. Future research should focus on building dual mechanisms that combine technological forensics and social validation to improve content credibility, particularly for attributing and explaining the trustworthiness of LLM-generated content (André et al., 2023; Kumarage and Liu, 2023). Another area of importance is addressing "technology-amplified cognitive biases" and their impact on information credibility perception and cognitive schema activation. For instance, The high fluency of LLM-generated content amplifies the halo effect (Augenstein et al., 2024); Multimodal synthetic content, such as deepfakes, enhances emotional infiltration, making misinformation harder to detect.

Interaction layer: strengthening cognitive alignment and belief intervention in social interactions. Integrating psychology and behavioral science insights can enable belief interventions using LLM-generated personalized and persuasive debunking content (Costello et al., 2024; Matz et al., 2024). For example, Employing dynamic intervention strategies, such as the "Friction Strategy," can suppress blind adherence and impulsive information sharing by increasing the cognitive processing cost of user decisions; LLMs could also provide real-time knowledge enhancement services for low-education user groups, helping mitigate the continued influence effect (Lewandowsky et al., 2012; Walter and Tukachinsky, 2020) (where misinformation persists even after being debunked) and the recognition gap driven by educational disparities (Afassinou, 2014; Hui et al., 2020). This approach fosters a synergistic effect between educational compensation and behavioral interventions.

Behavior layer: Enhancing robustness and adaptability in dynamic environments. To improve the adaptability of LLMs in dynamic environments, future research can leverage LLM-agent architectures that perform multi-phase reasoning chains (e.g., planning-execution-reflection), enabling zero-shot automatic feature annotation and latent rumor identification in real-time. Additionally, a more comprehensive rumor baseline assessment should be established to address the evolving characteristics of misinformation in complex environments.

By integrating beneficial insights from across disciplines into a rumor detection framework based on collective intelligence, this research line can promote the development of efficient and highly adaptable systems, pushing the boundaries of rumor detection methods.

B.3 Information Ecosystem Governance Under Multi-Multi-dimensionalraints

In the context of rumor governance, the synergistic governance of legal, ethical, and technological constraints emerges as a necessary approach.

Legal measures should focus on regulating data usage while ensuring privacy protection. Privacy-preserving technologies(such as differential privacy (Chua et al., 2024; Mai et al., 2023) and federated learning)(Kuang et al., 2024; Wu et al., 2024a; Ezzeldin et al., 2023)), combined with compliance frameworks(General Data Protection Regulation (GDPR) and the Artificial Intelligence Act (AI Act)

(Parliament, 2023)), enhance model performance while safeguarding data security. These measures serve as a foundation for responsible rumor detection and governance in the digital age.

LLMs have been shown to possess the capability of inferring psychological tendencies from user-generated texts (Peters and Matz, 2024; Perc et al., 2019), potentially influencing users' false memories (Chan et al., 2024; Acerbi and Stubbersfield, 2023). Furthermore, LLM-generated content could be weaponized for privacy infringements, cognitive attacks, and social media manipulation (Huang and Zhu, 2023; Park et al., 2024). To address these challenges, platforms, and developers should proactively disclose algorithm designs, ensure data sources and security measures, and establish transparent accountability chains to enhance transparency and responsibility allocation (Dwivedi et al., 2023).

At the social governance level, advancing multi-stakeholder collaborative mechanisms is essential. This involves building a governance ecosystem that includes developers, policymakers, and sociologists, aimed at enhancing the transparency and societal adaptability of LLM technologies and achieving a comprehensive balance between technological efficiency and societal impact (Hu et al., 2024; Tokayev, 2023), which ensures that the governance of false information can be effectively expanded in different social and technical environments.

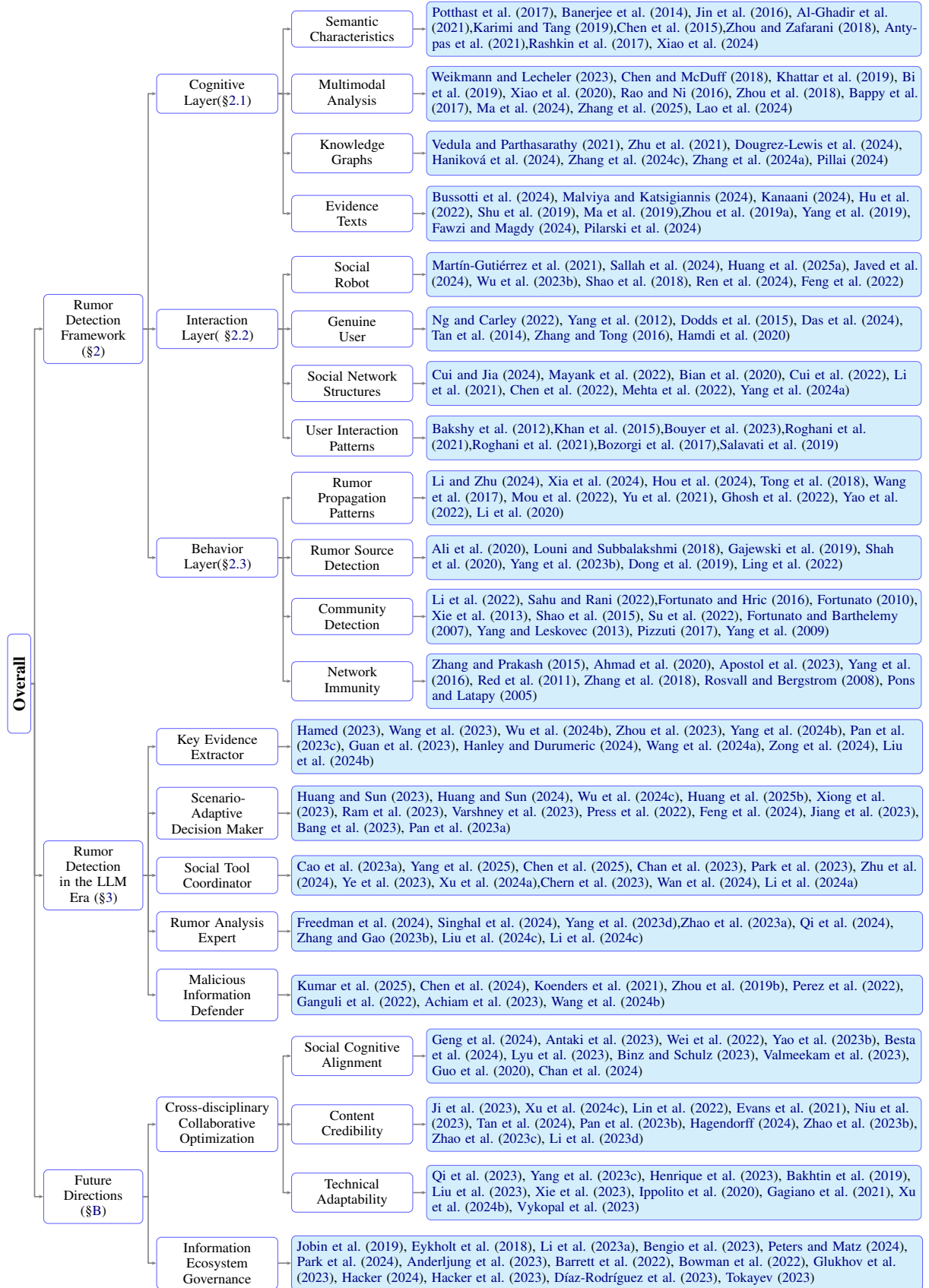


Figure 4: Classification of Rumor Detection Methods, Applications in the LLM Era, Challenges, and Future Directions under the CIB Framework