

AI AGENT IN SCI-TECH INTELLIGENCE ANALYSIS: CURRENT APPLICATIONS AND FUTURE TRENDS

Han Yang

School of Public Administration, Xiangtan University, Xiangtan 411105, Hunan, China.

Abstract: Objective: In the data-intelligent era, sci-tech intelligence analysis faces the dual challenges of information overload and insufficient depth in knowledge mining. Driven by large language models, AI Agents—as new autonomous actors capable of perception, planning, execution, and memory—are profoundly reshaping the intelligence workflow paradigm. This paper aims to systematically analyze the current applications, core capabilities, and development trends of AI Agents in sci-tech intelligence analysis, providing a reference for theoretical evolution and innovative practice in information science. Process: Using literature review and logical deduction, this paper first constructs an analytical framework encompassing four elements: "human-agent-information-technology." It then examines the core capabilities of AI Agents, from planning and memory to multi-agent collaboration. On this basis, it delves into their current applications in intelligence perception, organization, and generation. Finally, it discusses emerging trends such as human-AI collaboration paradigms and knowledge production transformation, along with potential challenges. Conclusion: AI Agents are driving the transformation of sci-tech intelligence analysis from "human-in-the-loop" to "human-on-the-loop," enabling intelligent and pipeline-based intelligence production processes. In the future, intelligence work will move toward a human-multi-agent collaborative paradigm, yet faces risks such as AI hallucinations and cognitive offloading, requiring urgent countermeasures at technical, literacy, and institutional levels.

Keywords: AI agent; Sci-tech intelligence analysis; Multi-agent collaboration; Intelligence paradigm

1 INTRODUCTION

In the data-intelligent era, global scientific and technological innovation has entered a period of vigorous activity, triggering profound changes in the sci-tech intelligence environment. International Data Corporation (IDC) projects that global data volume will reach 393.8 ZB by 2028[1]. The "data tsunami" formed by multi-source heterogeneous information—including scientific literature and other sources—has made the traditional intelligence workflow, centered on "manual retrieval-empirical judgment-linear analysis," increasingly inadequate for extracting high-value knowledge from information, resulting in a structural contradiction where "knowledge thirst" coexists with "information overload"[2]. This contradiction is driving a paradigm shift in sci-tech intelligence work.

Large language model (LLM)-powered AI Agents offer a technological pathway to address this predicament. AI Agents can autonomously perceive, plan, invoke tools, and execute complex operations[3]. Industry observers have designated 2025 as the "Year of AI Agents," with development progressing from content generation through tool-calling agents toward an explosive phase of multi-agent collaboration[4]. Meanwhile, the *Large Model Research Report (2026)* that the industry is moving toward a systematic reasoning era characterized by "model-architecture-scenario" synergy[5]. This technological progress marks a shift from AI as a passive responder to human commands to an active "task executor" that autonomously acts, continuously optimizes, and adapts to new environments—a development whose intelligence-enabling potential merits in-depth investigation.

How do AI Agents systematically reshape sci-tech intelligence analysis? What new trends and challenges do they engender? The traditional "human-information-technology" triadic framework treats technology merely as a tool. However, agents such as OpenClaw can autonomously search for information, analyze it, and generate intelligence products, thereby becoming independent actors in information behavior[3]. Therefore, this paper attempts to construct a four-element analytical perspective—"human-agent-information-technology"[3]—to more precisely dissect their intervention in and restructuring of the entire intelligence analysis process.

To address the research questions posed above, this paper first analyzes the three core modules of planning and memory, tool use, and multi-agent collaboration. Next, based on the "perception-organization-analysis-service" intelligence process chain, it systematically elaborates on current applications with case illustrations. It then draws on the latest 2026 reports and frontier literature to project development trends. Finally, it examines challenges at technical, organizational, and ethical levels and discusses coping strategies. This study aims to provide valuable reference for the development of library and information science in the intelligent era.

2 CORE CAPABILITIES: HOW DO AI AGENTS WORK?

An AI Agent is a composite system integrating large models, planning, memory, and tool-use capabilities. Its core working capabilities follow an evolutionary path from foundational large models to single agents and then to multi-agent collaboration[4]. Understanding these core capabilities is a prerequisite for comprehending their application in

intelligence.

2.1 Planning and Memory

The "brain" capabilities of AI Agents are primarily manifested in their planning and memory modules. The planning module endows AI Agents with the ability to structurally decompose complex intelligence problems and formulate optimal execution paths. Technically, planning capabilities mainly rely on the chain-of-thought (CoT) and its derivative techniques within large language models, guiding the model to produce multi-step continuous reasoning processes before arriving at conclusions[6-7]. Building on this, the ReAct framework further integrates reasoning and acting: the agent dynamically adjusts at each "think-act-observe" step. For major global decisions, the Plan-and-Execute framework generates a detailed execution plan, which is then strictly implemented by an "executor" agent[4], ensuring both comprehensive planning and coherent action. A more advanced planning capability is embodied in the introduction of the "Reflection" mechanism. After completing key steps, the agent can self-evaluate and critically examine its conclusions, promptly correcting any identified flaws or superior alternatives, thus forming a closed loop of "act-reflect-optimize." This enables the agent to iteratively improve without external supervision and is a key factor driving intelligence analysis from one-off deliverables toward continuous evolution[8].

The memory module endows the agent with "experience" and "knowledge." Short-term memory primarily maintains contextual coherence during current task execution, enabling the agent to accurately remember user requirements from several dialogue turns earlier, key findings from previous analysis steps, and its current stage. Long-term memory serves as the agent's "experience archive." Its core technological pathway is Retrieval-Augmented Generation (RAG), which essentially provides the large model with an externally and dynamically callable knowledge base. When executing tasks, the agent not only relies on knowledge internalized during model training but also retrieves the most relevant information fragments in real time from external databases or previously accumulated analysis cases, injecting them into the model's reasoning process to achieve self-evolution[8].

2.2 Tool Use

Tool use is a key feature distinguishing agents from simple language models[9]. In sci-tech intelligence analysis, agents can be authorized to invoke APIs of specialized academic databases, Python code interpreters for complex data computations, and other external tools to transcend the limitations of their own model training[10]. In early tool-use implementations, developers had to write dedicated code files to adapt to different agent frameworks, which constrained ecosystem development and application in the intelligence field. At the end of 2024, Anthropic introduced the Model Context Protocol (MCP)[11], establishing an open standard for connecting AI models with external data sources and tools. MCP also reduces the adaptation cost of tool development, enabling agents to flexibly invoke diverse tool resources. The *Cloud-Native Agents (Report)* indicates that cloud computing platforms are becoming the core carrier of the MCP protocol, providing elastic computing power and full lifecycle management for MCP Servers, and fostering tool marketplaces that accelerate ecosystem integration and scenario-based deployment[12]. This standardization process has far-reaching implications for building an open and shared intelligence analysis tool ecosystem.

2.3 Multi-Agent Collaboration

Given the limited cognitive depth and processing capacity of single AI Agents, Multi-Agent Systems (MAS) have emerged, marking a shift from "single-agent operations" to "legion collaboration" in AI applications[13]. In cognitively intensive intelligence analysis, MAS achieves discrete decomposition and specialized coordination of complex intelligence tasks through task decomposition and communication protocols. Specifically, complex tasks are first progressively decomposed by a manager agent, typically in hierarchical or hybrid organizational patterns[14], and then executed in parallel or sequentially by specialized worker agents. Throughout this process, communication and coordination mechanisms, such as global state graphs, help avoid information silos and redundant efforts[10,13]. Errors from a single agent can be cross-validated and promptly corrected by others, significantly improving the reliability and consistency of intelligence products[10].

3 CURRENT APPLICATIONS: HOW DO AI AGENTS EMPOWER INTELLIGENCE ANALYSIS?

The organic integration of AI Agents' core capabilities is driving a transformation of sci-tech intelligence analysis from linear processes to automated intelligence. From the perspective of information science, AI Agents deeply restructure the core links of traditional intelligence work. Specifically, intelligence perception shifts from passively responding to human commands to actively exploring user needs; knowledge organization transitions from manual linear processing to automated pipelines; and intelligence services evolve from "human-in-the-loop" to human-machine collaboration. Following the full intelligence analysis process of "perception-organization-analysis-service"[15], this section systematically explains the current applications of AI Agents in each core link, supported by case studies.

3.1 Intelligence Perception and Collection: From Passive Retrieval to Active Monitoring

In the intelligence perception and data collection stage, AI Agent applications enable a leap from "passive retrieval" to

"active" intelligent monitoring. Traditional intelligence collection relies heavily on analysts manually conducting keyword searches across various databases and websites, which is not only inefficient but also prone to missing cross-domain important clues due to cognitive inertia, making it difficult to detect latent technological mutations or unknown risks[16].

Currently, agent-based intelligent perception systems can continuously monitor and sense the broad information environment, fundamentally changing the *modus operandi* in intelligence analysis. First, at the demand perception level, they can autonomously perceive and interpret vague user needs, proactively clarifying and analyzing user intent through "reference interview agents" without waiting for precise human instructions[17]. Second, at the information retrieval level, they can invoke multi-source tools—including academic databases, patent databases, policy texts, and social media—for comprehensive scanning. More importantly, "multi-source intelligence management agents" can perform preliminary semantic correlation and cross-validation on perceived multimodal data, identifying anomalous information and automatically adjusting information collection strategies accordingly, thereby capturing weak signals of emerging technologies and enabling early warning of potential risks[18]. For example, a food security emergency intelligence system, powered by large-model agents, actively acquires early warning information from meteorological, logistics, and market entities, and combines it with historical data such as inventory and prices, as well as past experience, for comprehensive research and analysis, achieving risk pre-assessment[19]. This core capability of intelligence agents to autonomously perceive user needs and promptly monitor and analyze anomalous information significantly compresses the time for sensing intelligence risks, laying a foundation for subsequent intelligence analysis and decision-making.

3.2 Intelligence Organization and Analysis: From Manual to Automated

Intelligence organization and analysis constitute the core value-adding link of intelligence work, transforming raw data into knowledge and actionable intelligence. Under the traditional model, organization and analysis personnel must manually filter key entities from scientific literature, identify relationships between entities, construct corresponding knowledge graphs or roadmaps, and then rely on expert research to form final conclusions. This traditional process suffers from high costs, long cycles, heavy reliance on manual effort and experience, and difficulties in standardizing and reusing the resulting knowledge products.

The intervention of AI Agents reconstructs this process into a highly automated, orchestrated "production line." In this process, the planning, tool-use, and multi-agent collaboration capabilities of agents decompose each procedure into independent functional modules and connect them through task sequencing into an end-to-end automated pipeline[3]. First, "information organization agents" perform intelligent cleaning, entity extraction, and semantic analysis on multi-source heterogeneous data. For instance, a patent analysis system based on TRIZ theory and AI Agent collaboration can automatically extract technical contradiction information from patent texts in the hydrogen storage field[20], with multiple agents collaboratively completing contradiction identification, TRIZ engineering parameter mapping, and solution matching, ultimately constructing a knowledge base containing patent technical contradictions and solutions. This illustrates that agent applications not only greatly improve the efficiency and systematization of patent analysis but also achieve a value-added leap from shallow information extraction to innovative knowledge generation. Second, "intelligence analysis agents" undertake the core judgment work. They can leverage knowledge graphs for deep reasoning, decompose complex intelligence problems through chain-of-thought techniques, and simultaneously invoke multiple analytical tools—such as statistical analysis and social network analysis—for comprehensive collaborative judgment[1]. For example, when analyzing a complex strategic issue like the "China-US AI competition landscape," a multi-agent collaborative system can simulate in-depth discussion among an analysis team, with agents reaching consensus through multi-round dialogues of "proposing viewpoints-questioning-supplementing evidence-logical verification," ultimately forming systematic analytical conclusions[9].

3.3 Intelligence Generation and Service: From "Human-in-the-Loop" to "Human-on-the-Loop"

In the stage of intelligence product generation and service, AI Agent applications redefine human-machine roles, realizing a paradigm shift from "human-in-the-loop" to "human-on-the-loop." Under the traditional "human-in-the-loop" model, intelligence personnel serve as the sole drivers of the intelligence production process, personally participating in every link—from raw information filtering, through intelligence analysis, to final report editing. This results in intelligence products consuming substantial time, energy, and resources, while output quality is constrained by individual cognition and experience. CrowdFlower (2016) research shows that data scientists spend an average of 60% of their time on data cleaning and organization[21], leaving only 20% for core analytical activities. This statistic starkly reveals the structural imbalance in human resource allocation under the traditional model.

With the intervention of AI Agents, the new "human-on-the-loop" working model positions intelligence experts and agents as co-drivers of the intelligence production process, with experts elevated to "task definers" and "quality supervisors"[5,8]. Intelligence experts need only define task objectives, boundaries, and success criteria, design the logical framework for analysis, and intervene or review at key decision nodes; the remaining execution is delegated to a collaborative team of multiple agents. For example, Li Baiyang et al. built a multi-agent system for defense sci-tech intelligence research based on large models, deploying general expert modules, intelligence collection expert modules, intelligence analysis expert modules, and intelligence writing expert modules. In a simulated test task of "analyzing development trends of AI in military reconnaissance applications over the next five years," the system automatically generated an intelligence research

report containing multi-dimensional analyses—including technology breakthrough pathways and equipment intelligence trends—through collaborative reasoning and functional execution of multiple agents, and could output it in Word or Markdown format as required[13].

The *Artificial Intelligence Industry Development Research Report (2025)* points out that the effectiveness of large models in real application scenarios has become a key industry concern[22], and the transition from "technologically viable" to "service-ready" is achieved precisely through this new human-machine collaborative model. This model frees human energy from massive information processing, allowing focus on high-value intelligent decision-making and achieving genuine complementarity between human and machine strengths. In information science research, the establishment of the "human-on-the-loop" new paradigm enables human wisdom and effort to concentrate on higher-level strategic decisions and creative content work, while AI Agents undertake a large volume of automatable tasks. The complementary strengths of both constitute the core driving engine for the new generation of intelligence analysis and services.

4 DEVELOPMENT TRENDS

With the deepening application of AI Agent technologies, their impact on sci-tech intelligence analysis has transcended the technological tool level and is gradually evolving into a challenge to and reshaping of foundational information science theories. Drawing on the latest developments in 2026, the following trends are particularly salient.

4.1 Establishment of the "Human-Multi-Agent" Collaborative New Paradigm

"Human-multi-agent collaboration" is emerging as a new core paradigm for intelligence work. The essence of this trend is the restructuring of human-machine relationships in intelligence work—from a primary-secondary "human-led, AI-assisted" relationship toward a symbiotic partnership of "human-guided, agent-executed, collaborative decision-making"[23]. The latest *Report (2025)* also indicates that in 2025, "foundation super-models" such as GPT-5 and DeepSeek V3.2 have deeply integrated reasoning, coding, and agent core capabilities, with agents becoming a "factory default" feature of large models[22]. With this technological foundation, the "human-multi-agent" collaborative paradigm is further confirmed as the future direction of intelligence analysis. The report also explicitly states that MaaS platforms are transitioning from "general capability provision" to "scenario value closure," with the intrinsic driver being precisely this refined division of labor and collaboration between humans and agents. Under this trend, the systematic experience, value judgments, and creative thinking of human experts form a deep complementary and cognitive coupling with the massive information processing capabilities, continuous action capabilities, and multi-dimensional cross-validation capabilities of agent collectives. This combination not only improves the efficiency and quality of intelligence production but also redefines the core value and positioning of both humans and agents in intelligence analysis.

4.2 Expansion of Information Behavior Actors

Theoretically, the most profound impact of this trend lies in challenging the "information person" assumption—a core research object of information science. Classic theories such as Wilson's information behavior model and Kuhlthau's information search process model all presuppose humans as the sole active agents of information behavior[3]. Within these theoretical frameworks, information needs originate from human cognitive gaps, information collection is initiated by humans, and information acquisition depends on human cognitive levels—thus, in the intelligence domain, the intelligence production process is also highly dependent on the human subject. When AI Agents can autonomously perceive information needs, plan search strategies, and independently extract and analyze intelligence from massive data, they themselves become a new type of "information behavior actor." However, leveraging agents to complete tasks also raises concerns about cognitive offloading: when intelligence analysts no longer personally experience the intelligence production processes of searching and filtering, will their perception of information source characteristics and ability to discriminate information quality degrade? Research indicates a significant negative correlation between frequent AI tool use and critical thinking skills, with cognitive offloading playing a key mediating role[24]. The process of seeking information has intrinsic value independent of the outcome—it is not only the path to answers but also the process by which researchers construct problem awareness, form knowledge frameworks, and innovate thinking.

4.3 Transformation of Knowledge Production Models

AI Agents not only retrieve and integrate information and expand information behavior actors; they are also deeply participating in and even leading the knowledge production process. This signals a major transformation in knowledge production models. According to the *Artificial Intelligence Industry Development Research Report (2025)*[22], large models are transitioning from the "human data era" to the "experience era," enabling iterative learning through environmental interaction, task trial-and-error, and self-reflection. In the intelligence domain, this manifests as agent-orchestrated pipeline-based knowledge production, where the entire process—from hypothesis generation, experimental design, and data analysis to final conclusion formulation—can be highly automated. For example, the DeepResearch Skill built on the OpenClaw platform can automatically trigger upon recognizing a deep research need, connecting multiple specialized agent modules—including multi-source retrieval, literature review writing, and citation verification—into a complete knowledge production pipeline[3]. This means that in the production of a professional-grade sci-tech intelligence analysis report, the proportion of direct human contribution is gradually decreasing, while the proportion of

human-machine collaborative generation and verification is rising. When intelligence products increasingly manifest as hybrid human-machine outputs, how should intellectual contributions be defined? How should intelligence products generated by agents be quality-assessed and academically accredited? The Committee on Publication Ethics (COPE) has explicitly stated that AI tools cannot meet the fundamental requirements for authorship[25], but this norm primarily addresses scenarios where human authors use AI tools to assist creation, not the emerging scenario where agents autonomously complete the entire production process. This theoretical gap urgently requires systematic responses from the information science community.

In summary, the above three trends are not isolated but are interrelated. The human-agent collaborative paradigm redefines human-machine division of labor; the expansion of information behavior actors undermines the theoretical foundation that "humans are the sole actors"; and the transformation of knowledge production models blurs the boundaries between human and agent contributions to intelligence creation. All three indicate that information science needs to reconstruct its core theoretical system and research paradigms based on the four-element framework of "human-agent-information-technology."

5 CHALLENGES AND COPING STRATEGIES

Despite the promising prospects, large-scale application of AI Agents in sci-tech intelligence analysis is far from smooth. Their high autonomy brings not only leaps in efficiency and quality but also unprecedented challenges and risks, requiring us to maintain a clear-headed awareness and actively seek solutions.

5.1 Technical Reliability Challenges

The primary technical challenge is the "AI hallucination" problem. When reasoning about complex issues, the underlying large language model within an AI Agent may generate seemingly logically coherent but factually false information. Current research confirms that this problem remains unresolved. In the intelligence analysis pipeline, this issue is particularly fatal—a minor hallucination at the source can be repeatedly cited and amplified by downstream agents, leading to severe deviations in the final intelligence product. Moreover, the open interactive nature of AI Agents significantly expands their attack surface. Malicious actors can craft "prompt injection" attacks to induce agents to perform unauthorized operations, such as leaking sensitive data or deleting/modifying critical files in response to forged instructions. Even more lethal is "memory poisoning," which can implant false or malicious information into the agent's long-term memory, causing sustained cognitive contamination[4]. According to a Gartner survey, up to 74% of IT leaders interview AI Agents as a new enterprise security attack vector[4].

5.2 Organizational and Ethical Challenges

At the organizational application level, when core intelligence operations are delegated to agents, can intelligence workers still hone their information perception and critical thinking skills through hands-on practice? Once a path-dependent reliance on agents is formed, the last line of defense in intelligence work—human critical thinking and professional judgment—may be gradually eroded[3]. More importantly, the blurring of knowledge production subjectivity brings attribution dilemmas. When an intelligence product is collaboratively generated by multiple specialized agents and reviewed by a human expert, how should responsibility be reasonably allocated among AI developers, deploying institutions, and reviewers in the event of erroneous decisions? Existing academic ethics and legal liability frameworks appear powerless when facing non-human "agents" as knowledge producers[3]. Thus, throughout the entire AI-agent-based sci-tech intelligence analysis process, there exists a critical "accountability vacuum" that urgently needs to be filled.

5.3 Coping Strategies

To address the above challenges, a "three-pronged" strategy encompassing technical, literacy, and institutional dimensions should be constructed. First, at the technical level, AI-embedded trustworthiness mechanisms should be built, deploying fact-checking modules within the intelligence product production pipeline, using cross-validation among multiple modules to suppress AI hallucinations to some extent, while imposing minimum privilege restrictions on agents' file operations to reduce external attack surfaces. Second, at the personnel literacy level, the intelligence profession should strengthen the cultivation of higher-order critical thinking and professional decision-making capabilities among intelligence talents, ensuring they always play a vital role in intelligence analysis and preventing capability degradation due to over-reliance on agents. Finally, at the institutional level, clear accountability mechanisms should be established for AI applications in intelligence work. This includes developing standardized operating procedures, defining human responsibilities and accountabilities in the intelligence product generation process, mandating pauses and handover to human handling when significant risks may arise[26], and, for agents, establishing agent behavior log review processes to ensure that technological applications always operate within the tracks of the rule of law and ethics.

6 CONCLUSION

AI Agents, as a new quality productive force in the data-intelligent era, are embedding themselves into the entire process of sci-tech intelligence analysis with unprecedented breadth and depth. Fundamentally, AI Agents are not poised to replace

humans; rather, through automation and intelligence, they liberate intelligence workers from the quagmire of information overload, enabling them to focus on higher-level strategic thinking and value judgments, leveraging the synergistic work of humans and agents throughout the intelligence product production process. At the heart of this transformation lies the fundamental restructuring of human-machine relationships—moving toward the new paradigm of "human-multi-agent collaboration." Facing the development trends of AI Agents in intelligence analysis, the information science community should embrace modern technology with an active, open, and inclusive mindset, promptly reflect on and reshape its existing theoretical systems, while optimistically confronting the accompanying technical, organizational, and ethical challenges, seizing the new development opportunities of the era, and making good use of the important role of agents in intelligence analysis.

COMPETING INTERESTS

The authors have no relevant financial or non-financial interests to disclose.

REFERENCES

- [1] Li Guangjian, Wu Xuan, Pan Jiali. From General AI to Intelligence AI: Background, Framework, and Key Technologies. Library and Information Service, 2025, 69(20): 3-15. DOI: 10.13266/j.issn.0252-3116.2025.20.001.
- [2] Ren Huichao, Wang Xuefeng, Liu Yuqin. Practice of Intelligent Sci-Tech Intelligence Analysis System for Strategic Decision-Making. Information Studies: Theory & Application, 2022, 45(04): 27-34. DOI: 10.16353/j.cnki.1000-7490.2022.04.004.
- [3] Wang Shuyi, Xu Jie. The Impact and Implications of AI Agents Represented by OpenClaw on Information Science. Library & Information Knowledge, 2026: 1-11. <https://link.cnki.net/urlid/42.1085.G2.20260429.0932.002>.
- [4] China Academy of Industrial Internet. AI Agent Technology Development Report. 2026. <https://www.china-aii.com/u/cms/www/202601/AI%20Agent%E6%99%BA%E8%83%BD%E4%BD%93%E6%8A%80%E6%9C%AF%E5%8F%91%E5%B1%95%E6%8A%A5%E5%91%8A.pdf>.
- [5] China Academy of Information and Communications Technology (CAICT). Artificial Intelligence Industry Alliance. Research Report on Key Technologies and Application Practices of Large Model Inference Optimization. 2026-04-15. <https://www.caict.ac.cn/kxyj/qwfb/ztbg/202604/P020260415615308750699.pdf>.
- [6] Zhao Bang, Cao Shujin. Exploring Knowledge Mining with Generative AI Large Models Combined with Knowledge Bases and AI Agents. Library & Information Knowledge, 2025, 42(04): 88-101. DOI: 10.13366/j.dik.2025.04.088.
- [7] Li Hailiang, Lei Shuai, Zhao Xiang'an. Research on Intelligent Decomposition Approaches for Complex Sci-Tech Intelligence Problems in the Context of Large Models. Information Studies: Theory & Application, 2025, 48(11): 87-93. DOI: 10.16353/j.cnki.1000-7490.2025.11.011.
- [8] Geng Guotong, Lyu Yifei, Zhao Keran, et al. Analysis of a Grand Intelligence Research Model Based on the "System-Experiment Dual-Cycle". Information Studies: Theory & Application, 2025, 48(11): 7-12+19. DOI: 10.16353/j.cnki.1000-7490.2025.11.002.
- [9] Pu Yunqiang, Tang Chuan, Xu Jing, et al. Research on Multi-Agent Collaborative Systems for Sci-Tech Strategic Intelligence Analysis. Information Studies: Theory & Application, 2026, 49(03): 170-178. DOI: 10.16353/j.cnki.1000-7490.2026.03.019.
- [10] Chen Xi, Meng Xuyang. Deep Literature Retrieval Method Based on Multi-Agent Collaboration. Research on Library Science, 2026(03): 55-65. DOI: 10.15941/j.cnki.issn1001-0424.2026.03.010.
- [11] CAICT. Blue Paper on New Generation Intelligent Terminals: From "AI + Terminals" to AI Terminals. 2025. <https://www.caict.ac.cn/kxyj/qwfb/bps/202603/P020260316619476244269.pdf>.
- [12] Cloud Computing Open Source Industry Alliance. Cloud-Native Agents: AI Agent Technology and Application Research Report. 2025.
- [13] Li Baiyang, Zhang Xinyu, Zhao Yingxiao, et al. Construction and Application of a Multi-Agent System for Defense Sci-Tech Intelligence Research Based on Large Models. Information Studies: Theory & Application, 2025, 48(09): 178-184. DOI: 10.16353/j.cnki.1000-7490.2025.09.019.
- [14] Zhang Yidong. Ecological Intelligence Analysis Paradigm Based on Multi-Agent Collaboration. Information Studies: Theory & Application, 2025, 48(S1): 38-43.
- [15] Liu Xiwen, Sun Mengge, Fu Yun. From Management Information Systems to Intelligence Agents: Reshaping the Paradigm of Sci-Tech Intelligence Work. Library and Information Service, 2025, 69(17): 3-15. DOI: 10.13266/j.issn.0252-3116.2025.17.001.
- [16] Yao Changqing, Cheng Qikai, Wang Lijun, et al. Intelligent Intelligence Technology: Connotation, Boundaries, and System. Journal of the China Society for Scientific and Technical Information, 2025, 44(01): 1-9.
- [17] Ji Ting, Xu Lei, Zhou Gang. Research on the Reconstruction of Library Information Retrieval Service Models Driven by Generative AI. Library Development, 2025(06): 134-145. DOI: 10.19764/j.cnki.tsgjs.20251624.
- [18] Zhou Honglei, Zhang Haitao, Liu Yanhui, et al. Research on the Intelligence Response System for Major Emergencies Empowered by Large-Model Agents. Library & Information, 2025(01): 21-31.

- [19] Yin Yue, Zhang Haitao, Liu Yanhui, et al. Research on the Construction of a Food Security Emergency Intelligence System Based on Large-Model Agents. *Information Studies: Theory & Application*, 2025, 48(01): 20-30. DOI: 10.16353/j.cnki.1000-7490.2025.01.003.
- [20] Zhai Dongsheng, Du Ruize, Qu Qianwei. Method for Constructing Domain Knowledge Bases Based on TRIZ-AI Agent. *Journal of Intelligence*, 2026, 45(04): 158-167.
- [21] Liu Xiwen, Fu Yun, Wei Huanan. Intelligence Agents – A New Paradigm for Sci-Tech Intelligence Work Toward the "15th Five-Year Plan". *Journal of Agricultural Library and Information Science*, 2024, 36(12): 20-34. DOI: 10.13998/j.cnki.issn1002-1248.24-0666.
- [22] CAICT. Artificial Intelligence Industry Development Research Report. 2025. <https://www.caict.ac.cn/kxyj/qwfb/bps/202602/P020260202487301304903.pdf>.
- [23] Zhou Peng, Zhang Ji, Wang Shuoyu, et al. Research on a Dual-Subject Knowledge Assimilation Model Driven by Human-AI Interaction. *Library and Information Service*, 2026: 1-13. <https://link.cnki.net/urlid/11.1541.G2.20260428.1614.002>.
- [24] Gerlich M. AI Tools in Society: Impacts on Cognitive Offloading and the Future of Critical Thinking. *Societies*, 2025, 15(1): 6.
- [25] COPE Council. COPE Position – Authorship and AI – English. 2026-05-15. DOI: 10.24318/cCVRZBms.
- [26] Jiang Ning, Feng Wengang, Yuan Lining. Constructing a Theoretical Model for AI-Empowered National Security Intelligence Decision-Making. *Journal of Intelligence*, 2026, 45(02): 65-72.