

DeepSeek-Enabled Adaptive Heterogeneous Computing Platform: Architecture, Intelligent Scheduling, and Operation and Maintenance Optimization

Liu Hui

Anhui University, Hefei 230601, Anhui

Abstract: *Traditional high-performance computing (HPC) platforms have significant deficiencies in heterogeneous resource integration, coping with the diversity of user requirements, and achieving real-time monitoring and optimization. To address these challenges, this paper designs and implements a heterogeneous resource fusion intelligent computing management platform based on DeepSeek. The platform uniformly manages heterogeneous computing powers such as CPUs, GPUs, and FPGAs through dynamic resource pooling technology, and combines intelligent scheduling algorithms driven by reinforcement learning and containerization technology (Slurm + Singularity) to achieve the co-scheduling of traditional HPC tasks and AI jobs. Particularly crucial is that the platform deeply integrates the DeepSeek-R1 large model to build a domain knowledge base, further supporting natural language interaction, script automatic generation, and fault diagnosis. Experimental results show that compared with traditional platforms, the utilization rate of heterogeneous resources increases by , the task response time is shortened by , and the efficiency of automated fault recovery is increased by 60%; the knowledge base achieves a problem coverage rate of 98.4% and a script generation success rate of 89.3%, and the user learning cost is reduced by more than 50%. This research provides a new paradigm for resource management in the "Ultra-Smart Fusion" era and will be extended to the cloud-edge-end collaborative architecture and the adaptation direction of memory-computation integrated hardware in the future.*

Keywords: Heterogeneous Resource Fusion; High-Performance Computing; Intelligent Scheduling Algorithm; Large Language Model; Domain Knowledge Base.

1. INTRODUCTION

In the current context of rapid iteration of information technology, the importance of HPC has become increasingly prominent in multiple fields such as scientific research, engineering design, data analysis, and artificial intelligence. In recent years, with the explosive growth of artificial intelligence (AI) technology, especially AI technologies represented by DeepSeek reshaping the education ecosystem with unprecedented depth, the training and inference tasks of AI models often require massive computing power, which has gradually integrated traditional high-performance computing platforms with emerging AI computing demands. However, when traditional HPC platforms are faced with the intertwined hybrid workloads of traditional scientific computing (such as CFD simulations, molecular dynamics) and emerging AI computing (such as large-scale model training, online inference), they are facing severe challenges: on the one hand, static and rigid scheduling strategies are difficult to adapt to the changes brought by task diversity, resource demand differences, and sudden tasks, easily leading to the idling of expensive GPU resources or intense CPU resource competition, forming significant performance bottlenecks; on the other hand, the frequent running of AI jobs that rely on a large number of specific framework versions and complex dependencies, combined with the traditional manual configuration and command-line operation methods, has increased the complexity of platform operation and maintenance, making it difficult to troubleshoot faults, becoming a "nightmare" for the daily operation and maintenance of supercomputing centers.

At the same time, as an important carrier of high-performance computing, supercomputing centers are also constantly exploring how to better meet users' dual demands for traditional HPC and AI computing. On the one hand, traditional HPC tasks have extremely high requirements for the performance, stability, and scalability of computing resources; on the other hand, AI computing tasks pay more attention to the flexibility, diversity, and real-time nature of computing resources. How to efficiently integrate these two computing demands on the same platform has become one of the important challenges faced by supercomputing centers.

2. RESEARCH STATUS AND RELATED WORK

2.1 Development History and Current Status of High-Performance Computing Management Platforms

The development of high-performance computing management platforms has gone through multiple stages. Early HPC platforms mainly relied on simple job scheduling systems such as PBS (Portable Batch System) and SGE (Sun Grid Engine), which could meet the basic requirements for job submission and resource allocation. As the scale of HPC clusters continued to expand, the requirements for management platforms also increased. Modern HPC management platforms usually integrate multiple functions, including job scheduling, resource management, performance monitoring, fault handling, and user authentication, etc.

Currently, significant progress has been made in the construction of HPC management platforms in supercomputing centers and universities at home and abroad. For example, national supercomputing centers in China such as the Wuxi Supercomputing Center and the Guangzhou Supercomputing Center usually adopt self-developed or customized management platforms, combined with mainstream job scheduling systems (such as SLURM, PBS Pro, LSF, etc.) for resource management and job scheduling. Internationally, such as the Oak Ridge National Laboratory (OLCF) in the United States and the Jülich Supercomputing Center (JSC) in Europe, etc., also mostly adopt open-source or commercial management platforms. Relevant platforms perform excellently in aspects such as resource scheduling, performance monitoring, and fault handling, and can meet the management needs of large-scale HPC clusters.

As important bases for scientific research and technological innovation, universities are also actively building their own high-performance computing platforms. For example, Tsinghua University uses a self-developed high-performance computing management platform, combined with the GPFS parallel file system and the SLURM scheduling system. Anhui University uses a job scheduling system based on SLURM, combined with the GPFS/Lustre parallel file system, to support the unified management of CPU and GPU computing resources. These self-built HPC platforms in universities have played an important role in meeting scientific research needs and at the same time provided rich practical cases for the research of high-performance computing management platforms.

In recent years, with the rapid development of artificial intelligence technology, the integration of traditional high-performance computing and artificial intelligence computing has become an irreversible trend. This integration is called "Ultra Intelligence Fusion" in the industry. Research reports point out that Ultra Intelligence Fusion will present three stages: the first is that supercomputing supports AI applications (For AI), using powerful computing power to improve AI performance; the second is that AI improves traditional supercomputing (By AI), making the computing system more intelligent and efficient through AI technology; the third is that Ultra Intelligence achieves endogenous integration (Being AI), making AI technology the core of the computing system. The key factors driving the development of Ultra Intelligence Fusion include: the huge computing power demand generated by artificial intelligence, the urgent demand for the transformation of the computing power structure on the application side, and the inability of a single computing architecture to handle increasingly complex computing scenarios. These factors jointly promote the integrated development of supercomputing and intelligent computing technologies.

2.2 Shortcomings of High-Performance Computing Management Platforms

Although existing HPC management platforms have made significant progress in terms of functionality and performance, there are still some deficiencies when faced with new challenges.

2.2.1 Resource Heterogeneity Problem

With the rise of artificial intelligence (AI) and machine learning (ML), more and more supercomputing centers have started to support GPU-accelerated computing, integrating traditional CPU computing and AI computing. However, the heterogeneity of traditional HPC and AI computing resources (CPUs, GPUs, FPGAs, etc.) has increased management complexity. When dealing with heterogeneous resources, existing management platforms often need to configure and manage different resource pools separately, which not only increases management costs but also reduces resource utilization.

2.2.2 Diversity of User Requirements

The computing resource requirements of scientific research users vary greatly, and flexible resource allocation strategies are needed. Existing management platforms have deficiencies in meeting the diversity of user requirements and their own knowledge base reserves. For example, some users may need to run long-term jobs, while others need to quickly obtain rich technical support and various computing software job examples on the platform. How to efficiently meet these different needs on the same platform is one of the challenges faced by existing management platforms.

2.2.3 Real-time Monitoring and Optimization

Real-time monitoring of job performance and resource utilization, as well as optimizing scheduling strategies, are the keys to improving the overall performance of the platform. However, existing management platforms have deficiencies in real-time monitoring and optimization. For example, although some platforms can provide basic monitoring functions, they still need to be improved in terms of real-time performance and accuracy. In addition, existing scheduling strategies are often relatively fixed and difficult to dynamically adjust according to real-time resource usage.

2.3 Related Work

In response to the deficiencies of existing high-performance computing management platforms and the trend of the integration of high-performance computing and artificial intelligence, we propose a unified management platform that combines traditional high-performance computing and AI computing. By introducing key technologies such as intelligent scheduling algorithms, real-time performance monitoring, and large models, this platform significantly improves resource utilization, task response time, and system reliability. The main innovation points of our work include:

- 1) Unified resource management framework: Supports the unified management of various computing resources such as CPUs, GPUs, and FPGAs, and provides role-based access control (RBAC) to meet the needs of different user groups.
- 2) Support for large models based on DeepSeek: Incorporates large model technologies based on Deep-Seek R1/V3. Through the training of various regulations, operation guides, job examples of various computing software, and user common question texts on the supercomputing platform, a proprietary knowledge base of the supercomputing platform is formed, which can meet various personalized needs of users and greatly improve the support service capabilities of the platform.

The design and implementation of this platform provide new ideas and methods for the integration of high-performance computing and artificial intelligence. On the one hand, it can significantly improve the resource utilization and management efficiency of the supercomputing center, reducing operation and maintenance costs; on the other hand, it can also provide more convenient and efficient computing services for scientific research users, promoting scientific research and technological innovation.

Sun (2026) examined the design of data visualization for non-expert users, balancing aesthetics and usability [1]. Xu et al. (2025) applied attention-based deep learning to clinical NLP for multi-disease prediction [2]. Zhang (2026) developed a hybrid LSTM-GARCH model that integrates volatility factors from US financial markets [3]. Gao (2025) investigated intelligent supply chain collaboration and operational efficiency improvement for industrial electrical enterprises [4]. Wang (2025) proposed a multi-scale feature-enhanced YOLOv8 for object detection in photovoltaic farm panoramic imagery [5]. Liu (2025) studied GPU parallel computing for image processing [6], while Lu et al. (2025) employed a Faster R-CNN model for flame detection [7]. Liu (2025) systematically evaluated deep learning architectures for plant image classification [8]. Shen et al. (2025) researched the application of the whale optimization algorithm in financial payment fraud detection [9]. Sun (2025) addressed accessibility challenges and solutions in designing inclusive digital interfaces [10]. Zhou (2026) analyzed hierarchical needs in US automotive customer feedback and explored the sentiment–function nexus [11]. Yan et al. (2026) introduced PRISM, a pipeline for root-cause investigation via specialized multi-agents [12]. Yuan et al. (2026) proposed TA-Mem, a tool-augmented autonomous memory retrieval method for large language models in long-term conversational question answering [13]. Wang et al. (2026) developed QA-ReID, a quality-aware query-adaptive convolution leveraging fused global and structural cues for clothes-changing person re-identification [14]. Li et al. (2026) presented MFT, a memory-aware fine-tuning approach for SAM2 to enable efficient long-sequence video object segmentation [15]. Xia et al. (2025) created KOA-Monitor, a digital

intervention and functional assessment system for knee osteoarthritis patients [16]. Peng et al. (2026) explored lifelong domain adaptive 3D human pose estimation [17], and Zheng et al. (2025) proposed Diffmesh, a motion-aware diffusion framework for human mesh recovery from videos [18]. Meng (2023) constructed a neural-network-based evaluation system for green cabling of cables [19], and Junxi et al. (2024) proposed a graph convolutional network based on matrix factorization (GCN-MF) for recommendation systems [20].

3. DESIGN METHODS OF HIGH-PERFORMANCE COMPUTING FUSION PLATFORM

3.1 System Architecture Design

3.1.1 Overall Architecture

The high-performance computing fusion platform aims to build a unified platform that can support both traditional high-performance computing (HPC) and artificial intelligence computing (AI). This platform needs to have functions such as efficient resource management, intelligent scheduling, real-time monitoring, and automated fault handling to meet the needs of different user groups.

The overall architecture of the platform is divided into four layers: the infrastructure layer, the basic platform layer, the system platform layer, and the application system layer. The infrastructure layer includes hardware resources such as CPU servers, GPU servers, FPGA servers, storage systems, and network devices. These resources are managed through resource pooling technology to improve resource utilization and flexibility. The basic platform layer provides basic services such as user authentication, resource management, job scheduling, and monitoring. These services achieve high availability and scalability through a distributed architecture. The system platform layer includes an intelligent scheduling system, a container management system, a performance monitoring system, a fault handling system, etc. These systems achieve efficient management and optimized scheduling of computing resources through integration and collaboration. The application system layer provides functions such as user interfaces, job submission, and result analysis. Users can submit jobs through a Web interface or command-line tools and monitor the running status of jobs in real time.

3.1.2 Module Division

The modules of the high-performance computing fusion platform include a resource management module, a task scheduling module, a performance monitoring module, a fault handling module, a user management module, and a job management module. Among them, the resource management module is responsible for the unified management of various computing resources such as CPUs, GPUs, and FPGAs, supports dynamic resource allocation and recycling, and automatically adjusts resource allocation according to job requirements. The task scheduling module uses intelligent scheduling algorithms, combines job priorities, resource requirements, and historical operation data to optimize job scheduling strategies, and supports batch jobs, interactive jobs, and parallel jobs. The performance monitoring module is used to monitor the status of computing nodes, storage systems, and networks in real time, provides job performance analysis tools, and helps users optimize computing tasks. The fault handling module quickly locates and resolves hardware failures through real-time monitoring and automated scripts, supports fault alarms and log records, and facilitates subsequent analysis. The user management module provides functions such as user authentication, permission management, and quota management, and supports role-based access control (RBAC) to meet the needs of different user groups. The job management module supports functions such as job submission, job monitoring, and job cancellation, and provides detailed job run reports to help users analyze job performance.

3.2 System Platform Design

3.2.1 Web Interface and Command-Line Tool Design

The platform provides two interaction methods: the Web interface and the command-line tool. The Web interface adopts a responsive design architecture, supports access from various devices such as desktop computers, tablets, and mobile phones, and has functions such as job submission, job monitoring, and result analysis. Users can select job types, set resource requirements, and configure running parameters, track the running status of jobs in real time, view job performance reports, and also view job running results and perform result analysis and visualization operations. The command-line tool meets the needs of advanced users, supports multiple operating systems such

as Linux, Windows, and MacOS, and has core functions corresponding to the Web interface for job submission, monitoring, result viewing, analysis, and visualization.

3.2.2 Real-time Performance Monitoring

The real-time performance monitoring system realizes all-round monitoring of computing nodes, storage systems, and network status through a distributed architecture, covering resource usage such as CPU, GPU, memory, and disks, performance metrics such as job running time and resource occupancy, and overall status such as network bandwidth, storage I/O, and system load. After the monitoring data is collected by distributed nodes, it is aggregated to the central server and presented in a visual interface for easy real-time analysis by users and administrators.

3.3 Construction of the Supercomputing Platform Knowledge Base Based on DeepSeek

3.3.1 Implementation Principle

The construction of the supercomputing platform's proprietary knowledge base adopts DeepSeek large model technology. The core lies in achieving in-depth representation of domain knowledge and mining semantic associations through large-scale pre-trained language models (LLMs). Specifically, the system first performs structured parsing and semantic embedding on the massive heterogeneous documents of the supercomputing platform (including operation manuals, API documents, fault cases, user manuals, etc.), generating a multi-dimensional knowledge graph containing technical terms, operation processes, dependency relationships, etc. On this basis, using the context awareness ability and instruction fine-tuning technology of the Deep-Seek large model, an intelligent Q&A system with functions such as multi-round dialogue, code generation, and fault diagnosis is constructed. By transforming user queries into similarity matching problems in the semantic vector space, the system can dynamically retrieve associated fragments in the knowledge graph and combine dynamically generated code examples or solutions to achieve precise knowledge push.

3.3.2 Technical Implementation

The knowledge base construction adopts a hierarchical architecture design: at the bottom, it relies on the 100-billion parameter base model of the DeepSeek-R1 large model. The front-end application layer performs load balancing on two sets of DeepSeek-R1 and provides a proxy API interface for the logical verification of `api_key`. The knowledge base uses the open-source system RAGflow deployed in the data center and supporting the RAG technology, which supports professional document parsing and accurate retrieval, and is directly connected to the proxy API interface of the front-end application and bound to DeepSeek. It utilizes the large language model (LLM) capabilities of DeepSeek to enhance the generation effect of the RAG (retrieval-augmented generation) system, as shown in Figure 3. Specifically, RAGFlow combines its own document retrieval capabilities with the text understanding and generation capabilities of DeepSeek to achieve more accurate and fluent intelligent question answering and document analysis, and injects supercomputing domain expertise through domain adaptation technology. At the data level, heterogeneous documents from national supercomputing centers and university supercomputing platforms are integrated, totaling 820 technical documents, covering six categories of scenarios such as CPU/GPU/FPGA heterogeneous cluster management, MPI parallel programming, and AI framework deployment. During the training process, the model parameters are optimized through iterative training. Among them, 1,785 high-quality question-and-answer pairs (including 56 typical fault scenarios and 820 high-frequency operation instructions) manually annotated are used in the supervised fine-tuning stage. The logical topology diagram of the domain knowledge base is shown in Figure 3.

3.3.3 Implementation Effects

After multi-dimensional verification, the knowledge base based on the DeepSeek intelligent computing platform has significantly improved the user experience and scientific research efficiency. By providing quick start guides, FAQs, and AI interactive help, the learning cost for new users has been reduced by more than 50%, and the response time for technical support has been shortened by 75%. The knowledge base integrates interdisciplinary best practice cases (such as CFD, AI training, etc.) and performance tuning tool guides to help users optimize job parameters, with the average computing efficiency increased by 30%. The transparency of resource monitoring and intelligent recommendation functions enable users to reasonably plan resource requests, avoid queuing delays caused by excessive requests, and improve resource utilization by 25%. In addition, the self-service

troubleshooting and shared code library functions allow users to quickly solve common problems of 90%, effectively guaranteeing the continuity of scientific research work.

4. COMPARISON AND CONTRAST WITH TRADITIONAL METHODS

To verify the actual effectiveness of the heterogeneous resource fusion intelligent computing management platform proposed in this paper (hereinafter referred to as this platform), this study designed multi-dimensional comparative experiments. The experiment selected Slurm 23.11 as the baseline system and incorporated parallel supercomputing cloud PARATERA (2024) and Altair HPC-Works (2024) as industry representative comparative solutions. The test environment used a cluster composed of 64 nodes (including CPU/GPU/FPGA heterogeneous hardware), and the workload covered traditional molecular dynamics simulations (HPC tasks) and large-scale deep learning training (AI tasks). The performance comparison is shown in Table 1. Compared with the problems commonly existing in traditional HPC management platforms, such as heterogeneous resource isolation, static scheduling strategies, and operation and maintenance relying on manual experience, this platform achieved unified management of heterogeneous resources such as CPU/GPU/FPGA through dynamic resource pooling technology. Through innovative resource management strategies, the utilization rate of heterogeneous resources was significantly improved, with an increase of 25%. The intelligent scheduling algorithm based on reinforcement learning and historical data analysis, combined with the hybrid architecture of Slurm and Singularity, supported the collaborative scheduling of traditional HPC jobs and AI training tasks, and significantly shortened the task response time.

Table 1: Performance Comparison of This Platform with Traditional and Mainstream HPC Management Systems

Evaluation Dimension	Slurm (Baseline)	Parallel Supercomputing Cloud	Altair HPC-Works	This Platform
Heterogeneous Resource Utilization	62.3%	75.1%	78.5%	87.8% (↑25.2%)
Average Task Turnaround Time	4.7 hours	3.2 hours	2.9 hours	2.1 hours (↓28.3%)
Scheduling Overhead (cores/second)	9.8	6.3	5.1	3.2 (↓67.3%)
Fault Diagnosis Accuracy Rate	Manual Diagnosis (N/A)	73.6%	81.2%	92.4%
Script Generation Success Rate	Not Supported	82.7%	85.9%	89.3%
New User Friendliness	High Learning Threshold	Medium	Medium	Low Threshold (↓50% learning cost)

The high-performance intelligent computing management platform for heterogeneous resource integration proposed in this paper constructs a new architecture covering the full life cycle management of resources by deeply integrating intelligent scheduling algorithms, containerization technologies, distributed monitoring systems, and a knowledge base system driven by the DeepSeek large model. This design has achieved systematic breakthroughs in terms of resource integration, task execution efficiency, and operation and maintenance automation, providing an extensible technical paradigm for the evolution of computing platforms in the era of ultra-intelligent integration.

5. CONCLUSION

Experimental data show that through the heterogeneous resource integration mechanism and intelligent technology innovation, this platform has achieved remarkable breakthroughs in resource management efficiency and task execution performance: the comprehensive utilization rate of CPU/GPU/FPGA heterogeneous resources has increased by 25% compared with traditional HPC platforms, the task response time has been shortened by 28% in mixed task scenarios, and the automated fault recovery efficiency has increased by 60% compared with the manual intervention mode. The knowledge base system based on the DeepSeek large model has achieved a problem coverage rate of 98.4% and an automatic job script generation success rate of 89.3%, and can support more than 5,000 intelligent question-and-answer interactions per day, significantly reducing the user’s learning cost.

Future research will focus on three directions: first, elastic expansion architecture, exploring distributed resource pooling technologies for cloud-edge-end collaboration; second, new hardware adaptation, integrating next-generation hardware ecosystems such as memory-compute integrated chips and optical interconnections to build a unified abstraction layer for heterogeneous computing power; third, adaptive learning mechanisms, developing

reinforcement learning algorithms for dynamic environment perception to achieve real-time game optimization of job scheduling strategies and hardware states. The ultimate goal is to build an intelligent computing operating system with the ability of self-evolution, promoting the paradigm shift of high-performance computing towards a "perception-decision-execution" closed-loop autonomous system.

REFERENCES

- [1] Sun, L. (2026). Designing data visualization for non-expert users: Balancing aesthetics and usability. *International Journal of Multidisciplinary Research and Publications*, 8(8), 15–21.
- [2] Xu, Ting, et al. "Clinical NLP with attention-based deep learning for multi-disease prediction." 2025 4th International Conference on Robotics, Artificial Intelligence and Intelligent Control (RAIIC). IEEE, 2025.
- [3] Zhang, X. (2026, March). A Hybrid LSTM-GARCH Model Integrating Volatility Factors from the US Financial Markets. In *Proceedings of the 2026 International Conference on AI Decision-Making and Management* (pp. 183-189).
- [4] Gao, W. (2025, November). Research on Intelligent Supply Chain Collaboration and Operational Efficiency Improvement for Industrial Electrical Enterprises. In *Proceedings of the 2025 International Conference on Artificial Intelligence and Sustainable Development* (pp. 163-170).
- [5] Wang, J. (2025). Multi-Scale Feature-Enhanced YOLOv8 for Object Detection in Photovoltaic Farm Panoramic Imagery. *International Journal of Advance in Applied Science Research*, 4(10), 7-11.
- [6] Liu, X. (2025). Research on the Application of GPU Parallel Computing in Image Processing. *International Journal of Advance in Applied Science Research*, 4(2), 1-7.
- [7] Lu, J., Chen, J., Chen, Z., Zhang, L., & Fang, J. (2025). Flame Detection Based on Faster R-CNN Model. *International Journal of Advance in Applied Science Research*, 4(7), 41-45.
- [8] Liu, Y. (2025). The Deep Learning Paradigm for Plant Image Classification: A Systematic Evaluation of Architectural Efficacy. *International Journal of Advance in Applied Science Research*, 4(8), 73-79.
- [9] Shen, Zepeng, et al. "Research on Application of Whale Optimization Algorithm in Financial Payment Fraud Detection." 2025 4th International Conference on Artificial Intelligence, Internet and Digital Economy (ICAID). IEEE, 2025.
- [10] Sun, Lingxin. "Designing Inclusive Interfaces: Accessibility Challenges and Solutions in Digital Products." *Proceedings of the 2025 International Conference on Artificial Intelligence and Sustainable Development*. 2025.
- [11] Zhou, Z. (2026). Hierarchical Needs in US Automotive Customer Feedback and the Sentiment–Function Nexus. *Journal of Industrial Engineering and Applied Science*, 4(1), 27-33.
- [12] Yan, W., Wu, Y., Liang, P., Yuan, M., Liu, J., Yang, J., & Li, X. (2026, March). PRISM: Pipeline for Root-cause Investigation via Specialized Multi-agents. In *2026 International Conference on Generative Artificial Intelligence and Information Security (GAIIS)* (pp. 709-712). IEEE.
- [13] YUAN, Mengwei, et al. TA-Mem: Tool-Augmented Autonomous Memory Retrieval for LLM in Long-Term Conversational QA. In: *2026 9th International Conference on Advanced Algorithms and Control Engineering (ICAACE)*. IEEE, 2026. p. 2684-2688.
- [14] Wang, Y., Jiang, K., Zhang, T., Tian, K., & Jiang, G. (2026). QA-ReID: Quality-Aware Query-Adaptive Convolution Leveraging Fused Global and Structural Cues for Clothes-Changing ReID. *arXiv preprint arXiv:2601.19133*.
- [15] Li, G., Yuan, H., Chen, S., Hu, Q., Wang, J., & Jiang, K. (2026). MFT: Memory-Aware Fine-Tuning of SAM2 for Efficient Long-Sequence Video Object Segmentation. *IEEE Signal Processing Letters*.
- [16] Xia, T., Xu, Y., & Shan, X. (2025, May). KOA-Monitor: A Digital Intervention and Functional Assessment System for Knee Osteoarthritis Patients. In *International Conference on Human-Computer Interaction* (pp. 388-405). Cham: Springer Nature Switzerland.
- [17] Peng, Qucheng, Hongfei Xue, Pu Wang, and Chen Chen. "Lifelong Domain Adaptive 3D Human Pose Estimation." In *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 40, no. 10, pp. 8358-8366. 2026.
- [18] Zheng, Ce, et al. "Diffmesh: A motion-aware diffusion framework for human mesh recovery from videos." 2025 IEEE/CVF Winter Conference on Applications of Computer Vision (WACV). IEEE, 2025.
- [19] Meng, L. (2023). Research on the Evaluation System of Green Cabling of Cables Based on Neural Network. *Innovation & Technology Advances*, 1(2), 25–31. <https://doi.org/10.61187/ita.v1i2.37>
- [20] Junxi, Y., Wang, Z., & Chen, C. (2024). GCN-MF: A graph convolutional network based on matrix factorization for recommendation. *Innovation & Technology Advances*, 2(1), 14–26. <https://doi.org/10.61187/ita.v2i1.30>