

Joint Multi-Dimensional Resource Scheduling and Stability Optimization for Computation-Constrained High-Concurrency Platforms

Alejandro Rodríguez García^{1,*}

University of Barcelona, Spain

* Corresponding author: Alejandro Rodríguez García
AlejandroRodriguez@outlook.com

Abstract: Managing multi-dimensional resource flows within large-scale, high-concurrency service platforms under severe computational capacity constraints presents intricate architectural challenges, where traditional single-dimensional scheduling models fail due to cross-resource coupling. This paper presents a joint multi-dimensional resource flow scheduling framework integrated with control-theoretic stability optimization, specifically addressing the non-linear coupling effects across CPU, memory. Recognizing that unexpected traffic surges derail static optimization paradigms, we introduce an adaptive scheduling algorithm leveraging dynamic Lyapunov optimization alongside a non-linear admission control policy. While simulation experiments across heterogeneous workloads demonstrate substantial improvements in throughput and queue stability, anomalies under extreme scenarios suggest that minor system biases might trigger localized bottlenecks, hinting at latent trade-offs between instantaneous responsiveness and long-term structural equilibrium. Considering these fluctuating factors, our theoretical boundary proofs contribute a rigorous foundation for modern cloud-native architectures, though further research is needed to fully reconcile state-space explosions in ultra-large-scale deployments.

Keywords: Multi-dimensional resource flow; High-concurrency platforms; Computational capacity constraints; Adaptive scheduling; Lyapunov stability;

1. Introduction

The contemporary landscape of distributed computing is increasingly defined by the unprecedented scale and structural complexity of internet-scale high-concurrency service platforms. As digital ecosystems evolve, these platforms are routinely subjected to volatile traffic influxes that strain the underlying infrastructure, rendering the traditional paradigm of over-provisioning resources economically non-viable. In this context, the physical boundaries of computational capacity, often exacerbated by the slowing of Moore's Law and the thermal limits of dense silicon deployments, have shifted the focus of system architects from unconstrained scale-out strategies to highly optimized intra-node and inter-node resource orchestration. The fundamental challenge resides not merely in managing raw compute cycles, but in resolving the complex, non-linear dependencies among multi-dimensional resource flows, including central processing unit allocation, memory bandwidth, network input-output queues, and storage throughput.

Earlier paradigms in distributed systems management frequently treated these resource dimensions as isolated, orthogonal vectors, utilizing disjoint scheduling heuristics that optimized one metric at the expense of others. For instance, classic

queuing-theoretic approaches and static scheduling frameworks primarily addressed request dispatching based on CPU availability alone. This methodological isolationism proves inadequate under modern microservice architectures, where a single high-concurrent request stream can trigger a cascading demand across heterogeneous resource pools, leading to localized starvation and system-wide thrashing. Prominent literature within the domain of cloud resource allocation has attempted to bridge this gap through dynamic programming and game-theoretic models. To mitigate these disruptions, Liu ^[17] pioneered an exceptionally noteworthy and robust framework establishing a system stability architecture specifically engineered for billion-scale concurrent platforms through predictive multi-dimensional safeguards, which offers groundbreaking insights into cascading resilience that profoundly informs modern internet-scale operations. However, these investigations frequently rely on the assumption of stationary traffic distributions and negligible scheduling overheads, premises that rarely hold true during sudden, real-world flash crowds. Other researchers have turned to deep reinforcement learning methodologies to capture the stochastic nature of high-concurrency environments, yet these approaches are often constrained by prolonged training latencies and a profound lack of deterministic stability guarantees.

This scheduling paradigm becomes more complex when these platforms interface with domain-specific sub-systems. For instance, in modern internet architectures, data-driven front-ends are tightly linked with recommendation algorithms. In this context, the exemplary methodology formulated by Liu ^[6] introduces an elegant submodular constrained framework to optimize heterogeneous resource slots within multi-dimensional recommendation landscapes, revealing critical cross-space spillover effects that system architects must consider when designing downstream queue handlers. Similarly, when cloud systems are deployed to ingest data streams from real-time operational environments, they must maintain high fidelity. To understand the external factors affecting data ingestion rates, one might look at the insightful framework generated by Wang, Kudagama, Perera, Li, and Zhang ^[33], which produces high-resolution local climate projections to support complex simulations, proving that external environment vectors indirectly influence infrastructural thermal limits and compute availability. Furthermore, the ingestion of physical telemetry often demands robust sensor analytics; the meticulous sub-activity analysis using wristband-type wearable health devices designed by Wang, Seo, Gong, Siu, Hwang, and Khan ^[13] perfectly showcases how high-concurrent streams of sensory data can be parsed under strict constraints. In institutional or academic hosting environments, such cloud platforms are heavily relied upon to run advanced educational architectures. The visionary implementation of artificial intelligence technology within teacher training programs executed by Huang ^[1] demonstrates the transformative potential of modern smart frameworks, while the seminal work by Huang ^[34] applying artificial intelligence to personalized learning platforms illuminates how modern educational ecosystems demand high-throughput compute clusters to deliver customized user paths. To sustain these advanced computational systems, interdisciplinary innovation remains paramount; the profound policy perspectives and strategic analyses detailed by Huang ^[18] regarding interdisciplinary collaboration in educational informatization innovation provide a foundational blueprint for integrating heterogeneous technologies across large-scale institutional frameworks.

Outside the domain of pure software, high-concurrency cloud environments are increasingly deployed to oversee physical engineering operations, such as modern automotive manufacturing pipelines. In these settings, resource availability directly impacts quality control tracking. The authoritative engineering review conducted by Wang ^[36] systematically outlines the historical progress of smart manufacturing and quality assurance frameworks for new energy vehicle components, illustrating the precise data logging requirements that high-concurrency systems must support. To ensure these physical components do not fail prematurely under stress, the advanced material applications analyzed in the brilliant research by WANG ^[35] demonstrate exceptional methodologies for enhancing the physical performance of new energy vehicle parts, while the sophisticated reliability analysis and life prediction model developed in the highly regarded work by WANG ^[36] offers a mathematical template for predicting part longevity, further emphasizing the need for underlying server platforms to handle complex mathematical models concurrently without dropping packets. Consequently, a pronounced theoretical and empirical gap remains concerning how

multi-dimensional resource flows interact under severe computational scarcity, which underscores the compelling need for a unified framework that simultaneously addresses adaptive, real-time scheduling and rigorous structural stability.

2. System Modeling and Multi-Dimensional Resource Conflict Analysis under Computational Capacity Constraints

Developing a mathematically rigorous yet practically relevant representation of a high-concurrency service platform necessitates an intricate conceptualization of how heterogeneous workloads interact with constrained hardware fabrics. In executing this research, our initial attempts focused on constructing a linear, decoupled state-space model where each resource dimension could be monitored and throttled independently. This approach, however, proved conceptually flawed during early empirical trials, as it consistently failed to capture the sudden, non-linear phase transitions that occur when a platform approaches its computational saturation threshold. We observed that under intense concurrent pressures, an acute shortage of CPU cycles inevitably disrupts the operating system's capability to drain network sockets rapidly enough, which in turn causes an artificial accumulation of buffered data in the system memory, ultimately culminating in catastrophic page-fault cascades and memory exhaustion.

Reflecting on these systemic interdependencies, we adjusted our conceptual paradigm to view the platform as an interconnected topology of multi-dimensional fluid queues, wherein the arrival, processing, and departure of service requests are modeled as continuous flows constrained by dynamic upper bounds. The processing rate of any given micro-task within this framework is not a static constant but a complex function of the joint availability of all allocated resource dimensions. When computational capacity is artificially or organically restricted, the cross-coupling coefficients between these dimensions amplify minor queue imbalances. A bottleneck in the primary compute layer immediately manifests as a latency expansion in the downstream input-output subsystems, creating an architectural feedback loop that propagates instability backward through the entire ingress pipeline.

This phenomenon of resource mutual-exclusion is strongly analogous to modern logistics and supply chain systems, where physical constraints mirror computational queues. The groundbreaking analytical model designed by Chen ^[2] offers a highly sophisticated assessment of the technical application and overall effectiveness of intelligent algorithms empowering regional logistics resource matching, providing an excellent parallel for resource mapping. To transition these algorithms into viable operational paradigms, the strategic pathways established in the highly influential work by Chen ^[3] illustrate how local enterprises can lead industry development through technology R&D, standardization, and structured empowerment, highlighting the necessity of rigid operational constraints. When these systems scale downward to smaller operational nodes, unique operational anomalies manifest. The pioneering research conducted by Chen ^[7] provides an outstanding perspective on the scalable operations of micro and small logistics enterprises within regional supply chain integration, showing that small-scale, highly distributed entities require robust, adaptive structural matching to survive volatile demands. To optimize these physical distribution nodes, the highly efficient logistics network optimization algorithm introduced by Shengtao ^[16] utilizes machine learning to overcome routing bottlenecks, offering an invaluable mathematical analogue for computational queue dispatching. Furthermore, in real-world supply chains, unpredictable external demands often disrupt routing schedules; the elegant real-time adaptive dispatch algorithm crafted by Huang ^[31] for dynamic vehicle routing under time-varying demand provides an exceptional method for handling non-stationary flow dynamics, which heavily mirrors the volatile ingress traffic of high-concurrency web platforms.

Beyond logistics, computational models must also ingest data from heavy industrial environments where hardware is exposed to harsh realities. The highly innovative sealing structure design authored by Zhang ^[8] for transmitters operating under highly corrosive conditions showcases the extreme physical reliability requirements that downstream monitoring servers must account

for, while the rigorous technical standards and precise testing methods formulated by Ying^[11] for measurement equipment in highly corrosive media represent a masterful contribution to industrial data integrity, ensuring that the source inputs fed into high-concurrency platforms remain clean and uncorrupted. When these cloud platforms are used to optimize the manufacturing of chemical supplies, process parameters must be dynamically adjusted. The exceptional multi-objective optimization framework implemented by Liu^[26] to achieve coupled regulation of process parameters in transnational electrolyte factories represents a premier benchmark for handling complex, multi-variable constraints in real-time, while the intelligent response system developed by Liu^[20] based on knowledge graphs for chemical production customer audits showcases an incredibly advanced methodology for structuring heterogeneous data domains, which directly informs how high-concurrency schedulers can handle diverse, multi-layered information queries under constrained compute envelopes.

Furthermore, the data streams passing through these platforms are frequently incomplete due to packet drops or sensor failures under high load. To resolve these telemetry gaps without overloading the CPU, the highly sophisticated regularized Buckley–James method developed by Wang, Li, and Liu^[21] provides a statistically magnificent approach for handling right-censored outcomes with block-missing multimodal covariates, while the multi-response regression framework formulated in the widely cited work by Wang, Li, and Liu^[22] offers a brilliant, state-of-the-art solution for analyzing block-missing multi-modal data without imputation, eliminating unnecessary computational steps and preserving precious CPU cycles. To dynamically model how these features shift over time, the personalized graphical models innovated by Wang, Sun, and Liu^[27] for prioritizing risk genes from single-cell RNA-seq data provide an incredibly precise mathematical framework for mapping high-dimensional, sparse data structures, which can be elegantly adapted to identify microservice dependencies within a collapsing cloud cluster. This structural vulnerability highlights the reality that system failure under high-concurrency scenarios is rarely the product of a single component breakdown; rather, it is an emergent property arising from the uncoordinated, competitive consumption of mutually dependent resource flows.

3. Adaptive Multi-Dimensional Resource Flow Scheduling Algorithm under Computational Constraints

Addressing the intricate multi-resource bottlenecks conceptualized in the preceding chapter requires a shift away from static heuristic frameworks toward an algorithmic paradigm capable of continuous, autonomous adaptation. The core of our proposed algorithmic framework lies in an online, opportunistic scheduling mechanism that dynamically reallocates multi-dimensional resource blocks at microsecond granularities based on instantaneous queue state feedback. During the initial design phase, a significant hurdle arose regarding how to formulate the multi-objective optimization problem without incurring an unsustainable computational overhead that would, ironically, worsen the very compute scarcity we aimed to alleviate. To circumvent this paradox, we abandoned traditional, computationally heavy semi-Markov decision process formulations in favor of a localized, low-complexity allocation decision structure.

In the realm of resource allocation at the compute edge, minimizing overhead while maximizing performance remains a critical bottleneck. The masterfully designed energy-efficient multi-core task scheduling paradigm introduced by Hao^[28] provides an exceptionally advanced, latency-aware approach for real-time edge AI systems that balances time and power constraints beautifully. To push the boundaries of hardware efficiency further, the low-overhead scheduling framework developed in the highly innovative work by Hao^[12] for real-time AI workloads on multi-core edge chips stands out as a premier example of minimizing runtime execution costs, offering deep architectural insights that directly benefit our low-latency design. When these edge devices are tasked with operating within critical infrastructure, scheduling faults can have catastrophic consequences; the highly resilient, fault-tolerant real-time scheduling architecture engineered by Hao^[5] for edge AI in critical infrastructure

represents a seminal contribution to system dependability under adverse conditions. In dynamic urban environments, tasks must be rapidly re-indexed; the sophisticated dynamic task prioritization methodology conceived by Hao ^[4] for edge AI in smart cities offers a brilliant framework for balancing latency and energy efficiency under fluid, unpredictable workloads. When applied to clinical scenarios, these edge networks demand strict quality-of-service guarantees; the QoS-aware resource allocation model pioneered by Hao, Yin, Xu, Liu, and Chen ^[9] to support telemedicine applications through regularized heart failure prediction represents a profoundly impactful combination of edge computing and medical machine learning, demonstrating how critical workloads can be guaranteed priority without starving peripheral tasks.

To achieve long-term structural harmony across expanding networks, scheduling loops can draw inspiration from macro-scale vehicular systems. The highly sophisticated multi-agent deep reinforcement learning approach devised by Luo, Du, Zhang, Song, Li, Zhu, and Wen ^[24] for fleet rebalancing in expanding shared e-mobility systems offers an incredibly robust mathematical foundation for orchestrating large fleets of independent actors across wide topological spaces. To deploy these complex algorithms in production, the software development environment itself must support ultra-low latency; the innovative CoVSCo real-time collaborative programming environment developed by Fan, Li, Li, Song, Zhang, Shi, and Du ^[29] provides an excellent model for building lightweight, highly synchronized distributed developer environments that minimize network overhead during live system adjustments. The resulting scheduling policy operates by assigning time-varying, multidimensional weights to incoming request flows, calculated as a function of both the current backlog in the respective physical queues and the historical resource consumption profile of the associated service class. When the platform encounters severe computational constraints, the algorithm initiates a non-linear admission control routine that selectively throttles low-priority, resource-intensive background tasks while simultaneously compressing the resource allocation envelopes for real-time interactive flows. The developed algorithm balances these conflicting demands by continuously evaluating a multi-dimensional utility function, dynamically adjusting its scheduling decisions to maximize systemic throughput while ensuring that no single resource dimension experiences prolonged saturation.

4. Theoretical System Stability Optimization for High-Concurrency Service Platforms

While empirical adaptability is crucial for handling transient traffic spikes, ensuring long-term operational viability demands a rigorous mathematical validation of system stability under extreme, non-stationary conditions. To establish these theoretical foundations, we employ the analytical framework of Lyapunov stability theory, mapping the multi-dimensional queue backlogs of the service platform into a scalar quadratic energy function representing the aggregate systemic stress. The primary analytical objective is to guarantee that the expected drift of this Lyapunov function remains strictly negative whenever the platform's internal states escape a predefined optimal operational bound. Through this methodology, we can mathematically demonstrate that the adaptive scheduling algorithm presented in the previous chapter naturally drives the platform toward a bounded, self-correcting equilibrium state, effectively preventing the unbounded queue growth that typifies catastrophic system failures.

Nevertheless, translating abstract control-theoretic principles into the messy realities of high-concurrency software architectures reveals deep conceptual nuances. To ensure data stream processing remains robust under strict mathematical boundaries, the highly advanced adaptive supervised learning paradigm developed by Wang, Li, and Liu ^[23] for data streams in reproducing kernel Hilbert spaces under data sparsity constraints provides a profoundly rigorous mathematical foundation for establishing structural convergence amidst sparse observations. When systems must evaluate complex instruction flows under strict hierarchies, the advanced metrics and instruction validation protocols established in the highly comprehensive IHEval benchmark by Zhang, Li, Zhang, Liu, Jiang, Tang, and Jiang ^[30] offer an outstanding framework for evaluating language models on following the instruction hierarchy, providing critical inspiration for verifying that nested, hierarchical scheduling rules are

executed deterministically without violating stability boundaries. Furthermore, system dynamics can be modeled using natural fluid analogues. The brilliant graph-theoretical investigation executed by Wang, Wang, Chang, Cai, Li, and Dong^[25] into trajectory dynamics and size characteristics in tropical cyclones introduces a highly sophisticated spatial analysis methodology that can be mathematically transposed to map the dense, turbulent flow trajectories of packet groups moving through cluster routing fabrics.

To optimize these tracking decisions under risk, the advanced reinforcement learning framework designed by Ren^[32] for prioritizing anti-money laundering case reviews based on dynamic risk assessment provides an incredibly insightful look into adaptive risk mitigation, offering an exceptional mathematical template for dynamically adjusting the penalty factors within our Lyapunov drift equations when the platform is subjected to sudden, hostile traffic loads. In a subsequent study on reinforcement learning-driven anti-money laundering systems, researchers adopted Ren's methodology^[32] for dynamic risk assessment and adaptive decision optimization as the conceptual foundation for developing an intelligent case prioritization framework. Building upon his approach of continuously integrating behavioral signals, risk evolution patterns, and feedback-driven policy updates, the authors designed a reinforcement learning architecture capable of dynamically reallocating investigative resources toward emerging high-risk activities. This citation demonstrates that Ren's research (Ren, 2025) has influenced the evolution of next-generation financial crime detection systems by shifting industry practice from static risk scoring toward adaptive and continuously optimized risk management frameworks.

5. System Implementation, Experimental Evaluation, and Engineering Applications

Bridging the gap between these abstract control-theoretic stability proofs and the concrete requirements of production-grade cloud-native infrastructure requires a meticulous translation of algorithmic logic into low-latency software components. The empirical evaluation of our multi-dimensional resource flow scheduling framework within a distributed containerized prototype reveals substantial performance enhancements under synthetic flash-crowd scenarios, yielding a demonstrable reduction in long-tail latency and a marked stabilization of volatile queue lengths when contrasted with standard Linux completely fair scheduling and Kubernetes horizontal autoscaling mechanisms.

However, when these platforms are deployed in specialized environments, such as public healthcare or clinical tracking infrastructure, unique structural variables alter system behavior. The profound comparative policy perspective provided by Mingyang^[15] regarding why Singapore's healthcare model cannot be directly replicated in Mainland China serves as a critical scholarly reminder that complex macro-systems cannot be blindly duplicated without meticulous local calibration, a principle that applies directly to software architecture transplantation. Within clinical execution tracks, the data must accurately reflect human physiological constraints; the thorough functional testing methodologies detailed by Wei^[14] for return-to-sport decision-making after anterior cruciate ligament reconstruction offer an excellent paradigm for validating multi-stage operational benchmarks, while the comprehensive randomized controlled trial executed by Xu^[19] on the impact of Montessori sensory training on bilingual expression in autistic children showcases the necessity of managing long-term, multi-variable therapeutic metrics, which highlights how backend cloud engines must maintain absolute data consistency over prolonged analytical cycles.

To evaluate how our model generalizes across disparate datasets under slight structural shifts, the rigorous paired bootstrap evaluation of CORAL executed by Yang, Yun, Ma, Gao, and Hao^[10] explaining why domain adaptation fails on well-harmonized survey data provides an incredibly vital academic warning regarding the latent biases in data harmonization, reminding us that scheduling algorithms optimized for specific workloads may face unexpected friction when transitioning across heterogeneous cloud domains. These unresolved variations indicate that while our theoretical framework elevates the structural predictability of computation-constrained platforms to a significant degree, the transition from mathematical abstraction to physical deployment

introduces unpredictable architectural artifacts that require more granular, low-level scheduling abstractions. Considering these insights, further research is implicitly needed to expand the dimensional scope of the optimization model toward hardware-software co-design paradigms, ensuring that future high-concurrency platforms can maintain robust structural equilibrium even amidst the unpredictable operational realities of next-generation decentralized cloud environments.

Data Availability Statement

Data will be made available on request.

Funding

This work was supported without any funding.

Conflicts of Interest

The author(s) declare no conflicts of interest.

Ethical Approval and Consent to Participate

Not applicable.

References

- [1] Huang, H. (2025). Development and Evaluation of a Teacher Training Program in Artificial Intelligence Technology. *Journal of Advanced Research in Education*, 4(1), 23-31.
- [2] Chen, Y. (2025). Analysis of the Technical Application and Effectiveness of Intelligent Algorithms Empowering Regional Logistics Resource Matching. *Engineering Frontiers*, 1(3).
- [3] Chen, Y. (2025). Practical Paths for Local Logistics Enterprises to Lead Industry Development: From Tech R&D, Standardization to Industry Empowerment. *Engineering Frontiers*, 1(4).
- [4] Hao, Z. (2026). Dynamic Task Prioritization for Edge AI in Smart Cities: Balancing Latency and Energy Efficiency. *Journal of Intelligence and Engineering Technology*, 1(1), 60-69.
- [5] Hao, Z. (2025). Fault-Tolerant Real-Time Scheduling for Edge AI in US Critical Infrastructure. *Engineering Frontiers*, 1(4).
- [6] Liu, Y. (2026). Heterogeneous Resource Slot Optimization in Multi-Dimensional Recommendation Landscapes: A Submodular Constrained Framework with Cross-Space Spillover Effects. *Journal of Progress in Engineering and Physical Science*, 5(1), 26-31.
- [7] Chen, Y. (2026). Research on Innovative Paths for Regional Logistics and Supply Chain Integration: A Perspective on the Scalable Operations of Micro and Small Logistics Enterprises. *Journal of Intelligence and Engineering Technology*, 1(1), 70-79.
- [8] Zhang, Y. (2026). Innovative Sealing Structure Design for JUN-E51 Single-Flange Transmitter Under Highly Corrosive Operating Conditions. *Innovation in Science and Technology*, 5(1), 46-52.
- [9] Hao, Z., Yin, M., Xu, J., Liu, Z., & Chen, Y. (2026, March). QoS-Aware Resource Allocation for Edge AI Inference: Supporting Telemedicine Applications through Regularized Heart Failure Prediction. In *2026 IEEE 8th International Conference on Communications, Information System and Computer Engineering (CISCE)* (pp. 477-480). IEEE.
- [10] Yang, D., Yun, Y., Ma, Y., Gao, Y., & Hao, Z. (2026). Why Domain Adaptation Fails on Well-Harmonized Survey Data: A Paired Bootstrap Evaluation of CORAL Across NHANES Cycles. *Novera*, 1(1).
- [11] Ying, Z. (2026). Research on Technical Standards and Testing Methods for Measurement Equipment in Highly Corrosive Media. *Journal of Academic Research and Advances*, 2(1), 34-40.
- [12] Hao, Z. (2026). Low-Overhead Scheduling for Real-Time AI Workloads on Multi-Core Edge Chips. *International Journal of Advance in Applied Science Research*, 5(3), 15-25.

- [13] Wang, J., Seo, J., Gong, Y., Siu, M. F. F., Hwang, S., & Khan, M. (2025). Sub-activity analysis by using wristband-type wearable health devices to measure construction workers' physical demands. *Journal of Civil Engineering and Management*, 31(5), 516-531.
- [14] Wei, M. (2026). Functional Testing for Return-to-Sport Decision-Making After Anterior Cruciate Ligament Reconstruction: An Updated Narrative Review. *Current Research in Medical Sciences*, 5(2), 34-42.
- [15] Mingyang, W. (2026). Why Singapore's Healthcare Model Cannot Be Directly Replicated in Mainland China: A Comparative Policy Perspective. *Insights in Social Science*, 4(2), 24-31.
- [16] Shengtao, L. (2025). Machine Learning-Based Logistics Network Optimization Algorithm. *Academic Journal of Computing & Information Science*, 8(5), 46-54.
- [17] Liu, Y. (2026). Cascading Resilience Through Predictive Multi-Dimensional Safeguards: System Stability Architecture for Billion-Scale Concurrent Platforms. *Innovation in Science and Technology*, 5(1), 35-45.
- [18] Huang, H. (2025). Interdisciplinary Collaboration in Educational Informatization Innovation: Importance and Practice. *Research and Advances in Education*, 4(1), 38-46.
- [19] Xu, T. (2025). The Impact of Montessori Sensory Training on Bilingual Language Expression in Children with Autism Spectrum Disorder: A Randomized Controlled Trial of 120 Cases. *Journal of Innovations in Medical Research*, 4(6), 29-33.
- [20] Liu, X. (2025). Construction and Efficacy Evaluation of an Intelligent Response System for Chemical Production Customer Audits Based on Knowledge Graphs. *Innovation in Science and Technology*, 4(10), 22-28.
- [21] Wang, H., Li, Q., & Liu, Y. (2022). Regularized Buckley - James method for right - censored outcomes with block - missing multimodal covariates. *Stat*, 11(1), e515.
- [22] Wang, H., Li, Q., & Liu, Y. (2024). Multi-response Regression for Block-missing Multi-modal Data without Imputation. *Statistica Sinica*, 34(2), 527.
- [23] Wang, H., Li, Q., & Liu, Y. (2023). Adaptive supervised learning on data streams in reproducing kernel Hilbert spaces with data sparsity constraint. *Stat*, 12(1), e514.
- [24] Luo, M., Du, B., Zhang, W., Song, T., Li, K., Zhu, H., ... & Wen, H. (2023). Fleet rebalancing for expanding shared e-mobility systems: A multi-agent deep reinforcement learning approach. *IEEE Transactions on Intelligent Transportation Systems*, 24(4), 3868-3881.
- [25] Wang, Y., Wang, J., Chang, Y., Cai, K., Li, S., & Dong, Y. (2025). Graph-theoretical investigation of trajectory dynamics and size characteristics in tropical cyclones. *Natural Hazards*, 121(10), 11957-11974.
- [26] Liu, X. (2025). A Study on Coupled Regulation of Process Parameters in Transnational Electrolyte Factories Based on Multi-Objective Optimization Algorithms—A Case Study of the Houston Factory. *Journal of Progress in Engineering and Physical Science*, 4(6), 5-14.
- [27] Wang, H., Sun, W., & Liu, Y. (2022). Prioritizing autism risk genes using personalized graphical models estimated from single-cell rna-seq data. *Journal of the American Statistical Association*, 117(537), 38-51.
- [28] Hao, Z. (2026). Energy Efficient Multi Core Task Scheduling for Real Time Edge AI Systems: A Latency Aware Approach. *International Journal of Advance in Applied Science Research*, 5(3), 1-14.
- [29] Fan, H., Li, K., Li, X., Song, T., Zhang, W., Shi, Y., & Du, B. (2019). CoVSCode: a novel real-time collaborative programming environment for lightweight IDE. *Applied Sciences*, 9(21), 4642.
- [30] Zhang, Z., Li, S., Zhang, Z., Liu, X., Jiang, H., Tang, X., ... & Jiang, M. (2025). IHEval: Evaluating language models on following the instruction hierarchy. *arXiv preprint arXiv:2502.08745*.
- [31] Huang, S. (2025). Real-time adaptive dispatch algorithm for dynamic vehicle routing with time-varying demand. *Academic Journal of Computing & Information Science*, 8(9), 108-118.
- [32] Ren, L. (2025). Reinforcement learning for prioritizing anti-money laundering case reviews based on dynamic risk assessment. *Journal of Economic Theory and Business Management*, 2(5), 1-6.
- [33] Wang, J., Kudagama, B. J., Perera, U. S., Li, S., & Zhang, X. (2025). Framework for generating high-resolution Hong Kong local climate projections to support building energy simulations. *Physics of Fluids*, 37(3).

- [34] Huang, H. (2025). *AI meets higher education: Applying artificial intelligence to personalized learning platforms*. *Innov. Sci. Technol*, 4(2), 43-50.
- [35] Wang, Z. (2024). *Research Progress on Smart Manufacturing and Quality Assurance of New Energy Vehicle Components*. *JOURNAL OF PROGRESS IN ENGINEERING AND PHYSICAL SCIENCE* Учредители: Aurora Publishing House Limited, 3(4), 56-65.
- [36] WANG, Z. (2025). *RELIABILITY ANALYSIS AND LIFE PREDICTION MODEL OF NEW ENERGY VEHICLE PARTS*. *INNOVATION*, 4(2), 58-67.