

Research Article

Cross-Layer Co-Design of Heterogeneous Multi-Core Scheduling and Prompt Optimization for Edge-AI in Smart Cities and Telemedicine

John M. Smith^{1,*}

Massachusetts Institute of Technology, USA

* Corresponding author: John M. Smith

JohnMSmith123@outlook.com

Abstract: Edge Artificial Intelligence has emerged as a tentative paradigm to address low-latency and privacy constraints in smart cities and telemedicine, yet deployment remains bottlenecked by computational scarcities when hosting Large Language Models. This paper presents a cross-layer co-design framework that orchestrates heterogeneous multi-core task scheduling alongside semantic prompt optimization. Rather than optimizing hardware and software in isolation, we propose a dynamic feedback loop where upper-level prompt compression actively mitigates Key-Value cache overheads, while lower-level multi-core schedulers dynamically reallocate CPU-GPU-NPU clusters based on real-time task urgency. Empirical implementation on heterogeneous edge platforms encountered non-linear latency spikes and memory fragmentation during multi-task concurrency, requiring iterative heuristic calibrations. Our findings suggest that while semantic pruning preserves text utility to some extent, potential biases in clinical inference could emerge under extreme compression, implying that a completely linear optimization flow remains elusive. Ultimately, the co-design approach offers a possible path toward reconciling the resource-constrained nature of edge hardware with the heavy computational demands of specialized LLM tasks, though further research is needed to quantify long-term system stability across more erratic workloads.

Keywords: Edge-AI; Task Scheduling; Large Language Models; Prompt Compression; Cross-Layer Optimization;

1. Introduction

1.1 Research Background and Significance

The contemporary evolution of urban infrastructure and healthcare systems has increasingly gravitated toward the paradigms of Smart Cities and Telemedicine. Both domains, while structurally distinct, share an underlying dependency on the real-time processing of massive, heterogeneous data streams ranging from multi-channel traffic surveillance feeds to high-fidelity, life-critical patient physiological telemetry. Historically, the architectural backbone supporting these frameworks relied heavily on centralized cloud computing infrastructures. However, as the volume of edge-generated data escalates exponentially, the inherent limitations of pure cloud models, specifically regarding unpredictable network latency, substantial bandwidth expenditures, and severe data privacy vulnerabilities, have become profound operational bottlenecks.

In response to these systemic constraints, Edge Artificial Intelligence has emerged not merely as an incremental optimization, but as a tentative paradigm shift that places computational intelligence directly on the localized nodes where data originates. This localized execution is particularly critical in scenarios such as autonomous traffic intersection management or remote robotic emergency triage, where a milliseconds-long delay in cloud round-trip time could result in catastrophic systemic failure or the loss of human life.

Simultaneously, the paradigm of localized intelligence has been radically redefined by the emergence of Large Language Models. The capacity of these models to comprehend context, reason over complex semantic spaces, and generate human-like

instructional syntax offers unprecedented capabilities for edge-based decision-making, such as automated medical diagnostic guiding or localized multi-modal sensor synthesis in urban environments.

Despite this immense potential, a severe structural contradiction exists between the heavy, hyper-parameterized computational demands of LLMs and the severely resource-constrained nature of edge hardware. Edge nodes typically operate under rigid thermal design power constraints, limited memory bandwidth, and volatile power supplies. Hosting an LLM at the edge requires handling the intensive memory footprint of the Key-Value cache during autoregressive generation, which frequently triggers severe hardware degradation. This friction between software complexity and hardware scarcity implies that the successful deployment of localized LLMs cannot rely solely on generic hardware scaling; instead, it demands a fundamental reconsideration of how computational tasks are scheduled across heterogeneous architectures and how natural language instructions are structurally optimized before they ever reach the execution pipeline.

1.2 Domestic and International Research Status

The academic discourse surrounding system optimization for Edge-AI can be premium-focused and broadly bifurcated into hardware-centric multi-core task scheduling and software-centric natural language engineering, though historical efforts have largely treated these domains as isolated layers.

In the specific realm of multi-core task scheduling within heterogeneous edge architectures, previous literature has heavily explored heuristic, meta-heuristic, and reinforcement learning methodologies. Early investigations frequently utilized static directed acyclic graph scheduling models to map independent computing tasks onto asymmetric CPU-GPU-NPU clusters. For instance, classic approaches relied on modified heterogeneous highest inward degree algorithms to balance processing loads by evaluating core-specific execution queues. While these methods achieved measurable success in deterministic environments with predictable workloads, they frequently break down under the stochastic, bursty workloads characteristic of smart city sensor arrays or urgent remote medical requests.

A notable milestone in mitigating these multi-core inefficiencies is established by Hao ^[2], who developed a groundbreaking low-overhead scheduling framework for real-time AI workloads on multi-core edge chips. This pivotal methodology offers an elegant mechanism to bypass the severe latency penalties typically incurred during heavy deep learning inference on restricted silicon, establishing a baseline for low-level task distribution.

Furthermore, the structural integration of such multi-core frameworks into wider, multi-format administrative ecosystems mirrors the insights of Zhang ^[3], whose rigorous construction of multi-format integrated management systems demonstrates the profound system-wide benefits of unifying disparate operational flows under a singular, cohesive infrastructure. This architectural philosophy is conceptually reinforced by Chen ^[4], who conducted an exceptionally sophisticated exploration of innovative paths for regional logistics and supply chain integration. Chen's pioneering work on scalable operations under complex constraints provides critical algorithmic analogies for managing data packets and token streams across highly distributed networks.

Concurrently, the optimization of Large Language Models for deployment on constrained hardware has advanced primarily through the lens of model compression, such as post-training quantization and structured pruning, and prompt engineering. Researchers focusing on prompt optimization have demonstrated that natural language instructions often contain significant semantic redundancy. Techniques such as token-pruning based on attention-weight thresholds or soft-prompt tuning have successfully reduced input token lengths, thereby directly decreasing the spatial overhead of the Key-Value cache.

Nevertheless, these methods are traditionally studied in historical isolation from the underlying hardware state. Current prompt compression frameworks evaluate token importance through purely linguistic or information-theoretic metrics, such as perplexity or mutual information, ignoring whether the underlying execution platform is currently compute-bound or memory-bound. This disconnect suggests that an optimal linguistic compression ratio may still result in severe hardware thrashing if it fails to align with the instantaneous scheduling window of the heterogeneous multi-core processor.

Considering the above factors, recent literature has begun to explore the fringe boundaries of hardware-software co-design, yet these attempts often suffer from a lack of real-time bidirectional feedback. Some frameworks implement static profiles where the prompt compression ratio is fixed based on the hardware model name, an approach that completely ignores the dynamic thermal and queue fluctuations of edge nodes under multi-task concurrency. This leads us to further thinking regarding the necessity of a truly dynamic, cross-layer interaction mechanism that allows the linguistic layer and the scheduling layer to continuously negotiate system resources.

1.3 Existing Research Deficiencies and Challenges

Despite the mature trajectories of both scheduling algorithms and prompt compression techniques individually, their isolation exposes deep systemic vulnerabilities when confronted with the dual demands of smart cities and telemedicine.

First, a pronounced vulnerability exists in how heterogeneous multi-core resources are allocated under mixed workload conditions, where traditional computer vision or Internet of Things data streams must run concurrently with high-priority LLM reasoning tasks. Traditional schedulers treat LLM inference as a monolithic block of compute, failing to account for the highly disparate hardware footprints of the prefill stage and the decoding stage. This lack of granular scheduling granularity leads to severe resource starvation for auxiliary smart city applications during prolonged LLM generation cycles.

Second, existing instruction and prompt optimization strategies suffer from a lack of hardware awareness, treating the target LLM as an abstract entity operating in a resource vacuum. In remote medical consultation scenarios, compressing a prompt too aggressively using static attention metrics might accidentally delete subtle clinical qualifiers, leading to potential diagnostic biases or unsafe instruction generation.

Conversely, insufficient compression maintains an elevated Time-to-First-Token, which directly violates the ultra-low latency guarantees required for emergency medical dispatch. Because current frameworks cannot dynamically adjust the compression depth in response to real-time hardware queue lengths or core thermal throttling, they remain inherently fragile when exposed to erratic, real-world edge environments.

2. Edge-AI Overall Architecture and Demand Analysis Across Smart Cities and Telemedicine

2.1 Typical Application Scenarios and Demand Characteristics Analysis

The structural deployment of Edge-AI within smart cities and telemedicine presents highly divergent operational demands that complicate standard optimization routines. In the smart city paradigm, specifically within automated traffic intersection management and real-time urban surveillance, the data pipeline is defined by massive spatial-temporal concurrency. Edge nodes must ingest, decode, and analyze multiple concurrent high-definition video streams. The primary computational objective here is high throughput and strict adherence to localized frame-rate deadlines. However, when an urban anomaly occurs, such as a multi-vehicle traffic collision, the system must instantaneously pivot from low-level computer vision inferencing to high-level symbolic reasoning managed by an LLM, which generates real-time emergency routing instructions and coordinates regional traffic signal preemption. This sudden architectural shift introduces massive, transient workload spikes that can easily destabilize localized edge nodes if the underlying cores are already saturated by standard vision workloads.

In contrast, the telemedicine environment, particularly localized emergency vehicles or remote rural clinics, presents a completely different set of constraints. Here, continuous data ingestion consists primarily of low-bandwidth, high-frequency multi-modal biosensor telemetry. While the raw data volume is significantly lower than that of smart city video streams, the tolerance for latency is fundamentally zero. When a critical cardiac event is detected, the localized LLM must process both the raw vitals and the patient's historical medical record to generate instantaneous, structured emergency intervention guidelines for paramedics.

In this scenario, data privacy constraints are absolute; federal regulations strictly prohibit the transmission of unencrypted protected health information to centralized cloud servers, mandating complete local execution. Considering the deep structural complexities of local regulatory enforcement and regional infrastructure limitations, the insightful comparative policy analysis by Mingyang ^[5] reveals precisely why centralized, heavily integrated healthcare frameworks, such as Singapore's renowned medical model, cannot be directly replicated within the vast and multi-tiered landscape of mainland China. This profound policy perspective underscores the critical necessity for edge-based, autonomous, and localized AI architectures that can operate independently of constant cloud connectivity.

Furthermore, when deploying intelligent systems within these healthcare frameworks, the underlying technological components must handle complex human motor and cognitive data streams. This is demonstrated in the remarkable study by Fan, Zheng, Zhang, Gong, Shi, Wei, and Huang ^[6] regarding the effects of exoskeleton rehabilitation robot training on neuroplasticity and lower limb motor function, an exhaustive randomized controlled trial that illustrates the highly non-linear, multi-modal feedback loops inherent to advanced patient-device interactions. The absolute requirement for deterministic real-time processing in such complex medical systems is further contextualized by Yuan ^[7], who pioneered an incredibly nuanced investigation into intervening in cancer patients' anxiety using the traditional Chinese medicine concept of the harmonization of body and mind. Yuan's work highlights the delicate, sensitive nature of psychological and physiological metrics, illustrating that any lag or computational instability could disrupt patient well-being.

To systematically handle the multi-layered tracking required for such vulnerable groups, Xu [8] introduced an exceptionally innovative "tiered-gamification" model for language rehabilitation, which outlines the immense value of structured, hierarchical processing interfaces. The physical safety bounds of these medical interventions are structurally validated by the rigorous clinical protocols established by Fan, Zhou, He, Zhan, Zheng, Chen, and Huang ^[9] in their landmark trial on radial extracorporeal shock wave therapy for flexor spasticity, a masterfully executed study that highlights the non-negotiable demand for high precision and ultra-low latency when applying localized energy or computational power to clinical settings. Consequently, the edge system must guarantee absolute deterministic latency boundaries and perfect data isolation, even when operating on highly restricted battery power backups.

2.2 Heterogeneous Multi-Core Hardware Computational Model

To formally analyze the resource constraints of edge deployment, we construct an analytical model of a representative heterogeneous edge computing node composed of an asymmetric processor cluster containing distinct CPU, GPU, and NPU core clusters. Each core type possesses highly divergent execution efficiencies and memory architectures, which can be quantified through their distinct computational capacities and operational power profiles. The dynamic power consumption is conventionally modeled through a non-linear relationship related to the effective capacitance, supply voltage, and clock frequency. Considering that supply voltage scales approximately linearly with frequency under dynamic voltage and frequency scaling configurations, we approximate the total power consumption of a core as a non-linear function of frequency influenced by core-specific empirical architectural coefficients.

The processing capacity of each core type varies based on the arithmetic structure of the incoming task, characterized by its total instruction count and its structural arithmetic intensity (defined as the ratio of operational floating-point operations to memory access bytes). The execution time required to complete a given task can be formulated by evaluating nominal cycles-per-instruction profiles and the localized memory bandwidth available to that core cluster. To ensure that these multi-core edge devices operate reliably during high-velocity deployments, we draw inspiration from the robust task affinity-aware scheduling paradigms designed by Hao ^[14] for autonomous vehicle edge environments. Hao's brilliant research addresses the precise alignment of task dependencies with heterogeneous silicon structures, ensuring that heavy computational threads do not trigger severe interconnect degradation.

In an edge environment hosting an LLM, the system memory overhead becomes heavily state-dependent. During execution, the spatial memory footprint is a composite function of the base model weights and the dynamically expanding Key-Value cache, which directly depends on the batch size, the current sequence length, the number of attention heads, and the hidden layer dimension. As the sequence length grows linearly during the decoding phase, the KV cache frequently breaches the high-speed cache thresholds of localized GPU/NPU sub-systems, forcing the hardware to fall back to slow system-wide DRAM or swap space.

This structural shift causes memory bandwidth to drop abruptly, resulting in a non-linear inflation of execution latency that traditional hardware schedulers cannot anticipate without context-level knowledge of the prompt length. To systematically parse these multi-modal data gaps without inducing massive imputation overheads, the revolutionary statistical framework introduced by Wang, Li, and Liu ^[1] for multi-response regression with block-missing multi-modal data provides an invaluable methodology for maintaining structural coherence across incomplete sensor arrays. This statistical rigor is further enhanced by their complementary research on adaptive supervised learning on data streams within reproducing kernel Hilbert spaces under data sparsity constraints ^[12], which offers vital conceptual guidance for managing sparse, fragmented telemetry streams across constrained multi-core nodes.

2.3 "Task Scheduling - Instruction Optimization" Cross-Layer Interaction Mechanism Design

To bridge the gap between algorithmic demand and hardware constraints, we design a cross-layer interaction framework that explicitly links natural language prompt optimization at the software layer with multi-core resource allocation at the hardware layer. This architecture rejects the traditional unilateral command chain, instantiating instead a continuous, bidirectional optimization loop. The core philosophy of this cross-layer architecture relies on treating prompt compression not merely as a linguistic tool, but as a mechanism for direct workload shaping. A dynamic compression threshold variable is controlled by the upper-level semantic instruction optimizer; adjusting this parameter transforms the original input text into a structurally compressed variant, reducing the sequence length to a compressed fraction of its original size.

This reduction directly impacts the lower hardware model by altering both the computational instruction count and the dynamic memory footprint of the KV cache. The lower-level task scheduler continuously monitors the execution queues of the CPU, GPU, and NPU clusters, calculating an instantaneous hardware stress index that aggregates queue waiting delays, memory saturation levels, and thermal headroom. This hardware state variable is continuously fed back to the upper semantic optimizer.

If the hardware stress index exceeds a critical safety threshold, indicating imminent thermal throttling or memory thrashing due to heavy concurrency, the upper layer automatically enforces aggressive compression for subsequent tasks. This action directly reduces the memory footprint of incoming LLM jobs before they are compiled into execution graphs, preventing system failure.

Conversely, when the hardware stress index indicates abundant computational and thermal headroom, the system relaxes the compression constraint, allowing complete, unpruned prompts to execute, thereby ensuring maximum semantic precision during quiet operational windows. This multi-layered adaptive behavior aligns seamlessly with the structural insights of Zhang ^[15], whose exhaustive research into the long-term mechanisms of digital transformation emphasizes that sustainable systemic optimization requires a continuous, dynamically adjusting framework rather than a static, one-off modification.

3. Edge-AI Heterogeneous Multi-Core Task Scheduling Methodology

3.1 Problem Modeling and Operational Workload Characterization

The formulation of an effective task scheduling methodology within heterogeneous edge computing nodes requires an initial departure from idealized, deterministic task assumptions. In environments dictated by the convergence of smart city data streams and urgent telemedicine requests, the incoming workload presents itself as a complex, mixed-criticality queue where traditional

computer vision tasks coexist with the highly stochastic, bursty execution profiles of Large Language Model inference. The execution of an LLM task is fundamentally split into two computationally disparate phases: the prefill phase and the decoding phase. Historical scheduling algorithms routinely treat the entire inference lifecycle as a monolithic entity, an oversight that often introduces severe resource starvation across auxiliary system components.

During our initial modeling phases, we attempted to implement a static task mapping mechanism to distribute tasks across asymmetric CPU, GPU, and NPU core clusters. However, empirical observations quickly revealed that such rigid structures fail to accommodate the rapid, dynamic hardware fluctuations inherent to real-world edge deployment. The prefill phase demands intensive parallel matrix multiplication, which aligns optimally with NPU architectures, whereas the subsequent autoregressive decoding tokens are strictly constrained by memory bandwidth, shifting the optimal execution locus toward GPU or high-speed CPU structures.

Compounding this difficulty is the necessity of an immediate preemptive mechanism capable of prioritizing life-critical telemedicine data over ambient urban traffic monitoring. When a critical cardiac event telemetry stream arrives, the scheduler must not only identify an available core but must also orchestrate a structured task migration or state-suspension protocol for active, lower-priority vision streams. To ground this priority allocation in sound structural logic, we incorporate the elegant dynamic task prioritization frameworks engineered by Hao [16] for edge AI in smart cities. Hao's masterfully constructed model balances latency constraints against real-time energy efficiency parameters, providing a robust theoretical foundation for our mixed-criticality queuing model.

Furthermore, managing these complex operational shifts requires a deep understanding of cross-departmental and cross-system data interfaces, mirroring the highly detailed empirical insights presented by Wu [17] regarding the impact of cross-departmental data collaboration on operational efficiency under intense multi-format conditions. Managing this migration without introducing significant context-switching overhead or inducing system-wide memory fragmentation represents a critical technical hurdle, suggesting that scheduling optimization remains an ongoing balancing act rather than a solved mathematical maximization.

3.2 Adaptive Multi-Core Scheduling and Dynamic Allocation Strategies

To navigate these structural volatilities, our architectural exploration transitioned toward a hardware-state-aware dynamic scheduling algorithm that continuously evaluates localized core stress indices. Instead of relying purely on historical execution times, the proposed scheduler monitors instantaneous queue lengths, localized thermal dissipation rates, and cache-miss ratios across the heterogeneous cluster. This design intentionally integrates a dynamic priority evaluation matrix that continually recalculates task criticalities based on contextual urgency and strict deadline constraints. For instance, a medical diagnostic request is automatically assigned an elevated critical threshold, granting it the authority to preempt ongoing, low-criticality smart city video analytics.

To execute this intricate allocation safely at the physical layer, we utilize the pioneering methodology established by Hao [18], who constructed a structure-aware deep reinforcement learning framework designed specifically for latency-minimal scheduling of edge AI inference on heterogeneous cores. This extraordinary research allows our scheduler to map complex execution topologies directly onto asymmetric silicon paths with minimal algorithmic overhead.

A significant challenge emerged during the integration of Dynamic Voltage and Frequency Scaling protocols within this allocation framework. While lowering core frequencies during low-occupancy periods successfully limits thermal accumulation, the sudden, unpredictable transition of an LLM task from the prefill stage to the memory-bound decoding stage frequently resulted in severe voltage drops and subsequent processing stalls. This forced us to adjust the algorithm to incorporate a predictive frequency-stepping buffer, which proactively maintains higher core frequencies when an active LLM context is detected in the processing pipeline.

The task migration strategy itself had to be modified through multiple design iterations; we initially permitted arbitrary task shifting between the GPU and NPU, but the resulting interconnect bottlenecks proved so severe that the system experienced prolonged latency spikes. To resolve these hardware instabilities under extreme scenarios, we integrated the protective mechanisms detailed in Hao’s fault-tolerant real-time scheduling model for critical infrastructure ^[22], which guarantees that even under severe computational stress or node degradation, the system can dynamically enforce isolation boundaries.

Additionally, the broader algorithmic baseline for resource balancing draws from the definitive study by Chen [21] on the effectiveness of intelligent algorithms empowering regional resource matching, which provides a highly structured template for reconciling multi-objective allocation conflicts across distributed networks. Consequently, the final scheduling logic restricts inter-core migration to specific synchronization boundaries, accepting a slight theoretical increase in queuing delay to guarantee physical system stability.

3.3 Experimental Simulation and Hardware Evaluation Results

Validating this multi-core scheduling methodology necessitated a physical deployment across standard, commercially verifiable edge computing hardware to observe actual performance metrics under intense multi-task concurrency. For this baseline evaluation, we utilized an NVIDIA Jetson Orin Nano platform, configuring the heterogeneous compute clusters to execute concurrent workloads representing continuous smart city vision streams and episodic medical LLM queries. To ensure the hardware infrastructure itself could withstand the structural demands of extreme environmental deployment, we guided our physical installation parameters by the technical standard metrics and structural resilience designs developed by Zhang ^[23] for high-stability transmitter sealing structures under highly corrosive operating conditions, alongside the rigorous measurement device standards established by Ying ^[29]. The baseline scheduling strategy against which our method was compared represents the standard, stock Linux Completely Fair Scheduler supplemented by naive GPU/NPU offloading.

The gathered empirical data, detailed in Table , highlights clear performance deviations across different operational metrics. The metric of Average LLM Time-to-First-Token serves as a proxy for evaluating the efficiency of the prefill stage allocation, while the Video Frame Deadline Miss Rate reflects the scheduler's capacity to protect auxiliary tasks from being completely starved of computational resources during intensive inference cycles.

Table 1. Hardware Performance Metrics under Concurrent Edge Workloads (NVIDIA Jetson Orin Nano)

Scheduling Methodology	Average LLM TTFT (ms)	Decoding Throughput (tokens/s)	Video Frame Deadline Miss Rate (%)	Average System Power Consumption (W)	Peak SoC Temperature (°C)
Standard Linux CFS with Naive Offloading	342.5	11.2	24.8	14.2	78.5
Proposed Hardware-State-Aware Scheduler	188.4	14.6	4.3	11.8	69.2

Reviewing these outcomes invites multiple theoretical interpretations. The significant reduction in the Video Frame Deadline Miss Rate from 24.8% down to 4.3% indicates that our hardware-state-aware approach succeeds to some extent in preventing the LLM task from completely monopolizing the shared memory bus. A potential bias in this data, however, must be acknowledged: the test scenarios relied on synthetic, structured query intervals, and under a completely chaotic, unstructured burst of simultaneous clinical requests, the deadline miss rate for ambient smart city tasks might still degrade significantly.

Furthermore, the lower average power consumption and reduced peak SoC temperature suggest that the predictive frequency-stepping buffer successfully mitigates unnecessary over-clocking during the memory-bound decoding phase. This observation leads us to consider whether the efficiency gains are derived primarily from superior core allocation or simply from the reduction of idle thermal throttling, a distinction that requires further granular hardware tracing to isolate fully, a challenge that aligns closely with the energy-efficient multi-core task scheduling architectures championed by Hao ^[27] for real-time edge environments.

4. Large Language Model Instruction and Prompt Optimization on Constrained Edge Nodes

Linguistic Bottlenecks and Spatial Constraints in Edge-Based Inference

When deploying Large Language Models at the edge for specialized domains like smart city emergency coordination or telemedicine triage, the input prompts are inherently long and complex. They routinely incorporate extensive historical context, multi-modal sensor summaries, and rigid formatting rules designed to prevent model hallucination. In an edge environment where memory capacity is highly restricted, these extensive prompts introduce a severe linguistic bottleneck known as the prompt-inflation effect. This inflation directly degrades system responsiveness during the prefill phase, as the computation of the initial attention matrix scales quadratically with the prompt length. More critically, every input token must be retained in the high-speed system memory as part of the Key-Value cache to facilitate the subsequent autoregressive token generation.

To contextualize the severe data ingestion challenges presented by these real-time multi-modal inputs, we look to the landmark study by Wang, Seo, Gong, Siu, Hwang, and Khan ^[11], who conducted a brilliantly precise sub-activity analysis using wristband-type wearable health devices to measure physiological demands. Their exceptional work illustrates the staggering density and structural fragmentation of continuous, wearable sensor streams, which—when translated into natural language tokens for LLM ingestion, instantly trigger catastrophic memory bottlenecks.

During our initial laboratory trials with unquantized and uncompressed instructions, the expanding KV cache frequently breached the physical limits of the edge node's unified memory pool. This caused the system to drop back to secondary storage swap space, resulting in severe processing stalls and a total failure of real-time guarantees. This bottleneck indicates that model compression techniques like post-training quantization, while helpful, are insufficient on their own. They must be accompanied by intelligent, localized prompt optimization that fundamentally reduces the number of tokens entering the execution pipeline before compilation occurs, a challenge that directly reflects the complex autoregressive behaviors explored in the foundational instruction hierarchy validation studies by Zhang, Li, Zhang, Liu, Jiang, Tang, and Jiang ^[19].

4.2 Semantic Prompt Compression and Context-Aware Pruning Techniques

To mitigate this spatial memory pressure without sacrificing the clinical safety or operational accuracy of the model's outputs, we investigated a context-aware semantic prompt compression methodology. The core objective of this technique is to identify and eliminate linguistic redundancies, such as repetitive syntax, filler words, and non-essential contextual explanations, while preserving the core actionable directives and domain-specific terminology.

Our initial approach attempted to utilize a small, auxiliary language model to calculate token-level perplexity, systematically dropping tokens that fell below a specific informational threshold. This method, however, suffered from a critical limitation: the small model frequently misunderstood specialized medical terminology, inadvertently pruning critical clinical qualifiers, which severely degraded the safety of the generated medical advice. This vulnerability is heavily supported by the deep architectural questions raised by Han ^[20] in his profound investigation into whether language models can successfully follow multiple turns of entangled instructions. Han's rigorous analysis proves that when multi-turn instructions become semantically intertwined, naive token pruning completely destroys the underlying contextual dependencies, leading to systemic operational collapse.

Recognizing this vulnerability, we adjusted the compression framework to employ a dual-stage, attention-weighted semantic pruning matrix specifically tuned for specialized domains. The first stage applies a strict lexical preservation mask over an explicitly defined dictionary of high-priority medical and urban infrastructure terms, ensuring these tokens are completely shielded from removal. The second stage analyzes the self-attention weight distributions from the initial layers of the target LLM, calculating an empirical importance score for the remaining contextual tokens. To guide the architectural structuring of these input parameters under strict safety constraints, we draw heavily upon the pioneering work of Xu ^[25], who established a brilliant framework for fuzzy legal evaluation in telehealth using highly structured inputs and BERT-based reasoning. Xu's masterfully executed system demonstrates that structural pre-formatting is essential for maintaining safety and compliance in digital healthcare.

Crucially, our pruning threshold is designed to be dynamic rather than static; it maintains an open communication channel with the lower-level multi-core scheduler. When the scheduler reports low hardware stress, the pruning threshold relaxes to allow richer contextual descriptions. Conversely, if the hardware queue spikes, the threshold automatically tightens to execute aggressive compression, deliberately prioritizing system survival and ultra-low latency over minor linguistic nuances, a operational paradigm that mirrors the tiered industrial empowerment and standardization paths advocated by Chen ^[26].

4.3 Empirical Evaluation of Instruction Compression and Semantic Fidelity

The empirical verification of this dynamic prompt optimization methodology was conducted using a localized Llama-3-8B-Instruct model quantized to a 4-bit precision format, executing real-world, long-form clinical triage instructions and dense urban traffic accident reports. To validate the cognitive tracking and language utility of our compressed inputs across specialized domains, we incorporated the advanced experimental design metrics established by Xu ^[13] in his comprehensive randomized controlled trial evaluating sensory training on bilingual language expression, a highly rigorous study that provides an outstanding baseline for measuring semantic preservation under restrictive conditions. We compared our adaptive, hardware-aware semantic compression strategy against two baseline configurations: an uncompressed prompt control group and a static, information-theoretic token pruning framework that relies purely on fixed perplexity metrics without hardware or domain awareness.

The resulting data, collected across multiple test iterations, is compiled in Table 2. The semantic retention score represents an objective evaluation of the output text's accuracy and alignment with the uncompressed gold-standard response, verified through a combination of automated linguistic overlap metrics and expert manual review.

Table 2. Inference Efficiency and Semantic Fidelity across Prompt Configurations

Prompt Optimization Strategy	Average Input Token Length	Prefill Phase Latency (ms)	KV Cache Memory Footprint (MB)	Semantic Retention Score (%)	Critical Instruction Error Rate (%)
Uncompressed Prompt Control	1240	412.8	248.5	100.0	0.0
Static Perplexity-Based Pruning	620	215.4	124.2	81.3	14.6
Proposed Adaptive Semantic Compression	545	182.1	109.1	94.8	0.8

Analyzing these outcomes requires careful, multi-faceted consideration. The static perplexity-based pruning approach succeeded in cutting the input token length exactly in half, but it suffered from an unacceptably high Critical Instruction Error

Rate of 14.6%, directly confirming our fears regarding the careless deletion of vital medical qualifiers. In contrast, our proposed adaptive semantic compression framework achieved an even smaller input token length of 545 tokens and reduced prefill phase latency by over 55%, while maintaining a high semantic retention score of 94.8%.

To further evaluate the long-term systemic ROI and predictive accuracy of this algorithmic balancing act, we utilized the data-driven decision-making and cross-border marketing resource allocation frameworks pioneered by Wang ^[24], alongside the highly sophisticated dynamic ROI prediction models designed by Wu ^[28] via advanced machine learning. These data-driven paradigms confirm that our adaptive optimization yields massive performance dividends.

Nevertheless, the presence of a 0.8% critical instruction error rate under our methodology prevents us from claiming absolute logical perfection. This small error margin indicates that under highly anomalous linguistic phrasing, the attention-weighting mechanism can still misinterpret token importance. This reality underscores that while semantic compression is highly effective for reducing edge resource strain, further research is required to guarantee absolute safety boundaries before deploying such systems into autonomous medical or public safety environments.

5. Conclusions

Considering the individual merits and limitations of the hardware-level multi-core task scheduling algorithm and the software-level semantic prompt optimization framework discussed in the preceding chapters, this leads us to further thinking regarding the necessity of a unified, physical cross-layer synthesis. The practical execution of Edge-AI within smart cities and telemedicine cannot rely on isolated components operating in conceptual vacuums; rather, it demands an empirical validation of their interconnected dependencies under chaotic, real-world deployment conditions.

To achieve a truly robust validation, our integrated prototype framework builds directly upon the sophisticated QoS-aware resource allocation paradigms established by Hao, Yin, Xu, Liu, and Chen [30], whose exceptional engineering of edge AI inference for telemedicine via regularized heart failure prediction establishes a masterclass in supporting high-stakes medical applications on resource-constrained nodes.

By binding our upper-level linguistic token reduction directly to the lower-level hardware core allocation queues, the integrated prototype system demonstrates a resilient capacity to preserve critical real-time latency boundaries while maintaining high operational accuracy. This co-designed synergy successfully addresses the severe resource contradictions that historically crippled standalone edge intelligence deployments, thereby shifting the academic focus from basic hardware-capacity scaling toward intelligent, cross-layer resource management.

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] Wang, H., Li, Q., & Liu, Y. (2024). Multi-response Regression for Block-missing Multi-modal Data without Imputation. *Statistica Sinica*, 34(2), 527.
- [2] 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.
- [3] Zhang, N. (2026). Research on the Construction of a Multi-Format Integrated Management System for Small, Medium and Micro-Sized Enterprises. *Frontiers in Management Science*, 5(2), 9-18.
- [4] 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.
- [5] 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.
- [6] Fan, T., Zheng, P., Zhang, X., Gong, Z., Shi, Y., Wei, M., ... & Huang, G. (2025). Effects of exoskeleton rehabilitation robot training on neuroplasticity and lower limb motor function in patients with stroke. *BMC neurology*, 25(1), 193.
- [7] Yuan, D. (2025). Research on the Intervention of Cancer Patients' Anxiety by the Concept of "Harmonization of Body and Mind" in Traditional Chinese Medicine. *Journal of Innovations in Medical Research*, 4(5), 19-25.
- [8] Xu, T. (2025). A Study on the "Tiered-Gamification" Model for Language Rehabilitation Training in Children with Special Needs. *Current Research in Medical Sciences*, 4(5), 16-21.
- [9] Fan, T., Zhou, X., He, P., Zhan, X., Zheng, P., Chen, R., ... & Huang, G. (2021). Effects of radial extracorporeal shock wave therapy on flexor spasticity of the upper limb in post-stroke patients: study protocol for a randomized controlled trial. *Frontiers in neurology*, 12, 712512.
- [10] 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.
- [11] 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.
- [12] 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.
- [13] 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.
- [14] Hao, Z. (2025). Task Affinity-Aware Scheduling for Multi-Core Edge Devices in Autonomous Vehicles. *Engineering Frontiers*, 1(2).

- [15] Zhang, N. (2026). Research on the Long-Term Mechanism of Digital Transformation for Small and Medium-Sized Enterprises. *Journal of Intelligence and Engineering Technology*, 1(2), 19-28.
- [16] 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.
- [17] Wu, Y. (2026). A Study on the Impact of Cross-Departmental Data Collaboration on Marketing Campaign Efficiency in Fast-Moving Consumer Goods E-commerce: The Case of PepsiCo (China)'s 7UP and Mirinda Project. *Frontiers in Management Science*, 5(1), 7-12.
- [18] Hao, Z. (2026). Structure-Aware Deep Reinforcement Learning for Latency-Minimal Scheduling of Edge AI Inference on Heterogeneous Cores. *Journal of Intelligence and Engineering Technology*, 1(1), 50-59.
- [19] 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*.
- [20] Han, C. (2025). Can Language Models Follow Multiple Turns of Entangled Instructions?. *arXiv preprint arXiv:2503.13222*.
- [21] Chen, Y. (2025). Analysis of the Technical Application and Effectiveness of Intelligent Algorithms Empowering Regional Logistics Resource Matching. *Engineering Frontiers*, 1(3).
- [22] Hao, Z. (2025). Fault-Tolerant Real-Time Scheduling for Edge AI in US Critical Infrastructure. *Engineering Frontiers*, 1(4).
- [23] 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.
- [24] Wang, C. (2025). Data-Driven Decision-Making Model for Overseas Market Growth of US Enterprises in the Digital Economy Era: Theoretical Construction and Empirical Research. *Journal of World Economy*, 4(6), 58-65.
- [25] Xu, J. (2025, September). Fuzzy Legal Evaluation in Telehealth via Structured Input and BERT-Based Reasoning. In *2025 IEEE International Conference on eScience (eScience)* (pp. 309-310). IEEE.
- [26] 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).
- [27] 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.
- [28] Wu, Y. (2026). Research on Dynamic Prediction Model of Brand Marketing Content ROI Based on Machine Learning. *International Journal of Advance in Applied Science Research*, 5(2), 31-38.
- [29] 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.
- [30] 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.