

Intelligent Edge-Assisted Network Slice Orchestration for Ultra-Reliable Low-Latency Communications in 5G Networks

Zhouzou Cheng

Department of Computer Science, Binghamton University, Binghamton, NY, USA.
zhoumail@binghamton.edu

Hudson Sims

Department of Computer Science, University of Houston, Houston, TX, USA.
hudson.work@uh.edu

Kaijin Hao

School of Electrical Engineering and Computer Science, Oregon State University, Corvallis,
OR, USA.
kaijin1975@oregonstate.edu

Roy Reynolds

Department of Electrical Engineering and Computer Science, University of Missouri,
Columbia, MO, USA.
reynolds1981@missouri.edu

Abstract

The emergence of fifth-generation (5G) wireless systems has introduced a paradigm shift through network slicing, enabling the co-existence of diverse service verticals over a common physical infrastructure. Among the most demanding slice categories is ultra-reliable low-latency communication (URLLC), which necessitates deterministic performance guarantees for mission-critical applications such as autonomous driving, industrial automation, and tele-surgery. However, achieving end-to-end URLLC across geographically distributed edge sites requires a fundamental rethinking of slice orchestration mechanisms that extend beyond centralized resource management. This paper presents a comprehensive system-level analysis of intelligent edge-assisted network slice orchestration for URLLC in 5G networks. We examine the architectural decomposition of orchestration functions across core, edge, and far-edge layers, highlighting the structural trade-offs between centralized coordination and distributed autonomy. The discussion explores how edge intelligence, particularly through machine learning, can enable proactive slice adaptation, predictive fault management, and context-aware policy enforcement. Special attention is devoted to the governance of multi-domain resources, the interplay between computational and networking elements, and the resilience strategies required to sustain URLLC under dynamic conditions. We further evaluate deployment constraints, infrastructure scalability, sustainability considerations, and fairness implications when URLLC slices compete with other service types. The paper contributes a forward-looking perspective on how policy frameworks and standardization efforts can align technical orchestration with socioeconomic objectives. By integrating perspectives from systems engineering, artificial intelligence, and regulatory governance, this

work provides a holistic reference for designing and operating next-generation network slices that simultaneously satisfy ultra-reliability and low-latency requirements.

Keywords

5G network slicing; edge computing; ultra-reliable low-latency communication; slice orchestration; artificial intelligence; policy governance; resilience; sustainability.

1. Introduction

The fifth generation of mobile networks has been designed to accommodate a radically heterogeneous set of service requirements through the concept of network slicing, a technique that partitions a single physical network into multiple logical networks, each optimized for a distinct class of application demands [1]. Among the standardized slice types, ultra-reliable low-latency communication (URLLC) represents the most stringent category, targeting scenarios where both packet loss and delay must be bounded within extremely narrow margins. Use cases such as remote robotic surgery, cooperative autonomous driving, and industrial process control not only rely on minimal latency but also impose hard reliability constraints, often exceeding five-nines availability targets. The satisfaction of such requirements cannot be achieved solely through radio access improvements; it demands holistic orchestration spanning the radio access network, transport domain, and distributed computing resources. As edge computing platforms proliferate, the locus of orchestration intelligence is progressively shifting from centralized data centers toward the network edge, enabling decisions to be made closer to where data originates and where actuation must occur.

This structural shift introduces a complex set of systems-level challenges. Orchestrating URLLC slices across edge clusters involves reconciling the need for low-latency control loops with the overarching requirement for global resource efficiency and policy compliance. The distributed nature of edge sites, often operated by different administrative stakeholders, introduces multi-domain coordination difficulties that are absent in traditional single-provider cloud environments [2]. Furthermore, the integration of artificial intelligence into the orchestration plane holds promise for automating complex decision-making processes, yet the placement of intelligence across the edge-to-cloud continuum raises fundamental questions about model freshness, inference latency, and resilience to environmental non-stationarity. This paper undertakes an interdisciplinary examination of intelligent edge-assisted network slice orchestration for URLLC, focusing on architectural principles, system-level trade-offs, governance mechanisms, and the broader policy implications that shape the feasibility and fairness of such deployments.

2. Background and Architectural Context

Network slicing leverages software-defined networking (SDN) and network function virtualization (NFV) to construct isolated logical topologies with dedicated resource allocations and tailored control plane behaviors [2]. In the context of URLLC, a slice must guarantee deterministic end-to-end performance, which implies that both the network path and the associated compute functions must meet stringent service-level objectives. Traditional orchestration frameworks, such as those defined within the European Telecommunications Standards Institute (ETSI) NFV management and orchestration (MANO) stack, adopt a hierarchical structure where a centralized orchestrator maintains a global view of resources and coordinates domain-specific controllers. While this model ensures consistency, its responsiveness is limited by communication delays between the orchestrator and the

infrastructure, which can violate URLLC timing requirements when control decisions must be made at millisecond granularity.

The emergence of multi-access edge computing (MEC) provides an architectural vehicle for moving orchestration logic closer to the data plane. MEC platforms constitute an intermediate layer between centralized clouds and end devices, hosting both application-level services and infrastructure management functions at aggregation points such as base stations or local data centers [3]. By embedding orchestration competencies at the edge, the control loop latency for resource adjustment and fault recovery can be significantly reduced. However, distributing orchestration raises the question of how to partition decision-making authority between edge-local controllers and a central coordinating entity. The resulting architecture often takes the form of a federated hierarchy in which edge orchestrators manage a limited set of resources within their purview while reporting summary state information upward for inter-slice coordination and policy enforcement. This decomposition introduces trade-offs between local optimality and global efficiency, as edge orchestrators operating with incomplete information may make decisions that conflict with the objectives of other slices or administrative domains.

The URLLC slice further complicates this picture because of its sensitivity to both communication and computation dynamics. For instance, a predictive safety function in an autonomous vehicle may require a dedicated slice that spans multiple radio cells and edge nodes, each executing containerized inference models with strict deadlines. The edge orchestrator must coordinate not only the allocation of radio resource blocks and backhaul bandwidth but also the scheduling of graphics processing unit (GPU) time and memory within edge servers. This tight coupling between networking and computing resources necessitates an orchestration model that is natively cross-domain and capable of reasoning about joint resource states across traditionally siloed management systems.

3. System Architecture for Edge-Assisted Slice Orchestration

Designing an edge-assisted orchestration system for URLLC slices requires a careful articulation of functional components, their placement, and their interactions. A well-considered architecture distinguishes between the global slice manager, regional edge orchestrators, and local resource controllers. The global slice manager operates at a high aggregation level, responsible for slice admission control, inter-slice resource budgeting, and lifecycle management in accordance with the service provider's policy objectives [4]. It maintains a long-term view of resource utilization trends and customer demand patterns, enabling strategic decisions such as the pre-provisioning of compute capacity in regions expecting high URLLC demand surges. Conversely, the regional edge orchestrator operates with a short time horizon, focusing on real-time adaptation within its geographical scope. It manages virtualized network function (VNF) placement, scaling, and migration decisions, informed by telemetry data from the local radio and transport controllers.

A critical design dimension is the interface between the global and regional orchestrators. If the interface is too prescriptive, regional flexibility is crippled; if it is too loose, global policy consistency erodes. One promising approach involves intent-based interfaces where the global manager expresses high-level objectives and constraints, leaving the regional orchestrator to discover feasible local configurations that satisfy those intents [5]. For URLLC slices, an intent specification might encode maximum allowed latency, minimum reliability, and geographic coverage area without dictating which specific resources are to be used. The regional orchestrator then continuously resolves these intents into concrete actions through a combination of model-based optimization and learning-driven heuristics. This separation of

concerns enhances scalability by reducing the volume of detailed state information that must traverse the wide-area network, while preserving the ability to enforce consistent policies such as security isolation and energy consumption targets.

The local resource controllers constitute the lowest tier of this architecture, interacting directly with hypervisors, SDN switches, and radio schedulers. They expose abstract resource models to the regional orchestrator and execute the allocated configurations. In URLLC scenarios, local controllers incorporate fast-failure detection mechanisms that can trigger immediate countermeasures without waiting for higher-layer decisions. For example, if a local controller detects a wireless link degradation that threatens the reliability target of an active URLLC slice, it may autonomously reallocate a subset of radio resources using a pre-authorized fallback plan, simultaneously alerting the regional orchestrator to consider a longer-term topology adjustment. Such hierarchical delegation ensures that the system can meet the bounded response times required by URLLC while preserving overall coherence.

4. Intelligence at the Edge: AI-Driven Decision Making

The complexity of URLLC slice orchestration, characterized by high-dimensional state spaces, non-linear interactions between resources, and stochastic demand patterns, has motivated the integration of artificial intelligence into the orchestration plane. Machine learning techniques, particularly deep reinforcement learning (DRL), have shown potential for learning near-optimal policies in environments where analytical optimization becomes intractable [6]. When deployed at the edge, DRL agents can observe local telemetry streams and directly adjust slice resource allocations without requiring detailed system models. The application of these techniques to slice orchestration involves training agents to maximize cumulative rewards that reflect the satisfaction of service-level agreements (SLAs), energy efficiency, and the avoidance of slice degradation events.

A representative line of investigation has examined the use of proximal policy optimization (PPO) algorithms to assure quality of service in 5G network slicing [7]. The approach demonstrates that an agent can learn adaptive policies for adjusting resource boundaries between slices in response to fluctuating traffic patterns, achieving more reliable SLA compliance than static partitioning. When embedded within an edge orchestrator, such an agent operates on a localized view of the network state, which reduces state dimensionality and accelerates learning convergence. However, a purely localized learning paradigm risks myopic behavior, as an agent may optimize its local slice performance at the expense of cross-slice fairness or end-to-end application requirements that span multiple edges. Consequently, architectural mechanisms for federated or collaborative learning are being explored, wherein edge agents share policy parameters or value function approximations without exposing raw telemetry data, thereby balancing local optimization with global alignment.

Beyond reinforcement learning, other forms of edge intelligence contribute to URLLC reliability. Predictive analytics models, trained on historical infrastructure data, enable preemptive scaling of edge resources before latency spikes materialize. Anomaly detection algorithms can identify incipient hardware faults or security breaches, triggering proactive slice migration to healthy nodes. These models operate under strict inference latency constraints, which influences the selection of model architectures. Lightweight recurrent neural networks or transformer variants optimized for embedded deployment are often preferred over computationally heavy models. Furthermore, the non-stationarity inherent in mobile network environments demands continuous model retraining and validation pipelines.

This introduces a meta-orchestration challenge: the orchestration of the AI lifecycle itself, including data collection, training, model deployment, and monitoring, must be seamlessly integrated into the overall slice management framework without violating the very URLLC requirements it is meant to safeguard.

5. Reliability, Latency, and Resource Governance

The twin pillars of URLLC are reliability and latency, both of which must be satisfied jointly across the entire service chain. Reliability engineering in the context of sliced networks extends beyond traditional redundancy; it encompasses proactive fault prediction, fast rerouting, and stateful redundancy of VNF instances. At the edge, reliability strategies are constrained by the limited footprint of physical resources, which makes full one-to-one redundancy cost-prohibitive in many deployment scenarios. This limitation motivates shared protection schemes where backup resources are pooled across multiple slices and dynamically assigned based on failure probability estimates. Such pooling inherently conflicts with the isolation principle of slicing, as it introduces dependencies between slices that could propagate failures if not managed with rigorous access control and scheduling discipline.

Latency governance similarly requires a system-wide perspective that accounts for queuing delays at multiple hops, processing latencies at VNFs, and propagation delays over limited-length optical paths. Edge-assisted orchestration can mitigate queuing latencies by enforcing deterministic scheduling policies such as time-sensitive networking extensions adapted to virtualized environments. The orchestrator must co-schedule networking and computing resources to ensure that a packet traversing an edge node experiences not only bounded transmission delay but also predictable processing time within the server. This joint scheduling problem becomes tractable only when the orchestrator possesses fine-grained visibility into both network forwarding rules and compute thread allocations, which underscores the importance of breaking down organizational silos between network operations and information technology teams. From a governance standpoint, such convergence necessitates new role definitions and accountability frameworks that clearly attribute responsibilities for composite resource allocations.

The multi-domain nature of URLLC service delivery amplifies the complexity of resource governance. When a URLLC slice spans multiple administrative domains, such as a radio access provider, a transport network operator, and an edge data center operator, each stakeholder may have conflicting economic incentives and operational policies. Inter-domain orchestration must therefore be supported by robust negotiation protocols and policy enforcement points that can translate abstract SLAs into domain-specific configurations. Blockchain-based distributed ledgers have been proposed as mechanisms for establishing trust and accountability in these multi-stakeholder environments, but their consensus latency may conflict with URLLC control loop requirements. Alternative approaches relying on federated policy authorities that periodically synchronize domain-level status reports represent a pragmatic compromise between trustworthiness and speed.

6. Deployment, Infrastructure, and Scalability

Deploying edge-assisted orchestration for URLLC at scale demands a careful assessment of infrastructure constraints and lifecycle costs. Edge sites are characterized by heterogeneity in hardware capabilities, including varying levels of CPU, memory, and hardware acceleration support. Orchestrators must be able to reason about these heterogeneous capacities when placing slice components, ensuring that latency-critical VNFs are mapped to nodes with direct

access to data-plane forwarding accelerators or GPU resources. This placement problem is further complicated by the spatial locality requirements of URLLC applications, which may dictate that certain functions, such as safety-related perception models for autonomous vehicles, must reside within a specific geographic proximity to the served area. The orchestration logic must therefore incorporate location-awareness as a first-class dimension in its decision-making.

Scalability concerns also manifest in the control plane overhead generated by distributed orchestration. As the number of edge sites grows, the volume of telemetry data and the frequency of control interactions can overwhelm centralized elements if not carefully throttled. Hierarchical aggregation and data summarization techniques must be employed to compress operational data before it is relayed to global managers. At the same time, edge orchestrators must be capable of operating with degraded central connectivity, entering safe autonomous modes that guarantee baseline URLLC performance until connectivity is restored. Designing such graceful degradation modes is an often-underestimated aspect of edge orchestration that directly impacts the overall resilience of the system.

Another dimension of deployment concerns the onboarding and lifecycle management of AI models used in orchestration decisions. Models must be trained in environments representative of actual edge conditions, yet realistic training data may be scarce during early deployment phases. Transfer learning and simulation-to-reality domain adaptation techniques can mitigate this data scarcity, but they introduce additional validation overhead and potential for model errors to manifest under unseen conditions. The infrastructure must therefore support canary deployments and shadow testing of new orchestration policies in isolated subsets of the infrastructure, allowing operators to evaluate policy changes without jeopardizing live URLLC traffic. This requirement connects the technical orchestration framework to operational processes and skill sets, highlighting the need for organizational evolution alongside technological advancement.

7. Sustainability, Robustness, and Fairness

The orchestration of URLLC slices does not occur in a vacuum; it must align with broader objectives of environmental sustainability, long-term robustness, and equitable resource distribution among competing services. Energy consumption in edge infrastructure, driven by both computing and cooling demands, represents a growing concern as the density of edge deployments increases. Intelligent orchestration can contribute to sustainability by dynamically consolidating workloads during low-demand periods and powering down underutilized edge nodes. However, such consolidation must be executed with extreme caution for URLLC slices, because the deliberate reduction of active capacity may impair the ability to guarantee reliability during sudden demand spikes. This creates a tension between energy-saving policies and the hard performance guarantees required by mission-critical applications, demanding orchestration policies that can balance long-term environmental goals with instantaneous reliability imperatives [8].

Robustness refers to the system's capacity to maintain acceptable performance in the face of unexpected operational disturbances, including hardware failures, misconfigurations, and adversarial attacks. For URLLC slices, robustness demands not only redundancy but also architectural diversity that prevents correlated failures from simultaneously compromising multiple slices. Edge-assisted orchestration can enhance robustness by distributing failure recovery logic across multiple independent agents, reducing the risk of a single point of orchestration failure. However, distributed recovery introduces the possibility of oscillatory

behavior, where independent agents compete for the same backup resources after a major failure event. Avoiding such cascading instabilities requires careful design of backoff strategies and prioritization hierarchies, which must be evaluated under worst-case stress testing scenarios that go beyond typical operational profiles.

Fairness in resource allocation among slices is another critical governance dimension. When URLLC slices coexist with enhanced mobile broadband (eMBB) or massive machine-type communication (mMTC) slices, the orchestration framework must prevent URLLC demands from starving other service types, while still honoring the strict guarantees of URLLC. This multi-objective allocation problem cannot be reduced to simple priority scheduling because eMBB slices may also carry socially significant traffic such as public safety communications. Policy-driven approaches that assign relative weights to different slice classes and allow operators to define fairness criteria through tunable parameters offer a pragmatic path. The orchestration system can then optimize resource distribution subject to these fairness constraints, using lexicographic optimization or weighted proportional fairness formulations implemented via learning-based controllers. Ensuring that these fairness policies are transparent and auditable is essential for regulatory compliance and public trust in critical communication infrastructures.

8. Policy Implications and Standardization

The technical architecture of edge-assisted slice orchestration is deeply intertwined with regulatory frameworks and industry standardization efforts. Deterministic URLLC services will underpin an increasing number of safety-of-life applications, which may eventually attract regulation mandating minimum performance levels and accountability mechanisms similar to those imposed on utilities. This prospect elevates the importance of developing policy-aware orchestration frameworks that can produce verifiable evidence of compliance. For instance, an orchestration system might be required to generate tamper-proof logs demonstrating that a URLLC slice maintained its latency bound for a specific percentage of time over a regulatory period. Integrating such compliance reporting into the orchestration loop adds overhead but is increasingly seen as a non-negotiable requirement for sectors such as connected health and autonomous transport.

Standardization bodies, including the 3rd Generation Partnership Project (3GPP), ETSI, and the Internet Engineering Task Force (IETF), have made substantial progress in defining slice management interfaces and MEC service models. However, the domain of intelligent orchestration remains fragmented, with competing proposals for how AI components interact with standard MANO functions. Bridging this gap requires standardizing not just the static interfaces but also the dynamic interaction patterns, such as the handover of control between human operators and automated agents under degraded conditions. The notion of a shared responsibility model, which clearly delineates the scope of automated decision-making, is gaining traction as a governance framework that balances innovation with accountability. Moreover, interoperability between orchestration solutions from different vendors will be critical for end-to-end URLLC slices that traverse multi-vendor environments, making conformance testing and certification an important area for future policy development.

International policy coordination further influences the deployment of edge-based URLLC orchestration, particularly concerning spectrum allocation and data sovereignty. Edge nodes processing sensor data for URLLC applications may raise cross-border data flow issues that national regulations restrict. Orchestration frameworks must therefore incorporate data residency constraints into their placement logic, ensuring that slice components handling

regulated data operate only within permissible geopolitical boundaries. This intersection of technical orchestration with legal geography represents a fertile ground for interdisciplinary research, connecting systems engineering with international law and digital policy. The long-term viability of URLLC services will depend as much on the sophistication of these socio-technological governance mechanisms as on the underlying radio and computing hardware.

9. Conclusion

The realization of ultra-reliable low-latency communications in 5G networks represents one of the most challenging systems engineering undertakings of the current decade. Edge-assisted network slice orchestration emerges as a pivotal enabler, offering the architectural flexibility to distribute control intelligence across the core-to-edge continuum. This paper has examined the orchestration challenge from a system-level perspective, analyzing how the decomposition of orchestration functions across global, regional, and local tiers can reconcile centralized policy coherence with the stringent timing requirements of URLLC. We identified structural trade-offs in reliability engineering, highlighting the tension between resource isolation and shared protection, and discussed the integration of AI techniques, including deep reinforcement learning, to automate complex resource allocation decisions under uncertainty.

The analysis further revealed that sustainable and fair URLLC orchestration extends beyond technical metrics to encompass energy consumption, cross-slice equity, and resilience against correlated failures. Deployment at scale demands addressing infrastructure heterogeneity, control plane overhead, and the lifecycle management of AI models in production settings. Importantly, the paper underscored the inseparable link between technical design and policy frameworks, arguing that future orchestration platforms must embed compliance verification, data sovereignty awareness, and transparent fairness governance as integral features rather than afterthoughts. As 5G networks evolve toward 6G, the foundational principles discussed here will become even more critical, necessitating continued interdisciplinary collaboration among network engineers, AI researchers, policymakers, and legal scholars. By constructing orchestration systems that are intelligent yet governed, agile yet robust, and performant yet fair, the community can unlock the full societal potential of URLLC while safeguarding the trust and resilience upon which critical digital infrastructures depend.

References

1. Foukas, X., Patounas, G., Elmokashfi, A., & Marina, M. K. (2017). Network slicing in 5G: Survey and challenges. *IEEE Communications Magazine*, 55(5), 94–100.
2. Ordóñez-Lucena, J., Ameigeiras, P., Lopez, D., Ramos-Munoz, J. J., Lorca, J., & Folgueira, J. (2017). Network slicing for 5G with SDN/NFV: Concepts, architectures, and challenges. *IEEE Communications Magazine*, 55(5), 80–87.
3. Mao, Y., You, C., Zhang, J., Huang, K., & Letaief, K. B. (2017). A survey on mobile edge computing: The communication perspective. *IEEE Communications Surveys & Tutorials*, 19(4), 2322–2358.
4. Popovski, P., Nielsen, J. J., Stefanovic, C., de Carvalho, E., Strom, E. G., Trillingsgaard, K. F., & Sørensen, T. B. (2018). Wireless access for ultra-reliable low-latency communication: Principles and building blocks. *IEEE Network*, 32(2), 16–23.
5. Taleb, T., Samdanis, K., Mada, B., Flinck, H., Dutta, S., & Sabella, D. (2017). On multi-access edge computing: A survey of the emerging 5G network edge cloud architecture and orchestration. *IEEE Communications Surveys & Tutorials*, 19(3), 1657–1681.

6. Zhang, C., Patras, P., & Haddadi, H. (2019). Deep learning in mobile and wireless networking: A survey. *IEEE Communications Surveys & Tutorials*, 21(3), 2224–2287.
7. Li, Q. (2026). QoS Assurance Mechanism for 5G Network Slicing Based on the Deep Reinforcement Learning PPO Algorithm. *arXiv preprint arXiv:2605.03345*.
8. Afolabi, I., Taleb, T., Samdanis, K., Ksentini, A., & Flinck, H. (2018). Network slicing and softwarization: A survey on principles, enabling technologies, and solutions. *IEEE Communications Surveys & Tutorials*, 20(3), 2429–2453.
9. Bennis, M., & Debbah, M. (2018). Ultra-reliable and low-latency communication in 5G downlink: A machine learning approach. *IEEE Transactions on Communications*, 66(6), 2501–2514.
10. Simsek, M., Aijaz, A., Dohler, M., Sachs, J., & Fettweis, G. (2016). 5G-enabled tactile internet. *IEEE Journal on Selected Areas in Communications*, 34(3), 460–473.
11. Samdanis, K., Costa-Perez, X., & Sciancalepore, V. (2016). From network sharing to multi-tenancy: The 5G network slice broker. *IEEE Communications Magazine*, 54(7), 32–39.
12. Giust, F., Cominardi, L., & Bernardos, C. J. (2016). Multi-access edge computing: The driver behind the wheel of 5G-connected cars. *IEEE Network*, 30(5), 56–62.
13. Bega, D., Gramaglia, M., Banchs, A., Sciancalepore, V., Samdanis, K., & Costa-Perez, X. (2017). Optimising 5G infrastructure markets: The business of network slicing. In *Proceedings of IEEE INFOCOM 2017* (pp. 1–9). IEEE.
14. Yan, M., Feng, G., & Qin, S. (2019). A multi-agent reinforcement learning method for network slicing in wireless networks. *IEEE Transactions on Vehicular Technology*, 68(12), 12336–12349.
15. Kang, J., Xiong, Z., Niyato, D., Wang, P., & Kim, D. I. (2019). Software-defined wireless network slicing: Resource sharing and lifecycle management. *IEEE Communications Magazine*, 57(4), 68–74.
16. Meshkova, Z., Riihijarvi, J., & Mahonen, P. (2018). Policy-based coordination and management for 5G network slices. *IEEE Network*, 32(6), 130–137.
17. Xu, J., Ota, K., & Dong, M. (2018). Saving energy on 5G networks: A survey of green communication approaches. *IEEE Communications Surveys & Tutorials*, 20(3), 1979–2011.
18. Mijumbi, R., Serrat, J., Gorricho, J. L., Latre, S., Charalambides, M., & Lopez, D. (2016). Management and orchestration challenges in network functions virtualization. *IEEE Communications Magazine*, 54(1), 98–105.
19. Bousard, M., Papillon, D., Zahariev, A., Keslassy, I., & Duarte, O. (2018). 5G network slicing: A security and trust perspective. In *2018 IEEE 5G World Forum* (pp. 455–460). IEEE.
20. Yrjola, S., Ahokangas, P., & Matinmikko, M. (2019). Evaluation of recent spectrum sharing concepts from business and scalability perspectives. *Telecommunications Policy*, 43(4), 101–110.