

AI-NATIVE NETWORK SLICING FOR 6G NETWORKS

Wen Wu, Conghao Zhou, Mushu Li, Huaqing Wu, Haibo Zhou, Ning Zhang, Xuemin (Sherman) Shen, and Weihua Zhuang

ABSTRACT

With the global rollout of fifth generation (5G) networks, it is necessary to look beyond 5G and envision 6G networks. 6G networks are expected to have space-air-ground integrated networks, advanced network virtualization, and ubiquitous intelligence. This article presents an artificial intelligence (AI)-native network slicing architecture for 6G networks to enable the synergy of AI and network slicing, thereby facilitating intelligent network management and supporting emerging AI services. AI-based solutions are first discussed across the network slicing life cycle to intelligently manage network slices (i.e., *AI for slicing*). Then network slicing solutions are studied to support emerging AI services by constructing AI instances and performing efficient resource management (i.e., *slicing for AI*). Finally, a case study is presented, followed by a discussion of open research issues that are essential for AI-native network slicing in 6G networks.

INTRODUCTION

Compared to existing wireless networking including the fifth generation (5G), 6G is more than an improvement of key performance indicator (KPI) requirements, such as increased data rates, enhanced network capacity, and low latency. 6G networks are envisioned to have the following unique features. First, space networks, such as low Earth orbit (LEO) satellites, air networks (e.g., unmanned aerial vehicles, UAVs), and ground networks, such as cellular base stations (BSs), are integrated into a space-air-ground integrated network (SAGIN) to provide global coverage and on-demand services [1]. Second, resource virtualization using network slicing techniques and end user virtualization using digital twin techniques can facilitate *advanced network virtualization* to provide flexible network management [2, 3].¹ Third, intelligence penetrates every corner of networks, ranging from end users through the network edge to the remote cloud, which results in *ubiquitous intelligence*. A number of network nodes are endowed with built-in artificial intelligence (AI) functionalities, thereby not only facilitating intelligent network management but also fostering AI services (e.g., deep neural network based applications). Hence, 6G networks are expected to create a new wireless networking ecosystem that brings societal and economic benefits.

The 6G networks will support a diverse set of services with different quality of service (QoS) requirements, such as multisensory extended real-

ity and hologram video streaming. To support diversified services as established in 5G networks, network slicing is a potential approach to construct multiple logically isolated virtual networks (i.e., slices) for different services on top of the common physical network [4]. The QoS requirements of different services can be guaranteed via cost-effective slice management strategies ranging from preparation, planning, and operation phases in the network slicing life cycle.

Developing network slicing schemes faces many challenges in 6G networks due to their unique features. First, managing slices over space, air, and ground network segments in the SAGIN requires judicious coordination of heterogeneous network segments. Moreover, 6G networks need to support a variety of new services while satisfying their different and stringent QoS requirements, which further complicates slice management. Hence, it is paramount to develop intelligent slice management solutions in 6G networks. Second, fueled by powerful computing capability and advanced AI techniques, ubiquitous intelligence is fostering abundant AI services with new QoS requirements, such as data quality, inference accuracy, and training latency. Hence, it is necessary to construct customized network slices to support the emerging AI services in 6G networks.

In this article, we propose an *AI-native* network slicing architecture for 6G networks to facilitate intelligent network management while supporting emerging AI services. AI-native means that, as a built-in component in the network slicing architecture, AI exists not only in the software-defined networking (SDN) controller for managing network slices, but also in network slices as services for end users. Hence, the synergy of AI and network slicing in the proposed architecture is two-fold. On one hand, AI techniques can be applied to manage network slices, namely *AI for slicing*. The network slicing life cycle including preparation, planning, and operation phases is introduced, along with specifying AI-based solutions for each phase. In addition, the detailed procedure of information exchange among end users, access points, and the SDN controller is presented. On the other hand, network slicing can be applied to construct customized network slices for various AI services, namely *slicing for AI*. Potential approaches such as AI instance construction and efficient resource management for AI services are introduced.

The remainder of this article is organized as follows. Expected features of 6G networks are

¹ End user virtualization is to virtualize end users in network operation and management by characterizing users' behaviours and status (e.g., service demands and QoS satisfaction), which can be achieved using digital twin concepts.

discussed, and then the AI-native network slicing architecture is proposed. The basic ideas of AI for slicing and slicing for AI are presented, respectively. A case study is presented. The research directions are identified, followed by the conclusion.

AI-NATIVE NETWORK SLICING FOR 6G NETWORKS

NETWORK SLICING

Network slicing is an emerging technology to support diversified applications in a cost-effective manner [4, 5]. The concept of network slicing can be traced back to the late 1980s [6]. Nowadays, network slicing is a key technology in 5G networks, supported by network function virtualization (NFV) and SDN techniques. Specifically, NFV enables virtualized resources and network functions for flexible resource management, while SDN facilitates centralized network management for network optimization. In 5G networks, network slicing has been defined in the 3rd Generation Partnership Project (3GPP) Release 15 [7]. Moreover, in coming 6G networks, network slicing will continue to evolve and play an increasingly important role.

The basic idea of network slicing is to create multiple logically isolated network slices on top of the common physical infrastructure, which can achieve flexible and adaptive network management. Its benefits are three-fold:

1. Multi-tenancy: Multiple virtual networks can share the common physical infrastructure, thus reducing capital expenditures in the network deployment.
2. Service isolation: Multiple slices are constructed for different services via judicious resource management, such that service level agreements of different slices can be effectively guaranteed.
3. Flexibility: Network slicing can support flexible network management, as slices can be created, modified, or deleted on demand.

FEATURES OF 6G NETWORKS

From 5G to 6G, it is in general expected that KPI requirements will be increased by at least an order of magnitude. According to a recent white paper [8], the KPI requirements of 6G networks include 1 Tb/s peak data rate, 20–100 Gb/s user experienced data rate, 0.1 ms end-to-end latency, 10 million devices/km², and near 100 percent coverage. Such KPI requirements demand several candidate technologies, such as THz communications and AI [6]. The 3GPP working group will discuss 6G candidate techniques by the end of 2026, and the first 6G standard is expected to debut by 2030.

Distinguished from 5G networks, 6G networks have several features:

- SAGIN: While current ground networks provide good coverage in highly populated areas, 6G needs to provide universal coverage, including in rural, remote, and sparsely populated areas. To achieve this goal, 6G will exploit the altitude dimension. Space, air, and ground network segments are integrated into the SAGIN [1, 9], which can provide global coverage, facilitate on-demand services, and support high-rate low-delay services;
- Diversified services: Many services have stringent QoS requirements in different dimensions. Mobile

virtual reality (VR) and hologram video streaming applications require a high data rate; for example, the uplink data rate of mobile VR is up to 5 Gb/s. Other applications may require ultra-high reliability, such as autonomous driving, industrial control systems, and robot/UAV swarm; for example, the required reliability of autonomous driving is up to 99.999 percent [10]

- Ubiquitous intelligence: With caching capability, a large amount of data can be stored in the network. In addition, with the development of AI techniques, edge computing, and device computing, intelligence is pushed from the remote cloud to the network edge and end users. As such, AI will be integrated into 6G networks for intelligent network management by directly learning from extensive data in the network. Moreover, ubiquitous intelligence will foster a number of AI services in which AI is provided as services.

AI-NATIVE NETWORK SLICING ARCHITECTURE

These features impose new challenges on developing network slicing schemes for 6G networks. First, the SAGIN not only increases the number of integrated network segments but also introduces extra dynamics on network resource availability due to satellite mobility and UAV manoeuvrability. Network slicing schemes should accommodate the large-scale SAGIN while taking dynamic resource availability into account. Moreover, supporting diversified services with stringent QoS requirements further complicates network slicing scheme design. Second, ubiquitous intelligence facilitates many emerging AI services that will be prevalent in 6G networks. Different from conventional services, facilitating AI services requires multiple steps, including collecting high-quality data samples, training satisfactory AI models, and performing low-latency model inference, which should meet diverse QoS requirements. How to satisfy such diverse QoS requirements for AI services remains a challenging issue.

To address the above challenges, an AI-native network slicing architecture for 6G networks is presented. As shown in Fig. 1, the architecture aims to integrate SAGIN and ubiquitous intelligence and support diverse services with stringent QoS requirements. Compared to network slicing for 5G networks, the proposed architecture has two new characteristics. First, AI is integrated into SDN controllers to realize intelligent network slicing such that a number of network slices with stringent QoS requirements can be managed efficiently and cost-effectively via AI techniques, which is referred to as *AI for slicing*. Second, emerging AI services are supported by network slicing. In addition to network slices for conventional services, new network slices are constructed for AI services on top of the common physical infrastructure, which is referred to as *slicing for AI*.

Two types of SDN controllers are deployed in the proposed architecture. One is the centralized SDN controller located at the cloud, which is to manage network slices. The other is the local SDN controller located at access points, which is to schedule resources to end users within each network slice. The SDN controller in the following refers to the centralized SDN controller unless

With caching capability, a large amount of data can be stored in the network. In addition, with the development of AI techniques, edge computing, and device computing, intelligence is pushed from the remote cloud to the network edge and end users.

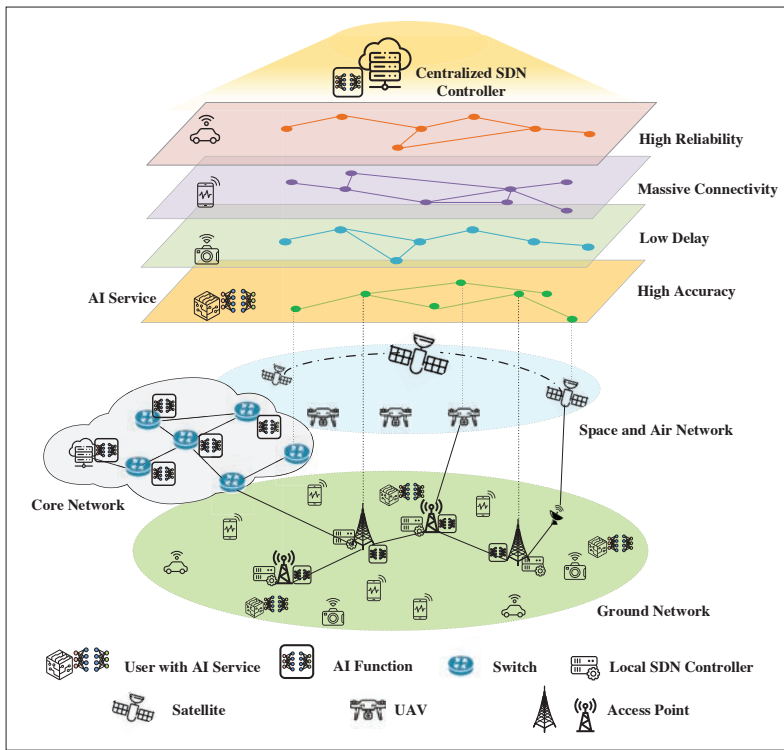


FIGURE 1. An illustration of the AI-native network slicing architecture for 6G networks.

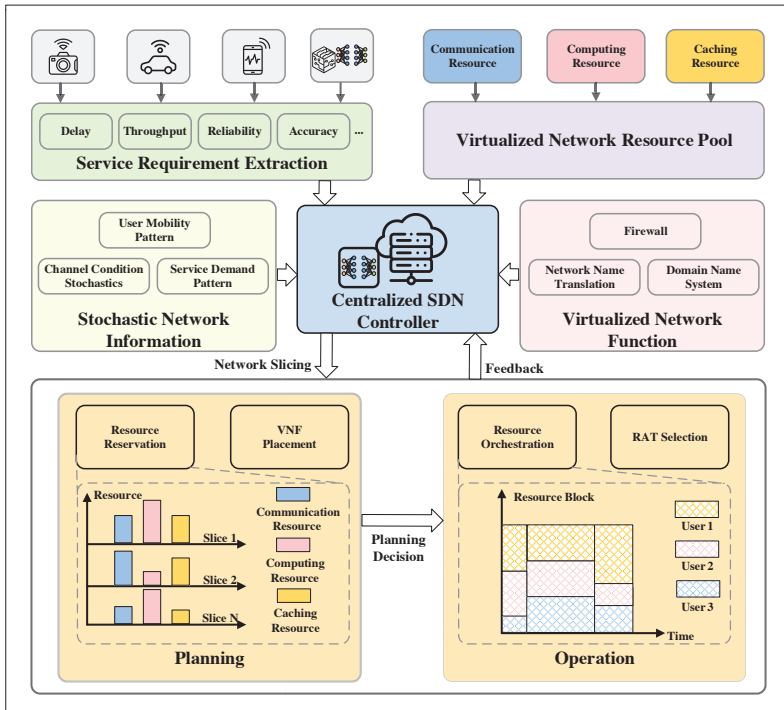


FIGURE 2. An illustrative example of the network slicing life cycle, including the preparation, planning, and operation phases.

otherwise stated. In the following, we illustrate the basic ideas of AI for slicing and slicing for AI, respectively.

AI FOR SLICING

In this section, we introduce the network slicing life cycle in three phases and then investigate potential AI solutions for each phase. Next,

the corresponding procedure of information exchange in AI for slicing is discussed.

NETWORK SLICING LIFE CYCLE

The network slicing life cycle consists of three phases: *preparation*, *planning*, and *operation*, as shown in Fig. 2. The centralized SDN controller is in charge of the preparation and planning phases, while the operation phase is coordinated by local SDN controllers.

Preparation Phase: This phase is to construct and configure network slices based on service requirements, data traffic, user information, and virtual network resource availability. To achieve the goal, the SDN controller conducts the following tasks:

- **Service requirement extraction:** This task is to classify services by extracting their QoS requirements, such as service delay, service priority, throughput, and reliability. 3GPP has standardized specific service/slice type values for classified services, such as enhanced mobile broadband, ultra-reliable low-latency communications, and massive machine-type communications services [7].
- **Network resource and function virtualization:** Network resources, such as communication, computing, and caching resources, are pooled into virtualized resource blocks via advanced resource virtualization techniques. Similarly, network functions, such as firewall, network name translation, and domain name system, are separated from dedicated hardware network functions into virtualized network functions (VNFs). Through virtualization, the SDN controller can flexibly manage network resources and functions.

Once these tasks are completed, the SDN controller can construct network slices for each admitted slice request.

Planning Phase: This phase aims to reserve network resources to slices for service provisioning. The planning phase operates over a long timescale. Time is partitioned into multiple planning periods (windows) for each slice. The duration of each planning window depends on service demand and network dynamics, whose value ranges from several minutes to several hours. To achieve the goal, the following two steps are conducted in the planning phase:

- **Service and network information collection:** Benefiting from the global control functionality of the SDN controller, extensive network information can be collected from underlying physical networks, such as service demands, stochastic channel conditions, and user mobility patterns. The collected information is utilized for the following resource reservation decision making.
- **Resource reservation:** At the beginning of each planning window, the SDN controller adjusts the amount of reserved network resources for each slice based on the monitored slice performance. The reserved virtualized network resources of each slice are mapped to the physical network. At the end of each planning window, some system information is fed back to the SDN controller, such as resource utilization, system performance, and service level agreement satisfaction. Based on the feedback information, the SDN controller can adjust resource reservation decisions to accommodate dynam-

ic network environments while guaranteeing QoS requirements.

Operation Phase: This phase is to schedule the service of a slice using the reserved resources for subscribed end users. The operation phase works over a much shorter timescale (e.g., 100 ms) than that of the planning phase. Specifically, under the coordination of the centralized SDN controller, local SDN controllers allocate network resources to end users in each slice according to their real-time data traffic. The operation decisions include selecting radio access technology (RAT), determining user association with specific radio access points, deciding proper protocol and associated parameters, and orchestrating resources among end users.

Roles of AI in Network Slicing: Although network slicing can facilitate service provisioning, managing a number of network slices incurs significant network management cost, especially in 6G networks. As shown in Fig. 3, AI-based network slicing is a potential solution in which AI plays different roles in different network slicing phases.

AI for preparation: In the preparation phase, AI needs to perform two tasks.

1. Service demand prediction: Based on historical data, service demand can be predicted via AI techniques, such as recurrent neural networks. Prior studies show that the service demand and resource usage of a slice can be accurately predicted [11]. The prediction results can be utilized for decision making in the planning phase.
2. Slice admission: The SDN controller admits slices to maximize network resource utilization considering resource availability and service demands. As the slice admission decision is binary, this problem is deemed as an integer optimization problem. In large-scale networks with complex resource availability distribution, conventional optimization solutions become complicated and intractable, while AI-based solutions have potential.

AI for planning: In the planning phase, AI can perform two tasks:

1. VNF placement: The SDN controller deploys VNFs to support services in the network. The resources allocated for VNFs should be dynamically adjusted for time-varying service demands to guarantee service delay requirements. Deep learning methods can be applied to enhance resource utilization in dynamic network environments.
2. Resource reservation: The SDN controller reserves resources for different slices based on their service demands. Since data traffic loads are time-varying, the resource reservation should be adaptive to dynamic real-time demands, which can be addressed via reinforcement learning (RL) methods, such as deep deterministic policy gradient (DDPG).

AI for operation: Two exemplary operation tasks are:

1. Resource orchestration: The reserved resources of a slice are allocated to end users. The decisions are determined based on real-time user mobility, service demands, and so on. To efficiently utilize resources, RL methods can be applied for dynamic resource orchestration;
2. RAT selection: To maximize system utility, an optimal RAT is selected among multiple can-

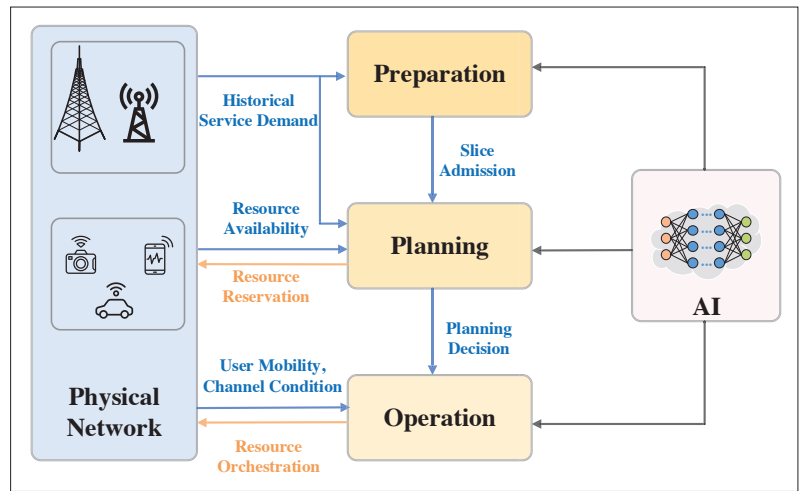


FIGURE 3. The considered AI-based network slicing solution in which AI plays different roles in the preparation, planning, and operation phases.

didate RATs for each end user. Due to user mobility, user-perceived service performance of an RAT is stochastic. This problem can be addressed by multi-armed bandit methods (e.g., contextual bandit).

PROCEDURE OF INFORMATION EXCHANGE

The AI for slicing procedure involves the information exchange among end users, access points, and the SDN controller. The procedure is illustrated in Fig. 4 with the following steps:

1. Access points collect user-level stochastic information, such as end users' service demand patterns, mobility patterns, and stochastic channel conditions.
2. Access points translate the user-level information into desired service-level information. For example, user density information can be obtained from processing user location information, and AI techniques can be used for such data abstraction, fusion, and analysis; . RNN can be applied for service demand prediction.
3. The processed service-level information is delivered to the SDN controller.
4. The SDN controller runs AI-based planning algorithms to make decisions based on the collected service-level information.
5. The determined planning decisions are sent back to all access points.
6. Access points enforce the received planning decisions (e.g., reserving network resources for corresponding slices).
7. End users in service report their real-time information to their associated access points, such as real-time service demands, channel conditions, and task data sizes.
8. Access points run the AI-based operation algorithm to allocate resources for end users based on real-time user-level information.
9. Service requests from end users are supported with the allocated network resources. For example, computation tasks can be offloaded to access points using communication resource and then processed using computing resources. For each operation slot within a planning window, Steps 7–9 are repeated;
10. Access points monitor slice performance in the network given the enforced planning decisions by measuring end users' satisfaction rates across

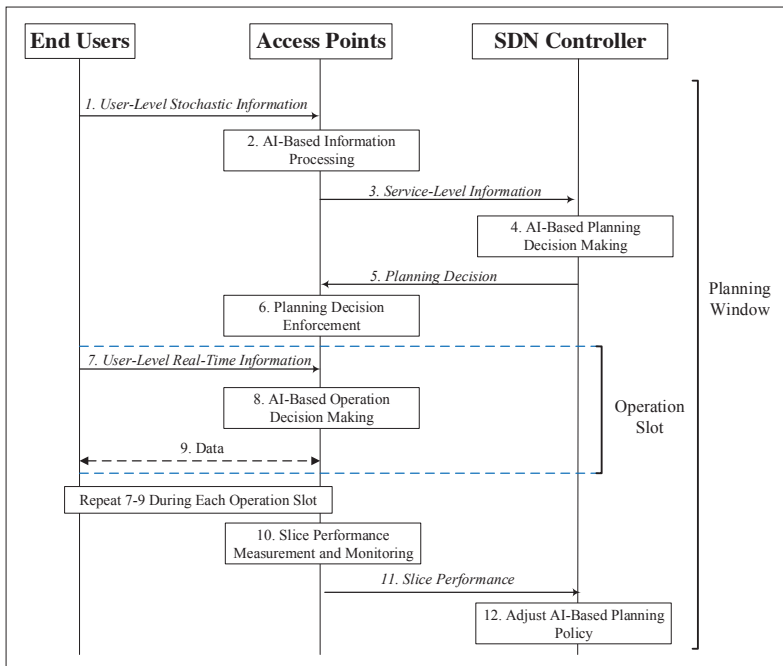


FIGURE 4. Procedure of information exchange in AI for slicing.

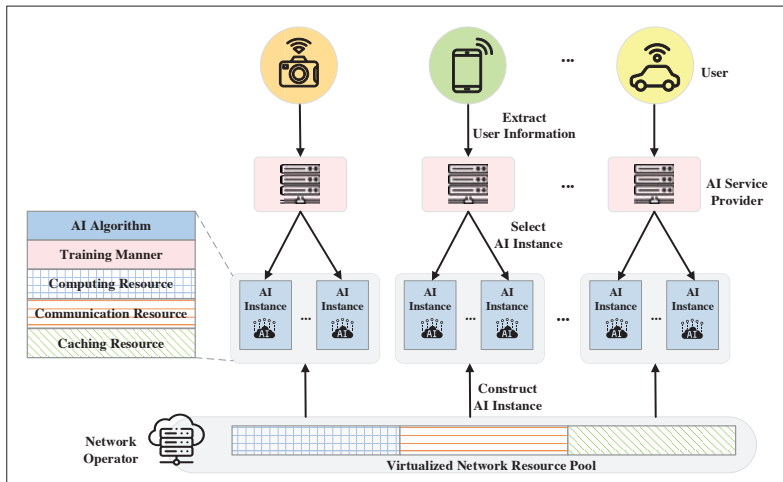


FIGURE 5. The conceptual AI instance management framework for AI services.

- all operation slots within a planning window.
11. Access points report network performance to the SDN controller.
 12. The SDN controller makes the planning decision for the next planning window and adjusts the planning policy based on the feedback information.

Note that in the preceding procedure, Steps 1–6 and Steps 10–12 are in the network planning phase, and Steps 7–9 are in the operation phase.

SLICING FOR AI

Slicing for AI utilizes network slicing to support AI services while satisfying QoS requirements. Potential solutions include constructing and selecting AI instances and efficient resource management in the AI service life cycle.

AI INSTANCE

There are diversified implementation options for supporting AI services. An AI service can be implemented via different kinds of algorithms, training manners, and network resource allocation. For

example, objective detection services can be implemented via ResNet32, Inception-v3, AlexNet, or VGG16 algorithms. Hence, the primary issue of supporting an AI service is to determine an appropriate implementation option in the network.

We introduce the concept of *AI instance* to address the issue, as shown in Fig. 5. An AI instance of an AI service represents an implementation option for an AI service. The basic idea is to construct multiple candidate implementation options and then select an appropriate one based on network environments. The procedure of the conceptual AI instance management framework consists of two steps:

1. AI instance construction: The network operator constructs multiple candidate AI instances for each AI service based on available network resources and service requirements. An AI instance may include:
 - The AI algorithm, which specifies the implementation algorithm and the corresponding neural network architecture
 - The training manner of the AI algorithm (e.g., centralized or distributed training)
 - The amount of required network resources
2. AI instance selection: In this step, the AI service provider selects an appropriate AI instance among candidate AI instances based on user service preference (e.g., privacy preservation preference). If an AI instance is selected, the AI service will be executed using the AI algorithm and the corresponding required amount of network resources by the AI instance. In summary, the idea of AI instance provides flexibility for AI service management.

The network operator can dynamically manage AI services via creating, modifying, and deleting candidate AI instances on demand. Note that AI instances are not limited to neural networks, and they can include other types of AI algorithms such as support vector machines and decision trees.

RESOURCE MANAGEMENT IN AI SERVICE LIFE CYCLE

Running AI services includes three stages: data collection, model training, and model inference (i.e., the AI service life cycle) [12, 13]. Specifically, *data collection* is to collect data via communication links, and the collected data can be stored in network edge servers. Based on the collected data, an AI model can be trained in the *model training* stage. The model training can be implemented in either a centralized or distributed manner. For example, multiple devices can work collaboratively to train a global model via federated learning. Next, well-trained AI models are deployed to execute specific computation tasks; this is referred to as *model inference*. The model inference can be performed in multiple ways. For example, device-edge collaborative inference approaches can allocate and process computation tasks at different network nodes to achieve low inference latency.

The performance of an AI service depends on all three stages of the AI service life cycle. For example, model inference accuracy depends on multiple factors, such as the quality of the collected data, the number of training iterations, and the approach of model inference. Meanwhile, all three stages consume multi-dimensional network resources. As a result, to optimize the performance

of AI services, network resources should be jointly allocated for these three stages. The reserved network resources in AI slices should be further allocated to these three stages to satisfy their corresponding QoS requirements.

CASE STUDY

In this section, a case study is provided on AI-assisted resource reservation, aimed at reducing long-term overall system cost.

CONSIDERED SCENARIO

We consider an air-ground integrated network for providing autonomous driving services to vehicles traversing a highway segment. For the considered highway segment with a length of 2 km, two BSs are uniformly deployed along the highway with a separation distance of 1 km, and one UAV is deployed in the centre hovering at a height of 100 m. When an autonomous vehicle is driving on the highway, extensive computation-intensive tasks must be processed. For prompt task processing, all access points are equipped with edge computing servers, and vehicles can associate with the nearest access point and upload their computation tasks. We consider a delay-sensitive autonomous driving service, object detection, whose delay requirement is 100 ms for road safety [14]. The data size and computation intensity of a task are set to 0.6 Mbit and 6×10^8 cycles, respectively. The service delay is characterized by the queuing theory since task arrivals are assumed to follow a Poisson process with rate $\lambda = 1$ packet/s. To guarantee the service delay requirement, a network slice is constructed in which spectrum and computing resources are reserved.

Resource reservation decisions are made to minimize the overall system cost considering vehicle traffic dynamics. The overall system cost is defined as $C = \sum_{t=1}^T (\omega_r C_r^t + \omega_s C_s^t + \omega_d C_d^t)$, which is a weighted summation of three cost components across all T planning windows.

1. Resource reservation cost C_r^t accounts for the amount of the reserved spectrum and computing resources at BSs at planning window t . The spectrum resource is allocated in a unit of sub-carrier of 5 MHz, and the computing resource is allocated in a unit of virtual machine (VM) instance with a processing rate of 10×10^9 cycles/s.
2. Slice reconfiguration cost C_s^t accounts for the difference between two consecutive resource reservation decisions [15].
3. Delay requirement violation penalty C_d^t refers to the penalty once service delay exceeds the delay requirement. These weight parameters are set to $\omega_r = 1$, $\omega_s = 20$, and $\omega_d = 200$, respectively. The planning window size is set to one hour.

We propose a DDPG-based solution to minimize the overall system cost [15]. In this solution, both actor and critic networks are fully connected neural networks with four layers, the numbers of neurons in two hidden layers are 128 and 64, respectively, and their learning rates are set to 2×10^{-4} and 2×10^{-3} , respectively. For performance comparison, we adopt an optimization-based solution, named *myopic resource reservation*, in which network resources are reserved to minimize the resource reservation cost at each planning window while satisfying the delay requirement.

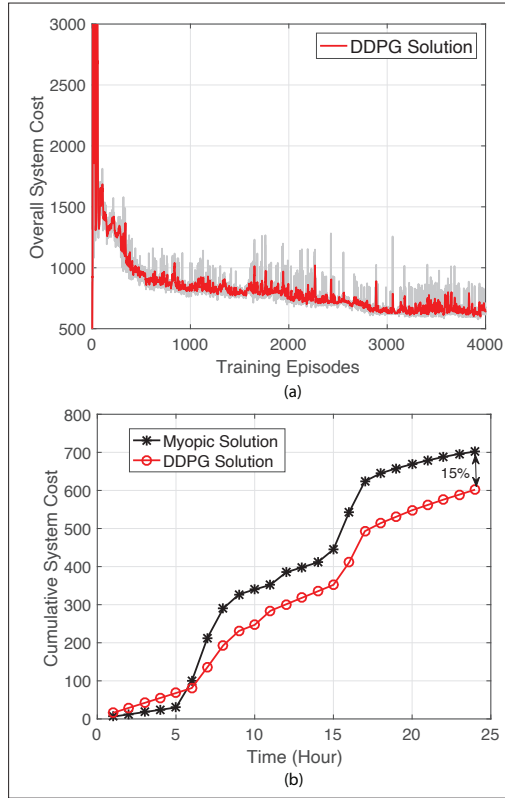


FIGURE 6. Performance evaluation of the proposed DDPG-based resource reservation solution: a) convergence performance; b) cumulative system cost within one day.

SIMULATION RESULTS

We evaluate the performance of the proposed DDPG-based solution based on real-world highway vehicle traffic flow trace collected by Alberta Transportation.² As shown in Fig. 6a, we first present the convergence performance of the proposed DDPG-based solution. A five-point moving average is applied to process raw simulation points to highlight the convergence trend (i.e., red curve). It can be seen that the DDPG-based resource reservation solution has converged after 4000 training episodes.

Next, as shown in Fig. 6b, the cumulative system cost within one day is presented. It can be observed that the DDPG-based solution can reduce the cumulative overall system cost within one day by around 15 percent as compared to the myopic solution. The reason is that the proposed DDPG-based solution is able to minimize the long-term overall system cost, while the myopic solution minimizes the short-term system cost, which incurs prohibitive slice reconfiguration cost due to frequent adjustment of network resource reservation in highly dynamic vehicular networks. The simulation results show that the proposed AI-based resource reservation solution can achieve a low system cost.

OPEN RESEARCH ISSUES

In the following, we discuss some open research issues pertaining to AI-native network slicing.

JOINT DESIGN OF NETWORK PLANNING AND OPERATION

Planning and operation are performed and coupled in two different timescales. Specifically, the

The cumulative system cost within one day is presented. It can be observed that the DDPG-based solution can reduce the cumulative overall system cost within one day by around 15 percent as compared to the myopic solution.

² Alberta Transportation; <http://www.transportation.alberta.ca/mapping/>.

The slicing for AI is to construct customized network slices to better accommodate various emerging AI services. To accelerate the pace of AI-native network slicing architecture development, extensive research efforts are required, such as in the identified research directions.

planning phase is performed in a long timescale (e.g., minute level) to reserve resources for different slices based on service demands, while the operation phase is performed in a short timescale (e.g., sub-second level) to allocate the reserved resources to on-demand users within each slice. Achieving the optimal network slicing performance requires a joint optimization design of planning and operation.

DATA MANAGEMENT FRAMEWORK

The cornerstone of AI-native network slicing is abundant data that can be used for AI model training. In 6G networks, data is widely distributed in the network. Due to limited communication resources, the cost of collecting a large amount of data cannot be neglected. In addition, the collected data needs to be processed to mine valuable information for network management. For example, abundant historical behaviour data from individual users can be analyzed to predict spatio-temporal service demand distributions. Hence, establishing a data management framework to collect and analyze data is necessary.

PREDICTION-EMPOWERED NETWORK SLICING

With the development of advanced AI technologies, the data traffic in the network can be predicted. How to effectively leverage the power of prediction for network slicing is an interesting topic. Since the prediction is imperfect, the prediction error may degrade the performance of network slicing. How to evaluate the impact of prediction errors on system performance and to develop corresponding solutions are important research issues.

CONCLUSION

In this article, we have proposed the AI-native network slicing architecture to facilitate intelligent network management and support AI services in 6G networks. The architecture aims to enable the synergy of AI and network slicing. AI for slicing helps reduce network management complexity, while adapting to dynamic network environments by exploiting the capability of AI in network slicing. Slicing for AI is to construct customized network slices to better accommodate various emerging AI services. To accelerate the pace of AI-native network slicing architecture development, extensive research efforts are required, such as in the identified research directions.

ACKNOWLEDGMENTS

This work was financially supported by the Natural Sciences and Engineering Research Council (NSERC) of Canada. The authors would like to thank Jie Gao, Qihao Li, and Kaige Qu for many valuable discussions and suggestions throughout the work.

REFERENCES

- [1] N. Zhang et al., "Software Defined Space-Air-Ground Integrated Vehicular Networks: Challenges and Solutions," *IEEE Commun. Mag.*, vol. 55, no. 7, July 2017, pp. 101–09.
- [2] X. Shen et al., "Holistic Network Virtualization and Pervasive Network Intelligence for 6G," *IEEE Commun. Surveys & Tutorials*, to be published, 2022. DOI: 10.1109/COMST.2021.3135829.
- [3] R. Minerva, G. M. Lee, and N. Crespi, "Digital Twin in the IoT Context: A Survey on Technical Features, Scenarios, and Architectural Models," *Proc. IEEE*, vol. 108, no. 10, Oct. 2020, pp. 1785–1824.

- [4] X. Shen et al., "AI-Assisted Network-Slicing Based Next-Generation Wireless Networks," *IEEE Open J. Vehic. Tech.*, vol. 1, no. 1, Jan. 2020, pp. 45–66.
- [5] W. Zhuang et al., "SDN/NFV-Empowered Future IoV with Enhanced Communication, Computing, and Caching," *Proc. IEEE*, vol. 108, no. 2, Feb. 2020, pp. 274–91.
- [6] X. You et al., "Towards 6G Wireless Communication Networks: Vision, Enabling Technologies, and New Paradigm Shifts," *Sci. China Info. Sci.*, vol. 64, no. 1, Jan. 2021, pp. 1–74.
- [7] A. Kaloxylis, "A Survey and an Analysis of Network Slicing in 5G Networks," *IEEE Commun. Stds. Mag.*, vol. 2, no. 1, Mar. 2018, pp. 60–65.
- [8] N. Rajatheva et al., "White Paper on Broadband Connectivity in 6G," arXiv:2004.14247, 2020; <https://arxiv.org/abs/2004.14247>.
- [9] C. Zhou et al., "Deep Reinforcement Learning for Delay-Oriented IoT Task Scheduling in Space-Air-Ground Integrated Network," *IEEE Trans. Wireless Commun.*, vol. 20, no. 2, Feb. 2021, pp. 911–25.
- [10] C. Campolo et al., "5G Network Slicing for Vehicle-to-Everything Services," *IEEE Wireless Commun.*, vol. 24, no. 6, Dec. 2017, pp. 38–45.
- [11] C. Gutterman et al., "RAN Resource Usage Prediction for a 5G Slice Broker," *Proc. ACM MobiHoc*, Catania, Italy, 2019.
- [12] I. Goodfellow, Y. Bengio, and A. Courville, *Deep Learning*, MIT Press, 2016.
- [13] M. Li et al., "Slicing-Based Artificial Intelligence Service Provisioning on the Network Edge: Balancing AI Service Performance and Resource Consumption of Data Management," *IEEE Vehic. Tech. Mag.*, vol. 16, no. 4, Dec. 2021, pp. 16–26.
- [14] S.-C. Lin et al., "The Architectural Implications of Autonomous Driving: Constraints and Acceleration," *Proc. ASPLOS*, 2018, pp. 751–66.
- [15] W. Wu et al., "Dynamic RAN Slicing for Service-Oriented Vehicular Networks via Constrained Learning," *IEEE JSAC*, vol. 39, no. 7, 2021, pp. 2076–89.

BIOGRAPHIES

WEN WU [S'13, M'20] (w77wu@uwaterloo.ca) earned his Ph.D. degree in electrical and computer engineering from the University of Waterloo, Ontario, Canada, in 2019. He received his B.E. degree in information engineering from South China University of Technology, Guangzhou, China, and his M.E. degree in electrical engineering from University of Science and Technology of China, Hefei, in 2012 and 2015, respectively. He worked as a postdoctoral fellow with the Department of Electrical and Computer Engineering, University of Waterloo. He is currently an associate researcher at Peng Cheng Laboratory, Shenzhen, China. His research interests include 6G networks, pervasive network intelligence, and network virtualization.

CONGHAI ZHOU [S'19] (c89zhou@uwaterloo.ca) received his B.E. degree from Northeastern University, Shenyang, China, in 2017 and his M.S. degree from the University of Illinois at Chicago in 2018. He is currently working toward a Ph.D. degree with the Department of Electrical and Computer Engineering, University of Waterloo. His research interests include space-air-ground integrated networks and machine learning in wireless networks.

MUSHU LI [S'17, M'21] (m475li@uwaterloo.ca) is a postdoctoral fellow with the Department of Electrical and Computer Engineering, University of Waterloo, where she also received her Ph.D. degree in electrical engineering in 2021. She received her M.A.Sc. degree from Ryerson University in 2017. She was the recipient of a Natural Science and Engineering Research Council of Canada Graduate Scholarship in 2018 and an Ontario Graduate Scholarship in 2015 and 2016, respectively. Her research interests include machine learning in wireless networks and network slicing.

HUAQING WU [S'15, M'21] (h272wu@uwaterloo.ca) received her Ph.D. degree from the University of Waterloo in 2021. She received her B.E. and M.E. degrees from Beijing University of Posts and Telecommunications, China, in 2014 and 2017, respectively. She received the prestigious Natural Sciences and Engineering Research Council of Canada (NSERC) Postdoctoral Fellowship Award in 2021. She is currently a postdoctoral research fellow at McMaster University, Ontario, Canada. Her current research interests include vehicular networks with emphasis on edge caching, wireless resource management, space-air-ground integrated networks, and application of artificial intelligence for wireless networks.

HAIBO ZHOU [M'14, SM'18] (haibozhou@nju.edu.cn) received his Ph.D. degree in information and communication engineering from Shanghai Jiao Tong University, China, in 2014. He is cur-

rently an associate professor with the School of Electronic Science and Engineering, Nanjing University, China. He was named a 2020 cross-field highly cited researcher. He was a recipient of the 2019 IEEE Communications Society Asia-Pacific Outstanding Young Researcher Award. He is currently an Associate Editor of the *IEEE Internet of Things Journal*, *IEEE Network*, and *IEEE Wireless Communications Letters*. His research interests include resource management in VANETs, 5G/B5G wireless networks, and SAGINs.

NING ZHANG [M'15, SM'18] (ning.zhang@uwindsor.ca) is currently an associate professor in the Department of Electrical and Computer Engineering at the University of Windsor, Canada. He received his Ph.D. degree in electrical and computer engineering from the University of Waterloo in 2015. His research interests include connected vehicles, mobile edge computing, wireless networking, and machine learning. He is a Highly Cited Researcher. He serves as an Associate Editor of the *IEEE Internet of Things Journal*, *IEEE Transactions on Cognitive Communications and Networking*, and the *IEEE Systems Journal*; and has served as a Guest Editor for several international publications, such as *IEEE Wireless Communications* and *IEEE Transactions on Industrial Informatics*.

XUEMIN (SHERMAN) SHEN [M'97, SM'02, F'09] (ssh@uwaterloo.ca) is currently a university professor with the Department of Electrical and Computer Engineering, University of Waterloo. His research focuses on network resource management, wireless network security, social networks, 5G and beyond, and vehicular ad hoc networks. He is a Canadian Academy of Engineering Fellow, a Royal Society of Canada Fellow, and a Chinese Academy of Engineering Foreign Fellow. He received the R.A. Fessenden Award in 2019 from IEEE, Canada; the James Evans Avant Garde Award in 2018 from the IEEE Vehicular Technology Society; and the Joseph LoCicero Award in 2015 and Education Award in 2017 from the IEEE Communications Society.

WEIHUA ZHUANG [M'93, SM'01, F'08] (wzhuang@uwaterloo.ca) has been with the Department of Electrical and Computer Engineering, University of Waterloo since 1993, where she is currently a professor and a Tier I Canada Research Chair in Wireless Communication Networks. She is a Fellow of the Royal Society of Canada, the Canadian Academy of Engineering, and the Engineering Institute of Canada. She is also an elected member of the Board of Governors and VP-Publications of the IEEE Vehicular Technology Society.