



REPUBLIQUE ALGERIENNE DEMOCRATIQUE ET POPULAIRE
MINISTRE DE L'ENSEIGNEMENT SUPERIEUR ET DE LA RECHERCHE SCIENTIFIQUE

UNIVERSITE ABOU-BEKR BELKAID – TLEMCE

THÈSE LMD

Présentée à :

FACULTE DES SCIENCES – DEPARTEMENT D'INFORMATIQUE

Pour l'obtention du diplôme de :

DOCTORAT

Spécialité: *Informatique Distribuée et Réseaux*

Par :

SAIDI Anas Nawfel

Sur le thème

**L'impact de la technologie 5G sur les performances des
communications dans les VANETs**

Soutenue publiquement le 04 Octobre 2025 à Tlemcen devant le jury composé de :

M ^r BENAMAR Abdelkrim	Professeur	Université de Tlemcen	Président
M ^r LEHSAINI Mohamed	Professeur	Université de Tlemcen	Directeur de thèse
M ^r MERAD BOUDIA Omar Rafik	Professeur	Université d'Oran 1	Examineur
M ^r BENDAOU D Fayssal	MCA	ESI de Sidi Belabbes	Examineur
M ^r HADJILA Mourad	Professeur	Université de Tlemcen	Examineur

*Laboratoire Systèmes et Technologies de l'Information et de la Communication (STIC)
BP 119, 13000 Tlemcen - Algérie*

Remerciements

Je tiens, avant toute chose, à remercier ALLAH TOUT PUISSANT de m'avoir donné la force, le courage, la volonté et la patience afin de réaliser ce travail de thèse.

J'exprime ma profonde gratitude à mon directeur de thèse, Professeur LEHSAINI Mohamed pour son esprit scientifique, sa pédagogie, sa disponibilité, sa patience, ses conseils et ses idées qui m'ont permis de mener à bien cette thèse.

Mes remerciements les plus sincères sont adressés au Professeur BENAMAR Abdelkrim de me faire l'honneur de s'intéresser à ce travail et d'avoir présidé le jury

J'exprime ensuite ma plus profonde gratitude aux Professeurs MERAD BOUDIA Omar Rafik et HADJILA Mourad, et au Docteur BENDAOUD Fayssal qui m'ont fait l'honneur d'accepter d'être les examinateurs de cette thèse.

Ce travail a été réalisé au sein du laboratoire Systèmes et Technologie de l'Information et de la Communication (STIC) de l'université de Tlemcen. Je souhaite également remercier tous membres de ce laboratoire et je souhaite beaucoup d'épanouissement à mes collègues dans leurs travaux.

Enfin, J'adresse mes sincères remerciements à toutes les personnes qui me sont très chères qui m'ont encouragé, soutenu durant toutes ces années, en commençant par ma famille : d'abord ma pensée va également vers mes très chers parents pour tous leurs sacrifices pour nous depuis toujours et pour m'avoir soutenu et encouragé dans mon choix de poursuites d'étude.

A vous tous, Merci !

Dédicaces

Je dédie ce travail à mes chers parents,

A la mémoire

SAIDI Anas Nawfel

Table des matières

Liste des figures	iv
Liste des tableaux	vi
Introduction Générale	1
1 Généralités sur le découpage réseau dans les réseaux véhiculaires 5G	4
1.1 Introduction	4
1.2 Généralités sur les VANETs	4
1.3 La technologie 5G	7
1.4 Découpage réseau	9
1.4.1 Définition et histoire	9
1.4.2 Le découpage réseau dans la 5G	9
1.4.3 Outils de découpage du réseau dans la 5G	11
a) L'architecture "Software Defined Networking (SDN)"	11
b) La virtualisation des fonctions réseau (NFV)	13
c) Cloud/Edge computing	16
1.4.4 Les types de découpage réseau dans la 5G	17
a) Le réseau central (CN)	17
b) Le réseau de transport	18
c) Le réseau radio "Radio Access Network (RAN)"	19
1.4.5 Les principes du découpage réseau	20
1.4.6 Cas d'utilisation génériques dans le découpage de réseau 5G	21
1.5 Découpage réseau 5G V2X	21
1.5.1 Les types de communications dans les réseaux véhiculaires	21
1.5.2 Les divers services V2X avec leurs exigences	22
1.5.3 Motivations de la transition vers le découpage réseau 5G V2X	26
1.5.4 Tranches de réseau 5G V2X	26
1.6 Défis de l'adaptation du découpage réseau pour les communications V2X	27
1.6.1 Gestion de la mobilité	28
1.6.2 Chaînage et placement de fonctions	29
1.6.3 Gestion des ressources	29
1.6.4 Déchargement/planification des tâches	30
1.6.5 Problèmes de la sécurité	31
1.7 Conclusion	32
2 État de l'art sur les approches de découpage réseau dans l'environnement V2X 5G	33

2.1	Introduction	33
2.2	Critères essentiels de gestion des tranches de réseau	33
2.3	Taxonomie des solutions utilisées dans la gestion des ressources dans le découpage réseau	34
2.3.1	Les approches d'allocation dans le découpage réseau	34
a)	Tranchage strict / la réservation complète des ressources	34
b)	L'approche sans réservation	36
c)	L'approche de réservation partielle des ressources	37
2.3.2	Les types d'allocation des ressources	38
a)	Gestion des ressources à une seule couche	38
b)	Gestion des ressources à deux couches	39
2.3.3	Techniques utilisées pour l'allocation des ressources dans le découpage de réseau	41
a)	Modèles basés sur la théorie des jeux	41
b)	Modèles de prédiction	43
c)	Modèles basés sur des heuristiques et des métaheuristiques	46
d)	Modèles basés sur l'optimisation mathématique	48
2.4	Discussion	49
2.5	Conclusion	53
3	Mécanisme de partage des ressources dans un environnement de découpage réseau V2X	54
3.1	Introduction	54
3.2	Motivation	54
3.3	Contribution	55
3.4	L'apprentissage par renforcement	56
3.4.1	Processus de décision Markovien	56
3.4.2	Exploration/Exploitation	58
3.4.3	Apprentissage en ligne/hors ligne	59
3.4.4	Algorithmes basés sur un modèle/sans modèle	59
3.4.5	L'algorithme Q-Learning	59
3.4.6	Deep Q-Learning	61
3.5	Contribution : Mécanisme proposé pour le partage des ressources entre les tranches V2X	62
3.5.1	Description générale de déclenchement du mécanisme proposé	62
3.5.2	Éléments du processus de décision Markovien	63
3.5.3	Réduire l'espace d'actions total	66
a)	Classification des actions préliminaires	66
b)	Élimination des actions irréalisables	70
3.5.4	Scénario de partage de ressources	71
3.6	Évaluation des performances	73
3.6.1	Environnement du travail	73
3.6.2	Description du scénario	74
3.6.3	Approches d'analyse comparative	74
3.6.4	Métriques	75
3.6.5	Taux d'utilisation des ressources	78

3.6.6	Probabilité d'abandon du transfert d'appel	79
3.6.7	Probabilité de blocage des nouveaux appels	80
3.6.8	Récompenses obtenues par différentes stratégies	81
3.6.9	Impact des paramètres de priorités	82
3.6.10	Évaluation de l'agent DQL durant l'apprentissage	84
3.7	Conclusion	85
4	Mécanisme de dissémination des informations du trafic véhiculaire dans les VANETs	86
4.1	Introduction	86
4.2	Contexte	86
4.3	Travaux connexes	87
4.4	Contribution : Le protocole ERGF	88
4.4.1	Dissémination d'informations sur le trafic routier	89
4.4.2	Création du paquet RTIP	91
4.4.3	Dissémination du paquet RTIP	91
4.4.4	Améliorations proposées du mécanisme de dissémination sur le trafic routier	92
	a) Première limitation dans le protocole PBRP	94
	b) Première amélioration	96
	c) Deuxième limitation présentée dans le protocole PBRP	96
	d) Deuxième amélioration	97
4.5	Évaluation des performances	99
4.5.1	Environnement de simulations	99
4.5.2	Résultats de simulation	100
	a) Taux de livraison de paquets	100
	b) Délai de bout en bout	100
	c) Overhead du routage	102
4.6	Conclusion	103
	Conclusions et Perspectives	104
	Bibliographie	107

Liste des figures

1.1	Types de communications dans les VANETs	5
1.2	Exigences de spécifications de la technologie 5G	8
1.3	Cycle de vie des tranches dans V2X 5G	11
1.4	Passage vers une architecture à base de SDN	12
1.5	L'architecture du SDN	13
1.6	Passage vers une virtualisation des fonctions réseau	14
1.7	Les modèles de Cloud	16
1.8	Les tranches génériques	22
1.9	Les modes de communications V2X	23
1.10	Un scénario de partage de l'information sur les obstacles	24
1.11	Architecture générale de la conduite à distance	24
1.12	Les tranches dans V2X	28
1.13	Le modèle de découpage de bout en bout	29
2.1	Classification des travaux de gestion des ressources dans le découpage réseau	35
2.2	Les types de gestion des ressources dans V2X 5G	38
3.1	Organigramme de déclenchement du mécanisme de partage	56
3.2	Interaction entre l'agent et l'environnement	57
3.3	La difference entre Q-Learning et Deep Q-Learning	62
3.4	Illustration du mécanisme de partage des ressources	63
3.5	Espace des actions préliminaires pour la $k^{ième}$ tranche (type 1)	68
3.6	Espace des actions préliminaires pour la $k^{ième}$ tranche (type 2)	69
3.7	Illustration d'un exemple scenario	72
3.8	Taux d'utilisation des ressources	78
3.9	Probabilité d'abandon du transfert	79
3.10	Probabilité de blocage des nouveaux appels	81
3.11	Récompenses obtenues ave différentes approches	82
3.12	Impact du paramètre "type priority"	83
3.13	Impact du paramètre "slice priority"	84
3.14	Entraînement de l'agent DQL	84
4.1	Dissemination des informations sur le trafic routier	89
4.2	Paquet d'information sur le trafic routier	90
4.3	Schéma de transmission des informations sur le trafic routier	91
4.4	Premier scénario : Partie 1	94
4.5	Premier scénario : Partie 2	95
4.6	Limitation 1 : Scénario idéal vs scénario actuel	96

4.7	Deuxième scénario	98
4.8	Limitation 2 : Scénario idéal vs scénario actuel	98
4.9	Taux de livraison des paquets dans les protocoles EGyTAR, PBRP, ERGF	101
4.10	Délai de bout en bout dans les protocoles EGyTAR, PBRP, ERGF	101
4.11	Nombre de paquets RTIP dans les protocoles PBRP et ERGF	102
4.12	Nombres de paquets RTIP diffusés par les RBNs dans les protocoles PBRP et ERGF	103

Liste des tableaux

2.1	Récapitulatif sur les solutions proposées pour la gestion des ressources dans le découpage réseau 5G	51
3.1	Paramètres du réseau	73
3.2	Paramètres de formation des agents	74
4.1	Paramètres de simulation	100

Glossaire

- 3D** Three Dimensional. 23
- 3GPP** 3rd Generation Partnership Project. 6
- 5G** Fifth Generation Mobile Network. 4
- 5GCAR** 5G Automotive Communication. 22
- AMF** Access and Mobility Function. 18
- AUSF** Authentication Server Function. 27
- BBU** Baseband Unit. 19
- BS** Base Station. 19
- BSR** Buffer Status Report. 41
- BSS** Business Support System. 15
- CAMs** Cooperative Awareness Messages. 23
- CN** Core Network. 14
- CNN** Convolutional Neural Network. 44
- CoCA** Collision Cooperative Avoidance. 25
- CPE** Collective Perception of the Environment. 25
- CSMA/CA** Carrier Sense Multiple Access with Collision Avoidance. 6
- DDoS** Distributed Denial of Service. 31
- DDPG** Deep Deterministic Policy Gradient. 46
- DL** Downlink. 24
- DLNN** Deep Learning Neural Network. 44
- DPI** Deep Packet Inspection. 14
- DQL** Deep Q-Learning. 45, 61

DRL Deep Reinforcement Learning. 3

DSR Dynamic Resource Sharing. 47

DSRC Dedicated Short Range Communication. 5

EDCA Enhanced Distributed Channel Access. 6

EGyTAR Enhanced Greedy Traffic Aware Routing. 88

eMBB Enhanced Mobile Broadband. 21

ERGF Efficient Routing Greedy Forwarding. 86

EtrA Emergency Trajectory Alignment. 25

ETSI Institut européen des normes de télécommunications. 15

ETX Expected transmission count. 87

eV2X Enhanced Vehicle to Everything. 26

FANS Fixed Access Network Sharing. 18

FoRCES Forwarding and Control Element Separation. 13

GBR Guaranteed Bit Rate. 39

GL Group Leader. 90

GPS Global Positioning System. 6

GyTAR Greedy Traffic Aware Routing. 88

H2H Human To Human. 41

IA intelligence artificielle. 7

IaaS Infrastructure as a Service. 16

ICT Information and Communication Technology. 26

IoT Internet des objets. 7, 48

IoV Internet of Vehicles. 1, 45

ITS Intelligent Transportation Systems. 1, 21

KKT Karush-Kuhn-Tucker. 48

KPI Key Performance Indicators. 44

LSGO Link State Aware Geographic Opportunistic. 87

LSTM Long Short Term Memory. 43

LTE Long Term Evolution. 1, 7

M2M Machine To Machine. 41

MAC Media Access Control. 19

MANETs Mobile Adhoc NETworks. 86

MBR Minimum Bit Rate. 45

MDP processus de décision Markovien. 56, 63

ML Machine Learning. 43

mMTC Massive Machine Type Communication. 21

MNOs Mobile Network Operators. 26

MVNO Mobile Virtual Network Operator. 19

NAT64 Network Address Translation. 14

NFV Network Function Virtualization. 13

NFV-MANO Gestion et orchestration du NFV. 14

NGMN Next Generation Mobile Network. 9

NN Neural Network. 43

NVS Network Virtualization Substrate. 41

OBU OnBoard Unit. 6

ONF Open Networking Foundation. 11

PaaS Platform as a Service. 16

PBRP Partial backwards routing protocol. 86, 87

PCF Policy Control Function. 18, 27

PDCCP Packet Data Convergence Protocol. 19

PDR Packet Delivery Ratio. 99

PGW Packet Data Network Gateway. 14

PL Linear Programming. 48

PNL Nonlinear Programming. 48

PRR Packet Reception Rate. 47

PSO Particle Swarm Optimization. 36

QoE qualité d'expérience. 7

QoS qualité de service. 1, 7

RAC Radio Admission Control. 39

RAN Radio Access Network. i, 2, 19

RBN RTIP Broadcast Node. 89

RIN RTIP Initiator Node. 89

RIT Road Information Table. 91

RL Reinforcement Learning. 56

RLC Radio Link Control. 19

RRC Radio Resource Control. 19

RRHs Remote Radio Heads. 19

RSUs Road-Side Units. 5, 88

RTIP Road Traffic Information Packet. 89

SaaS Software as a Service. 16

SAC Slice Admission Control. 39

SBI SouthBound. 13

SDN Software Defined Networking. 9

SFC Service Function Chain. 29

SGW Serving Gateway. 14

SINR Signal to Interference plus Noise Ratio. 47

SLA Service Level Agreement. 20

SMDP Semi Markov Decision Process. 39

SMF Session Management Function. 17

SSMS Sensor and Status Map Sharing. 25

SSR Service Satisfaction Rate. 46

- TCP** Transmission Control Protocol. 14
- TN** Transport Network. 18
- UDM** Unified Data Management. 27
- UIT-T** International Telecommunication Union - Telecommunication Standardization Sector. 4
- UL** Uplink. 24
- UPF** User Plane Function. 18
- URLLC** Ultra Reliable Low Latency Communication. 21
- V2I** Vehicle-to-Infrastructure. 21
- V2N** Vehicle-to-Network. 22
- V2P** Vehicle-to-Pedestrian. 22
- V2V** Vehicle-to-Vehicle. 21
- V2X** Vehicle to Everything. 1, 21
- VaD** Video Data Sharing for Automated Driving. 25
- VANETs** Vehicular Ad hoc NETworks. 1, 4
- VNF** Virtual Network Function. 14
- VPN** Virtual Private Network. 14, 18
- VUE** Vehicle User Equipment. 46

Résumé

Avec l'évolution et l'émergence de nouveaux services visant à améliorer l'expérience des passagers de véhicules, la communication entre diverses entités telles que les véhicules, les unités routières, les systèmes de gestion du trafic et les applications cloud a considérablement augmenté le volume d'informations échangées. Cela a entraîné une plus grande complexité dans la gestion des diverses exigences de qualité de service, notamment dans la partie communication, au sein des réseaux Vehicular-to-Everything (V2X). En conséquence, la technologie 5G est considérée comme une solution indispensable pour prendre en charge efficacement les communications V2X grâce au paradigme de découpage de réseau. Le découpage du réseau permet à la technologie 5G de gérer divers services avec des exigences variées au sein d'une infrastructure partagée, en créant des réseaux virtuels personnalisés isolés, appelés tranches. Cela a conduit les chercheurs à proposer l'intégration du découpage réseau dans les VANETs afin de tirer parti de ses avantages. Cependant, ce paradigme nécessite une certaine adaptation pour s'intégrer efficacement avec les VANETs, qui présentent des spécificités uniques, telles que la mobilité élevée. Parmi les domaines nécessitant une adaptation, la gestion des ressources entre les tranches véhiculaires est particulièrement cruciale, notamment dans la partie radio où la bande passante est une ressource limitée. Dans le cadre de cette thèse, l'objectif est de relever les défis liés à la gestion des ressources dans l'environnement 5G V2X avec le paradigme du découpage réseau. Pour réaliser cet objectif, nous avons proposé un mécanisme de partage des ressources entre les tranches véhiculaires, basé sur l'apprentissage par renforcement profond. Les résultats de simulations de cette contribution ont été comparés à d'autres travaux pertinents en matière de performances et ils ont illustré une certaine supériorité par rapport à ces travaux. Une autre contribution abordée dans le cadre de cette thèse concerne l'amélioration du mécanisme de routage dans les VANETs afin de permettre une dissémination plus efficace et fiable des informations routières.

Mot clés : V2X, 5G, Découpage réseau, Gestion des ressources, Apprentissage par renforcement, VANETs, Protocoles de routage.

Abstract

With the evolution and emergence of new services aimed at improving the experience of vehicle passengers, communication between various entities such as vehicles, roadside units, traffic management systems and cloud applications has considerably increased the volume of information exchanged. This has led to greater complexity in managing the various quality of service requirements, particularly on the communication side, within Vehicular-to-Everything (V2X) networks. As a result, 5G technology is seen as an indispensable solution to efficiently support V2X communications through the network slicing paradigm. Network slicing enables 5G technology to manage various services with varying requirements within a shared infrastructure, by creating isolated custom virtual networks, called slices. This has led researchers to propose the integration of network slicing into VANETs to take advantage of its benefits. However, this paradigm requires some adaptation to integrate effectively with VANETs, which have unique features such as high mobility. Among the areas requiring adaptation, resource management between vehicular slices is particularly crucial, especially in the radio part where bandwidth is a limited resource. In this thesis, the goal is to address the challenges related to resource management in the 5G V2X environment with the network slicing paradigm. To achieve this goal, we proposed a resource sharing mechanism between vehicular slices based on deep reinforcement learning. The simulation results of this contribution have been compared with other relevant works in terms of performance and have shown a certain superiority over these works. Another contribution addressed in this thesis concerns the improvement of the routing mechanism in VANETs in order to allow a more efficient and reliable dissemination of traffic information.

Keywords: V2X, 5G, Network slicing, Resource management, Reinforcement learning, VANETs, Routing protocols.

ملخص

مع تطور وظهور خدمات جديدة تهدف إلى تحسين تجربة ركاب المركبات، أدى التواصل بين مختلف الكيانات مثل المركبات ووحدات الطرق وأنظمة إدارة المرور وتطبيقات السحابة إلى زيادة كبيرة في حجم المعلومات المتبادلة. وقد أدى هذا إلى مزيد من التعقيد في إدارة متطلبات جودة الخدمة المختلفة، وخاصة في جزء الاتصالات، ضمن شبكات المركبات إلى كل شيء (V2X). لذلك، تعتبر تقنية الجيل الخامس 5G حلاً لا غنى عنه لدعم اتصالات V2X بكفاءة من خلال نموذج تقطيع الشبكة. يتيح تقطيع الشبكة لتقنية الجيل الخامس إدارة خدمات متنوعة ذات متطلبات مختلفة ضمن بنية تحتية مشتركة من خلال إنشاء شبكات افتراضية معزولة ومخصصة، تسمى الشرائح. وقد دفع هذا الباحثين إلى اقتراح دمج تقسيم الشبكة في شبكات VANETs من أجل الاستفادة من فوائدها. ومع ذلك، يتطلب هذا النموذج بعض التكيف ليتكامل بشكل فعال مع شبكات VANETs، التي تتمتع بخصائص فريدة، مثل القدرة العالية على الحركة. ومن بين المجالات التي تتطلب التكيف، تعد إدارة الموارد بين شرائح المركبات أمراً بالغ الأهمية، وخاصة في الجزء اللاسلكي حيث يكون النطاق الترددي مورداً محدوداً. في هذه الأطروحة، الهدف هو معالجة تحديات إدارة الموارد في بيئة 5G V2X مع نموذج تقطيع الشبكة. ولتحقيق هذا الهدف، اقترحنا آلية لتقاسم الموارد بين شرائح المركبات، استناداً إلى التعلم التعزيزي العميق. تمت مقارنة نتائج المحاكاة لهذه المساهمة مع أعمال الأداء الأخرى ذات الصلة وأظهرت بعض التفوق على هذه الأعمال. تتضمن مساهمة أخرى تناولتها هذه الأطروحة تحسين آلية التوجيه في شبكات الطرق السريعة من أجل السماح بنشر معلومات الطرق بشكل أكثر كفاءة وموثوقية.

الكلمات المفتاحية : V2X، 5G، تقطيع الشبكة، إدارة الموارد، التعلم التعزيزي، شبكات VANETs، بروتوكولات التوجيه.

Introduction Générale

Introduction Générale

Le secteur automobile a toujours captivé l'attention, tant dans le domaine académique qu'industriel, créant ainsi une demande continue pour améliorer la sécurité routière en réduisant les accidents, en anticipant les conditions de circulation dangereuses et météorologiques, ainsi qu'en optimisant l'expérience utilisateur, notamment en termes de confort et de plaisir de conduite. Cela inclut également des efforts pour diminuer les embouteillages et la pollution. Cet objectif a été concrétisé par divers projets visant à déployer des systèmes de transport intelligents "Intelligent Transportation Systems (ITS)", répondant ainsi aux besoins croissants du secteur automobile. Dans ce contexte, la mise en place d'une infrastructure de transport fiable, capable de fournir des informations précieuses en matière de sécurité et de services de confort, nécessite des réseaux de communication efficaces. Ces réseaux doivent être en mesure de répondre aux exigences croissantes en termes de performance. Parmi les premiers réseaux mis en place et au cœur des recherches en développement, les réseaux ad-hoc de véhicules "Vehicular Ad hoc NETWORKS (VANETs)" constituent une catégorie spécifique de réseaux ad hoc basés sur des technologies de communication telles que le 802.11p, où les véhicules servent de nœuds mobiles. Ces réseaux offrent aux utilisateurs divers services, tels que des informations sur le trafic routier. Cependant, la naissance de nouveaux besoins, qui se transforment en services avec des exigences de qualité de service "qualité de service (QoS)" très difficiles à satisfaire avec les technologies ad hoc, combinée à la quantité considérable de données générée par le nombre croissant de véhicules connectés, pose un défi majeur. En effet, les VANETs sont insuffisants pour déployer et gérer efficacement tous ces services.

L'évolution des réseaux cellulaires a été révolutionnée de manière significative au cours des dernières décennies, avec l'émergence de nouvelles générations de réseaux cellulaires, chacune étant plus puissante que la précédente. Dans ce contexte, le réseau "Long Term Evolution (LTE)" de quatrième génération (4G) a prouvé son efficacité et désormais intégré dans notre quotidien. La quatrième génération a particulièrement marqué des avancées dans la communication véhiculaire. En effet, la 4G présente certains avantages par rapport aux VANETs, tels qu'une bande passante plus élevée, des vitesses de données supérieures et une latence réduite. Ces caractéristiques permettent une communication plus rapide et plus fiable entre les véhicules et les infrastructures, facilitant ainsi la transmission d'informations essentielles. Cependant, avec l'émergence de l'Internet des véhicules ("Internet of Vehicles (IoV)"), basé sur les communications "Vehicle to Everything (V2X)", qui permettent aux véhicules d'interagir en temps réel avec leur environnement, ainsi que de nouvelles applications exigeant une latence ultra-faible et une capacité extrêmement élevée, telles que le contrôle à distance ou les véhicules autonomes, la coexistence de tous ces services dans une architecture unifiée comme celle de la 4G est devenue de plus en plus complexe.

La nouvelle cinquième génération de communications sans fil (5G) a été introduite dans l'environnement véhiculaire comme solution ultime pour surmonter les limites techniques des communications dans les réseaux sans fil ad hoc et la 4G, en offrant une meilleure capacité, une latence ultra-faible et une fiabilité accrue, répondant ainsi principalement aux exigences liées à la coexistence de tous les services véhiculaires dans une seule infrastructure.

Problématique et contributions

La technologie 5G repose principalement sur le découpage réseau (network slicing), un paradigme qui permet de découper l'architecture réseau originale en plusieurs tranches logiques distinctes où chaque tranche étant dédiée à un service spécifique. Ce principe de découper le réseau physique en plusieurs tranches permet à chacune d'elles de se gérer de façon indépendante et autonome par rapport aux autres, en contrôlant son propre fonctionnement, notamment en matière d'allocation des ressources, de gestion de la mobilité et de mécanismes de sécurité, offrant ainsi une flexibilité optimale pour répondre aux besoins spécifiques de chaque service. Ce concept s'avère être une solution idéale pour assurer la coexistence des communications V2X et de leurs divers services. Le découpage réseau 5G V2X présente de nombreux défis, notamment en ce qui concerne la gestion des ressources, la gestion de la mobilité et la sécurité des tranches V2X. Ces défis soulignent la nécessité accrue de recherches dans ce domaine, en particulier sur la gestion des ressources dans un environnement hautement dynamique. Les travaux de recherche dans cet axe s'appuient sur l'analyse de plusieurs indicateurs, tels que l'isolation entre les tranches et la satisfaction des exigences de QoS pour chaque tranche.

Le découpage réseau entre les tranches V2X dans la partie "RAN" constitue l'un des principaux défis du déploiement des tranches de réseau, particulièrement en ce qui concerne la répartition des ressources entre les tranches, étant donné leur rareté. Une conception optimale de ce partage est donc essentielle. Rappelons que l'un des principes fondamentaux du découpage réseau est de garantir une isolation entre les tranches. Un aspect clé de cette isolation réside dans la gestion des ressources, chaque tranche disposant de ses ressources dédiées ainsi que d'une gestion personnalisée intra-tranche de ces ressources. Des solutions basées sur des approches comme le découpage strict des ressources entre les tranches peuvent garantir ce principe d'isolation. Cependant, dans un environnement très dynamique tel que celui des communications véhiculaires, ce manque de flexibilité entraînera une utilisation sous-optimale des ressources du système. L'intégration d'une amélioration via le partage des ressources entre les tranches est essentielle. Néanmoins, les solutions existantes rencontrent des difficultés pour trouver un équilibre entre l'isolation et l'efficacité de l'utilisation des ressources dans l'environnement véhiculaire.

Dans ce contexte, cette thèse présente une solution axée sur la gestion des ressources de la bande passante dans le RAN, au sein d'un environnement 5G V2X reposant sur le découpage réseau, afin d'adresser les défis de gestion des ressources abordés précédemment. La solution proposée est un mécanisme de partage des ressources entre les tranches 5G V2X, basé sur l'apprentissage par renforcement profond. Ce mécanisme fonctionne à un niveau supérieur au contrôle d'admission des tranches. Il permet d'accepter une demande de connexion en empruntant des ressources à d'autres tranches, tout en veillant

à protéger les tranches prêtes contre une surcharge due à ce processus de partage. En outre, une deuxième contribution est apportée, visant à améliorer le mécanisme de routage dans les environnements VANETs. Cette contribution cible spécifiquement la résolution des limitations rencontrées par les mécanismes de diffusion des informations sur le trafic routier dans certains scénarios critiques.

Organisation de la thèse

La structure de ce manuscrit de thèse est comme suit :

- Le premier chapitre propose une introduction aux réseaux ad hoc véhiculaires et à leurs services, en explorant leur évolution, notamment à travers les technologies de communication telles que la 5G, ainsi que les défis liés à la mise en œuvre des services véhiculaires dans un environnement 5G de découpage réseau, avec un focus particulier sur la gestion des ressources entre les services véhiculaires. Ce chapitre sert de préambule, posant les bases de l'ensemble de la thèse.
- Le deuxième chapitre de cette thèse se concentre sur l'analyse des travaux existants, en mettant particulièrement l'accent sur l'aspect central de notre recherche : la gestion des ressources dans le cadre du découpage réseau 5G, classée selon différents aspects. Cette étude nous a permis de tirer parti des atouts de ces travaux et de proposer de nouvelles solutions adaptées à l'environnement 5G V2X et au découpage réseau.
- Le troisième chapitre présente notre première contribution, qui consiste à la proposition d'un nouveau mécanisme de partage des ressources basé sur l'apprentissage par renforcement profond "Deep Reinforcement Learning (DRL)". De plus, nous avons évalué les performances du mécanisme proposé selon plusieurs métriques et les comparer à celles obtenues par d'autres approches pour illustrer les bénéfices de notre approche par rapport aux autres.
- Le quatrième chapitre présente notre deuxième contribution, qui consiste à la proposition d'améliorations de la partie dédiée à la dissémination des informations sur le trafic routier dans le protocole de routage proposé, qui représente une composante importante, responsable de fournir une vision sur l'état du réseau routier. De plus, nous avons évalué les performances du protocole de routage proposé pour les VANETs et les comparé à d'autres protocoles de routage, en utilisant plusieurs métriques d'évaluation.

Finalement, une conclusion générale clôture ce manuscrit et rappelle les principales contributions tout au long de ce travail de thèse. Elle présente également les perspectives sous-jacentes aux différentes problématiques abordées dans cette thèse. Les travaux contenus dans ce document ont été publiés dans deux journaux et une conférence.

Chapitre 1

Généralités sur le découpage réseau dans les réseaux véhiculaires 5G

Chapitre 1

Généralités sur le découpage réseau dans les réseaux véhiculaires 5G

1.1 Introduction

L'évolution rapide des diverses exigences des services connectés a rendu nécessaire l'adoption de technologies capables de répondre à cette complexité croissante. Contrairement aux générations précédentes, la "Fifth Generation Mobile Network (5G)" apporte la capacité de s'adapter à un large éventail d'applications, où les demandes variées des différents services ne peuvent être satisfaites par une approche réseau unique. Pour répondre à ce défi, la 5G a été renforcée par le paradigme du découpage réseau. Les communications dans les réseaux ad hoc véhiculaires "VANETs", par exemple, présentent une grande diversité en termes d'exigences de service, ce qui rend indispensable l'utilisation du découpage réseau. Ainsi, un nouveau domaine de recherche a émergé dans le secteur des transports : l'application du paradigme du découpage réseau dans les réseaux véhiculaires 5G.

Dans ce chapitre, nous commencerons par présenter les réseaux VANETs. Ensuite, nous aborderons les réseaux cellulaires, notamment la technologie 5G. Nous détaillerons ensuite le paradigme du découpage réseau utilisé dans la 5G, qui est au cœur de cette thèse.

1.2 Généralités sur les VANETs

Dans le monde entier, le nombre de personnes utilisant les systèmes de transport continue d'augmenter, ce qui engendre une forte croissance du nombre de véhicules en circulation. Cette augmentation du nombre de véhicules sur les routes entraîne de nombreuses conséquences, telles que la congestion et les accidents mortels. Pour améliorer la sécurité routière, divers paradigmes de réseaux de véhicules ont été proposés, parmi lesquels les réseaux ad hoc de véhicules (VANETs). Le concept de réseau VANET a été introduit pour la première fois en 2003, lors de la conférence de normalisation des communications automobiles de l'Union Internationale des Télécommunications - secteur de la normalisation des télécommunications "International Telecommunication Union - Telecommunication Standardization Sector (UIT-T)" dans le but de construire des systèmes de transport routier plus sûrs avec des fonctionnalités telles que la gestion du trafic, tout en offrant le confort et le plaisir aux conducteurs ainsi qu'aux passagers sur les routes

publiques. Les messages échangés entre les entités VANET peuvent être classés en deux catégories, à savoir :

- **Message lié à la sécurité** : Cette catégorie se divise en deux types de messages. Le premier est le message beacon, qui fournit des informations cruciales telles que l'identité, la position, la vitesse et la direction du véhicule. Diffusé régulièrement, ce message facilite l'identification des véhicules voisins et est essentiel pour les protocoles de routage et de sécurité. Le deuxième type est le message d'alerte ou d'urgence, envoyé pour informer les autres véhicules sur les conditions routières telles que des accidents, des congestions, des conditions météorologiques défavorables ou le passage des véhicules de secours. Ce type de message contribue à améliorer la circulation et la sécurité routière en donnant aux conducteurs plus de temps pour réagir face aux situations d'urgence.
- **Message à valeur ajoutée** : Ce type de message vise à améliorer le confort des usagers de la route en fournissant des informations telles que la localisation des restaurants ou des hôtels, ainsi que des données multimédias.

Le réseau VANET se compose de véhicules à haute mobilité et de stations disposées de manière sporadique le long des routes. Chaque véhicule et station est équipé de dispositifs de communication sans fil et, éventuellement, de dispositifs de détection. Les VANETs, qui se concentraient initialement sur les communications entre véhicules, s'étendent désormais à d'autres types de communications impliquant les véhicules et les infrastructures routières "Road-Side Units (RSUs)".

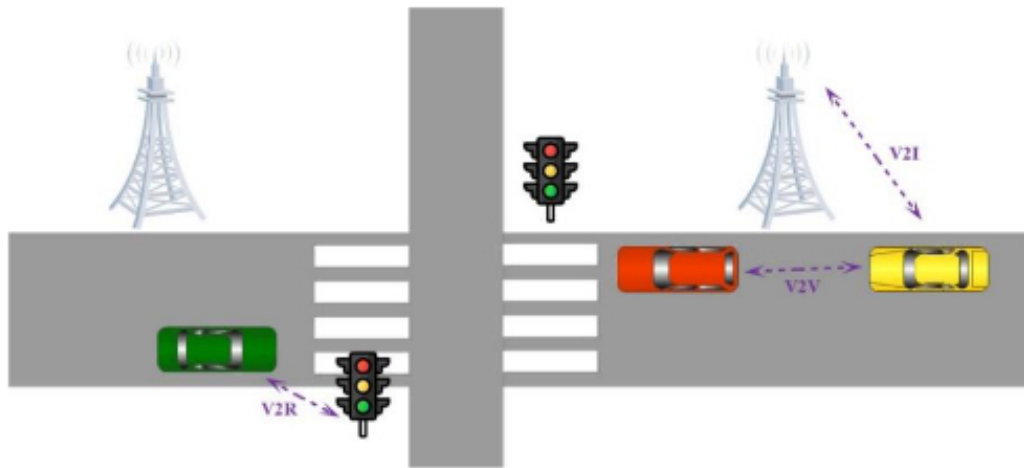


FIGURE 1.1 – Types de communications dans les VANETs

Pour établir les communications entre ces différentes entités de l'infrastructure du réseau VANET, diverses technologies ont été appliquées pour offrir divers services et augmenter la portée des communications et la bande passante. Ainsi, une norme de communication appelée Dedicated Short Range Communication (DSRC) (Dedicated Short Range Communication) a été adoptée, avec une couche physique basée sur la norme IEEE 802.11a. Plus tard, l'IEEE s'est inspirée de cette norme pour créer la norme 802.11p actuellement utilisée, qui définit principalement les services de sécurité et le format des messages.

Pour permettre l'échange d'informations entre les entités dans l'environnement VANET, les unités de base impliquées dans la communication sont : l'unité de bord de Véhicule "OnBoard Unit (OBU)" et l'unité de bord de route (RSU).

- **L'unité de bord de Véhicule (OBU)** : Il s'agit d'un dispositif électronique installé à l'intérieur des véhicules, comprenant un processeur, un système de positionnement global "Global Positioning System (GPS)", une mémoire vive et des capteurs. Il utilise les réseaux VANETs pour échanger des informations avec d'autres OBUs et RSUs.
- **L'unité de bord de route (RSU)** : Les RSUs sont des unités fixes situées le long de la route et conçues pour offrir une connectivité et un support d'information aux véhicules, y compris des alertes de sécurité et des mises à jour sur le trafic.

Les VANETs, basés sur Les normes IEEE 802.11 a/b/g/p offrent des fonctionnalités précieuses pour les communications entre les entités d'environnement VANETs. Comme ils permettent une communication directe entre les points de terminaison source et de destination, ils peuvent fonctionner indépendamment de la couverture réseau par un réseau d'infrastructure, fonctionnant de manière entièrement décentralisée, aident parfois à prévenir les accidents et à améliorer la qualité routière, mais leur efficacité n'est pas toujours garantie. Les couches physiques et MAC du DSRC ont subi des modifications pour s'adapter au mode ad hoc de l'environnement véhiculaire. Le DSRC a été confronté à de nombreux défis dans la prise en charge des communications véhiculaires, notamment les suivants :

- L'inconvénient majeur des normes IEEE 802.11 est que ses performances de débit et de délai se détériorent rapidement à mesure que la charge du réseau augmente, par exemple en cas de densité d'utilisateurs élevée. Ce problème découle de la stratégie d'accès au support basée sur les conflits "Carrier Sense Multiple Access with Collision Avoidance (CSMA/CA)" centrale à "Enhanced Distributed Channel Access (EDCA)", qui adhère au principe "écouter avant de parler" [1].
- En plus, les VANETs ne peuvent pas assurer une détection en temps réel des conditions routières ou maintenir une connectivité fiable en raison de la forte mobilité des véhicules, et une portée de communication réduite qui est généralement de quelques centaines de 100 mètres.
- IEEE802.11p n'est pas adapté aux applications à données élevées et à latence ultra faible en raison de sa surcharge de communication importante, qui consomme de la bande passante et augmente la latence via les retransmissions d'erreurs [2].
- IEEE802.11p est toujours confronté à des problèmes de livraison des paquets peu fiable dans les réseaux à grande échelle en raison des limitations du spectre radio disponible.

Les communautés de normalisation, comme le "3rd Generation Partnership Project (3GPP)", proposent l'intégration des technologies cellulaires dans les communications véhiculaires.

1.3 La technologie 5G

L'évolution des réseaux sans fil a été remarquablement transformatrice au cours des dernières décennies, portée par l'émergence de générations d'appareils sans fil de plus en plus puissantes. Cette progression a été principalement alimentée par la demande d'une qualité d'expérience (QoE) améliorée, une augmentation continue du nombre d'appareils utilisateurs et une utilisation croissante des données [3]. Le réseau "LTE" de quatrième génération (4G) s'est imposé comme un élément fondamental de la vie quotidienne, offrant des avantages significatifs par rapport à son prédécesseur, le réseau mobile de troisième génération (3G). Malgré le succès de la 4G, l'avènement de l'Internet des objets "Internet des objets (IoT)" et la prolifération de nouvelles applications exigeant une latence ultra faible et une capacité élevée ont compliqué la gestion des ressources réseau. Cette complexité résulte de la nécessité de répondre à un grand nombre de demandes tout en prenant en charge des exigences hétérogènes en matière de QoS. Par conséquent, ces défis ont mis en évidence les limites de la 4G et la nécessité d'une infrastructure réseau plus avancée. La cinquième génération (5G) est conçue pour répondre aux limites techniques de la 4G. Cette technologie offre une amélioration substantielle de la QoS perçue par les utilisateurs par rapport au réseau 4G existant et, surtout, introduit de nouvelles opportunités pour l'industrie verticale. Le réseau 5G devrait offrir des débits de données allant de 1 à 10 Gbps, soit près de dix fois le débit de données maximal théorique de 150 Mbps offert par les réseaux 4G LTE [4]. Avec ces débits de données, la 5G fournira des services de haute qualité avec une qualité de service assurée pour l'utilisateur final et des expériences de haut débit mobile véritablement omniprésentes et illimitées, même dans les zones très fréquentées telles que les stades, les voitures, les trains, les concerts ou les centres commerciaux, via des terminaux améliorés avec capacités d'intelligence artificielle (IA) [5]. La 5G met l'accent sur plusieurs aspects clés, illustrés dans la figure 1.2, notamment [4,6] :

- **Une latence plus faible** : La latence est une mesure cruciale dans la définition du réseau 5G. Il mesure le temps nécessaire aux données pour voyager entre un appareil 5G et la station de base. La 5G s'efforce de réduire la latence à 1 milliseconde seulement, car elle est conçue pour prendre en charge les systèmes vitaux, les services avec une tolérance de retard minimale et les applications en temps réel.
- **Une plus grande capacité réseau** : Le réseau 5G vise à augmenter sa capacité. Il a également l'intention de prendre en charge une grande variété de types d'appareils avec une connectivité omniprésente. De plus, la 5G cherche à offrir une expérience utilisateur ininterrompue en maintenant une approche centrée sur l'utilisateur grâce à une connectivité complète.
- **Une bande passante et des vitesses de données plus élevées** : Le réseau 5G améliorera considérablement les connexions à haut débit, permettant une transmission et une réception rapides des données. Son objectif est d'atteindre des vitesses de téléchargement allant jusqu'à 10 gigabits par seconde, soit plus de dix fois plus rapides que celles offertes par la 4G.
- **La disponibilité** : La disponibilité est définie comme le pourcentage d'utilisateurs ou de liens de communication dans une zone géographique spécifique qui répondent

aux exigences de QoE. Le contrôle des véhicules autonomes nécessite une disponibilité de 99,999

- **La sécurité** : Mesurer la sécurité d'une communication donnée est assez difficile. Une méthode pour le quantifier consiste à estimer le temps nécessaire pour qu'une entité non autorisée puisse accéder à l'information.
- **La densité du volume de trafic** : La densité du volume de trafic est définie comme la quantité totale de trafic échangé par tous les appareils de la zone considérée pendant une période de temps prédéfinie divisée par la taille de la zone. Par exemple, pour une ville intelligente, la densité de volume de trafic devrait être de 700 Gbps/km².
- **La Fiabilité** : La fiabilité est généralement définie comme la probabilité que des données soient transmises avec succès d'un émetteur à un récepteur avant qu'un certain délai n'expire. La fiabilité doit être de 99% pour les applications liées à la sécurité.
- **La Consommation d'énergie** : La consommation d'énergie est généralement définie comme l'énergie par bit d'information (particulièrement pertinente dans les environnements urbains) et comme la puissance par unité de surface (souvent pertinente dans les environnements suburbains/ruraux). Par exemple, pour les communications d'urgence, l'autonomie de la batterie doit être d'une semaine.
- **Densité de connexion** : La densité de connexion est le rapport entre le nombre d'appareils actifs ou d'utilisateurs présents simultanément dans une zone définie sur une période de temps spécifiée et la taille de cette zone. Par exemple, devant les stades, la densité de connexion doit être 200 000 utilisateurs par km².

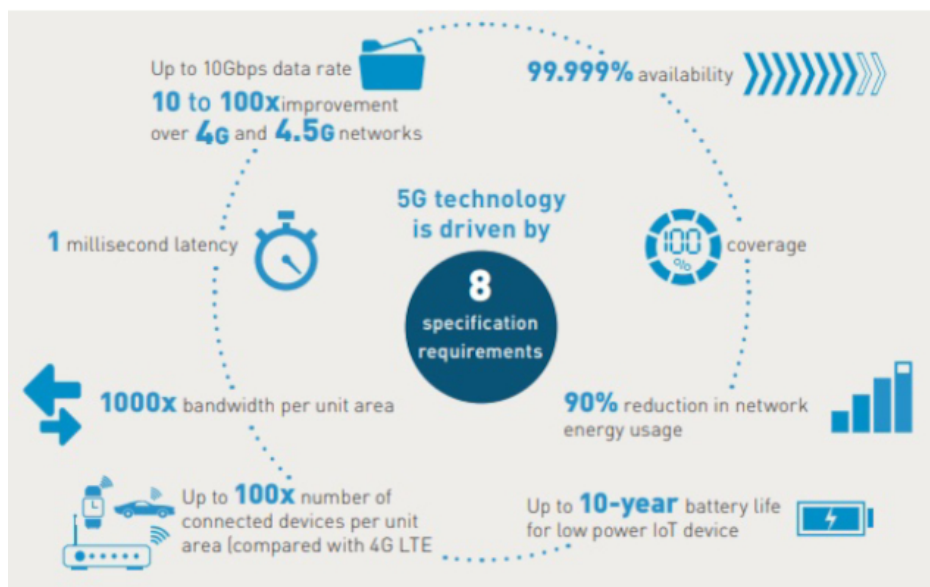


FIGURE 1.2 – Exigences de spécifications de la technologie 5G

[7]

1.4 Découpage réseau

Dans cette section, nous présentons l'historique et quelques définitions concernant le découpage réseau puis nous mettons l'accent sur le découpage réseau dans la 5G.

1.4.1 Définition et histoire

Le concept de découpage de réseau trouve ses racines dans les années 1960, fortement influencé par les principes de virtualisation, notamment avec la conception du premier système d'exploitation d'IBM (CP-40), prenant en charge le partage du temps et la mémoire virtuelle. Ce système permettait à jusqu'à quinze utilisateurs de travailler simultanément, chacun accédant indépendamment à un ensemble distinct de ressources matérielles et logicielles [8]. Cela a marqué le début de la virtualisation des réseaux, où des entités virtuelles pouvaient être créées à partir d'entités physiques. La vision était d'étendre les systèmes virtuels sur diverses ressources réseau, infrastructures informatiques et périphériques de stockage. Dans le début des années 1980, la virtualisation des réseaux a gagné du terrain dans les centres de données, permettant des connexions sécurisées et aux performances contrôlées entre des sites distants via Internet. À la fin des années 1980, l'idée des réseaux superposés est apparue, dans laquelle des nœuds de réseau connectés via des liens logiques formaient des réseaux virtuels sur une infrastructure physique partagée. Ces réseaux superposés constituaient une des premières formes de découpage de réseau, combinant les ressources de divers domaines administratifs tout en garantissant la QoS aux utilisateurs finaux. Cependant, ils manquaient d'automatisation et de programmabilité dans les contrôles réseau. Les années 1990 et le début des années 2000 ont vu la proposition de réseaux actifs et programmables, dans lesquels les systèmes d'exploitation de nœuds fournissaient des cadres de contrôle des ressources [9]. En 2009, le modèle d'architecture réseau "Software Defined Networking (SDN)" a permis aux chercheurs de mener des expériences au sein d'une tranche de réseau d'un réseau de campus, en utilisant la programmabilité via des interfaces ouvertes. Cette évolution a jeté les bases du découpage de réseau moderne [10].

1.4.2 Le découpage réseau dans la 5G

L'architecture de l'infrastructure du réseau mobile des générations précédentes de type "the one-size-fits-all", telle que dans la 4G, est insuffisante pour répondre aux diverses exigences des nouvelles applications 5G, ce qui a conduit à la conception d'une nouvelle architecture [11]. La nouvelle architecture doit offrir une plateforme flexible, s'adaptant à un large éventail de cas d'utilisation plutôt que de se concentrer sur une seule "application tueuse 5G". Il doit être adaptable à divers besoins, en particulier ceux des industries verticales telles que l'automobile, l'énergie et la fabrication, leur permettant d'obtenir des solutions personnalisées utilisant une infrastructure réseau partagée [4]. La nouvelle génération des réseaux mobiles "Next Generation Mobile Network (NGMN)" a été la première qui a introduit le découpage du réseau 5G comme nouveau paradigme pour répondre aux diverses exigences des nouvelles applications 5G [12]. Cette génération "NGMN" a défini une tranche de réseau comme un réseau logique de bout en bout fonctionnant sur une infrastructure physique sous-jacente partagée. Il est configuré pour répondre aux besoins variés de cas d'utilisation spécifiques concernant les fonctionnalités telles que la sécurité

et la prise en charge de la mobilité et les performances de livraison (y compris la latence, la fiabilité et le débit), mutuellement isolées, dispose d'un contrôle et d'une gestion indépendants et peut être créé sur demande. Le découpage de réseau vise à permettre la création de différents types de tranches pour diverses applications sur la même infrastructure. Cette création de tranches de réseau est réalisée en allouant des portions des ressources de l'infrastructure physique partagée à chaque tranche. Une tranche de réseau peut inclure des composants inter-domaines provenant de différents domaines au sein de la même ou de différentes administrations, ou des composants pertinents pour le réseau d'accès, le réseau de transport, le réseau central et le réseau périphérique [9]. Les tranches de réseau sont autonomes, isolées l'une de l'autre, étables et programmables pour prendre en charge le multi-service et la multi-location. Outre la génération NGMN, d'autres organismes de normalisation ont exploré la définition du découpage de réseau sous différents angles. L'Union Internationale des Télécommunications (UIT) considère le découpage du réseau comme un concept fondamental de la softwarisation du réseau, facilitant les partitions de réseau logiques isolées constituées de plusieurs ressources virtuelles. Ces tranches sont isolées et équipées d'un plan de contrôle et de données programmables [13].

Le 3GPP définit le découpage de réseau comme une technologie qui permet aux opérateurs de créer des réseaux sur mesure pour fournir des solutions optimisées pour divers scénarios de marché avec des exigences diverses, telles que la fonctionnalité, les performances et l'isolation [14]. Avec une gestion du cycle de vie de tranche en quatre étapes [15, 16] comme montre la figure 1.3 :

1) La phase de préparation de l'environnement réseau : Cette phase est la première dans la vie d'une tranche durant laquelle, les tâches suivantes sont réalisées :

- Création et vérification des tranches réseau selon les exigences/besoins du client.
- Préparation de l'environnement réseau nécessaire pour faciliter le cycle de vie de la tranche réseau.
- Planification de la capacité de la tranche réseau.
- Intégration des tranches réseau.
- Évaluation des besoins de la tranche réseau.
- Identification et préparation des exigences supplémentaires pour le réseau.

2) La phase d'activation : Durant cette phase, l'allocation des ressources à la tranche réseau est confirmée, qu'elles soient partagées ou dédiées.

3) La phase d'exécution : Dans cette phase, la tranche est prête à accepter les demandes de connexion de ces utilisateurs et à les servir.

4) La phase de terminaison : Cette phase consiste à désactiver la tranche réseau et à libérer les ressources allouées.

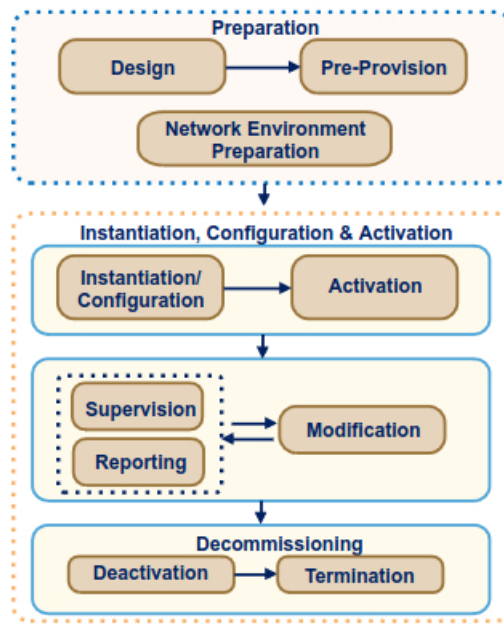


FIGURE 1.3 – Cycle de vie des tranches dans V2X 5G
[16]

1.4.3 Outils de découpage du réseau dans la 5G

a) L'architecture "Software Defined Networking (SDN)"

L'infrastructure réseau traditionnelle est basée sur un modèle dans lequel les routeurs intègrent à la fois une logique de contrôle et de données; ce qui signifie que chaque périphérique réseau dispose d'un plan de contrôle pour construire une table de transfert et d'un plan de données qui transfère le trafic sur la base des règles de la table de transfert, conduisant à une gestion décentralisée et complexe.

L'architecture SDN est un paradigme qui introduit de la flexibilité dans les réseaux 5G, permettant une orchestration et un contrôle intelligents, précis et à l'échelle du réseau, des applications et des services [10, 17]. Dans une architecture basée sur SDN comme montre la figure 1.4, les fonctionnalités de contrôle sont séparées du réseau de transmission de données, ce qui signifie que l'architecture SDN sépare le processus de transfert des paquets réseau (plan de données) du processus de routage (plan de contrôle), ce qui donne lieu à un plan de contrôle centralisé pour effectuer les fonctionnalités de la couche de contrôle. Ce plan de contrôle centralisé peut être représenté par plusieurs contrôleurs. Ces contrôleurs sont considérés comme le cerveau du réseau SDN, intégrant toute l'intelligence.

Ce contrôle centralisé, avec une vue globale sur l'ensemble du réseau, se traduira par une gestion plus facile des périphériques réseau. Cela fournit des capacités pour répondre rapidement aux évolutions des conditions du réseau, aux exigences commerciales, à la dynamique du marché et aux besoins des utilisateurs finaux. La scalabilité peut constituer un inconvénient, car le contrôleur centralisé peut rencontrer des inefficacités. Pour "Open Networking Foundation (ONF)" [18], l'architecture SDN est représentée dans les trois couches qui sont généralement définies comme suit :

1. Couche d'application (plan d'application) :

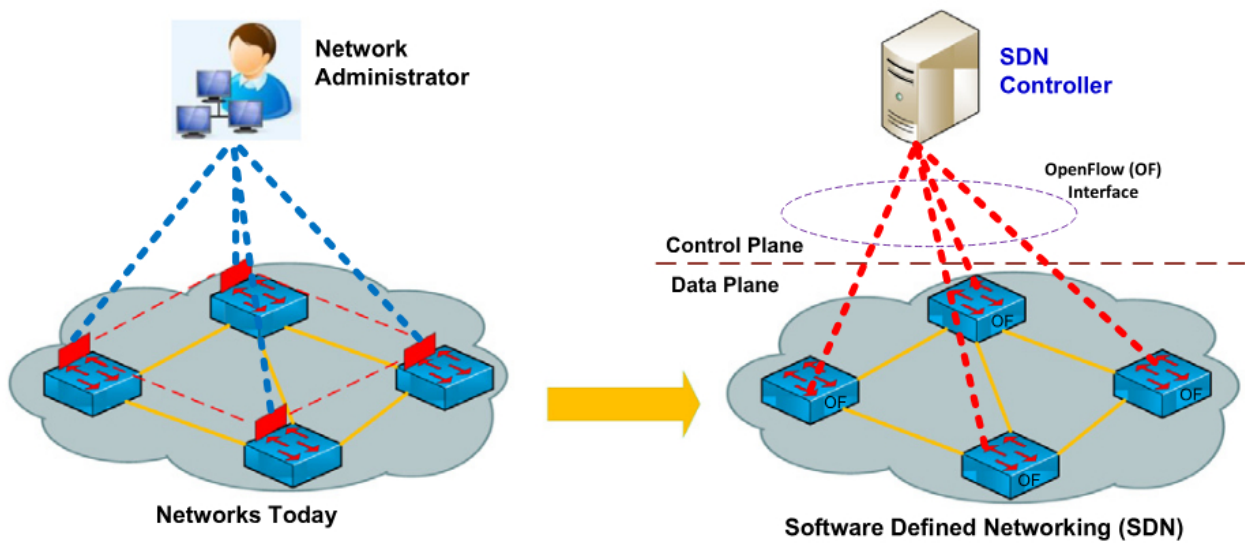


FIGURE 1.4 – Passage vers une architecture à base de SDN
[9]

- Cette couche comprend les applications et les services qui s'exécutent sur l'infrastructure SDN.
- Ces applications interagissent avec le contrôleur SDN via des APIs pour demander des services et des configurations réseau.
- Ces applications incluent la gestion de réseau, les outils d'ingénierie du trafic, les applications de sécurité, etc.

2. Couche de contrôle (plan de contrôle) :

- La couche de contrôle héberge le contrôleur SDN, qui est responsable de la gestion et du contrôle du réseau.
- Traduit les exigences du plan d'application SDN au plan de données.
- Il résume les fonctions de transfert du réseau et fournit une vue et un contrôle centralisés sur les périphériques réseau.
- Le contrôleur communique avec les périphériques réseau (commutateurs, routeurs) dans le plan de données pour programmer des comportements de transfert basés sur les politiques et exigences du réseau.

3. Couche de données (plan de données) :

- La couche de données, également appelée plan de données, est constituée de périphériques réseau tels que des commutateurs et des routeurs.
- Ces appareils transmettent les paquets selon les règles et configurations reçues du contrôleur SDN.

Dans l'architecture SDN, chaque couche communique avec sa couche adjacente à l'aide d'interfaces spécifiques comme montre la figure 1.5 :

1. SDN SouthBound Interface : Est l'interface définie entre un contrôleur SDN et les éléments du plan de données. Elle permet un contrôle programmatique des opérations de transfert, des rapports de statistiques et des notifications d'événements. Cette interface permet au contrôleur SDN d'effectuer des ajustements dynamiques en fonction des exigences du réseau en temps réel.
2. SDN NorthBound Interface : Est l'interface reliant les applications SDN aux contrôleurs SDN. Elle permet aux applications SDN d'accéder à des représentations de réseau abstraites et d'articuler directement les comportements et les exigences du réseau.

SDN permet la programmabilité directe du contrôle du réseau via des interfaces South-Bound (SBI) standardisées telles que OpFlex [9], Forwarding and Control Element Separation (FoRCES) [19] et OpenFlow [20]. Ces interfaces spécifient comment la communication s'effectue entre les périphériques de transfert du plan de données et les éléments des plans de contrôle et de gestion.

La fondation pour réseaux ouverts "ONF" [21] propose une analyse détaillée du paradigme SDN pour le découpage des réseaux 5G. Dans ce cadre, le contrôleur SDN est chargé de gérer les tranches de réseau via des règles ou politiques prédéfinies. Il permet la création de contextes serveur et client et facilite l'installation de leurs politiques respectives [22]. Plus précisément, le contrôleur SDN maintient un contexte client de tranche réseau, ce qui lui permet de superviser dynamiquement les tranches réseau. Le contrôleur régit ces tranches et effectue l'orchestration des ressources dans le contexte du serveur.

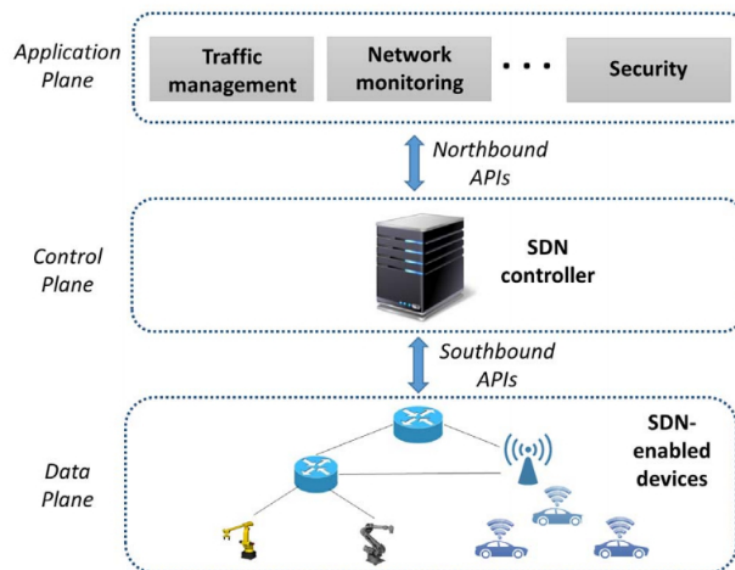


FIGURE 1.5 – L'architecture du SDN [23]

b) La virtualisation des fonctions réseau (NFV)

La virtualisation des fonctions réseau "Network Function Virtualization (NFV)" [24] est conçue pour créer une architecture permettant de virtualiser les fonctionnalités sur du

matériel standard tel que les pare-feux, les optimiseurs "Transmission Control Protocol (TCP)", "Network Address Translation (NAT64)", "Virtual Private Network (VPN)", et "Deep Packet Inspection (DPI)". La NFV envisage de déployer ces fonctions de réseau virtuel sur du matériel standard, passant d'une infrastructure centrée sur le matériel à une infrastructure centrée sur le logiciel. Avec la NFV, les fonctions réseau peuvent être facilement déployées et allouées dynamiquement sur différents périphériques matériels du réseau. De plus, les ressources réseau peuvent être affectées efficacement aux fonctions de réseau virtuel "Virtual Network Function (VNF)". Ces VNF peuvent être liées ou intégrées en tant que composants pour fournir un service de communication réseau optimal. Une fonction virtuelle peut consister en une ou plusieurs machines virtuelles fonctionnant sur divers appareils tels que des serveurs, des nœuds de calcul de pointe, une infrastructure cloud et des commutateurs réseau [25] comme c'est illustré par la figure 1.6. Cette approche permet l'instanciation de fonctions de réseau virtuel sans avoir besoin d'ajouter de nouveau matériel.

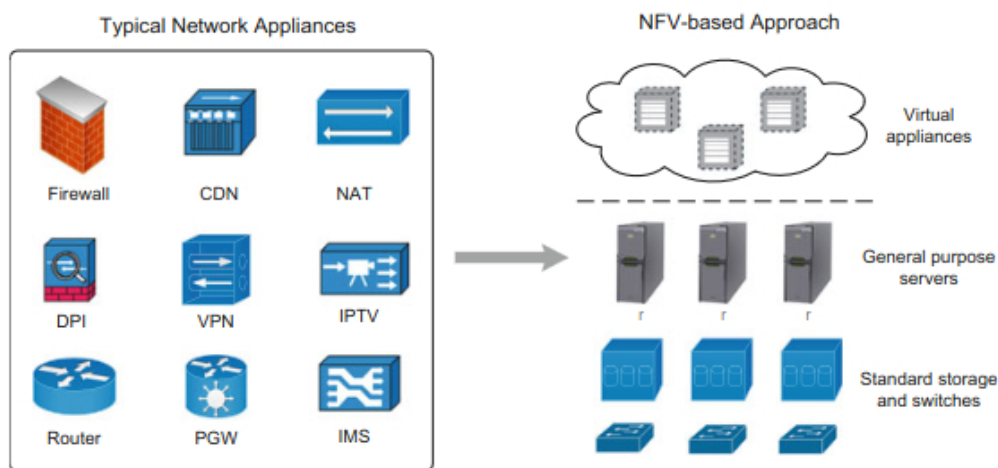


FIGURE 1.6 – Passage vers une virtualisation des fonctions réseau [25]

Étant donné que le logiciel est découplé du matériel, les deux peuvent exécuter diverses fonctions à des moments différents. Cela offre une flexibilité significative aux opérateurs pour introduire de nouveaux services innovants en utilisant la même plate-forme matérielle, avec une fourniture de services dynamique qui peut être adaptée en fonction des demandes des clients en faisant évoluer de manière dynamique les performances de NFV. Dans le cadre de la 5G, un groupe de VNF peuvent être regroupés pour simplifier la gestion. Par exemple, le Serving Gateway (SGW) et le Packet Data Network Gateway (PGW) d'un réseau central "Core Network (CN)" 4G peuvent être regroupés en un seul package, ou ils peuvent être divisés en blocs fonctionnels plus petits pour améliorer la réutilisations et réduire les temps de réponse [26].

L'architecture NFV est définie par trois composants principaux :

1. Gestion et orchestration du NFV (NFV-MANO) : NFV-MANO est spécifiquement conçu pour orchestrer et gérer le cycle de vie des ressources physiques ou virtuelles qui sous-tendent la virtualisation de l'infrastructure et la gestion des fonctions de

réseau virtuel (VNF). Il se compose de trois composants principaux : l'orchestrateur, les gestionnaires VNF et les gestionnaires d'infrastructure virtualisée. Le concept de NFV dans les infrastructures des opérateurs a été initialement étudié par l'Institut européen des normes de télécommunications (ETSI) [18]. Cette exploration visait principalement à relever les défis liés aux services flexibles et agiles et à établir une plate-forme facilitant la monétisation future du réseau.

2. L'infrastructure NFV (NFVI) : NFVI englobe à la fois des ressources matérielles et logicielles, notamment du matériel informatique, de stockage et de réseau. L'architecture NFV peut bénéficier d'une couche de virtualisation existante, comme un hyperviseur, qui exécute des fonctions standards pour extraire les ressources matérielles et les allouer aux fonctions de réseau virtuel (VNF).
3. Les services : Un service est un ensemble des VNFs, qui peuvent être déployées dans une ou plusieurs machines virtuelles. La création, la configuration, la surveillance et la gestion des performances de la fonction de réseau virtuel sont gérées par l'EMS (Element Management System). Un EMS est également chargé de fournir des informations essentielles au Operational Support System (OSS). L'intégration entre l'OSS et le Business Support System (BSS) permet aux fournisseurs de déployer et de gérer efficacement plusieurs services de communication de bout en bout.

Dans le contexte du découpage de réseau 5G, NFV agit comme un outil de découpage d'un réseau sans fil, simplifiant la création et la gestion de tranches en virtualisant certaines fonctions et en les exécutant de manière centralisée, plutôt que de s'appuyer sur du matériel propriétaire [24]. Le cadre de virtualisation des fonctions réseau permet la gestion des fonctions de réseau virtuel et du chaînage de services. L'architecture est composée de trois composants principaux [18] :

1. Le gestionnaire de fonctions de réseau virtuel : Est responsable de la gestion du cycle de vie des instances des fonctions de tranche du réseau.
2. Le gestionnaire d'infrastructure virtuelle : Effectue la gestion et le contrôle des ressources de l'infrastructure de virtualisation des fonctions réseau de l'opérateur réseau.
3. L'orchestrateur de virtualisation des fonctions réseau : Gère l'orchestration des ressources réseau, valide et autorise les demandes du gestionnaire de fonctions réseau virtuel pour les ressources d'infrastructure de virtualisation des fonctions réseau et gère le cycle de vie de la tranche réseau.

Une tranche réseau peut être décomposée en un ensemble de fonctions réseau virtualisées, qui sont ensuite exécutées sur des serveurs situés à différents emplacements. Cette approche permet aux fonctions réseau virtualisées d'être facilement créées, déplacées ou détruites n'importe où et à tout moment, offrant ainsi une flexibilité et réduisant les coûts pour l'opérateur réseau [24].

c) Cloud/Edge computing

Le cloud computing est un paradigme qui permet d'accéder et d'utiliser des ressources informatiques diverses telles que des serveurs, du stockage, des bases de données, des réseaux, des logiciels, etc., en ligne [27].

Le cloud peut être réalisé selon trois types de modèles de livraison comme c'est illustré par la figure 1.7 :

1. Infrastructure as a Service (IaaS) : Loue des ressources informatiques virtualisées en tant que service sur Internet. Exemples de IAAS incluent Cisco Metapod, and Google Compute Engine.
2. Platform as a Service (PaaS) : Il permet de développer, exécuter et gérer différentes applications sans avoir à gérer la complexité de la création et de la maintenance de l'infrastructure cloud. Exemples de PAAS incluent Apprenda, Windows Azure, les services Web Amazon.
3. Software as a Service (SaaS) : Ce type des services donne la possibilité aux utilisateurs d'accéder à divers logiciels à la demande sans avoir besoin de les installer sur un ordinateur personnel local. Le logiciel est hébergé sur le cloud et disponible pour les utilisateurs, qui ne le paient qu'en fonction de leur utilisation, et ce logiciel peut être accessible quel que soit l'emplacement. Des exemples de SaaS incluent Dropbox, Gmail, Microsoft Office 365 et Facebook.

Pour les applications temps réel et gourmandes en ressources, les applications sensibles à la latence telles que la réalité augmentée et la e-santé trouvent le cloud moins adapté à leurs besoins en raison de problèmes de latence. L'Edge Computing résout ce problème en rapprochant les ressources cloud de la périphérie du réseau.

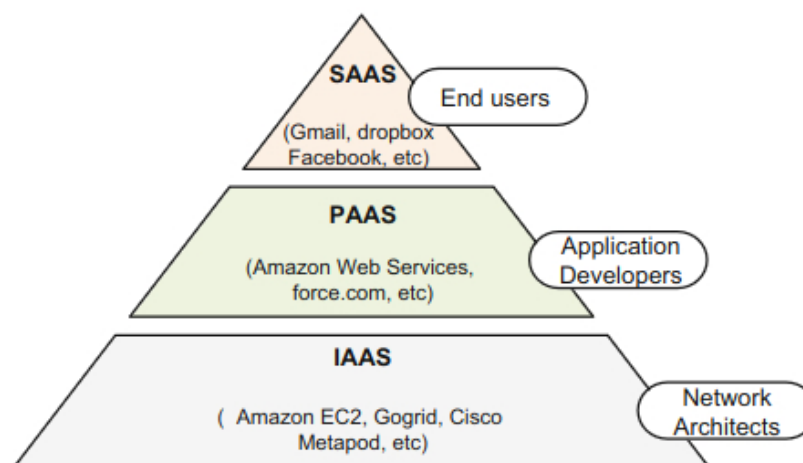


FIGURE 1.7 – Les modèles de Cloud
[25]

1.4.4 Les types de découpage réseau dans la 5G

Une tranche de réseau de bout en bout (E2E) dans les réseaux mobiles 5G s'étend sur la partie réseau central (CN), le réseau du transport et la partie radio. Du point de vue du 3GPP [28], les solutions de découpage de réseau sont classées en trois ensembles :

- L'ensemble *A* comprend un RAN partagé avec des tranches de réseau central entièrement dédiées, conduisant à une gestion indépendante des abonnements, des sessions et de la mobilité pour chaque tranche.
- L'ensemble *B* utilise également un RAN commun, mais partage la gestion des abonnements et de la mobilité entre les tranches, tandis que la gestion des sessions reste dédiée à chaque tranche.
- L'ensemble *C* envisage un RAN entièrement partagé et un plan de contrôle de réseau central commun, avec des plans utilisateur du réseau central affectés à des tranches spécifiques.

a) Le réseau central (CN)

L'architecture du réseau central définie par les générations précédentes n'est pas adaptée aux nouvelles exigences du réseau 5G, qui incluent l'élasticité, scalabilité et la flexibilité, ce qui nécessite des adaptations pour répondre à ces demandes. Cependant, en raison des contraintes variables des différentes tranches réseau déployées, une infrastructure de CN unique pourrait ne pas convenir aux besoins de toutes les tranches [29]. Par exemple, pour les services à faible latence tels que les véhicules autonomes, il est peu pratique de les accueillir si les fonctions chargées de la gestion de la mobilité sont situées proche du RAN. Pour cela, il est devenu crucial de diviser certaines fonctions du CN en plusieurs sous-fonctions ou de regrouper certaines fonctions sur un même serveur, en fonction des exigences de chaque type de tranche. Par exemple, les entités de réseau LTE MME, S-GW et P-GW ont été remplacées par des fonctions telles que la gestion de la mobilité, le routage et le transfert, le contrôle d'accès et la tarification [30]. En appliquant NFV, ces fonctions réseau plus granulaires sont virtualisées, ce qui leur permet d'être déployées à différents endroits du réseau [31, 32].

Le modèle 3GPP a remodelé l'entité EPC, le réseau central de la technologie 4G, en fonctions réseau plus granulaires, à savoir le nouveau CN [33]. Les principales fonctions de réseau définies dans le nouveau CN sont les suivantes :

1. Fonction de gestion de session "Session Management Function (SMF)" : La fonction de gestion de session (SMF) dans la 5G est responsable de la gestion des sessions, y compris l'établissement, la modification et la fin des sessions lorsqu'elles ne sont plus nécessaires ou lorsque l'utilisateur se déconnecte du réseau.
2. Gestion unifiée des données (UDM) : C'est une fonction qui est similaire à l'entité HSS de LTE, responsable de la gestion des données et des profils des abonnés, ce qui est important pour authentifier les utilisateurs et stocker les détails de l'abonnement.

3. Fonction de gestion des accès et de la mobilité "Access and Mobility Function (AMF)" : Est responsable de la gestion de l'accès et de la mobilité des équipements utilisateur (UE). Cette fonction est importante pour garantir une connectivité transparente et une prise en charge de la mobilité pour les utilisateurs sur l'ensemble du réseau, et intègre la fonctionnalité de tranche de réseau dans le cadre de son ensemble de fonctions de base.
4. Policy Control Function (PCF) : Cette fonction gère la qualité de service et assure l'application des politiques de contrôle pour garantir que les exigences des utilisateurs de chaque tranche, ce qui correspond au PCRF en LTE.
5. User Plane Function (UPF) : La fonction User Plane Function (UPF) est responsable du routage et du transfert des paquets, de l'inspection des paquets et de la gestion de la QoS. Elle joue un rôle crucial dans le découpage du réseau, permettant la création de tranches de réseau adaptées aux besoins des différentes applications et services, tout en garantissant une transmission de données efficace et fiable.

De plus, la composition de ces fonctions d'une certaine manière permettrait d'adapter le CN 5G à une tranche de réseau particulière. Certaines fonctions du réseau peuvent être partagées entre différentes tranches pour réduire la charge de signalisation et simplifier la gestion des tranches du réseau en partageant les processus de gestion de la mobilité.

b) Le réseau de transport

Le réseau de transport "Transport Network (TN)" assure la connectivité entre les composants RAN désagrégés et le réseau central. Le volume croissant du trafic présentera des défis non seulement pour les liaisons d'accès sans fil, mais également pour le réseau de transport. La désagrégation du RAN et les exigences émergentes des services nécessitent une architecture flexible pour les réseaux de transport [34]. Par conséquent, le secteur de normalisation des télécommunications de l'Union internationale des télécommunications (UIT-T) a proposé une architecture de référence pour le contrôle SDN des réseaux de transport [35]. Grâce au découpage de réseau dans les réseaux de transport, les flux de trafic peuvent être différenciés ou isolés les uns des autres, consacrant ainsi des ressources à une instance de tranche de réseau. Cette configuration applique des exigences spécifiques aux tranches tout en tirant parti de la flexibilité du contrôle SDN. Dans [36], les auteurs ont proposé une approche basée sur le SDN, qui permet l'acquisition d'une vue globale de l'allocation des ressources du réseau, et le calcul des chemins de transmission. Cependant, cette approche introduit une surcharge importante du plan de contrôle pour la planification à mesure que la complexité du réseau augmente et peut ne pas optimiser efficacement l'utilisation des ressources avec l'allocation de chemin. La mise en place du partage de réseaux d'accès fixe "Fixed Access Network Sharing (FANS)", qui consiste à partager le réseau physique en plusieurs tranches de réseau virtuel et à les allouer à différentes tranches, basé sur les technologies de virtualisation de réseau [37]. Deux modèles de solutions ont été proposés basés sur le partage du réseau d'accès fixe, qui sont les suivants [26] :

- Introduction du découpage réseau au niveau de la gestion via la configuration d'un réseau privé virtuel "VPN".

- Nœud d'accès virtuel (VAN), dans lequel les ressources sont réparties en plusieurs instances, avec des états de port virtuel à physique gérés par un mappage dans le système de gestion de l'opérateur.

c) Le réseau radio "RAN"

L'avènement de la technologie 5G entraîne une gamme diversifiée d'applications et de services, chacun ayant des exigences de performances uniques. Pour gérer efficacement ces différentes demandes, il existe également un besoin crucial de découper la partie réseau d'accès radio "RAN" du réseau. Les principaux défis pour la virtualisation des infrastructures résident dans le RAN. Les solutions qui pré-attribuent des parties de spectre distinctes aux instances de stations de base "Base Station (BS)" virtuelles (tranches) sont faciles à mettre en œuvre et garantissent l'isolation des ressources radio. Cependant, ils entraînent une utilisation inefficace des ressources radio. L'approche alternative c'est le partage de spectre dynamique et à granularité fine pour éviter cette limitation, ce qui en fait une option plus souhaitable [29]. En fait, chaque tranche du réseau ne peut pas ajuster la quantité de ressources (pRB) allouées à lui, même si elles ne sont pas utilisées. Plus précisément, les pRB au sein du spectre partagé sont gérés par un ordonnanceur centralisé, qui est un planificateur inter-slices. Cet inter-slices planificateur alloue des ressources aux tranches en fonction d'une politique prédéterminée et d'autres critères. Après que chaque tranche a reçu ses ressources dédiées à ce niveau haut, le deuxième niveau consiste à allouer ces ressources aux utilisateurs appartenant à cette tranche, ce qui est une allocation intra-tranche.

Les concepts de SDN et NFV sont principalement motivés par le besoin de flexibilité, permettant des progrès substantiels vers l'avenir dans la partie radio. Cependant, les deux technologies ont également été étendues aux architectures RAN. Avec la virtualisation RAN, le paradigme RaaS [38] va au-delà du simple partage de ressources radio entre les locataires de tranches (par exemple, les "Mobile Virtual Network Operator (MVNO)"). Il permet la création dynamique d'instances RAN virtuelles, chacune avec un ensemble personnalisé de fonctions de contrôle virtualisées (par exemple, planification, gestion de la mobilité) adaptées aux exigences spécifiques d'une tranche ou d'un service. Cette approche garantit également l'isolement entre les différentes tranches (instances virtuelles RAN). Avec ce paradigme, chaque tranche a accès à un pourcentage d'espace dédié pRB et ses propres instances des fonctions, qui sont : "Media Access Control (MAC)", "Packet Data Convergence Protocol (PDCP)", "Radio Resource Control (RRC)" et "Radio Link Control (RLC)". Avec cette approche de la softwarisation BS du RAN, nous pouvons déplacer certaines fonctions du RAN pour qu'elles s'exécutent sur des plates-formes logicielles distantes, comme les Clouds. Le concept Cloud-RAN connaît un développement important [39]. Dans C-RAN, les fonctions du RAN sont réparties entre les Remote Radio Heads (RRHs), qui gèrent l'équipement d'antenne et l'accès radio, et la Baseband Unit (BBU), qui est hébergée dans le cloud. La BBU gère les tâches de traitement en bande de base telles que la modulation et la démodulation du signal, et se connecte aux RRHs via des liaisons à haut débit pour traiter les signaux radio.

1.4.5 Les principes du découpage réseau

Le découpage du réseau dans la 5^{ème} génération doit respecter les principes fondamentaux suivants :

1. **Isolation des tranches** : L'isolation des tranches est une exigence essentielle dans le découpage du réseau, garantissant les performances sécurisées et fiables des différents locataires. Elle permet aux tranches de chaque locataire de fonctionner indépendamment des autres, empêchant ainsi tout accès non autorisé ou toute modification des tranches appartenant à différents locataires. L'isolation est un aspect fondamental de la virtualisation des fonctions réseau, qui inclut également la possibilité d'appliquer des restrictions sur l'utilisation des ressources réseau [25]. L'isolement peut être mis en œuvre via : (i) des ressources partagées utilisant la virtualisation pour la séparation, (ii) une utilisation dédiée de ressources physiques distinctes et (iii) un partage de ressources régi par des politiques spécifiques qui décrivent les droits d'accès des locataires individuels [26]. Cette capacité garantit que les ressources sont partagées équitablement entre les différents locataires, garantissant ainsi des performances optimales pour chacun.
2. **Élasticité** : L'élasticité fait référence à la capacité des tranches du réseau à s'adapter aux conditions changeantes afin de maintenir les garanties "Service Level Agreement (SLA)". L'adaptation peut prendre en considération des facteurs tels que le nombre d'utilisateurs du service, la mobilité des utilisateurs, les conditions radio et les performances globales du réseau. Après avoir pris en compte ces facteurs, l'élasticité peut être obtenue en réaffectant les fonctions du réseau virtuel, en ajustant l'allocation des ressources (augmentation ou diminution) et en reconfigurant les fonctionnalités des éléments de contrôle et de données [40].
3. **Personnalisation de bout en bout** : Garantit que les ressources partagées allouées à différentes tranches sont utilisées efficacement, depuis les fournisseurs de services jusqu'à l'utilisateur final. Une tranche de réseau de bout en bout intègre diverses ressources provenant de différentes parties du réseau, notamment le réseau central (CN), le réseau de transport, le réseau d'accès radio (RAN) et le Cloud. Pour permettre la personnalisation, la virtualisation des fonctions réseau, la mise en réseau définie par logiciel et l'orchestration du réseau peuvent être exploitées pour fournir une allocation des ressources de manière flexible et agile.
4. **Automatisation** : Cela signifie créer automatiquement une tranche de réseau, sans intervention manuelle ni accords contractuels fixes, en fonction des exigences souhaitées telles que le SLA, la latence, le flux et la durée [40].
5. **Programmabilité** : Ce principe permet de contrôler les ressources réseau et cloud de la tranche allouée via des API ouvertes, facilitant l'élasticité des ressources et la personnalisation des services à la demande [40].

1.4.6 Cas d'utilisation génériques dans le découpage de réseau 5G

Dans les réseaux 5G, divers types d'applications auront des exigences différentes en matière de QoS, notamment les applications liées aux véhicules autonomes, à la réalité augmentée, aux industries intelligentes, aux maisons intelligentes et aux villes intelligentes. Les besoins de ces applications peuvent être classés en trois types généraux de services. Cela signifie que chaque service peut avoir sa propre tranche dédiée à ces utilisateurs, initialisé à partir de ces tranches génériques. Les trois tranches génériques dans la 5G sont [4, 25] :

- **Enhanced Mobile Broadband (eMBB)** : les tranches de ce type seront dédiées au support de nouvelles applications nécessitant des débits de données plus élevés comme la réalité augmentée ou virtuelle, ou le streaming vidéo ultra haute définition. Ce débit de données élevé doit également être maintenu en mobilité, même lors de déplacements à grande vitesse en voiture ou en train.
- **Ultra Reliable Low Latency Communication (URLLC)** : ce type de tranche donne la priorité à la disponibilité, à la latence et à la fiabilité plutôt qu'au débit de données pour les applications 5G présentant des exigences extrêmes, telles que des liaisons de communication ultra-fiables à faible latence pour les applications d'urgence et de fabrication industrielle.
- **Massive Machine Type Communication (mMTC)** : cette tranche de réseau 5G doit fournir une densité de connexion la plus élevée pour assurer une connectivité sans fil à des dizaines de milliards d'appareils connectés au réseau, en donnant la priorité à une connectivité évolutive pour le nombre croissant d'appareils IoT tels que dans les villes intelligentes [14], à une transmission efficace de petites charges utiles, à une couverture étendue et à une pénétration profonde des débits de données.

1.5 Découpage réseau 5G V2X

1.5.1 Les types de communications dans les réseaux véhiculaires

Dans les communications "V2X", les véhicules utilisent quatre modes de communication pour faciliter les échanges entre eux et avec les infrastructures comme montre la figure 1.9 :

- **Vehicle-to-Vehicle (V2V)** : Les applications V2V prévoient que les véhicules à proximité échangeront des informations telles que la vitesse, la direction et l'état du freinage. Les véhicules peuvent communiquer directement entre eux ou via des véhicules intermédiaires.
- **Vehicle-to-Infrastructure (V2I)** : V2I permet la transmission de messages entre les véhicules et les infrastructures fixes de systèmes de transport intelligents ("ITS"), appelées RSUs. V2I soutient l'efficacité du trafic en permettant l'échange d'informations sur le marquage des voies, les panneaux de signalisation et les feux de

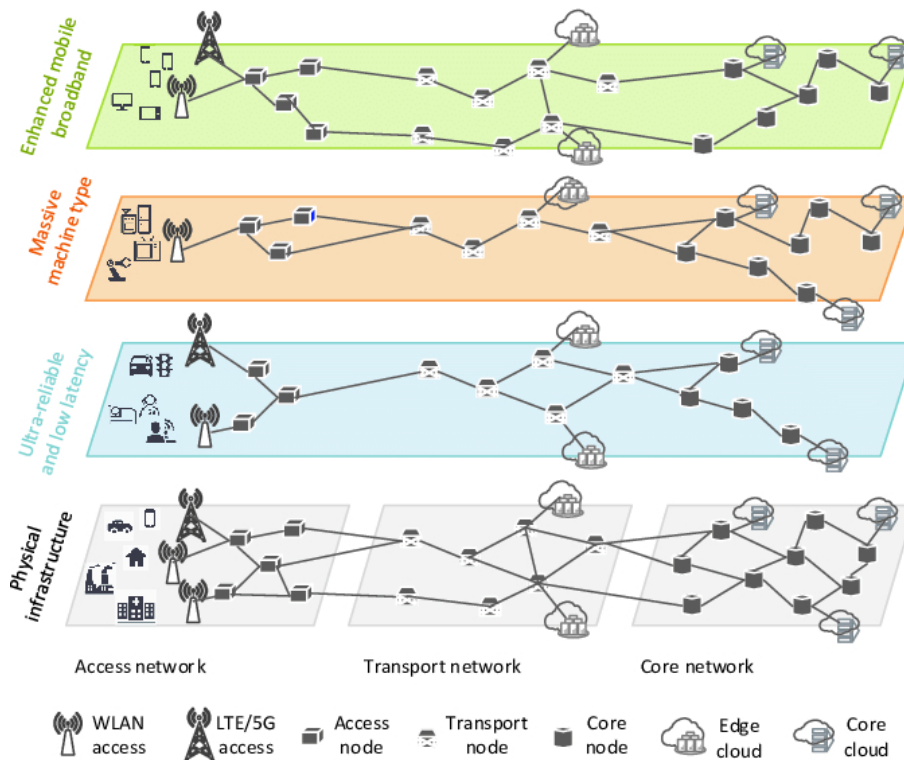


FIGURE 1.8 – Les tranches génériques [41]

circulation. De plus, V2I améliore la sécurité routière en fournissant des messages urgents contenant des informations liées à la sécurité, telles que les accidents de la route et la présence de chantiers de construction.

- **Vehicle-to-Network (V2N) :** V2N permet aux véhicules de communiquer avec un serveur prenant en charge les applications V2N, appelé serveur d'applications V2X. Ce serveur fournit des informations telles que la répartition du trafic, les détails des routes et les mises à jour des services.
- **Vehicle-to-Pedestrian (V2P) :** Le V2P facilite l'échange de messages entre les véhicules et les piétons, qui utilisent leur téléphone ou d'autres appareils sans fil portables pour envoyer et recevoir des informations.

1.5.2 Les divers services V2X avec leurs exigences

Les progrès attendus dans le domaine de la conduite autonome et des véhicules connectés introduiront de nombreuses applications de véhicule à tout (V2X), englobant à la fois des services liés à la sécurité et des services de divertissement. Cela se traduira par un large éventail d'exigences de QoS au sein des réseaux de véhicules [42, 43], et avec le réseau de la 5G qui promet de répondre aux diverses exigences de QoS pour les services V2X. L'initiative de recherche et d'innovation automobile de communication de cinquième génération "5G Automotive Communication (5GCAR)" propose cinq catégories de cas d'utilisation :

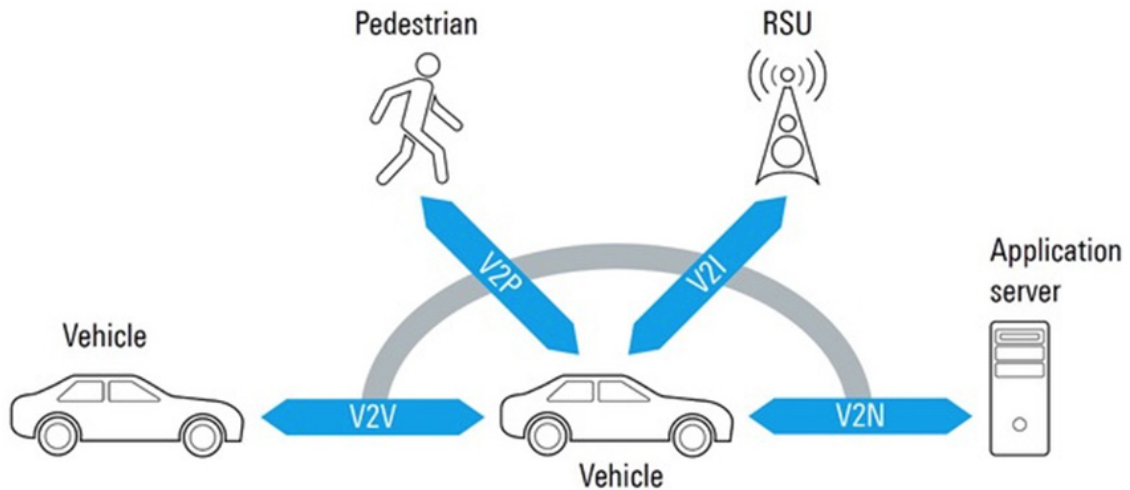


FIGURE 1.9 – Les modes de communications V2X

manœuvres coopératives, perception coopérative, sécurité coopérative, navigation autonome et conduite à distance [44]. Par ailleurs, le projet de la troisième génération "3rd Generation Partnership Project (3GPP)" présente 25 cas d'usage [45]. Ceux-ci sont divisés en quatre catégories principales, ainsi qu'un groupe de cas d'utilisation générale et un autre dédié à la QoS des véhicules. Pour mieux comprendre les défis des communications automobiles, nous examinons les groupes de cas d'utilisation les plus importants. Des informations supplémentaires et des exigences pour chaque cas d'utilisation sont également examinées.

1. **PLATON DE VÉHICULE** : Ce cas d'utilisation permet la formation dynamique d'un convoi de véhicules voyageant ensemble, dans le but d'économiser du carburant et de réduire le nombre de conducteurs, réduire les embouteillages, améliorer la sécurité routière. Il y parvient en échangeant des informations entre les véhicules du peloton, telles que la sélection d'itinéraire, le freinage, l'accélération et le remplacement du véhicule de tête [28]. Cela peut être accompli en échangeant au moins 30 messages de sensibilisation coopérative (Cooperative Awareness Messages (CAMs)) par seconde, avec des tailles de message allant d'environ 50 à 1 200 octets et atteignant une fiabilité de 90 %, en utilisant deux modes de communication V2V et V2I [14, 46]. La figure 1.10 illustre un scénario de partage de l'information entre véhicules en présence d'obstacles.

Par exemple, dans le travail [47], un système prototype 5G-V2X centré sur la conduite autonome collaborative. Ce système est créé à l'aide d'une plate-forme expérimentale de réseau d'accès radio de nouvelle génération, d'un peloton de conduite coopératif et d'un serveur MEC (Mobile Edge Computing) qui fournit des services de cartographie "Three Dimensional (3D)" dynamiques haute définition. Une autre travail dans [48], suggère un système de gestion dynamique de peloton de véhicules autonomes basé sur 5G-V2X communications.

2. **CONDUITE À DISTANCE** : Ce groupe de cas d'utilisation se concentre sur le contrôle à distance de la conduite des véhicules, que ce soit par des humains ou

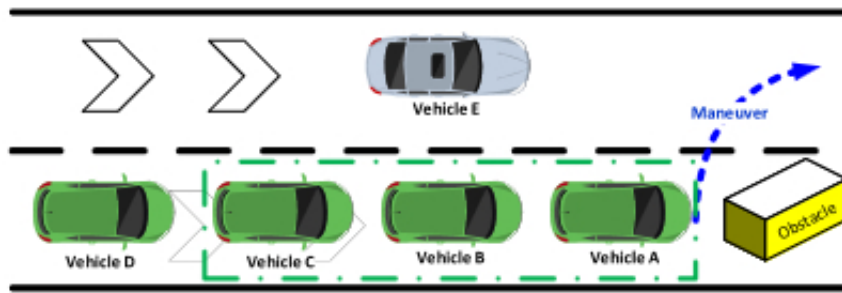


FIGURE 1.10 – Un scénario de partage de l'information sur les obstacles [43]

par des applications cloud/edge computing, via des diffusions vidéo en direct et une capacité d'exécuter des commandes en temps réel sur le véhicule. Cela peut se produire dans des situations où des problèmes inattendus empêchent le véhicule de rouler de manière autonome ou dans des conditions dangereuses où il est nécessaire d'éviter la présence physique du conducteur dans le véhicule [49]. Pour prendre en charge ces services, avec des véhicules atteignant des vitesses allant jusqu'à 250 km/h, les communications V2X, en particulier le mode V2N doit garantir des débits de données allant jusqu'à 1 Mbps sur la liaison descendante "Downlink (DL)" pour la réception par le véhicule des messages de contrôle et de commande liés aux applications, et 20-25 Mbps sur la liaison montante "Uplink (UL)" pour la vidéo et données des capteurs envoyées par le véhicule. De plus, il nécessite une fiabilité de 99,999 % et une latence de bout en bout de 5 ms entre le serveur d'applications V2X et le véhicule. L'étude menée par les chercheurs dans [43] recommande que, pendant la phase initiale, la conduite à distance soit limitée à des routes spécifiques (réseau adapté pour fournir des services logistiques spécialisés) en raison des difficultés à assurer une conduite à distance fiable sur une zone de réseau plus large en particulier dans des conditions de charge élevée ou d'interférences où il peut ne pas être possible de satisfaire à l'exigence de fiabilité de 99,999 %. La figure 1.11 illustre l'architecture générale de la conduite à distance proposée dans [43].

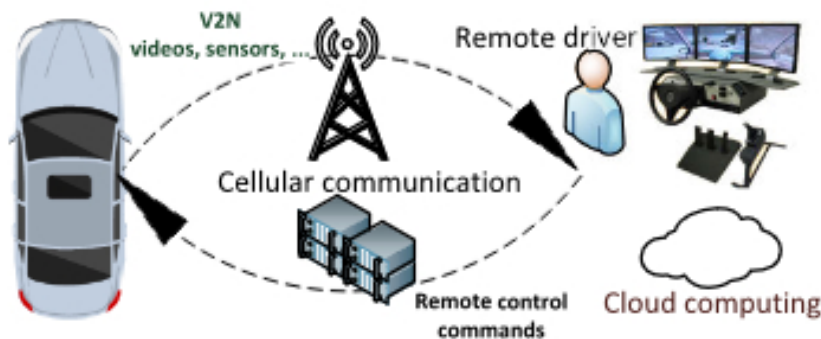


FIGURE 1.11 – Architecture générale de la conduite à distance [43]

3. CAPTEURS ÉTENDUS : Ce groupe est classé en trois cas d'utilisation :

- (a) partage de capteurs et de cartes d'état "Sensor and Status Map Sharing (SSMS)",
- (b) perception collective de l'environnement "Collective Perception of the Environment (CPE)",
- (c) partage de données vidéo pour la conduite automatisée "Video Data Sharing for Automated Driving (VaD)".

De cette manière, les véhicules peuvent améliorer leur conscience environnementale au-delà des capacités de leurs propres capteurs grâce aux informations partagées par des entités externes telles que les RSUs, les dispositifs pour piétons et les serveurs d'applications V2X. Les exigences de performances pour le groupe de capteurs étendu varient selon le cas d'utilisation et incluent une latence allant de 3 ms à 1 000 ms, une taille de charge utile minimale de 1 600 octets et une fiabilité de 90 % à 99,999% [43].

4. CONDUITE AVANCÉE : Cela permet aux véhicules de coordonner leurs trajectoires ou leurs manœuvres en partageant leurs intentions de conduite avec les véhicules à proximité, améliorant ainsi la sécurité. Nous pouvons classer cela en multiple cas d'utilisation :

- Évitement coopératif des collisions "Collision Cooperative Avoidance (CoCA)".
- Partage d'informations pour une conduite automatisée limitée.
- Partage d'informations pour une conduite entièrement automatisée.
- Alignement de trajectoire d'urgence "Emergency Trajectory Alignment (EtrA)".
- Fourniture d'informations sur la sécurité aux intersections pour les zones urbaines.

Ces cas d'utilisation ont des exigences spécifiques, notamment une latence inférieure à 10 ms, une fiabilité supérieure à 99,99%, un débit de 10 Mbps et une taille de message allant jusqu'à 2 Ko.

5. Service d'information sur l'état du véhicule : Un constructeur automobile ou un centre de diagnostic automobile peut obtenir périodiquement des mises à jour de l'état des véhicules en mode V2N pour un diagnostic à distance. De même, les applications de gestion de flotte peuvent utiliser ces données [50]. Ce service n'a pas d'exigences spécifiques concernant la latence et la fiabilité.

6. Service de divertissement : La navigation sur le Web, l'accès aux réseaux sociaux et d'autres options de divertissement embarquées sont considérées comme essentielles pour les véhicules modernes et deviendront de plus en plus importantes avec l'essor des véhicules autonomes. Ces services utilisent la communication V2N et nécessitent une bande passante de 0,5 Mb/s pour la navigation Web et jusqu'à 15 Mb/s pour la vidéo HD.

1.5.3 Motivations de la transition vers le découpage réseau 5G V2X

Le secteur automobile est indéniablement un moteur clé des systèmes 5G. La communication véhicule-vers-tout (V2X) et son amélioration "Enhanced Vehicle to Everything (eV2X)" ont été intégrées dans la version 14 de la technologie LTE de 3GPP et incluses parmi les cas d'utilisation de la 5G. L'intérêt des opérateurs de réseaux mobiles "Mobile Network Operators (MNOs)" pour le secteur automobile est également démontré par la création récente de la 5G Automotive Association. Cette association rassemble de grands constructeurs automobiles et des entreprises de technologies de l'information et des communications "Information and Communication Technology (ICT)" pour promouvoir des solutions interopérables pour le V2X cellulaire. Les technologies de communications cellulaires sont des concurrents sérieux pour répondre aux besoins des communications automobiles en raison de leur couverture étendue, de leur sécurité robuste, de leurs services mobiles et de leur capacité réseau élevée. Bien que la technologie cellulaire LTE ait intégré la communication V2X et sa version améliorée (eV2X), les réseaux LTE actuels ont encore du mal à prendre en charge efficacement les diverses exigences de QoS des services V2X en raison de leur architecture "taille unique". Par conséquent, ces technologies, notamment les variantes IEEE 802.11 de VANETs et la technologie LTE, ne peuvent pas répondre à des exigences aussi strictes. Ce défi a poussé les communautés de recherche et industrielles à développer des spécifications pour une nouvelle technologie avec une voie d'évolution claire et évolutive vers le système de cinquième génération (5G) avec son paradigme : "network slicing".

1.5.4 Tranches de réseau 5G V2X

Le terme V2X englobe un large éventail de cas d'utilisation, caractérisés par un ou plusieurs locataires et des exigences de service variables. Ceux-ci incluent tout, depuis un véhicule autonome naviguant dans une ville intelligente, jusqu'au streaming vidéo HD sur un système d'infodivertissement embarqué, en passant par les systèmes avancés de navigation embarqués en temps réel. Ces caractéristiques distinctes, variées et complexes des services V2X rendent difficile leur mappage direct sur des types de tranches standards prenant en charge les services eMBB, URLLC et mMTC, ou leur intégration dans une seule tranche V2X [49, 50]. Cela implique la nécessité de concevoir des tranches de réseau sur mesure pour divers services V2X.

Les tranches proposées pour contenir tous les V2X services proposés incluent celles pour la conduite autonome, la conduite téléopérée, l'infodivertissement des véhicules, ainsi que le diagnostic et la gestion à distance des véhicules. La figure 1.12 montre les différents types de tranches (slices) dans V2X.

- **La tranche de conduite autonome :** S'appuie sur les communications de type V2V à très faible latence comme mode de connexion RAT principal et sur des fonctions RAN/CN supplémentaires, telles que l'allocation de ressources contrôlée par le réseau sur l'interface PC5 (dans l'eNodeB) et la mobilité, l'authentification, l'autorisation et la gestion des abonnements. De plus, un échange vidéo/données fiable et à faible latence doit être pris en charge avec un AS V2X, déployé à la

périphérie du réseau, pour aider les véhicules à traiter des cartes 3D de la zone environnante et à créer une vision augmentée au-delà de leur perception visuelle.

- **La tranche de conduite téléopérée :** La tranche supportant la conduite téléopérée doit garantir une latence ultra faible et une connectivité de bout en bout très fiable entre le véhicule contrôlé et l'opérateur distant, qui est généralement situé en dehors du réseau central. Contrairement à la tranche de conduite autonome, ces services seront limités à quelques véhicules et activés uniquement dans des circonstances spécifiques, ce qui entraînera une légère charge sur les entités du réseau central.
- **La tranche d'infodivertissement pour véhicule :** Cette tranche doit garantir des données et des services d'infodivertissement à haut débit via un AS d'infodivertissement V2X. Ce cas d'utilisation est essentiel dans les environnements automobiles, car des activités telles que la navigation sur le Web, l'accès aux réseaux sociaux, le téléchargement de fichiers et le streaming vidéo HD pour les passagers deviennent de plus en plus courantes.
- **La tranche de diagnostic et de gestion à distance du véhicule :** Elle doit être capable de prendre en charge l'échange de petites quantités de données basse fréquence entre de nombreux véhicules et des serveurs distants à l'extérieur du CN. Pour cela, les fonctionnalités UP doivent gérer plusieurs interactions et les fonctionnalités CP doivent être instanciées en conséquence.

Toutes ces tranches V2X peuvent être instanciées et orchestrées par différents opérateurs pour fournir différents types de services à leurs véhicules. Un ensemble de fonctionnalités communes a été identifié qui devraient être fournies et éventuellement acceptées par tous les opérateurs pour une tranche V2X afin de simplifier la gestion de l'itinérance des tranches, comme les fonctionnalités AMF qui doivent être disponibles quel que soit le service V2X spécifique proposé pour accompagner la mobilité des véhicules. Pour éviter la surcharge de l'AMF et l'augmentation de la latence lors des procédures de sélection de tranche et de l'étape de connexion. Avec d'autres fonctionnalités peuvent être liées dynamiquement et différents comportements peuvent être configurés de manière flexible pour répondre aux exigences de performances spécifiques d'une catégorie de service V2X particulière [52]. "Authentication Server Function (AUSF)", "PCF" et "Unified Data Management (UDM)" doivent être déployés en tant que NFs, communs à toutes les tranches V2X.

1.6 Défis de l'adaptation du découpage réseau pour les communications V2X

Le découpage réseau, étant un paradigme de réseau encore relativement nouveau, soulève des questions quant à son adaptation aux environnements des véhicules dans les réseaux 5G, qui restent inexplorées. Nous présentons ici quelques défis auxquels l'adaptation de l'environnement véhiculaire devra faire face dans le découpage de réseau 5G.

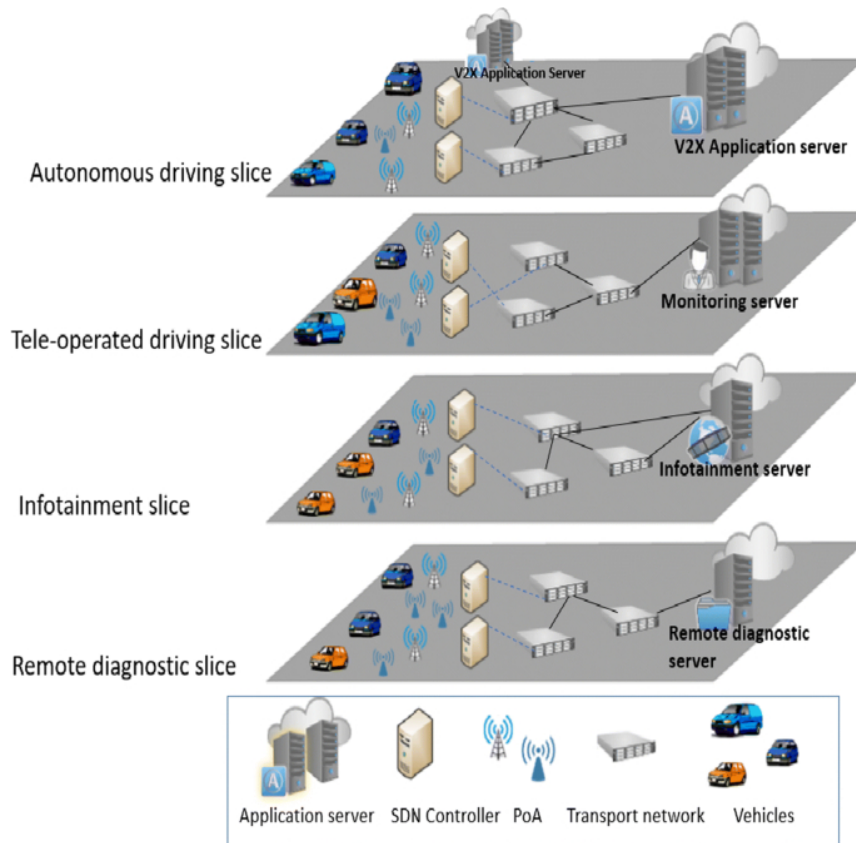


FIGURE 1.12 – Les tranches dans V2X
[51]

1.6.1 Gestion de la mobilité

La gestion de la mobilité dans les réseaux de communication antérieurs est passée de la gestion de simples scénarios de transfert à la gestion de situations complexes, où la mobilité doit gérer non seulement les transferts horizontaux traditionnels (c'est-à-dire les transferts entre cellules/stations de base) et les transferts verticaux (c'est-à-dire les transferts inter-technologies), mais également entre différentes tranches (c'est-à-dire les transferts entre tranches). Ces tranches de réseau ont des exigences uniques en matière de mobilité, de latence et de fiabilité, retards du réseau. Par exemple, le nombre de transferts horizontaux déclenchés peut varier en fonction de la vitesse des véhicules où les véhicules à grande vitesse provoquant de nombreux transferts en peu de temps et les véhicules à basse vitesse en provoquant peu [11]. En revanche, les applications IoT nécessitent des communications dont la plupart n'ont pas de mobilité, donc pas de transfert horizontal. Cela rend les tranches réseau V2X assez unique en termes de gestion de la mobilité de ses services. La nouvelle approche de gestion de la mobilité dans l'environnement 5G véhiculaire basé sur le découpage réseau doit prendre en compte les trois types de transferts ainsi que la grande mobilité des véhicules.

1.6.2 Chaînage et placement de fonctions

Chaque tranche de réseau représente un réseau de bout en bout (E2E) logiquement indépendant, comprenant un ensemble de fonctions réseau (NF), notamment des VNFs après leur virtualisation, ainsi que des ressources associées, adaptées pour fournir des services à la demande pour des scénarios commerciaux spécifiques. Généralement, ces services impliquent une séquence de fonctions réseau disposées dans un ordre spécifique, appelée chaîne de fonctions de service "Service Function Chain (SFC)" [53]. Chaque fonction réseau au sein du SFC est assurée par des nœuds de réseau désignés, par exemple des machines virtuelles (VMs) où chaque VM porte des VNFs spécifiques. Le placement de ces VNFs et l'enchaînement dynamique de services sont cruciaux pour répondre aux diverses exigences de QoS, telles qu'une faible latence et un débit élevé, essentiels pour des applications allant de la conduite autonome à la gestion du trafic en temps réel. La figure 1.13 illustre le modèle de découpage de bout en bout comme présenté dans [54].

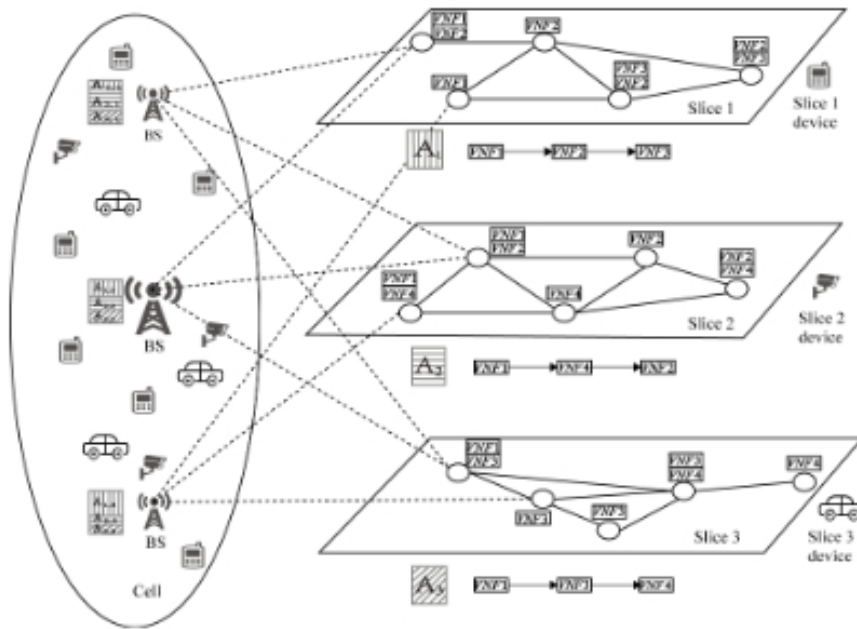


FIGURE 1.13 – Le modèle de découpage de bout en bout [54]

1.6.3 Gestion des ressources

La gestion efficace des ressources est nécessaire pour répondre aux besoins de performances dans les communications V2X. Elle implique la gestion des ressources au sein de plusieurs parties du réseau, notamment le réseau d'accès radio, le réseau central, ainsi que les processus de déchargement et de planification des tâches. Dans ce qui suit, nous présentons les deux mode de gestion de ressources du réseau :

1. Gestion des ressources du réseau d'accès radio : La gestion des ressources dans la partie RAN, est considérée comme l'aspect le plus crucial de la gestion des ressources

en raison de leur caractère limité et de leur distribution géographique. Il existe trois types différents de ressources : les ressources spectrales, les ressources de mise en cache et les ressources de calcul "Edge-Computing".

- **Les ressources spectrales :** Le spectre sans fil dans les RAN présente le défi le plus important à relever en raison de sa rareté. De plus, compte tenu de la grande mobilité des véhicules et de leurs diverses exigences en matière de QoS, une stratégie d'allocation de ressources dédiée à l'environnement de découpage V2X sera essentielle.
 - **Les ressources de mise en cache :** La mise en cache est utilisée pour stocker temporairement le contenu populaire dans un cache de l'élément BS, permettant ainsi une latence plus faible et une utilisation réduite du backhaul [25]. Un modèle de mise en cache adapté aux différents besoins de contenu des différentes tranches V2X constitue une tâche difficile.
 - **Les ressources de calcul en edge :** Les applications V2X des différentes tranches devront effectuer de nombreuses tâches de calcul, notamment en périphérie du réseau. Des approches efficaces de gestion des ressources de calcul seront cruciales pour permettre une variété de services en temps réel. Une option consiste à adopter ces approches spécifiquement pour les exigences des services V2X.
2. Gestion des ressources du réseau central : Dans le réseau central, le découpage implique l'allocation et la gestion dynamiques des ressources de calcul, de stockage et de réseau pour satisfaire les besoins variés des différentes tranches du réseau. Chaque tranche fonctionne comme un réseau virtuel indépendant, personnalisé pour des applications spécifiques telles que les véhicules autonomes. Les VNFs essentiels sont déployés dans différents nœuds du réseau central pour prendre en charge une tranche de réseau en fonction du service spécifique auquel elle s'adresse. Ces nœuds sont interconnectés pour permettre la communication entre les fonctions réseau instanciées. Ces VNFs peuvent être mis à l'échelle selon les besoins, en fonction des changements dans les exigences de service et de performances. La gestion de la coexistence des tranches de réseau et la garantie que les ressources nécessaires sont disponibles pour leurs fonctions, en termes de ressources de calcul dans chaque nœud, ainsi que de ressources de communication entre les fonctions, seront une tâche difficile.

1.6.4 Déchargement/planification des tâches

Avec l'introduction de nouveaux services V2X, une charge importante est également apparue concernant les ressources de calcul requises pour gérer les tâches de ces applications, et ces tâches ont des exigences différentes en termes de délai, de bande passante, de fiabilité et de puissance de traitement. Le DSRC étant une méthode de communication typique basée sur la compétition, il n'est pas adapté à plusieurs tâches en raison du problème d'évolutivité. Par conséquent, la plupart de ces tâches ont été traitées localement au sein du véhicule [55]. Avec l'émergence de nouveaux services V2X, la gestion de ces tâches localement au sein du véhicule est devenue quasiment impossible en raison de la complexité accrue et des besoins en ressources. La technologie 5G résout ces inconvénients

en permettant le traitement de diverses tâches provenant de différentes tranches V2X à trois emplacements distincts : localement dans le véhicule, en périphérie et dans le cloud distant, grâce à ses capacités de communication améliorées. Cela introduit de nouveaux défis dans le développement de solutions permettant de gérer efficacement le déchargement et la planification de ces tâches dans l'environnement de découpage 5G V2X.

1.6.5 Problèmes de la sécurité

Les réseaux mobiles sont constamment la cible d'un nombre croissant d'attaques visant leurs divers services et applications, notamment la finance, la santé, l'éducation, les villes intelligentes et l'e-gouvernement. Le découpage de réseau (NS) est envisagé comme un paradigme crucial dans la 5G. Cependant, NS introduit de nouvelles vulnérabilités et vecteurs d'attaques tels que la sécurité du cycle de vie, la sécurité intra-tranche, la sécurité inter-tranche et la sécurité du courtier de tranches. Ces nouvelles vulnérabilités et vecteurs d'attaques seront exposés aux attaques potentielles telles que le suivi de localisation, la fraude, les fuites de données et les attaques par canal secondaire sur différentes tranches du réseau. Nous explorerons certains de ces risques dans cette section [16] :

1. **Attaques de localisation des véhicules** : Le suivi de la localisation des véhicules dans l'environnement de découpage du réseau 5G peut présenter des risques importants. Une fonction de réseau virtuel malveillant pourrait obtenir un ticket d'autorisation et envoyer une demande au serveur de suivi de localisation en utilisant l'identifiant d'une véhicule cible d'une tranche [16]. De plus, les défauts de conception des NF hybrides et la compromission des fonctions de périphérie du réseau pourraient également être exploitée pour suivre illégalement la localisation des véhicules.
2. **Fuite de données** : Étant donné que les tranches V2X peuvent partager plusieurs fonctions réseau pour réduire la complexité et que ces tranches peuvent appartenir à différents secteurs verticaux et propriétaires, de nouveaux types d'attaques pourraient exploiter ces fonctions partagées [56]. De telles attaques pourraient potentiellement accéder aux données de véhicules situés dans d'autres tranches, compromettant ainsi leurs informations. Cela nécessite des recherches plus approfondies pour explorer les mécanismes de confiance avancés, tels que les modèles de confiance et de réputation, spécifiquement adaptés à l'environnement de découpage du réseau 5G V2X.
3. **Épuisement des ressources de sécurité dans d'autres tranches** : Les attaques de ce type, telles que les attaques de déni de service distribué "Distributed Denial of Service (DDoS)", exploitent le manque d'isolation entre les tranches. L'attaquant cible une tranche dont la sécurité est plus faible en épuisant les ressources de fonctions réseau partagées avec les autres tranches, telles que la puissance de traitement, la mémoire et le matériel, qui sont communes à plusieurs tranches. Cela représente un risque important, notamment pour les tranches qui gèrent des services critiques tels que les véhicules autonomes.

1.7 Conclusion

Dans ce chapitre, nous avons retracé l'évolution des réseaux véhiculaires, en mettant en lumière la transition des réseaux ad hoc vers les réseaux cellulaires, en commençant par les VANETs et en détaillant leurs caractéristiques. Ensuite, nous avons exploré la technologie 5G et ses améliorations en termes de communications. Nous avons également abordé le découpage réseau, le paradigme responsable de la coexistence des différentes tranches sur la même infrastructure, ainsi que ces facilitateurs technologiques telles que les architectures SDN et NFV. Le découpage réseau en 5G, comme toute nouvelle technologie, doit être adapté à l'environnement véhiculaire. Dans ce contexte, nous avons détaillé les problèmes de cette adaptation, en identifiant les exigences et les domaines qui représentent des défis.

Chapitre 2

État de l'art sur les approches de découpage réseau dans l'environnement V2X 5G

Chapitre 2

État de l’art sur les approches de découpage réseau dans l’environnement V2X 5G

2.1 Introduction

Le succès du déploiement de paradigme de découpage réseau dans l’environnement véhiculaire dépend de plusieurs critères, l’un d’eux étant la proposition de solutions dans le domaine de la gestion des ressources afin de permettre la coexistence de ces tranches V2X sur une même infrastructure partagée. L’allocation des ressources disponibles entre les tranches V2X est une tâche très complexe. En effet, au niveau du cœur de réseau (CN), l’extension du réseau par l’ajout de ressources matérielles supplémentaires peut améliorer les performances du réseau ou ajouter davantage de capacité de calcul en périphérie/nuage (Edge/Cloud) pour traiter les tâches générées par les services des véhicules. En revanche, le découpage du réseau d’accès radio (RAN) est confronté à une rareté des ressources. Étant donné que la bande passante est une ressource limitée, le découpage du RAN doit gérer cette contrainte de manière efficace en optimisant l’utilisation des ressources tout en assurant une haute isolation entre les tranches. Dans ce chapitre, nous allons analyser et examiner les travaux existants dans le domaine de la gestion des ressources, plus précisément la gestion de la bande passante dans la partie radio de la technologie 5G.

2.2 Critères essentiels de gestion des tranches de réseau

Dans cette partie, nous présentons les critères fondamentaux qui guident l’allocation et la gestion efficaces des ressources dans le découpage du réseau. Chaque aspect joue un rôle primordial dans l’efficacité et la performance des tranches de réseau.

- **Isolation entre les tranches :** L’isolation est la capacité de limiter l’impact des variations de la charge de trafic ou de la qualité de la connexion d’une tranche de réseau sur la qualité de service des autres tranches, même si elles partagent les

mêmes ressources. Cela signifie que toute fluctuation des ressources allouées à une tranche ne devrait pas affecter le nombre de ressources allouées aux autres tranches.

- **Qualité de service** : Cet aspect représente un des critères les plus importants pour évaluer l'efficacité de l'allocation des ressources. L'environnement de découpage réseau propose de gérer plusieurs services avec différents niveaux de QoS, tels que le débit, la latence, etc.
- **Personnalisation** : Chaque tranche peut avoir des exigences de QoS différentes et des modèles de tarification distincts tout en s'efforçant de répondre aux exigences des utilisateurs. Il est donc primordial que chaque tranche ait la capacité de mettre en œuvre sa propre stratégie de planification et de contrôle d'admission sur mesure pour atteindre ses objectifs.
- **Flexibilité** : Pour s'adapter aux fluctuations de la demande de communication sur chaque tranche de réseau, les tranches doivent avoir la possibilité de demander des ressources supplémentaires ou de libérer celles qui ne sont pas utilisées. Cette stratégie garantit que les ressources sont utilisées efficacement, permettant au réseau de s'adapter aux différents besoins de chaque tranche.
- **Équité** : L'équité dans la gestion des ressources de découpage du réseau fonctionne à deux niveaux : la répartition équitable des ressources entre les différentes tranches du réseau et l'allocation équitable des ressources entre les utilisateurs au sein de chaque tranche. Garantir l'équité est un facteur important pour empêcher qu'une seule tranche ne monopolise les ressources, ce qui priverait d'autres tranches des ressources nécessaires. De même, l'équité garantit que les utilisateurs finaux de chaque tranche reçoivent une part équitable des ressources.

2.3 Taxonomie des solutions utilisées dans la gestion des ressources dans le découpage réseau

Dans cette section, nous proposons une classification des solutions de gestion des ressources dans le découpage réseau, en se basant sur divers aspects tels que l'approche utilisée, les types d'allocation, et la technique employée.

2.3.1 Les approches d'allocation dans le découpage réseau

Plusieurs approches de base ont été proposées dans la littérature pour définir comment les ressources sont allouées entre les tranches tout au long de leur cycle de vie, telles que la réservation complète des ressources, l'approche sans réservation, et l'approche de réservation partielle des ressources.

a) Tranchage strict / la réservation complète des ressources

Cette approche est basée sur une réservation complète, ce qui signifie qu'une portion fixe des ressources globales est allouée exclusivement à chaque tranche du système au moment de sa création, et maintenue tout au long de sa durée de vie. Dans ce contexte, aucune

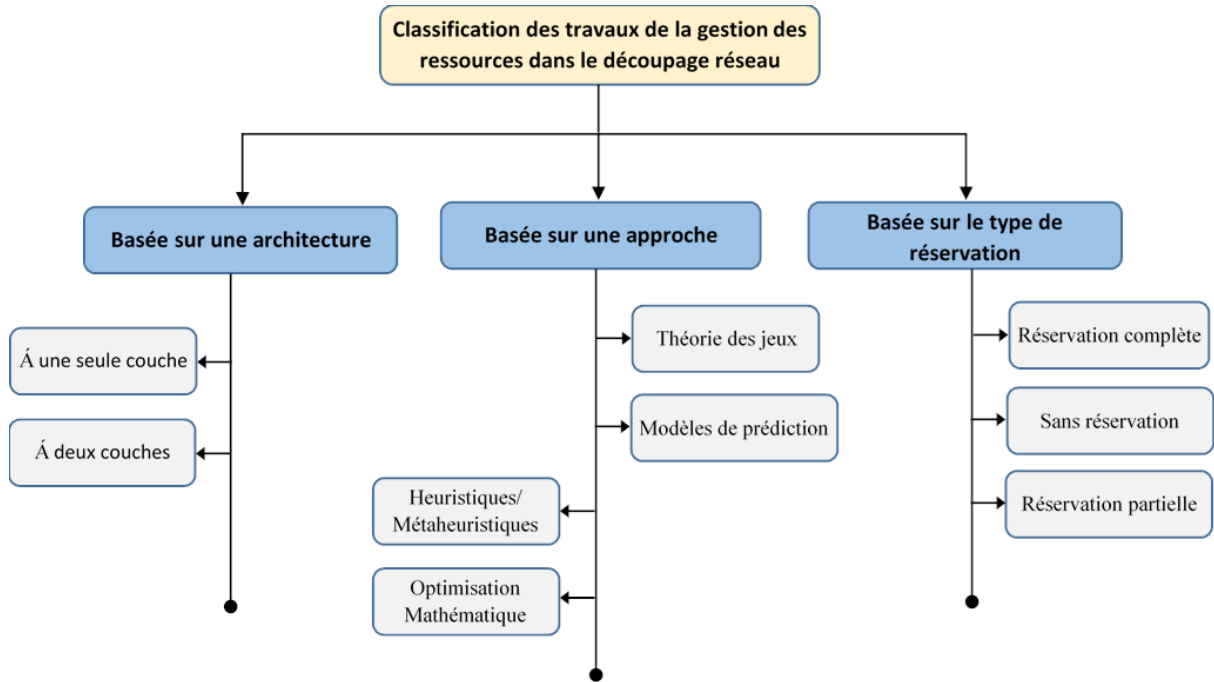


FIGURE 2.1 – Classification des travaux de gestion des ressources dans le découpage réseau

tranche n'interfère avec les autres. Il en résulte une isolation totale entre les tranches de l'infrastructure. Cette isolation garantira des performances stables pour chaque tranche. Cependant, cette isolation totale entre les tranches créera une limitation en termes de flexibilité. Les ressources dédiées à chaque tranche peuvent ne pas être utilisées à 100% pendant sa durée de vie, et comme cette approche est basée sur une isolation totale, ces ressources non utilisées ne peuvent pas être réallouées aux tranches qui en ont besoin.

Raza et al. [57] proposent une politique de contrôle d'admission de tranche basée sur l'apprentissage par renforcement, qui est divisée en deux catégories de services avec des contraintes de QoS différentes : les tranches faiblement prioritaires avec des exigences de latence non strictes et les tranches hautement prioritaires avec des exigences de latence strictes. L'objectif est d'identifier la stratégie optimale qui maximise le profit du système en acceptant le nombre maximum de demandes des tranches tout en évitant toute dégradation de la QoS.

Zhang et Wong [58] proposent une solution pour la gestion des ressources à deux niveau modélisée comme un problème stochastique, visant à maximiser le profit des propriétaires d'une tranche tout en répondant aux exigences de QoS des utilisateurs. La solution propose un algorithme inter-tranche pour réserver de manière optimale les ressources sur une base à long terme pour la tranche de réseau, et un algorithme intra-tranche pour l'allocation à court terme de ces ressources réservées, basé sur les statistiques de trafic des utilisateurs.

Il existe une partie des solutions qui ont choisi de tirer parti de la sous-utilisation des ressources des tranches tout au long de leur durée de vie dans le cadre d'un découpage strict. Ces travaux proposent d'activer un mécanisme de partage des ressources entre les tranches, permettant ainsi de satisfaire les besoins des tranches surchargées en réaffectant les ressources inutilisées des autres tranches.

Le travail présenté dans [51] propose deux schémas différents de partage de ressources V2X en utilisant des jeux de négociation pour réajuster les ressources d'une tranche surchargée lorsqu'une demande de partage est reçue de la part d'une tranche surchargée. Le premier schéma se concentre uniquement sur la maximisation de l'utilité de la tranche surchargée, tout en respectant la qualité de service de la demande de transfert, avec chaque tranche prêteuse offrant une part significative de ses ressources pour fournir la totalité des ressources demandées pour l'appel de transfert. Le deuxième schéma vise à maximiser la somme des utilités des tranches prêteuses et de la tranche surchargée. Ce schéma cherche à proposer un compromis entre les ressources nécessaires pour l'appel de transfert dans la tranche surchargée et les ressources disponibles dans les tranches prêteuses. Dans ce travail, les auteurs ont ensuite proposé de résoudre le premier problème formulé en utilisant un algorithme heuristique à faible complexité, et le second à l'aide de l'algorithme d'optimisation par essaims de particules "Particle Swarm Optimization (PSO)".

Dans [59], les auteurs ont proposé une solution appelée une approche appelée "SoftSlice" pour le partage des ressources radio, dans laquelle le schéma modélise le trafic des utilisateurs dans les tranches de réseau comme un processus de Markov. Ce mécanisme introduit un module de politique de partage des ressources basé sur un accord de partage prédéfini, qui détermine comment les ressources radio inutilisées sont partagées entre les différentes tranches de réseau et comment chaque tranche moins chargée y contribue. Ce module est activé lorsqu'une des tranches devient surchargée.

Maule et al. [60] ont proposé un mécanisme de partage dynamique des ressources entre les tranches basé sur un seuil fixe. Ce schéma vise à transférer progressivement les ressources radio inutilisées des tranches sous-chargées, appelées esclaves, vers une tranche surchargée, appelée maître, jusqu'à ce que l'accord de niveau de service du maître soit garanti ou que la tranche esclave atteigne un seuil de non-partage.

b) L'approche sans réservation

Dans cette approche, toutes les ressources du système sont partagées entre toutes les tranches sans aucune réservation spécifique pour les tranches individuelles. Plutôt que de recevoir une partie fixe de ressources, chaque tranche se voit allouer dynamiquement des ressources après chaque période de temps, en fonction de ses besoins actuels. Cette méthode flexible permet aux tranches de demander et d'utiliser des ressources selon leurs besoins, sans être limitées par une allocation prédéfinie. En conséquence, les ressources sont utilisées plus efficacement et peuvent s'adapter aux demandes variables des différentes tranches, garantissant ainsi des performances et une flexibilité optimale au sein du réseau. Une des lacunes de cette approche réside dans la complexité de l'allocation répétitive des ressources à chaque période de temps entre les tranches. De plus, l'isolation des tranches pose un problème majeur, car elle n'est pas garantie tout au long du cycle de vie de chaque tranche.

Dans [61], Khan et al. proposent une approche combinée pour la sélection de la qualité vidéo et l'allocation des ressources visant à maximiser la qualité d'expérience (QoE) en tirant parti de la dynamique de la file d'attente et de l'état des canaux des appareils véhiculaires tout en garantissant une lecture vidéo transparente. La solution proposée fonctionne en deux parties : premièrement, les véhicules sont divisés en plusieurs tranches de réseau à l'aide d'un algorithme de clustering. Dans la deuxième partie, la planification

des véhicules et la sélection de la qualité sont formulées sous la forme d'un problème d'optimisation stochastique, qui est résolu à l'aide de la méthode de dérive de Lyapunov [62].

Un modèle de contrôle d'admission pour deux types d'utilisateurs basé sur leurs distributions de débits : i) les utilisateurs homogènes et ii) les utilisateurs hétérogènes dans les réseaux 5G est présenté dans [63]. Le modèle proposé cherche à maximiser le nombre d'utilisateurs cohérents qui peuvent être admis, même si cela nécessite d'autoriser une légère dégradation de la qualité de service. Il comprend une analyse du nombre d'utilisateurs pouvant être hébergés tout en maintenant un tarif cohérent pour un groupe homogène. Sur la base de cette analyse, l'algorithme décide d'admettre ou non un nouvel utilisateur arrivant dans la cellule.

c) L'approche de réservation partielle des ressources

L'approche de réservation partielle des ressources divise l'allocation des ressources entre les tranches en deux périodes. Initialement, les ressources du système dans son ensemble sont divisées en deux parties : une partie est désignée comme pool partagé, tandis que l'autre est divisée en parties plus petites. Lorsqu'une tranche est créée, une petite partie des ressources de la deuxième partie lui est allouée. Au cours de la deuxième période, qui se produit de manière répétée tout au long de la durée de vie de la tranche, celle-ci peut demander des ressources supplémentaires au pool partagé. Cette méthode combine les limitations des deux approches : l'approche sans réservation, avec la complexité de l'allocation répétitive des ressources de pool partagé à chaque période de temps, et celle de tranchage strict, avec le manque de flexibilité, ce qui signifie que même les petites portions de ressources réservées peuvent ne pas être entièrement utilisées tout au long du cycle de vie des tranches.

Les auteurs de [64] présentent un schéma de réservation partielle de ressources conçu pour les réseaux LTE. Ce dispositif garantit un niveau minimum de ressources, appelé part réservée à l'opérateur, pour chaque tranche de trafic, qu'il s'agisse d'un débit garanti ou non garanti. Les ressources système restantes, après avoir pris en compte les parties réservées de toutes les tranches, sont réparties entre les tranches à l'aide d'un mécanisme de contrôle d'admission qui suit une politique de gestion "premier arrivé, premier servi".

Dans [65], Kalil et al. proposent un algorithme itératif qui fonctionne en deux étapes séquentielles. La première étape alloue directement les ressources aux utilisateurs afin de maximiser l'utilité globale, définie pour refléter le niveau de satisfaction des utilisateurs. Ensuite, les ressources restantes sont attribuées aux tranches les moins satisfaites, où elles sont réparties entre les utilisateurs. L'étude considère deux types de tranches : la première vise à garantir l'équité entre ses utilisateurs, tandis que la seconde cherche à maximiser son débit, chacune en fonction de sa politique de planification personnalisée.

Dans [66], les auteurs proposent un système dynamique avec une réservation partielle des ressources appelé "CellSlice", qui inclut un algorithme d'adaptation basé sur les retours pour partager les ressources inutilisées à partir de tranches sous-utilisées sans compromettre l'exigence d'isolation. Ce système est conçu pour gérer la charge fluctuante des différentes tranches causée par la mobilité de l'utilisateur et les conditions variables des canaux.

2.3.2 Les types d'allocation des ressources

Dans cette section, nous explorons deux types distincts de la gestion des ressources dans le découpage de réseau : la gestion des ressources à une seule couche et la gestion des ressources à deux couches. La figure 2.2 illustre les types de gestion de ressources dans le découpage réseau.

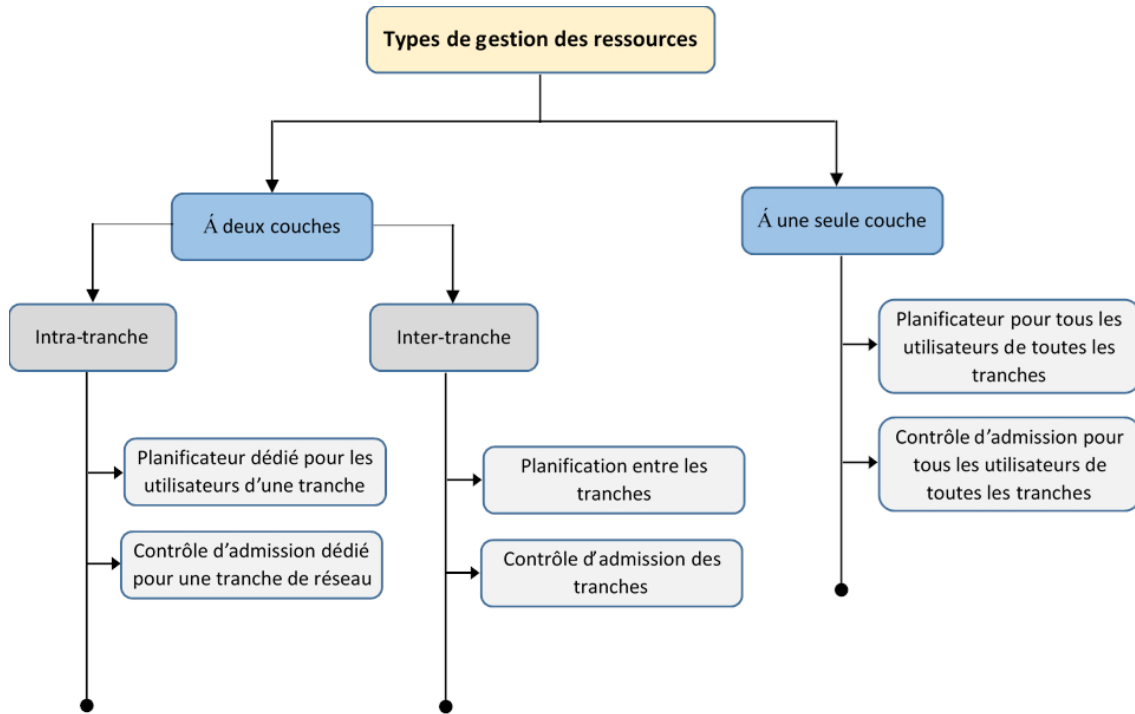


FIGURE 2.2 – Les types de gestion des ressources dans V2X 5G

a) Gestion des ressources à une seule couche

La méthode la plus simple pour gérer les ressources dans le découpage réseau. En effet, elle permet de gérer simultanément tous les utilisateurs de différentes tranches de réseau. Ce qui signifie que tous les utilisateurs des différentes tranches, avec leurs exigences variées en termes de qualité de service, seront traités de la même manière. Cela ajoute également à la complexité de traiter tous les utilisateurs des tranches de réseau. Cette approche monocouche peut être appliquée à différents niveaux, tels que la planification des ressources et le contrôle d'admission des utilisateurs. Une des limitations de cette approche est que l'isolation et la personnalisation des tranches deviennent plus difficiles.

- **La planification monocouche :** L'ordonnancement monocouche est une méthode d'ordonnancement des ressources qui traite l'allocation des ressources dans une seule dimension ou couche. Les ressources sont directement allouées aux utilisateurs de chaque tranche. Diverses approches d'ordonnancement de référence ont été proposées dans la littérature pour être utilisées comme un planificateur monocouche pour toutes les tranches [67, 68], telles que le proportionnel équitable (PF), le tourniquet (RR), l'ordonnancement par priorité, l'ordonnancement basé sur la QoS, et l'ordonnancement à délai minimal.

- **Le contrôle d'admission monocouche** : Dans ce type de contrôle d'admission, chaque utilisateur souhaitant se connecter à une tranche doit soumettre une demande de connexion au contrôle d'admission monocouche, qui est responsable de traiter les requêtes de toutes les tranches du système.

Les auteurs de [69] présentent un cadre d'optimisation pour le contrôle d'admission radio "Radio Admission Control (RAC)" destiné à tous les utilisateurs du système, dans un scénario 5G multi-locataires et multi-services, en modélisant le problème à l'aide d'un processus de décision semi-Markovien "Semi Markov Decision Process (SMDP)". La méthode d'itération sur la valeur proposée vise à résoudre ce problème basé sur le SMDP. Ce mécanisme a pour objectif de réguler l'acceptation des nouvelles demandes de service à débit garanti "Guaranteed Bit Rate (GBR)" provenant de différents locataires, afin de respecter leur accord de niveau de service (SLA) tout en tenant compte de la probabilité de congestion des cellules du système. Cela permet d'éviter les situations où la somme des débits garantis de tous les utilisateurs acceptés dépasse la capacité de la cellule.

Un mécanisme de planification des ressources à un seul niveau entre les utilisateurs de toutes les tranches de différentes priorités est présenté dans [70]. L'algorithme propose une modification de l'algorithme ordinaire de planification des paquets, tel que l'algorithme de proportionnalité équitable, afin de garantir une isolation des ressources en limitant le nombre maximal de blocs de ressources (RBs) alloués à chaque tranche.

Dans [71], un algorithme de planification des ressources à un seul niveau basé sur la programmation dynamique pour la communication véhicule-à-tout (V2X) est proposé. Cet algorithme vise principalement à améliorer l'expérience utilisateur pour les trois tranches V2X hétérogènes dédiées aux trois services : services de sécurité de base, services de sécurité avancés, et services d'efficacité du trafic, tout en maximisant l'efficacité spectrale.

b) Gestion des ressources à deux couches

En général, l'approche la plus prometteuse semble consister à séparer la gestion des ressources pour les tranches de celle pour les utilisateurs finaux, ce qui implique un système de gestion des ressources à deux niveaux. Une couche gère la gestion des ressources pour chaque tranche, tandis que l'autre couche, spécifique à chaque tranche, gère les ressources des utilisateurs individuels. Cette approche réduit la complexité de la gestion des ressources en la divisant en deux parties distinctes. Chaque tranche peut gérer ses propres ressources de manière indépendante, permettant une allocation des ressources plus personnalisée.

1. La gestion des ressources entre les tranches : La gestion des ressources inter-tranches représente la partie supérieure où les ressources du système sont d'abord allouées aux différentes tranches. Cette répartition peut se baser sur plusieurs critères, tels que la bande passante, le nombre de ressources, les accords de niveau de service (SLA), les interférences, la charge de trafic ou une combinaison de ces éléments. L'inter-tranches (inter-slicing) est appliquée dans différents niveaux, tels que le contrôle d'admission des tranches et la planification des ressources entre les tranches.

- **Le contrôle d'admission des tranches "Slice Admission Control (SAC)"** : À ce niveau, la responsabilité d'accepter ou de rejeter une tranche est mise en

œuvre. Le dilemme de choisir quelle tranche de réseau à accepter ou à rejeter est très complexe. Le propriétaire de l'infrastructure (InfProv) vise à maximiser l'utilisation des ressources du réseau en acceptant autant de tranches de réseau que possible pour des raisons commerciales. Cependant, les ressources du réseau sont limitées et les exigences des tranches de réseau en termes de QoS doivent être satisfaites. Ce contrôle d'admission est généralement utilisé dans des situations où une demande de connexion d'une tranche nécessite un certain nombre fixe de ressources, qui sera fournie pendant toute la durée de vie de la tranche.

- **Le planificateur entre les tranches :** Ce type de gestion des ressources est généralement appliqué dans les situations où les tranches sont créées sans qu'une quantité de ressources soit garantie tout au long de leur cycle de vie. Cela signifie qu'après chaque période, la tranche de réseau communiquera ses besoins en termes de ressources à satisfaire pour la période suivante. Après avoir reçu les demandes de ressources de toutes les tranches existantes dans l'infrastructure, le planificateur inter-tranche décide de la manière dont les ressources seront allouées entre les tranches.
2. La gestion des ressources entre les utilisateurs dans chaque tranche : Après que les ressources de l'infrastructure ont été allouées à chaque tranche du système par le niveau supérieur de gestion des ressources inter-tranches, il incombe à la gestion des ressources intra-tranche de chaque tranche de gérer ces ressources pour les utilisateurs finaux de cette tranche. Cela permet de personnaliser l'allocation des ressources en fonction des besoins spécifiques des services et des applications desservis par chaque tranche. L'inter-tranche (intra-slicing) est appliquée dans différents niveaux, tels qu'un contrôle d'admission personnalisé pour les utilisateurs de chaque tranche et la planification des ressources personnalisée entre les utilisateurs finaux.
- **Le contrôle d'admission des utilisateurs d'une tranche :** À ce niveau, chaque tranche peut avoir sa propre fonction de contrôle d'admission personnalisée, adaptée à son SLA ou à ses exigences en terme de QoS. Cela signifie que chaque utilisateur souhaitant se connecter à une tranche envoie une demande au contrôle d'admission de cette tranche, qui décide d'accepter ou de rejeter la demande de connexion. Si les ressources disponibles sont suffisantes pour répondre aux besoins de l'utilisateur sans compromettre la performance des autres utilisateurs déjà connectés, la demande est acceptée, sinon elle est rejetée ou mise en attente.
 - **Le planificateur pour les utilisateurs d'une tranche :** La planification des ressources à deux couches est réalisée en deux phases qui fonctionnent sur des échelles de temps différentes. La première échelle de temps est celle de la planification inter-tranche. La période entre les deux planifications inter-tranche, la plus large, est divisée en plusieurs petites périodes dédiées à la planification intra-tranche.

De nombreuses recherches ont exploré des approches basées sur l'allocation de ressources à deux couches. Dans le travail [72], les auteurs proposent une nouvelle méthode basée sur la planification des ressources à deux niveaux pour mettre en œuvre

le découpage de réseau dans le RAN 5G. L'algorithme d'allocation inter-tranches, appelé Resource Manager (RM), est chargé de mapper les ressources virtuelles aux ressources physiques, chaque tranche disposant d'une quantité prédéfinie de ressources. L'allocation intra-tranche est réalisée par un algorithme de planification personnalisé pour répondre aux besoins spécifiques de chaque tranche.

Une solution hybride d'isolation basée sur une enchère combinatoire hiérarchique à deux niveaux, avec la capacité de personnalisation intra-tranche est présentée dans [73]. Dans ce travail, l'allocation des ressources se fait en deux parties. Dans la première partie, un certain nombre de ressources est réservé pour chaque tranche en fonction de ses accords de service prédéterminés, tandis que dans la deuxième partie, le reste des ressources est alloué pour améliorer l'utilisation efficace de la bande passante.

Dans [74], deux schémas de découpage des ressources sont proposés pour gérer la coexistence des services "Machine To Machine (M2M)" et "Human To Human (H2H)". Ces solutions visent à isoler la planification des ressources des services M2M des autres services existants en réservant des ressources dédiées à sa communication lors de la prochaine période de découpage. Cette réservation est basée sur les rapports CQI et la quantité de données mises en file d'attente dans le tampon, qui est communiquée à l'eNB à l'aide des rapports d'état du tampon "Buffer Status Report (BSR)". La durée de la période de découpage est ajustée en fonction des conditions changeantes, telles que les variations de la demande, les conditions des canaux et les situations d'urgence.

Dans [75], un nouveau framework appelé "Network Virtualization Substrate (NVS)" est proposé, conçu pour répondre à trois exigences clés : l'isolation, la personnalisation et l'utilisation efficace des ressources. Le framework proposé introduit deux fonctionnalités principales : (i) l'introduction d'un planificateur de tranches et d'un algorithme de contrôle d'admission prouvés comme optimaux qui prennent en charge la coexistence de tranches avec des réservations basées sur la bande passante et d'autres sur les ressources au sein de la même infrastructure, et (ii) la possibilité de planification de flux personnalisée au sein de la station de base par tranche.

2.3.3 Techniques utilisées pour l'allocation des ressources dans le découpage de réseau

Dans le contexte du découpage de réseau (network slicing), plusieurs solutions basées sur différents modèles ont été développées pour gérer et allouer efficacement les ressources. Ces modèles sont conçus pour relever les défis uniques du découpage de réseau, tels que la gestion dynamique, la flexibilité et l'efficacité des ressources.

a) Modèles basés sur la théorie des jeux

La théorie des jeux, discipline axée sur l'étude des interactions et de la prise de décision entre individus ou entités, constitue depuis longtemps un outil analytique fondamental dans la recherche en sciences sociales. Son application s'est désormais étendue bien au-delà de son domaine d'origine, trouvant une pertinence significative dans divers domaines, notamment la gestion des ressources. Ces dernières années, les défis posés par le découpage réseau dans la gestion des ressources ont suscité un intérêt accru pour l'utilisation de la théorie des jeux, où elle est applicable aux situations dans lesquelles les décideurs sont en

compétition pour une ressource limitée.

Dans le scénario de découpage du réseau 5G, il y a trois participants typiques dans les jeux de la gestion de ressources : l'opérateur de réseau (ou fournisseur d'infrastructure), le locataire d'une tranche de réseau et l'utilisateur d'une tranche de réseau. Ce paragraphe décrit divers modèles et méthodes de la théorie des jeux appliqués à la gestion des ressources, en mettant l'accent sur des objectifs multiples. La théorie des jeux peut être classée en deux types : coopérative et non coopérative.

1. Modèles théoriques des jeux coopératifs : Dans les modèles de théorie des jeux coopératifs, les joueurs peuvent collaborer pour atteindre des objectifs communs grâce à un comportement coopératif et parvenir à des accords sur la manière de partager et de distribuer les gains collectifs résultant de leur coopération. Ces modèles se concentrent sur les mécanismes permettant de répartir équitablement et efficacement les récompenses des efforts conjoints, garantissant que les bénéfices sont répartis d'une manière qui reflète la contribution de chaque acteur et maintient la stabilité au sein de l'accord.

Un algorithme de découpage des ressources basé sur des jeux de faillite a été proposé dans [76] et amélioré dans [77]. L'algorithme traite les situations dans lesquelles les ressources totales demandées par toutes les tranches dépassent les ressources système disponibles. Il envisage de servir trois tranches de réseau distinctes : eMBB, mMTC et URLLC, tout en prenant en compte une fonction utilitaire qui reflète les exigences de qualité de service des utilisateurs de chaque tranche et garantit l'équité dans le partage des ressources entre les tranches.

2. Modèles théoriques des jeux non coopératifs : Dans les modèles non coopératifs, le jeu se concentre principalement sur l'analyse et la modélisation du comportement concurrentiel des participants et des décisions stratégiques qui découlent de leurs interactions. Dans ce contexte, chaque joueur choisit indépendamment sa stratégie dans le but soit d'améliorer sa propre performance, soit de minimiser ses pertes. Par conséquent, des accords contraignants entre les entités de jeux ne sont pas possibles.

Dans [78], les auteurs ont proposé une structure hiérarchique à trois niveaux contenant le niveau fournisseur d'infrastructure (InP), le niveau des opérateurs de réseaux mobiles virtuels (MVNO) et le niveau utilisateur. Cette solution repose sur un jeu de Stackelberg pour l'allocation des fréquences et de la puissance entre les MVNOs, où chacun étant propriétaire d'une tranche dédiée à ses utilisateurs. Cette allocation se déroule en cinq étapes. Tout d'abord, le fournisseur d'infrastructure (InP) détermine les prix unitaires des ressources en fréquence et en puissance et les annonce à l'ensemble des MVNOs. Ensuite, chaque MVNO fixe les prix de revente unitaires des ressources en fréquence et en puissance pour ses utilisateurs. Par la suite, chaque utilisateur décide de la quantité de ressources qu'il souhaite acheter auprès du propriétaire de sa tranche. Une fois cette décision prise, chaque MVNO collecte les montants d'achat auprès de ses utilisateurs. Enfin, l'InP alloue les ressources en fréquence et en puissance nécessaires aux tranches, qui seront ensuite distribuées aux utilisateurs.

Caballero et al. [79] ont proposé un framework de découpage de réseau (NES) pour gérer l'allocation des ressources en analysant les interactions de plusieurs tranches non

coopératives, où chacune visant à maximiser sa propre utilité réseau. La gestion des ressources dans ce framework comprend trois modules : le premier module dédié pour le contrôle d'admission avec deux politiques : une conservatrice et une agressive, le deuxième module pour l'allocation de pondération conçu pour maximiser l'utilité de chaque tranche, et le dernier module qui met en œuvre une stratégie pour abandonner les utilisateurs lorsque les garanties de taux ne peuvent pas être respectées.

Un modèle d'enchères à deux niveaux est proposé dans [80] pour l'allocation des ressources. L'objectif du modèle proposé est de traiter à la fois les aspects techniques et économiques de l'allocation des ressources sur le plan technique et le plan économique. Sur le plan technique, les enchères permettent d'améliorer l'efficacité de l'utilisation des ressources, tandis que sur le plan économique, elles sont adaptées pour maximiser les revenus du fournisseur d'infrastructure. Ce modèle d'enchères se présente en deux niveaux d'allocation : le niveau supérieur et le niveau inférieur. Dans le niveau supérieur, l'enchère se déroule entre le fournisseur d'infrastructure (InP), qui est l'enchérisseur, et le propriétaire de la tranche, qui est l'enchéreur. Une fois l'allocation effectuée à ce niveau, au niveau inférieur, le propriétaire de la tranche devient l'enchéreur, et les utilisateurs sont les enchérisseurs.

b) Modèles de prédiction

Les modèles de prédiction sont utilisés pour estimer de manière appropriée ou optimale l'allocation des ressources en s'appuyant sur des connaissances empiriques ou des données historiques. Ces modèles peuvent prédire en se basant sur divers facteurs, tels que le taux d'arrivée total des utilisateurs au sein de chaque tranche de réseau, le trafic de données dans une tranche spécifique, ainsi que les demandes de création ou de modification de tranches de réseau. Les méthodes de prédiction de base reposent souvent sur l'utilisation directe de données empiriques ou sur des modèles probabilistes bien établis pour générer ces prédictions. Parmi ces modèles de prédiction, on peut citer :

1) Modèles d'apprentissage automatique : Il s'agit d'un modèle qui permet d'extraire des connaissances sur divers problèmes à partir de grands ensembles de données en créant un modèle d'apprentissage automatique "Machine Learning (ML)" et en lui fournissant des entrées telles que des ensembles de données. Ce type de modèle nécessite une période d'apprentissage pour atteindre le résultat souhaité, et cette période dépendant de la complexité de l'algorithme et de la taille de l'ensemble de données. Ces modèles ont été largement utilisés pour prédire divers facteurs, tels que la demande future de tranches de réseau sur la base de données d'utilisation historiques, en analysant les modèles de comportement des utilisateurs de chaque tranche et de leurs trafics.

Dans [81], un algorithme d'apprentissage automatique, en particulier un modèle de réseau neurones "Neural Network (NN)", est proposé pour l'allocation de tranches. Le modèle NN est utilisé pour fournir des décisions d'admission, en particulier dans des conditions de réseau où aucune politique optimale n'a été calculée en raison d'un environnement complexes. L'objectif du travail est d'améliorer l'efficacité de l'utilisation des ressources, de garantir l'équité envers les fournisseurs de services, de maximiser les revenus des propriétaires de réseaux et de gérer la complexité.

Dans [82], un réseau récurrent "Long Short Term Memory (LSTM)" est proposé pour

optimiser le processus de réservation en minimisant à la fois les sur et sous-réservations effectuées par le propriétaire de la tranche auprès des opérateurs de réseaux mobiles (MNO). La surréservation entraîne le paiement des ressources inutilisées, tandis que la sous-réservation peut entraîner des difficultés dans le maintien de la qualité de service pour les utilisateurs. Cette approche vise à garantir une QoS de haute qualité pour la demande future de trafic tout en minimisant les coûts. Le modèle est entraîné sur un ensemble de données de trafic cellulaire collectées à partir de données réelles dans la ville de Shanghai.

Dans [83], un algorithme Random Forest (RF) et des réseaux neuronaux d'apprentissage profond "Deep Learning Neural Network (DLNN)" sont utilisés pour prédire la charge du réseau sur chaque tranche en analysant les modèles de trafic. Le modèle proposé apprend à l'aide d'un ensemble de données contenant des indicateurs clés de performance "Key Performance Indicators (KPI)" pour chaque tranche réseau. Son but est de sélectionner correctement la tranche de réseau appropriée pour chaque demande et de prévoir avec précision l'allocation des ressources pour chaque tranche, en veillant à ce que des ressources suffisantes soient allouées en fonction des prévisions de trafic.

Khan et al. [84] ont proposé un modèle hybride d'apprentissage profond combinant un réseau neuronal convolutif "Convolutional Neural Network (CNN)" pour la gestion des ressources et un réseau de mémoire à long terme (LSTM) pour le traitement des informations statistiques telles que l'équilibrage de charge et les taux d'erreur. Le but de la solution proposée est de prévenir les pannes et d'éviter une surcharge dans les tranches.

2) Modèles basés sur l'expérience : Ces modèles se concentrent sur l'apprentissage par interaction avec un environnement, en prenant des mesures appropriées en fonction de l'état actuel de l'environnement et en recevant des récompenses ou des pénalités pour leurs actions. La récompense/pénalité détermine l'efficacité de l'action entreprise. Ces modèles sont souvent utilisés pour apprendre à atteindre un objectif dans des environnements incertains et potentiellement complexes. En fait, les défis actuels en matière de gestion des ressources dans les environnements de découpage de réseau sont dynamiques et très complexes, ce qui rend ces modèles particulièrement précieux. Plusieurs travaux de gestion des ressources en découpage proposés dans la littérature s'appuient sur le modèle d'apprentissage par renforcement [85].

Dans [86], une approche d'apprentissage par renforcement hors ligne, spécifiquement le Q-learning combiné à une heuristique de faible complexité, est introduite pour mieux prendre en charge la coexistence de deux types de tranches 5G : les tranches eMBB qui nécessitent un débit garanti et les tranches V2X qui sont très sensibles à la latence. L'objectif de la solution de découpage RAN proposée est de déterminer le rapport de découpage optimal entre ces deux services afin de maximiser l'utilisation des ressources tout en minimisant la probabilité de pannes.

Le travail dans [87], les auteurs ont proposé une approche sensible à la demande basée sur le nombre de paquets arrivant dans chaque tranche dans une fenêtre de temps spécifique. Cette approche utilise l'apprentissage en profondeur pour améliorer l'efficacité et la flexibilité du découpage du réseau tout en partageant la bande passante agrégée entre les différentes tranches du système. La solution prend en compte deux SLA clés : le débit et la latence.

Les auteurs de [88] ont présenté une solution formulée comme un modèle de processus

décisionnel de Markov où l'état du système est représenté par un vecteur contenant le nombre de chaque type de tranche acceptée. En appliquant trois modèles d'apprentissage par renforcement (Q-Learning, Deep Q-learning et Regret Matching), ils prédisent les actions optimales (rejeter/accepter) à chaque état afin de maximiser les revenus du fournisseur d'infrastructure lors du contrôle d'admission des tranches. Le modèle génère des revenus en acceptant des demandes de connexion provenant de nouvelles tranches, mais encourt une pénalité lorsqu'il accepte une demande sans disposer de ressources suffisantes pour répondre aux besoins en ressources de la tranche admise.

Un framework appelé Hap-SliceR [89] pour l'allocation de ressources dans les communications haptiques dans le découpage du réseau 5G. Ce framework donne la priorité à la qualité de service de la tranche haptique, en se concentrant spécifiquement sur la latence, étant donné que les sessions de liaison montante et descendante sont couplées dans les communications haptiques en raison de l'échange bidirectionnel d'informations haptiques. Hap-SliceR vise à déterminer la stratégie de découpage optimale en utilisant une approche d'apprentissage par renforcement plus précisément Q-apprentissage, en donnant la priorité à la QoS de la tranche haptique par rapport à celle des autres tranches.

Les auteurs de [90] ont mis en œuvre une entité de gestion dynamique des ressources utilisant l'apprentissage par renforcement profond. La solution proposée considère trois types de tranches : une tranche pour un débit constant "Constant Bit Rate (CBR)", une autre tranche pour un débit minimum "Minimum Bit Rate (MBR)" et la dernière tranche dédiée pour meilleur effort "Best Effort (BE)", chacune représentant différentes exigences de service en termes d'indicateurs de performance clés (KPI). L'état de l'environnement est représenté sous la forme d'une matrice des paramètres de contrôle, et des KPIs tels que le débit moyen, le taux d'admission et le taux d'abandon. Les actions impliquent l'ajustement des paramètres de contrôle tels que les taux dans le planificateur et les seuils d'admission pour chaque tranche. L'objectif de cette solution est de respecter l'accord de niveau de service (SLA) pour chaque tranche en optimisant les paramètres de planification et de contrôle d'admission.

Dans [91], une stratégie d'allocation de ressources de découpage de réseau de bout en bout (E2E) est introduite. Cette solution, basée sur le Q-apprentissage profond "Deep Q-Learning (DQL)", prend en compte à la fois les tranches du réseau d'accès radio et les tranches du réseau central, ajustant dynamiquement l'allocation des ressources en réponse aux changements environnementaux. L'état est représenté comme un vecteur contenant deux types d'informations : la probabilité d'accès réussi du côté accès pour chaque tranche et le ratio d'utilisateurs obtenant un accès E2E par rapport à ceux tentant d'accéder du côté accès. Le but de cette solution est de satisfaire les contraintes des deux types de tranches, notamment les contraintes de débit et de délai.

Les auteurs dans [92] ont proposé un framework innovant de découpage des ressources radio pour décharger les tâches des véhicules de service URLLC via un réseau de découpage 5G. Le problème d'optimisation est modélisé comme un processus de décision de Markov, où l'état du système à un instant précis "t" se compose de cinq composants : le nombre de tâches, les ressources de calcul dans le serveur MEC, le serveur de convergence, le serveur cloud, et la bande passante disponible. Un agent d'apprentissage par renforcement profond est formé pour gérer le découpage des ressources de bande passante, facilitant ainsi le déchargement des tâches tout en garantissant que les exigences de délai sont respectées.

Une solution de découpage dynamique du RAN pour les services Internet des véhicules

"IoV" avec différentes exigences de QoS est proposée dans [93]. Dans cette approche, une tranche est dédiée aux services sensibles au retard et l'autre aux services tolérants au retard. L'algorithme de gradient politique déterministe profond "Deep Deterministic Policy Gradient (DDPG)" proposé alloue dynamiquement le spectre radio et les ressources de calculs en fonction de la densité du trafic automobile. En raison de la complexité du problème, il est divisé en deux sous-problèmes distincts : un sous-problème d'allocation des ressources et un sous-problème de répartition de la charge de travail.

Un mécanisme d'allocation à deux niveaux basé sur un algorithme d'apprentissage par renforcement profond (DRL) est proposé dans [94]. Ce mécanisme vise à trouver une solution adéquate pour maximiser la qualité d'expérience (QoE) et l'efficacité spectrale (SE) des tranches en réalisant une planification des blocs de ressources (RB) et une allocation de puissance à petite échelle temporelle, tout en tenant compte de la dynamique du trafic des tranches sur une plus grande échelle temporelle.

Dans [95], un algorithme d'apprentissage par renforcement profond basé sur la méthode de l'acteur-critique est présenté pour allouer et gérer des ressources limitées entre différentes tranches, spécifiquement conçues pour les applications automobiles. Cette approche crée un environnement qui simule le trafic en temps réel pour chaque type de service et entraîne l'algorithme pour améliorer la qualité de service de chaque service au sein du réseau. L'environnement est surveillé via deux paramètres : le nombre d'équipements utilisateur du véhicule "Vehicle User Equipment (VUE)" demandant des ressources à la tranche "k" au pas temps actuel, et le nombre de RB alloués à la tranche "k" au pas de temps précédent. L'agent vise à maximiser les récompenses à long terme en prenant des mesures qui aboutissent à une répartition optimale des RB entre les tranches en fonction de leurs demandes de trafic.

Une nouvelle stratégie d'allocation de découpage RAN basée sur l'apprentissage automatique est proposée dans [96]. Elle consiste à établir une politique d'allocation de découpage RAN à travers un réseau d'apprentissage par renforcement "Advantage Actor Critic (A2C)", entraîné en utilisant les caractéristiques temporelles de la mobilité des utilisateurs, prétraitées par des réseaux (LSTM). L'objectif de la stratégie proposée est de maximiser l'efficacité spectrale tout en garantissant le taux de satisfaction des services "Service Satisfaction Rate (SSR)" des différentes tranches du réseau.

Dans [97], un modèle de découpage de réseau hiérarchique à trois niveaux, basé sur l'apprentissage par renforcement profond fédéré (FDRL) et les techniques de gradient de politique déterministe profond "DDPG", tout en préservant la confidentialité. L'objectif de ce modèle proposé est d'optimiser les opérations inter-tranches et la gestion des ressources en maximisant un indicateur d'utilité, qui donne la priorité à l'allocation des ressources pour les tâches déchargées selon leur urgence et importance, tout en minimisant la latence pour les applications critiques dans les réseaux VANETs.

c) Modèles basés sur des heuristiques et des métaheuristiques

1) Modèles basés sur des heuristiques : Les solutions heuristiques sont simples et faciles à mettre en œuvre, conçues pour résoudre des problèmes spécifiques rapidement. Leur objectif principal est de trouver une solution satisfaisante plutôt qu'optimale, en mettant l'accent sur la rapidité et la complexité réduite de la résolution des problèmes.

Un mécanisme de partage des ressources pour résoudre la congestion du trafic de données et la surcharge dans le découpage de réseau 5G est présenté dans [98]. Cette

étude propose deux modes de fonctionnement différents, appelés schéma de partage des ressources statique (SSR) et schéma de partage des ressources dynamique "Dynamic Resource Sharing (DSR)". L'objectif de cette solution est de réaffecter équitablement les ressources partagées disponibles du réseau 5G grâce à un algorithme basé sur une heuristique, tout en protégeant une tranche 5G pour qu'elle ne surcharge pas les autres tranches 5G.

Les auteurs dans [99] ont proposé une solution d'allocation de ressources inter-tranches utilisant un algorithme heuristique adaptatif, avec deux objectifs principaux : i) déterminer le nombre de canaux dédié à chaque tranche, et ii) sélectionner les positions de canaux appropriées pour chaque tranche.

Les auteurs dans [100] ont proposé un modèle de communication basé sur le découpage du réseau pour les communications V2X. Ce modèle divise le trafic automobile en deux tranches logiques : l'une dédiée à la conduite autonome pour l'échange de messages de sécurité, et l'autre à l'infodivertissement, fournissant un streaming vidéo dans un scénario d'autoroute à plusieurs voies. L'approche utilise une stratégie de relais dans laquelle les véhicules dotés d'un faible rapport signal/interférence plus bruit "Signal to Interference plus Noise Ratio (SINR)" pour le streaming vidéo s'appuient sur des véhicules dotés de liaisons V2V et V2I de haute qualité pour améliorer les performances du système. Le but de l'approche proposée est d'augmenter le débit de réception des paquets "Packet Reception Rate (PRR)" dans chaque tranche.

2) Modèles basés sur des métaheuristiques : Ces modèles reposent sur une approche heuristique ou une combinaison de plusieurs heuristiques pour générer un ensemble de solutions initiales, puis cherchent à les améliorer. Cependant, la qualité des résultats obtenus n'est pas toujours garantie par ces méthodes, ce qui pousse la plupart des chercheurs à démontrer la qualité de leurs solutions en développant des bornes inférieures suffisamment bonnes pour les comparer à leurs résultats. Toutefois, obtenir de bonnes bornes nécessite souvent un temps de calcul important, ce qui constitue une limite notable. Les métaheuristiques sont plus flexibles et peuvent être adaptées à différents types de problèmes. Parmi les méthodes les plus fréquentes sont les algorithmes génétiques.

Dans [101], un optimiseur en ligne basé sur des algorithmes génétiques a été proposé pour le contrôle d'admission de tranches hétérogènes, visant à maximiser l'utilité à long terme du réseau tout en garantissant une grande évolutivité. Cette approche ne nécessite aucune connaissance préalable des modèles de trafic et d'utilité. Les stratégies de découpage ont été codées sous forme de séquences binaires pour gérer les demandes et le mécanisme de décision. La fonction d'utilité du système dépend de celle de chaque tranche acceptée, laquelle est définie de manière flexible selon le point de vue du propriétaire de la tranche, en fonction de certains indicateurs clés de performance (KPI) spécifiés, tels que le débit du réseau, la latence moyenne ou la fiabilité du réseau.

Les auteurs de [102] ont introduit une nouvelle méthode utilisant un algorithme génétique à encodage réel pour la gestion des ressources inter-tranches. L'objectif de ce travail est de réduire le nombre d'utilisateurs n'ayant pas obtenu le débit requis, tout en minimisant l'isolation des ressources radio inter-tranches. Pour cela, ils ont défini un nouveau concept appelé Facteur d'Isolation, destiné à estimer le degré d'isolation des ressources radio inter-tranches causée par l'interférence entre cellules des tranches situées dans différentes cellules. Chaque solution candidate à ce problème d'allocation de ressources dans

un système composé de plusieurs cellules est codée sous forme d'un tableau de valeurs de paramètres, appelé chromosome. Chaque chromosome est composé de plusieurs gènes, représentés sous forme de vecteur, qui indique le schéma d'allocation des ressources des tranches dans une cellule spécifique.

d) Modèles basés sur l'optimisation mathématique

Ces modèles tels que la programmation linéaire "Linear Programming (PL)" et non linéaire "Nonlinear Programming (PNL)" sont largement exploités dans la gestion des ressources dans le découpage réseau 5G. Ils visent à trouver la meilleure solution possible selon un ou plusieurs critères d'optimisation. Ces modèles permettent de modéliser les contraintes et les objectifs complexes du système.

Dans [103], Les auteurs ont introduit un ordonnanceur flexible à deux couches pour le découpage du RAN avec une faible complexité, prenant en compte diverses contraintes : les exigences de QoS des utilisateurs pour chaque tranche, la priorité, l'isolation et l'équité entre les tranches. Le problème résultant de ces contraintes a été étendu de la méthode des multiplicateurs de Lagrange aux conditions de Karush-Kuhn-Tucker (KKT) [104], où les contraintes sont des fonctions linéaires, l'espace de solution est un ensemble convexe, et la fonction d'optimisation est concave. Ces KKT conditions sont suffisantes pour trouver une solution optimale. L'objectif de cette approche de découpage est de trouver un compromis adéquat entre ces contraintes.

Dans [105], un algorithme d'allocation dynamique pour le découpage de la bande passante et l'allocation de ressources visant la coexistence de deux tranches est proposé : une pour l'IoT et l'autre pour le streaming vidéo, en utilisant la méthode d'optimisation de Lyapunov [62]. Le problème est formulé sous deux contraintes principales : minimiser la dépense totale d'énergie de tous les dispositifs IoT tout en garantissant que toutes les données générées par ces dispositifs soient transmises vers le serveur backend IoT, connecté à une station de base via Internet, dans un délai imparti. De plus, la tranche dédiée au streaming vidéo vise à maximiser la qualité de l'expérience (QoE) pour tous les utilisateurs. L'algorithme proposé opère à deux échelles de temps différentes : une échelle de temps longue pour le découpage des ressources entre les deux tranches, et une échelle de temps courte pour l'allocation des ressources, la planification des dispositifs IoT et l'allocation de la puissance de transmission, ainsi que pour la prise de décision concernant la qualité de la vidéo pour le service de streaming.

Dans [106], une solution est proposée pour allouer des ressources entre différentes tranches au sein d'une infrastructure à capacité partagée, où chaque tranche a des exigences distinctes en matière de QoS et de SLA. La stratégie de découpage est formulée à partir d'un problème de programmation convexe pour gérer l'allocation des ressources et d'une politique de contrôle d'admission qui utilise la préemption des ressources entre les tranches. De plus, un algorithme basé sur la méthode de projection de gradient est appliqué pour résoudre le problème d'optimisation. Les principaux objectifs de l'approche de découpage des ressources proposée sont de garantir une utilisation efficace des ressources, de maintenir l'équité et d'isoler les performances entre les tranches.

Dans [107], les auteurs introduisent un algorithme conçu pour améliorer le débit total de l'opérateur de réseau cellulaire tout en garantissant un certain niveau d'équité pour les utilisateurs mobiles des réseaux 5G. L'algorithme vise à réaffecter équitablement la bande passante inutilisée en termes d'un débit fixe pour améliorer la qualité de l'expérience pour

les utilisateurs ayant des exigences de tarif fixe, après avoir satisfait à ces exigences.

2.4 Discussion

Comme nous l'avons vu plus haut, différentes solutions ont été proposées pour la gestion des ressources dans le contexte du découpage réseau. Nous allons maintenant discuter de la possibilité d'adapter ces approches à l'environnement véhiculaire.

Les approches basées sur la réservation fixe ont démontré leur utilité en offrant une isolation totale dans la gestion des ressources, en particulier pour les tranches ayant un trafic prévisible qui change rarement, comme les tranches IoT. Cependant, cette option devient une limitation dans un environnement tel que l'environnement véhiculaire, où le trafic est en constant changement en raison de la mobilité des véhicules. En fait, dans la littérature, les solutions basées sur l'approche de réservation totale ont été adaptées à l'environnement véhiculaire en ajoutant un mécanisme de partage entre les tranches pour surmonter les fluctuations des ressources. Certains chercheurs ont travaillé sur le contrôle d'admission en déclenchant un mécanisme de partage chaque fois qu'une demande de connexion provenant d'un véhicule est rejetée par sa tranche, d'autres chercheurs ont proposé de permettre aux tranches de réseau de partager leurs ressources pour une période.

La plupart de ces solutions ont reconnu l'importance de s'appuyer sur deux niveaux d'allocation pour permettre une personnalisation pour chaque tranche tout en réduisant la complexité en gérant chaque niveau séparément.

Les solutions basées sur la non-réservation, contrairement à la réservation fixe, ont démontré leur flexibilité dans la gestion des fluctuations des ressources. De nombreuses solutions de ce type ont été adaptées à l'environnement véhiculaire en proposant des méthodes de complexité réduite, car l'allocation est effectuée de manière répétée à intervalles réguliers et rapprochés, que ce soit entre les tranches ou entre les utilisateurs finaux, ce qui souligne l'importance de recourir à des techniques pour trouver des solutions à ce type de problème avec une complexité réduite. La plupart de ces solutions n'ont pas réussi à fournir un niveau suffisant d'isolation et d'équité entre les tranches.

Les solutions basées sur la réservation partielle des ressources ont été proposées comme une approche hybride combinant les deux approches précédentes. Bien qu'elles héritent à la fois des limites et des points positifs de ces approches, elles nécessitent encore une adaptation spécifique au réseau véhiculaire, notamment en ce qui concerne le choix du nombre de ressources à allouer spécifiquement à chaque tranche au début de son cycle de vie, ainsi que la part qui sera partagée en concurrence entre toutes les tranches coexistantes.

Les solutions proposées pour les environnements non-véhiculaires présentent une limitation commune lorsqu'elles sont appliquées directement dans l'environnement V2X : elles manquent la flexibilité à gérer la forte mobilité des véhicules, ce qui rend la demande de ressources très fluctuante au fil du temps. Ces solutions peuvent présenter une complexité très élevée pour prendre une décision d'allocation de ressources.

En analysant les différents travaux proposés utilisant des techniques telles que les approches heuristiques, méta-heuristiques, la théorie des jeux et d'autres. L'apprentissage par renforcement a montré une efficacité accrue dans la prise de décision concernant le partage des ressources. Contrairement aux approches, qui reposent souvent sur des hypothèses prédéfinies ou nécessitent une modélisation précise des environnements, les modèles d'apprentissage par renforcement ont la capacité de s'adapter dynamiquement aux varia-

tions continues de l’environnement. Ils se montrent particulièrement efficaces même face à des situations imprévues, grâce à leur interaction directe avec l’environnement, qui leur permet un apprentissage en temps réel tout en réduisant la complexité computationnelle. En effet, nous avons besoin d’une solution qui permet une gestion des ressources dans l’environnement du découpage réseau V2X avec une flexibilité qui garantit une utilisation efficace des ressources. Cette solution doit également prendre en compte l’isolation existante entre les tranches tout en offrant une allocation personnalisée pour chaque tranche. Le tableau 2.1 présente un récapitulatif sur les travaux susmentionnés.

TABLE 2.1 – Récapitulatif sur les solutions proposées pour la gestion des ressources dans le découpage réseau 5G

Approche	Solution	Cas d'utilisation	Niveau de gestion des ressources	Techniques utilisées	Partage des ressources entre les tranches
Réservation complète des ressources	[51]	Environnement de découpage V2X	Contrôle d'admission des utilisateurs	Jeux de négociation (La théorie des jeux)	Oui
	[59]	Type de tranches non spécifié	Planification des ressources	Un accord de partage prédéfini	Oui
	[98]			Algorithme heuristique	Oui
	[108]			Programmation Convexe	Oui
	[69]		Contrôle d'admission des utilisateurs	La méthode d'itération sur les valeurs	Oui
	[60]	Les trois tranches génériques	Planification des ressources	Migration des ressources radio non utilisées en utilisant un seuil fixe	Oui
	[109]	Type de tranches non spécifié	Contrôle d'admission des tranches	Algorithme génétique (AG)	Non
	[57]			L'apprentissage par renforcement	Non
	[81]			L'apprentissage automatique	Non
	[82]			Le modèle LSTM	Non
	[83]	Les trois tranches génériques	Planification des ressources	Un algorithme Random Forest (RF) et des réseaux neuronaux d'apprentissage profond (DLNN)	Non
	[84]			Un réseau neuronal convolutif (CNN) + Le modèle LSTM	Non
	[58]	Type de tranches non spécifié	Planification des ressources	Les techniques de Branch-and-Bound et de dualité primale relaxée	Non
	[79]		Planification des ressources + Contrôle d'admission des utilisateurs	The Fisher Market solution (La théorie des jeux)	Oui
	[72]	Les trois tranches génériques	Planification des ressources	Allocation des ressources à deux niveaux fixes	Non
	[26]		Contrôle d'admission des tranches	Méthode Q-Learning + Méthode Deep Q-Learning + Méthode de Regret Matching	Non
Réservation partielle des ressources	[64]	Type de tranches non spécifié	Contrôle d'admission des utilisateurs	Le premier arrivé, premier servi pour la partie partagée non réservée	La partie non réservée est partagée
	[65]		Planification des ressources	Un algorithme itératif à faible complexité	
	[73]			Une enchère combinatoire hiérarchique (La théorie des jeux)	
	[66]		Planification des ressources + Contrôle d'admission des utilisateurs	Algorithme d'adaptation basé sur les retours	
Sans réservation des ressources	[99]	Type de tranches non spécifié	Planification des ressources	Algorithme heuristique	Aucun partage nécessaire
	[110]			Algorithme génétique	
	[74]	Une tranche pour M2M devices + Une tranche pour H2H devices		Algorithme heuristique	
	[71]	Environnement de découpage V2X		Programmation dynamique	

Approche	Solution	Cas d'utilisation	Niveau de gestion des ressources	Techniques utilisées	Partage des ressources entre les tranches
	[103]	Les trois tranches génériques		Méthode des multiplicateurs de Lagrange	
	[61]	Environnement de découpage V2X		Méthode d'optimisation de Lyapunov	
	[100]			Une approche de relais basée sur le rapport signal-sur-interférence-plus-bruit (SINR)	
	[86]			Méthode Q-Learning	
	[92]			Méthode Deep Q-Learning	
	[93]			Méthode Deep Q-Learning	
	[95]			Le modèle Actor-Critic (Deep Q-Learning)	
	[97]			L'apprentissage par renforcement profond fédéré	
	[76, 77]	Type de tranches non spécifié		Jeu de faillite (La théorie des jeux)	
	[94]			Méthode Deep Q-Learning	
	[96]			Advantage Actor Critic (A2C) + Long Court Terme Mémoire (LSTM)	
	[105]	Une tranche IoT + Une tranche pour video streaming		Méthode d'optimisation de Lyapunov	
	[106]	Type de tranches non spécifié	- Planification des ressources - Contrôle d'admission des utilisateurs	Réallocation des ressources basée sur le rapport signal-sur-interférence-plus-bruit (SINR)	
	[63]		Contrôle d'admission des utilisateurs	Approche basée sur le SINR des utilisateurs	
	[75]		- Planification des ressources - Contrôle d'admission des utilisateurs	Algorithme heuristique	
	[78]		Planification des ressources	Jeu de Stackelberg	
	[80]			Jeu d'enchères	
	[87]	Une tranche URLLC + Une tranche pour video + Une tranche pour VoLTE	Planification des ressources	Méthode Deep Q-Learning	
	[89]	Une tranche pour les communications haptiques + Une tranche eMBB		Méthode Q-Learning	
	[90]	Une tranche à débit constant (CBR) + tranche à débit minimum (MBR) + tranche à meilleur effort (BE)	- Planification des ressources - Contrôle d'admission des utilisateurs	Méthode Deep Q-Learning	
	[91]	Une tranche eMBB + Une tranche URLLC	Planification des ressources	Méthode Deep Q-Learning	

2.5 Conclusion

Le découpage réseau est un domaine de recherche émergent qui a vu le jour pour répondre aux différentes exigences des nouveaux services dans le cadre des réseaux 5G. Ce domaine est accompagné de nombreux défis, parmi lesquels figure le problème de la gestion des ressources entre les différentes tranches de réseau créées par le découpage réseau. L'une des parties du réseau où les ressources sont particulièrement limitées est le réseau d'accès radio (RAN), notamment en termes de ressources en bande passante. Contrairement à d'autres parties du réseau où il est possible d'ajouter plus de matériel pour augmenter les ressources disponibles, le RAN ne peut pas simplement être étendu de cette manière, ce qui rend la gestion des ressources dans cette zone plus complexe. L'environnement véhiculaire, avec son aspect V2X, qui cherche à tirer parti des avantages du découpage réseau en s'adaptant à ses services, héritera également des défis de gestion des ressources.

Dans ce chapitre, nous avons présenté les différentes solutions de gestion des ressources proposées dans le découpage réseau 5G. Cet état de l'art est articulé autour d'une classification de ces solutions selon certains aspects, nous a permis de constater qu'une solution adaptée à la communication V2X dans le cadre du découpage réseau est essentielle, car elle permet d'utiliser les ressources de manière efficace tout en assurant une qualité de service élevée pour chaque service, avec une isolation renforcée.

Les méthodes existantes de gestion des ressources ne peuvent pas être appliquées directement avec succès pour gérer les tranches V2X, une adaptation à l'environnement véhiculaire est nécessaire, car celui-ci présente des besoins différents et des contraintes spécifiques. Nous avons également étudié différentes méthodes proposées pour les tranches V2X ainsi que leurs limites. Les principales limitations incluent l'efficacité dans l'utilisation des ressources, la complexité accrue due à la forte mobilité, et l'isolation des tranches. Ces observations ont conduit à la proposition d'une contribution basée sur la réservation fixe qui se concentre sur le partage des ressources entre les tranches V2X. Les détails de cette contribution seront présentés en profondeur dans le chapitre suivant.

Chapitre 3

Mécanisme de partage des ressources dans un environnement de découpage réseau V2X

Chapitre 3

Mécanisme de partage des ressources dans un environnement de découpage réseau V2X

3.1 Introduction

L'adaptation du découpage réseau à l'environnement véhiculaire est encore en phase d'évolution, ce qui implique qu'il est primordial de prêter attention aux défis de gestion des ressources. Plus précisément, la partie RAN du réseau représente un défi majeur comparé aux autres parties, en raison de la rareté de la bande passante. Parmi les approches de gestion de bande passante, on trouve l'approche "strict slicing", qui garantit une isolation totale mais manque de flexibilité, nécessitant ainsi des améliorations pour être adaptée à l'environnement véhiculaire.

Dans ce chapitre, nous présentons notre première contribution [111, 112], qui permet une flexibilité dans la gestion des ressources en déclenchant un mécanisme de partage de ressources entre les tranches V2X sous-utilisées et surchargées. La solution proposée a été développée en utilisant le simulateur OMNET++, évaluée et comparée à certains autres solutions. Notre solution a permis d'améliorer l'utilisation des ressources en réduisant la probabilité de blocage des nouveaux appels ainsi que la probabilité de coupure lors du passage d'une cellule à une autre, tout en maintenant une isolation élevée entre les tranches.

3.2 Motivation

Comme nous l'avons vu dans le chapitre précédent, de nombreux travaux ont été proposés pour la gestion des ressources dans le découpage réseau, à différents niveaux de gestion à l'image du contrôle d'admission entre les utilisateurs ou les tranches, en se basant sur diverses approches, avec utilisation de plusieurs techniques telles que l'apprentissage par renforcement, tout en intégrant ces solutions dans des environnements variés tels que le découpage de réseau véhiculaire. L'approche de découpage strict des ressources entre les tranches "strict slicing" garantit une isolation totale, protégeant ainsi chaque tranche de toute fluctuation provenant des autres tranches. Cependant, cette isolation totale devient

un obstacle pour une gestion des ressources avec une efficacité, surtout dans les scénarios où certaines tranches ne parviennent pas à utiliser 100% de leurs ressources, alors que d'autres sont surchargées. En ce qui concerne l'environnement véhiculaire, celui-ci présente une particularité, à savoir une mobilité élevée, ce qui entraîne de fortes fluctuations dans les demandes de ressources des différentes tranches au cours de leur cycle de vie. De ce point de vue, on peut constater que l'approche strict a besoin d'une adaptation à cet environnement.

Plusieurs chercheurs ont proposé une adaptation via le partage des ressources à travers une planification des ressources. Cela signifie qu'après chaque période, une nouvelle décision de partage est prise entre les tranches. Cependant, cela entraîne une complexité de calcul, car il devient nécessaire de trouver une solution à chaque période de temps, ces périodes étant très courtes. Le déclenchement du processus de partage uniquement lorsque cela est nécessaire semble être une meilleure option. Cela signifie que lorsqu'une demande de connexion est rejetée dans une tranche via son contrôle d'admission, nous déclenchons le processus de partage pour accepter cette demande de connexion. Beaucoup de ces chercheurs se sont concentrés sur l'amélioration des performances, telles que l'utilisation des ressources, sans accorder trop d'importance à l'isolation. Cela revient à négliger l'un des principes fondamentaux du découpage réseau, notamment en pénalisant les actions de partage qui dégradent l'isolation entre les tranches.

L'apprentissage par renforcement a suscité l'intérêt de nombreux chercheurs en raison de son aptitude à résoudre des problèmes de prise de décision, en particulier dans des environnements dynamiques et incertains, ce qui en fait de l'apprentissage par renforcement une approche efficace pour résoudre la problématique du partage des ressources dans un environnement aussi dynamique que celui du découpage réseau V2X.

3.3 Contribution

Sur la base de toutes les considérations ci-dessus, les principaux apports de nos contributions sont résumés comme suit :

- Ce travail propose un modèle de Q-learning profond (DQL) pour optimiser la politique de prise de décision de partage entre les tranches dans le scénario 5G V2X de découpage réseau.
- Contrairement au travail basé sur la planification des ressources, où une nouvelle décision de découpage pour l'allocation des ressources est prise après chaque période, notre solution n'est déclenchée que dans les situations où il y a une tranche surchargée, comme illustré dans la figure 3.1.
- La solution est proposée avec une flexibilité obtenue en faisant simplement varier quelques paramètres, la priorité entre les nouvelles demandes d'appels et demandes d'appels de transfert (handover), et entre tranches surchargées.
- Contrairement aux approches de partage de ressources proposées pour l'allocation de ressources fixes décrites dans la littérature, notre approche considère les décisions de partage qui conduisent à une dégradation des performances dans des tranches sous-chargées comme de mauvaises décisions, même si elles conduisent à une meilleure utilisation des ressources.

- L'aspect personnalisation de l'environnement de découpage est garanti puisque notre approche n'interfère pas dans la politique de contrôle d'admission du slice. La solution que nous proposons fonctionne à un niveau élevé.

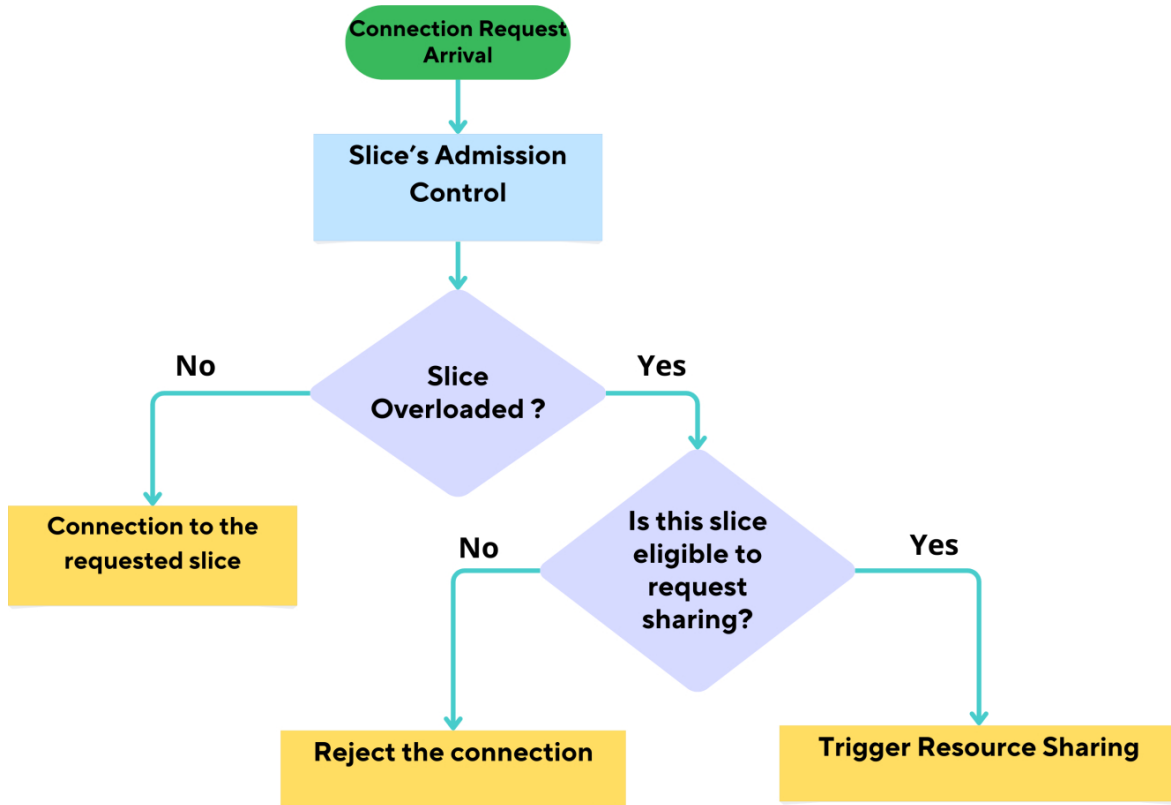


FIGURE 3.1 – Organigramme de déclenchement du mécanisme de partage

3.4 L'apprentissage par renforcement

L'apprentissage par renforcement "Reinforcement Learning (RL)" est situé entre l'apprentissage supervisé et l'apprentissage non supervisé. Le paradigme de l'apprentissage par renforcement traite de l'apprentissage dans le cadre des problèmes de prise de décision séquentielle dans lesquels le "feedback" est limité. RL vise à trouver une solution à un problème en permettant à l'apprenant/décideur que nous appelons l'agent d'apprendre comment se comporter dans un environnement, où le seul retour d'information consiste en un signal de récompense scalaire. Il peut traiter des scénarios plus difficiles pour lesquels aucune connaissance préalable du problème n'est disponible. Dans ce contexte, l'algorithme doit interagir et expérimenter avec l'environnement pour apprendre à optimiser son comportement. Le processus d'apprentissage est piloté par les retours "feedback" évaluatifs reçus de l'environnement.

3.4.1 Processus de décision Markovien

Pour formuler les interactions, les éléments du problème d'apprentissage par renforcement peuvent être modélisés à l'aide du modèle "*processus de décision Markovien (MDP)*".

Les MDPs sont devenus un formalisme fondamental pour la planification en théorie de la décision, l'apprentissage par renforcement et d'autres problèmes d'apprentissage dans les domaines stochastiques. En général, un modèle MDP se compose à travers le 5-uplet suivant $M = \langle S, A, P(s'|s, a), R, \gamma \rangle$ [113]. La figure 3.2 illustre l'interaction entre l'agent et l'environnement.

- **L'état** : L'ensemble des états environnementaux S est défini comme l'ensemble fini $\{s_1, \dots, s_N\}$ où la taille de l'espace d'état est N , c'est-à-dire $|S| = N$. Un état est le composant responsable de fournir à l'agent une représentation de l'environnement à chaque moment de l'interaction.
- **L'action** : L'ensemble des actions A est défini comme l'ensemble fini $\{a_1, \dots, a_K\}$ où la taille de l'espace d'action est K , c'est-à-dire $|A| = K$. L'action peut être utilisée pour contrôler l'état du système. L'ensemble des actions pouvant être appliquées dans un état particulier ($s \in S$), est noté $A(s)$, où $A(s) \subseteq A$. Dans certains systèmes, toutes les actions ne peuvent pas être appliquées dans chaque état.
- **La fonction de transition** : $P(s'|s, a)$ indique la probabilité qu'une application de l'action ($a \in A$) sous l'état ($s \in S$) à l'emplacement " t " conduit à l'état ($s' \in S$) à l'emplacement " $t + 1$ ".
- **La récompense** : $R(s, a)$ est une récompense immédiate après avoir effectué l'action " a " sous l'état " s ". La fonction de récompense est une partie importante du modèle MDP qui spécifie implicitement l'objectif de l'apprentissage.
- **Un facteur d'actualisation** : $\gamma \in [0, 1]$, qui reflète l'importance décroissante des récompenses futures par rapport à celles du présent.

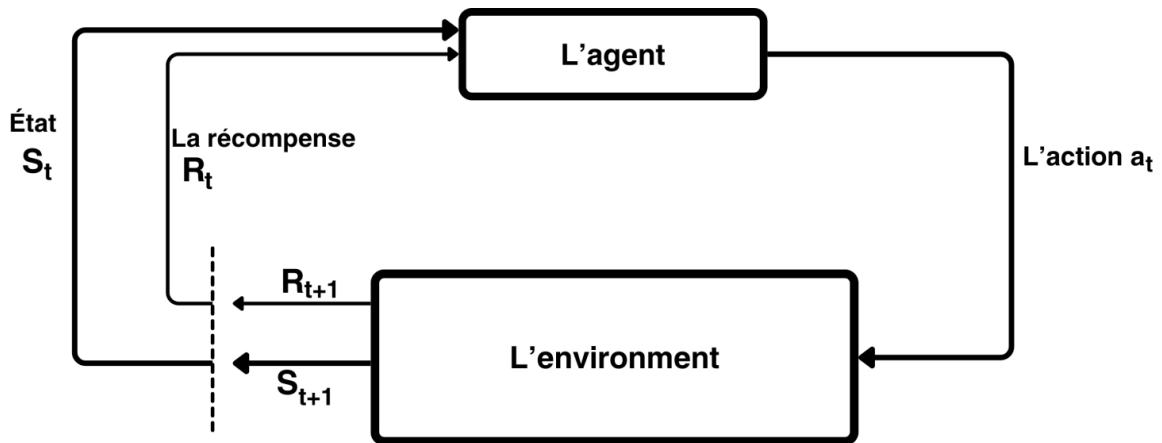


FIGURE 3.2 – Interaction entre l'agent et l'environnement

Presque tous les algorithmes d'apprentissage par renforcement impliquent l'estimation de fonctions de valeur où il s'agit de fonctions d'états (ou de paires état-action) qui évaluent dans quelle mesure il est bénéfique pour l'agent d'être dans un état particulier (ou de reprendre une action spécifique dans un état particulier). Généralement, le but d'un processus de décision Markovien (MDP) est de déterminer une politique $a = \pi(s)$ qui

sélectionne une action "a" dans l'état "s" afin de maximiser la fonction de valeur. Si l'agent ne se préoccupait que de la récompense immédiate, un critère d'optimalité simple serait d'optimiser $E[r_t]$. Rappelons qu'une politique, π , est une fonction qui associe à chaque état ($s \in S$) et action ($a \in A(s)$) la probabilité $\pi(a|s)$ de choisir l'action "a" lorsqu'on se trouve dans l'état "s". Les MDPs déterminent la politique optimale en apprenant des fonctions de valeur, qui estiment à quel point il est avantageux pour l'agent d'être dans un état spécifique ou d'entreprendre une action particulière dans cet état. La mesure du "à quel point" est définie selon un critère d'optimalité, généralement basé sur le rendement attendu. Les fonctions de valeur sont associées à des politiques spécifiques [113]. La valeur d'un état "s" sous la politique π , notée $V^\pi(s)$, représente le retour attendu lorsqu'on commence dans l'état "s" et qu'on suit ensuite la politique π . Cette fonction de valeur est définie comme la récompense cumulative escomptée et actualisée, selon l'équation de Bellman [113].

$$V^\pi(\hat{s}) = \mathbb{E}_\pi \left[\sum_{k=0}^{\infty} \gamma^k R(s^{(k)}, \pi(s^{(k)})) \mid s^{(0)} = \hat{s} \right] \quad (3.1)$$

$$V^\pi(\hat{s}) = \mathbb{E}_\pi \left[R(\hat{s}, \pi(\hat{s})) + \gamma \sum_{s' \in S} P(s' \mid \hat{s}, \pi(\hat{s})) V^\pi(s') \right] \quad (3.2)$$

où $E_\pi[\cdot]$ désigne la valeur attendue d'une variable aléatoire étant donné que l'agent suit la politique π .

3.4.2 Exploration/Exploitation

Parmi les défis dans l'apprentissage par renforcement, contrairement à d'autres types d'apprentissage, consiste à gérer l'équilibre entre l'exploration et l'exploitation. Pour obtenir une récompense cumulée élevée, l'agent doit choisir dans chaque état des actions qu'il a précédemment essayées et jugées efficaces pour générer des récompenses. Cependant, pour identifier ces actions réussies, l'agent doit également expérimenter des actions qu'il n'a jamais essayées auparavant. L'agent est confronté à un dilemme : il doit exploiter ses connaissances actuelles pour obtenir des récompenses, mais il doit également explorer de nouvelles actions pour améliorer sa prise de décision future. S'appuyer uniquement sur l'exploration ou l'exploitation aboutirait à l'échec. Par conséquent, l'agent doit appliquer une série d'actions et donner de la priorité progressivement à celles qui lui semblent les plus prometteuses. Dans un environnement stochastique, chaque action doit être testée plusieurs fois pour obtenir une estimation fiable de sa récompense attendue. Il existe plusieurs méthodes simples pour équilibrer l'exploration et l'exploitation, comme la politique $\epsilon - greedy$. Dans cette approche, l'apprenant sélectionne sa meilleure action dans l'état actuel avec une probabilité de $(1 - \epsilon)$ ou choisit une action aléatoire avec une probabilité de ϵ . Une autre méthode est la stratégie d'exploration *softmax*, qui est légèrement plus complexe que la stratégie $\epsilon - greedy$. Dans cette stratégie, la sélection des actions reste aléatoire, mais les probabilités de sélection sont pondérées en fonction des valeurs relatives "Q-values" des actions.

3.4.3 Apprentissage en ligne/hors ligne

Dans l'apprentissage en ligne, les actions sont appliquées directement à l'instance du problème, permettant à l'agent d'apprendre des interactions en temps réel. Cependant, cette approche peut entraîner des problèmes de performances importants car lors de la phase d'exploration, l'agent peut prendre des décisions incorrectes ou non testées à certains moments. D'autre part, l'apprentissage hors ligne consiste à utiliser un simulateur de l'environnement pour générer un grand nombre d'exemples de formation de manière sûre et efficace. L'agent formé peut ensuite être déployé sur le problème réel. Le choix entre ces approches dépend de la nature du problème, en particulier lorsque la simulation du problème réel à des fins de formation est difficile.

3.4.4 Algorithmes basés sur un modèle/sans modèle

Les méthodes basées sur un modèle fonctionnent sous l'hypothèse qu'un modèle du MDP est connu à l'avance. Ce modèle est utilisé pour calculer des fonctions de valeur et des politiques à l'aide de l'équation de Bellman, dans le but de déterminer des fonctions de valeur d'état. Lorsqu'un modèle est disponible, ces fonctions de valeur peuvent ensuite être utilisées pour sélectionner les actions optimales. À la différence des méthodes basées sur un modèle, les méthodes sans modèle s'appuient sur une interaction directe avec l'environnement, telle qu'une simulation, plutôt que sur la connaissance d'un modèle parfait de la politique. À travers cette interaction, elles produisent des échantillons de transitions d'état et de récompenses, qui sont ensuite utilisés pour estimer les fonctions de valeur "état-action". Dans ce contexte, l'agent doit explorer le MDP pour recueillir les informations requises.

3.4.5 L'algorithme Q-Learning

L'un des modèles classiques d'apprentissage par renforcement est le Q-learning, qui appartient à la catégorie sans modèle, hors ligne. Le concept de base du Q-learning consiste à estimer progressivement la valeur de chaque paire "état-action", connue sous le nom de "Q-valeur", sur la base d'un feedback sous forme de récompenses et de la fonction de Q-valeur de l'agent. Ces estimations sont stockées dans une Q-table. Ce processus est appelé exploration de l'environnement inconnu. La fonction de valeur peut généralement être mise à jour de deux manières : la mise à jour de Monte-Carlo et la mise à jour de différence temporelle (TD). Dans une mise à jour Monte-Carlo, l'agent révisé son estimation pour une paire "état-action" en faisant la moyenne des retours de plusieurs épisodes. Dans l'approche TD, la mise à jour affine les estimations de valeur en comparant les estimations obtenues lors de deux épisodes consécutifs, améliorant ainsi progressivement la précision de l'approximation de la Q-valeur. L'algorithme Q-Learning met à jour ses Q-valeurs à l'aide de la méthode de mise à jour par différence temporelle. Il est important de noter que la Q-table est initialement définie sur une certaine valeur et est mise à jour avec de nouvelles Q-valeurs après chaque épisode en utilisant l'équation (3.3) [113].

$$Q_{\text{new}}(s_t, a_t) = (1 - \alpha)Q(s_t, a_t) + \alpha \left(r_{t+1} + \gamma \max_{a \in A} Q(s_{t+1}, a) \right) \quad (3.3)$$

où :

- $Q(s_t, a_t)$: est l'ancienne Q-valeur.
- $Q_{\text{new}}(s_t, a_t)$: est la nouvelle valeur obtenue après mise à jour de l'ancienne.
- α : est le taux d'apprentissage.
- γ : est le facteur qui indique l'importance des récompenses futures.
- r_{t+1} : est une récompense reçue de l'action a_t .
- $\max_{a \in A} Q(s_{t+1}, a)$: est l'estimation de l'action future optimale.

L'algorithme Q-Learning comprend généralement trois étapes principales :

1. L'agent sélectionne une action "a" dans un état donné "s" en fonction d'une politique, telle que ϵ -greedy.
2. L'agent reçoit alors une récompense $R(s, a)$ de l'environnement, et l'état passe à l'état suivant s_{t+1} .
3. Enfin, l'agent met à jour la fonction de Q-valeur en utilisant une approche de différence temporelle comme décrit par l'équation 3.3.

L'algorithme (1) représente les étapes de l'algorithme Q-Learning.

Algorithm 1 L'algorithme Q-Learning

Require: γ and α

- 1: Initialize Q arbitrarily (e.g., $Q(s, a) = 0$ for all $s \in S$ and $a \in A$)
 - 2: **for** each episode **do**
 - 3: s is initialized as the starting state
 - 4: **repeat**
 - 5: choose an action $a \in A(s)$ based on an exploration strategy (e.g., ϵ -greedy)
 - 6: perform action a
 - 7: observe the new state s_{t+1} and received reward r
 - 8: calculate the new Q value through equation (3.3)
 - 9: store the new Q value in the Q table
 - 10: $s := s_{t+1}$
 - 11: **until** s_{t+1} is a goal state
 - 12: **end for**
-

Un inconvénient majeur du l'apprentissage Q-Learning est que lorsque l'espace d'état devient très grand, la taille de la Q-table peut croître de façon exponentielle, ce qui augmente considérablement le temps d'apprentissage "training". L'agent a besoin de plus de temps pour explorer tous les états possibles. Dans de tels cas, la Q-table nécessite non seulement un stockage important, mais rend également peu pratique l'exploration rapide de toutes les paires "état-action". À mesure que le nombre d'états augmente, la mémoire nécessaire au stockage et à la mise à jour de la table augmente également. De plus, le temps nécessaire pour analyser chaque état et générer la Q-table nécessaire devient irréaliste. En raison de cette "malédiction de la dimensionnalité", les méthodes d'approximation des fonctions deviennent une alternative plus attrayante.

3.4.6 Deep Q-Learning

Le DQL utilise des réseaux de neurones pour approximer la fonction $Q(s, a)$ comme $Q(s, a; \theta)$, où θ est un vecteur contenant tous les poids et les biais du réseau neuronal et paramétrise Q . θ est le paramètre qui est mis à jour pendant l'apprentissage. À chaque étape de temps " t ", le DQL prend l'état s_t comme entrée et produit $Q(s_t, a; \theta)$ pour tout $(a \in A)$, où A est l'ensemble des actions possibles. En effet, chaque nœud de sortie représente la Q-valeur d'une action particulière. L'une des lacunes du modèle DQL lors de la phase d'entraînement est son instabilité, principalement en raison de la combinaison avec une fonction de Q-valeurs non linéaires (réseau de neurones). Pour améliorer la stabilité pendant l'entraînement, deux améliorations ont été ajoutées :

- **Reprise d'expérience (Experience Replay)** [114] : Pour éviter la corrélation entre les échantillons consécutifs de l'environnement, l'agent sauvegarde ses expériences passées représentées par le tuple $(s_t, a_t, s'_t, R(s_t, a_t))$ à l'épisode " t " dans un ensemble de données $D_t = \{e_1, \dots, e_t\}$. L'agent sélectionne ensuite au hasard un mini-lot d'expériences à partir de cet ensemble de données pour mettre à jour le réseau neuronal de valeur $Q(s, a, \theta)$. Cela permet à l'agent d'apprendre à partir des mêmes expériences plusieurs fois.
- **Clonage de réseau** [114] : Il est suggéré de créer deux réseaux de neurones (clonage de réseau). Le premier réseau, appelé réseau neuronal d'évaluation, est mis à jour à chaque étape de formation $Q(s, a; \theta)$, tandis que le deuxième réseau, appelé réseau neuronal cible $Q(s, a; \theta^-)$, est mis à jour toutes les " C " étapes en le remplaçant par le réseau neuronal d'évaluation. Les résultats de simulation montrent que cette approche de clonage de réseau améliore la stabilité de l'apprentissage.

Les expériences entre l'agent et l'environnement ($e_t = \langle s_t, a_t, s_{t+1}, R(s_t, a_t) \rangle$) peuvent être stockées dans un buffer de réexamen à chaque étape de temps " t " (un ensemble de données $D_t = \{e_1, \dots, e_t\}$) pour éviter la corrélation entre les échantillons consécutifs de l'environnement. À chaque itération d'entraînement, nous choisissons un mini-lot aléatoire des expériences pour entraîner le réseau. Contrairement à l'algorithme Q-Learning, où pendant la phase d'entraînement, la Q-valeur d'une paire "*état-action*" est mise à jour directement, dans l'algorithme Deep Q-Learning, nous mettons à jour la fonction de perte comme s'est exprimée dans l'équation 3.4 qui compare notre prédiction de la Q-valeur avec la Q-cible, et nous utilisons la descente de gradient pour ajuster les poids de notre réseau de neurones afin d'améliorer l'approximation de nos Q-valeurs [114].

$$L_i(\theta_i) = \mathbb{E}_{(s,a,r,s') \sim U(D)} \left[\left(r + \gamma \max_{a'} Q(s', a'; \theta_i^-) - Q(s, a; \theta_i) \right)^2 \right] \quad (3.4)$$

où γ représente le facteur de mise à jour définissant l'horizon futur de l'agent, θ_i sont les paramètres du réseau " Q " à l'étape " i ", et θ_i^- sont les paramètres du réseau nécessaires pour calculer la cible à l'étape " i ".

L'agent de Deep Q-Learning peut être entraîné en minimisant une séquence de fonctions de perte $L_i(\theta_i)$ qui est mise à jour à chaque itération " i ". De plus, le modèle d'apprentissage DQL peut entraîner sa politique en arrière-plan et la stocker dans le réseau de neurones, ce qu'on appelle le Q-Learning hors ligne. L'utilisation d'une stratégie d'apprentissage en ligne, c'est-à-dire l'application de la configuration directement sur le réseau

réel, peut entraîner une dégradation significative des performances en raison de l'exploration effectuée par l'agent à un moment donné. Ainsi, la principale composante de la solution DQL est le processus d'entraînement qui optimise l'ensemble des paramètres du réseau θ .

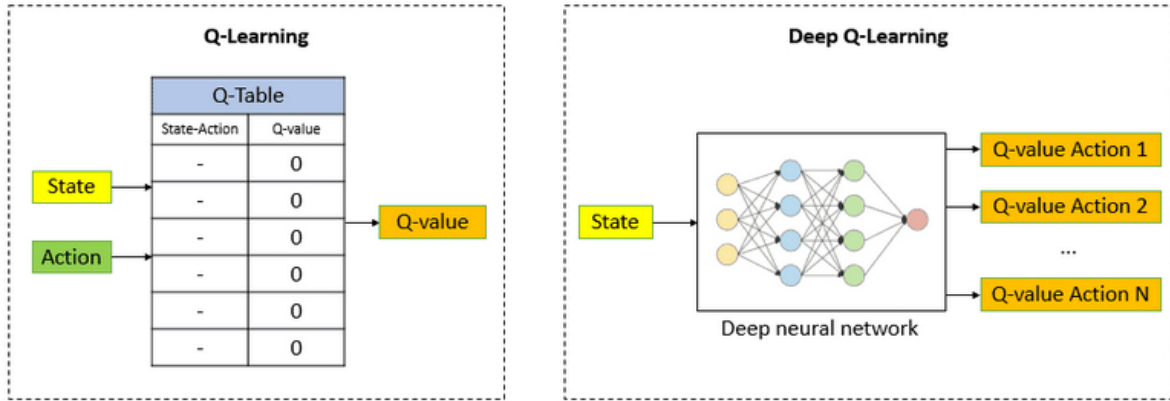


FIGURE 3.3 – La difference entre Q-Learning et Deep Q-Learning

3.5 Contribution : Mécanisme proposé pour le partage des ressources entre les tranches V2X

Dans cette section, nous expliquons comment le Deep Q-Learning est appliqué dans notre problème de partage de ressources destiné pour l'environnement de découpage V2X. Le mécanisme de partage de ressources basé sur DQL est illustré dans la figure 3.4.

3.5.1 Description générale de déclenchement du mécanisme proposé

Dans notre contribution, lors de la création des quatre tranches V2X, chaque tranche reçoit ses ressources dédiées tout au long de son cycle de vie en se basant sur le modèle "strict slicing", où chaque tranche dispose de sa propre politique de contrôle d'admission pour accepter les demandes de connexion. Étant donné que l'environnement véhiculaire est connu pour sa forte mobilité dynamique, cela entraîne de grandes fluctuations dans le trafic de données au sein des tranches, ce qui peut entraîner une surcharge de certaines tranches et une sous-utilisation des autres. Dans cette situation, la politique de contrôle d'admission de la tranche surchargée commencera à rejeter les demandes de connexion, alors que des ressources sont disponibles dans d'autres tranches, ce qui sera considéré comme inefficace en termes d'utilisation des ressources. Dans ce contexte, après le rejet de la connexion, nous tenterons d'accepter cette demande en déclenchant notre mécanisme de partage. Une demande de partage sera alors créée et envoyée à notre module de partage. Cette demande de partage contiendra le type de tranches surchargées et le type de demande (nouvel appel ou transfert d'appel). Notre module de partage vérifiera si la tranche surchargée a la permission de demander des ressources au près des autres tranches. Si c'est le cas,

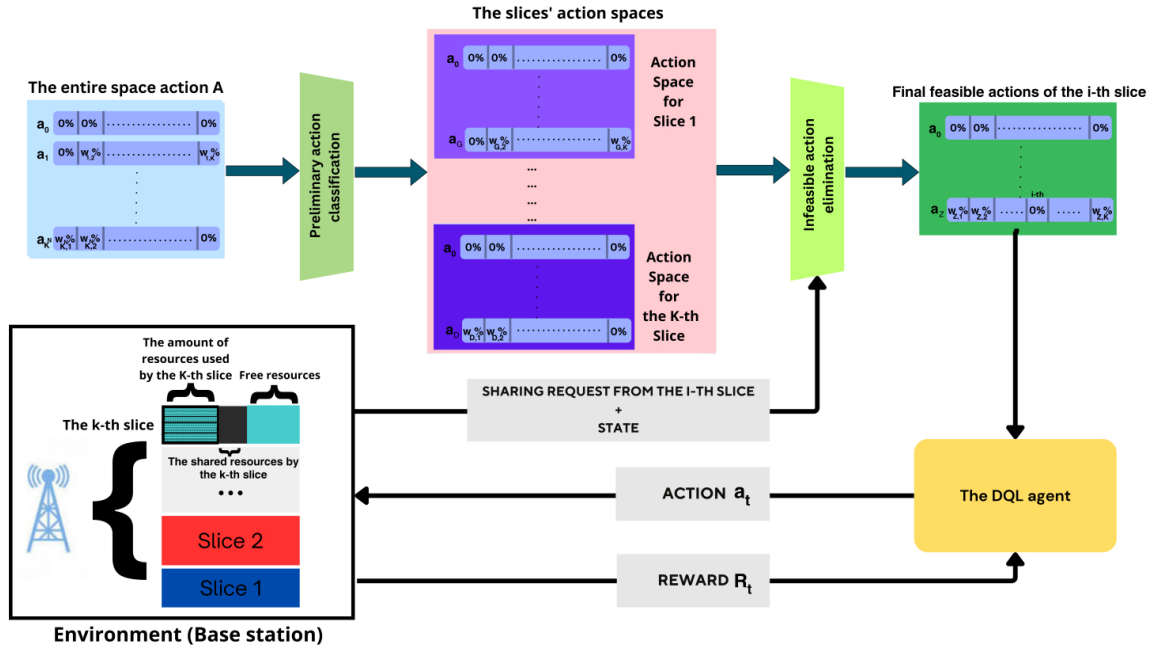


FIGURE 3.4 – Illustration du mécanisme de partage des ressources

la demande de partage sera transférée à l'agent DQL et ce dernier prendra la décision concernant la demande de partage, en choisissant soit de la rejeter, soit de l'accepter. En cas d'acceptation, l'agent déterminera le nombre de ressources que chaque tranche sous-utilisée devra contribuer pour satisfaire cette demande de partage, tout en tenant compte de certaines conditions. L'algorithme (2) illustre comment le partage des ressources est déclenché depuis l'algorithme de contrôle d'admission de tranche jusqu'à l'agent DQL.

3.5.2 Éléments du processus de décision Markovien

Pour appliquer l'apprentissage par renforcement à notre problème, nous devons identifier les trois éléments du MDP : l'état, l'action et la récompense. Ces éléments importants sont définis comme suit :

- **Action** : Chaque action (a_x) est un ensemble de pourcentages :
 $a_x = [w_1\%, w_2\%, w_3\%, \dots, w_K\%]$ où $w_{k,x}$ représente le pourcentage des ressources dédiées de la $k^{ième}$ tranche offertes pour accepter une demande d'une tranche surchargée par l'action (a_x). Si l'on désigne l'ensemble des actions possibles par A , alors la taille de A (c'est-à-dire le nombre d'actions possibles) est N^K , où K est le nombre de tranches du réseau et N est le nombre de pourcentages possibles que la tranche peut offrir pour le partage. Par exemple, $a_0 = [0\%, 0\%, 0\%, \dots, 0\%]$ est une action qui signifie le rejet de la demande.
- **État** : Pour notre système, nous supposons qu'un état $s = (r_1, r_2, \dots, r_K, b, m)$ est composé de trois types d'informations :
 - r_k représente la somme de tous les pourcentages des ressources partagées par la $k^{ième}$ tranche qui ont été offertes à toutes les demandes acceptées.

Algorithm 2 Pseudo-code of Resource Sharing Trigger

Require:

- Connexion_Request_Queue : The queue of the connection requests received for the slice k .
- Sharing_Request_Queue : The queue of the sharing requests received by the sharing module.
- Allowed_Slice_List : The list of slices that have permission to request sharing from other slices.

```

1: // Slice's admission control module
2: while Connexion_Request_Queue is not empty do
3:   Connexion_Request request = Connexion_Request_Queue.front()
4:   Connexion_Request_Queue.pop()
5:   Decision decision = Admission_control_policy_k.Decision(request)
6:   if decision is rejection then
7:     Create sharing_request_packet
8:     sharing_request_packet.SetSliceID(getSliceID())
9:     sharing_request_packet.SetRequestType(request.getType())           {Hando-
       ver/NewCall}
10:    Send sharing_request_packet to Sharing_Module
11:  end if
12: end while
13: // Sharing module
14: Sharing_Module receives sharing_request_packet
15: if sharing_request_packet.GetSliceID() exists in Allowed_Slice_List then
16:   Sharing_Request_Queue.put(sharing_request_packet)
17: else
18:   Remove sharing_request_packet
19: end if
20: while Sharing_Request_Queue is not empty do
21:   Sharing_Request_Packet request_packet = Sharing_Request_Queue.front()
22:   Sharing_Request_Queue.pop()
23:   Send request_packet to DQL_Agent_Module
24: end while

```

- “ b ” est une valeur pouvant être égale à 1, 2, 3 ou 4 pour indiquer le type de tranche (tranche de conduite autonome, tranche de conduite téléguidée, tranche de diagnostic et gestion à distance des véhicules ou tranche d’infodivertissement véhicule) de la dernière demande acceptée.
- “ m ” est une valeur pouvant être égale à 1 ou 2 pour indiquer le type de demande (demande de transfert ou nouvelle demande) de la dernière demande acceptée.

où $r_k = 0$, $b = 0$, $m = 0$ dans l’état initial S_0 .

Les différentes transitions du système se produisent lorsqu’une nouvelle demande de partage arrive d’une tranche surchargée, qu’il s’agisse d’une demande de type

"handover" ou d'une nouvelle demande de connexion, nécessitant une décision, ou lorsqu'une demande acceptée quitte le système fonctionne de la manière suivante. En cas d'acceptation, $w_{k,x}$ les ressources partagées par la tranche "k" dans l'action a_x sont ajoutées à r_k , le pourcentage des ressources déjà partagées par la tranche "k", et "b" reçoit la valeur qui indique le type de tranche, et "m" reçoit la valeur qui indique le type de demande. Il est important de mentionner que les ressources partagées sont retournées à leurs tranches d'origine à la fin de la communication.

- **Récompense** : Dans le processus de Markov, l'objectif est de trouver un chemin qui maximise la récompense, qui est une partie importante du MDP spécifiant implicitement l'objectif de l'apprentissage. L'objectif de notre solution est de maximiser l'utilisation des ressources en offrant des ressources aux tranches surchargées sans affecter les performances des tranches sous-utilisées. Il est important de mentionner que les récompenses sont dans la gamme de $[C, \alpha_{\max} \cdot p_{\max} \cdot R_{\max}]$, où C est la pénalité, $\alpha_{\max}$ est la priorité maximale donnée à une tranche surchargée, $p_{\max}$ est la priorité maximale donnée à un type de demande, et $R_{\max}$ est la récompense positive maximale. La récompense est obtenue à la suite de la décision de partage a_x prise lors de l'arrivée d'une demande de partage de la tranche surchargée. Les différentes conditions de la partie récompense sont expliquées comme suit :

1. Le premier cas est lorsque l'agent DQL décide de rejeter directement la demande de partage, qu'il s'agisse d'une demande de handover/nouvelle d'une connexion de la tranche surchargée, via l'action a_0 (ligne 1), ce qui signifie qu'aucune ressource n'est partagée pour accepter la demande. La récompense dans ce cas est $R = 0$.
2. Dans l'autre cas, la décision de l'agent DQL est de choisir une action autre que a_0 . Tout d'abord, chaque tranche offrira la quantité minimale de ressources entre ses ressources disponibles et le pourcentage demandé par l'action. La quantité de ressources offertes par l'action (a_i) est calculée comme suit :

$$x_i = \sum_{k=1}^K \min(B_k \cdot w_{i,k}, B_k - U_{k,t} - (B_k \cdot r_k)) \quad (3.5)$$

B_k est le nombre de ressources dédiées à la tranche "k", $U_{k,t}$ est le nombre de ressources utilisées par la tranche "k" à l'instant "t". Ensuite, nous vérifierons que le nombre de ressources offertes par l'action n'est pas inférieure à celle demandée dans la demande de partage reçue (ligne 4). Sinon, l'action sera interprétée comme un rejet, puisque la demande de partage ne sera pas acceptée avec moins de ressources, et la récompense sera une pénalité $R = C$ afin d'éviter ces actions considérées comme mauvaises (ligne 5), car l'agent a la possibilité, via l'espace d'actions de chaque état, de rejeter directement la demande de partage en utilisant l'action a_0 ou de choisir une autre action pouvant offrir la quantité demandée de ressources. Ainsi, si l'action offre la quantité demandée de ressources, la demande de partage sera acceptée et l'environnement passera à l'état suivant (ligne 7). Ensuite, pendant le nouvel état, nous vérifions que chaque rejet dans la tranche "k" qui a offert des ressources ($w_{x,k} > 0$) n'est pas une conséquence de ses ressources partagées par l'action a_x (lignes 13-14). Par conséquent, si tel est le cas, la récompense positive R est donnée selon l'équation (3.6).

$$R = \alpha_k \times p_j \times R_{\max} \quad (3.6)$$

Cette récompense dépend de la priorité du type de demande accepté et de la tranche surchargée. Sinon, la récompense sera une pénalité $R = C$ (ligne 15), afin d'éviter que toute dégradation des performances des tranches sous-utilisées ne se produise en raison du partage. Nous pouvons résumer la fonction de récompense comme s'est exprimée par l'équation (3.7). L'algorithme (3) illustre la fonction de récompense.

$$R = \begin{cases} 0 & \text{If } a_x = a_0 \\ R_{\max} \cdot p_j \cdot \alpha_k & \text{If } (M_k = x_i \text{ And No rejection is caused in the lender slices by the action } a_x) \\ C & \text{Otherwise} \end{cases} \quad (3.7)$$

3.5.3 Réduire l'espace d'actions total

Dans cette partie, nous désignons l'ensemble des actions possibles par A , où la taille de A (le nombre total d'actions possibles) est N^k . Cela indique que chaque action correspond à une combinaison de N pourcentages possibles qu'une tranche peut proposer en partage, K étant le nombre de tranches dans le réseau. L'objectif principal de la réduction de l'espace d'actions est d'éliminer les décisions de partage irréalisables. Par conséquent, cette approche peut améliorer la probabilité de prendre la décision de partage optimale. La méthode proposée pour la réduction de l'espace d'actions se compose de deux parties :

a) Classification des actions préliminaires

Dans cette étape, nous divisons l'ensemble de l'espace d'actions en des sous-espaces d'actions distincts pour chaque tranche. Nous évaluons si une action peut répondre à la demande de ressources d'une tranche. Si c'est le cas, nous incluons cette action dans l'espace d'actions correspondant à la tranche. Dans la première partie, nous commençons par calculer le nombre de ressources fournies par chaque action dans l'espace d'actions global comme s'est exprimée par l'équation (3.8).

$$x_{\text{offered}} = \sum_{k=1}^K B_k \cdot w_{i,k} \quad (3.8)$$

Puis dans une deuxième partie, nous évaluons l'adéquation de chaque action créée à être ajoutée à l'espace d'actions d'une tranche. Pour ce faire, nous considérons deux types de tranches :

1. **Le premier type** : Lorsqu'une tranche atteint le nombre maximum de connexions $N_{\max}$ qu'il peut accepter simultanément et a zéro ressource restante comme montre l'équation (3.9).

$$B_k - \sum_{n=1}^{N_{\max}} b_k = 0 \quad (3.9)$$

où :

Algorithm 3 Pseudo-code of Reward Function

Require:

- S_n : the actual state of the environment
- a_x : the action received from the agent
- $R_{x,n}$: the generated reward from the action a_x for state n
- B_k : the total amount of resources dedicated to slice k
- M : the amount of resources requested
- U_k : the actual amount of used resources for slice k

```

1: if  $a_x = a_0$  then
2:   {In case of direct rejection ,the environment stays in the state  $S_n$  with  $R_{x,n} = 0$  }
3: else
4:   {In case  $a_x \neq a_0$ , calculate  $x_i$  through equation 3.5}
5:   if  $x_i < M$  then
6:      $R_{x,n} = C$  {The offered resources are less than the requested amount, the action
       is interpreted as rejection. The environment stays in state  $S_n$ }
7:   else
8:      $S_n = S_{n+1}$  {Move to the new state by updating the percentage of shared resources
        $r_k$  for each slice and other values  $b, m, n$ }
9:     Calculate  $R_{x,n}$  according to equation 3.6 {Initialize the reward at the start of the
       new state}
10:    while new state do
11:      {Observe the new state}
12:      for each slice  $k$  with  $r_k > 0$  do
13:        {For each slice sharing some of its resources}
14:        if admission_control_policy_k( $B_k - U_k - (B_k \cdot r_k)$ ) rejects a connection
          request in slice  $k$  then
15:          if admission_control_policy_k( $B_k - U_k - (B_k \cdot (r_k - w_{x,k}))$ ) could accept
            this connection request in slice  $k$  then
16:             $R_{x,n} = C$  {The  $k$ -th lender slice connection request was rejected due to
              the resources  $w_{x,k}$  offered by the agent in action  $a_x$ }
17:          end if
18:        end if
19:      end for
20:    end while
21:  end if
22: end if

```

$$N_{\max} = \frac{B_k}{b_k}, \quad N_{\max} \in \mathbb{N} \quad (3.10)$$

b_k représente le débit de données dédié à chaque connexion acceptée dans la tranche k .

Dans ce scénario, lorsqu'une tranche "k" est surchargée, elle n'a plus de ressources à s'offrir à elle-même. Par conséquent, nous n'incluons dans l'espace d'actions de la tranche "k" que les actions qui suggèrent que la tranche surchargée s'offre un pour-

centage de ressources égal à zéro, c'est-à-dire où $w_{x,k} = 0\%$, en utilisant l'équation 3.11 (voir la figure 3.5).

$$A_k = \{a_x \in A_k \mid w_{x,k} = 0 \wedge x_{\text{offered}} = x_{\text{requested}}\} \quad (3.11)$$

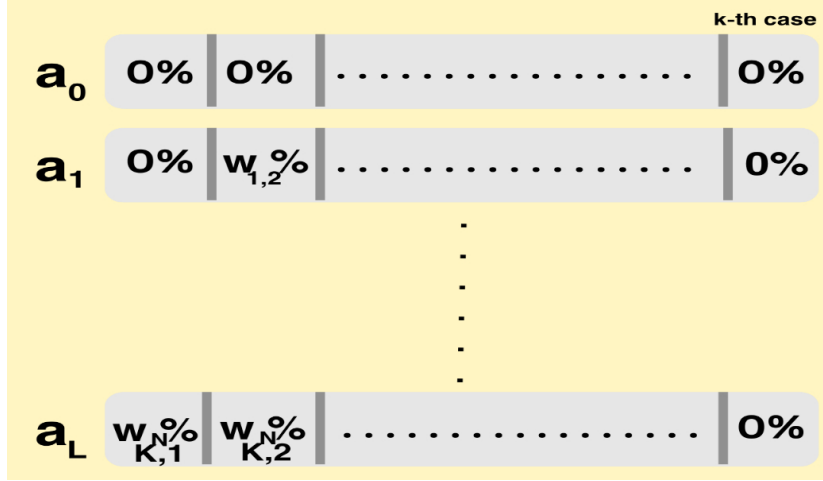


FIGURE 3.5 – Espace des actions préliminaires pour la $k^{\text{ième}}$ tranche (type 1)

2. **Le deuxième type :** Lorsqu'une tranche atteint le nombre maximum de connexions qu'elle peut les accepter simultanément, mais elle dispose encore de quelques ressources disponibles, bien qu'insuffisantes pour accepter une nouvelle connexion (voir l'équation 3.12).

$$B_k - \sum_{n=1}^{N_{\max}} b_k > 0 \quad (3.12)$$

Dans ce scénario, afin d'utiliser les ressources restantes disponibles dans la tranche surchargée et d'éviter leur gaspillage, nous incluons non seulement les actions proposant que la tranche surchargée s'attribue 0 % de ressources, mais également les actions suggérant qu'elle s'attribue toutes les ressources disponibles. Ces ressources disponibles représentent un pourcentage de toutes les ressources dédiées à cette tranche, comme montre la figure 3.6. Tout d'abord, nous devons calculer le nombre de ressources disponibles selon l'équation (3.13).

$$\text{Res}_{\text{available}} = B_k - \sum_{n=1}^{N_{\max}} b_k \quad (3.13)$$

Après cela, nous pouvons calculer le pourcentage de ressources disponibles en utilisant l'équation (3.14).

$$\text{Per}_{\text{available}} = \frac{\text{Res}_{\text{available}} \times 100}{B_k} \quad (3.14)$$

Ensuite, nous utilisons l'équation 3.15 pour gérer la classification.

$$A_k = \{a_x \in A_k \mid (w_{x,k} = 0 \vee w_{x,k} = \text{Per}_{\text{available}}) \wedge (x_{\text{offered}} = x_{\text{requested}})\} \quad (3.15)$$

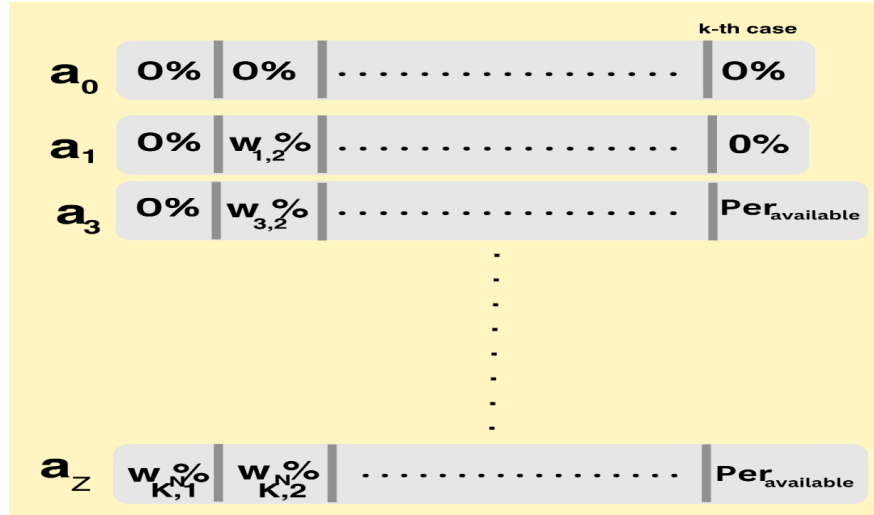


FIGURE 3.6 – Espace des actions préliminaires pour la $k^{\text{ième}}$ tranche (type 2)

Il est important de noter que ce processus est exécuté une seule fois, uniquement avant le début du processus d'entraînement de l'agent. L'algorithme 4 illustre la fonction de classification préliminaire des actions.

Algorithm 4 Pseudo-code of Preliminary Action Classification

Require:

- B_k : Total amount of resource dedicated to the slice k
- A : The set of all actions
- A_i : The set of all possible actions for the overloaded slice i at state S_0 , initialized with action a_0
- M_i : Required resources

```

1: for each  $a_x \in A$  do
2:   Sum = 0
3:   for  $i \leftarrow 1$  to  $K$  do
4:     Sum +=  $a_x[i] \cdot B_k$ 
5:   end for
6:   for  $i \leftarrow 1$  to  $K$  do
7:     if Sum ==  $M_i$  then
8:       if (slice  $i$  is type 1 and  $a_x[i] = 0$ ) or (slice  $i$  is type 2 and ( $a_x[i] = 0$  or
           $a_x[i] = \text{Per}_{\text{available}}$ )) then
9:         Add action  $a_x$  to the action space  $A_i$ 
10:      end if
11:    end if
12:  end for
13: end for

```

b) Élimination des actions irréalisables

Dans cette étape, nous visons à réduire l'espace d'actions en supprimant les décisions de partage qui ne sont pas faisables, c'est-à-dire que toutes les actions qui ne peuvent pas être appliquées dans chaque état. La méthode proposée pour l'élimination des actions se compose des étapes suivantes :

1. Considérons $A_k = \{a_1, \dots, a_n\}$ comme l'espace d'actions de la $k^{ième}$ tranche, qui représente la tranche surchargée de laquelle nous avons reçu la demande de partage.
2. Nous évaluons la faisabilité d'une décision de partage ($a_x \in A_k$) en vérifiant si $w_{x,k}$ dépasse le seuil de $(100 - r_k)$ (voir l'équation 3.16).

$$A_k = \{a_x \in A_k \mid w_{x,k} \leq 100 - r_k\} \quad (3.16)$$

3. Ensuite, si la tranche " k " de type 2 est la tranche surchargée demandant le partage et que son $r_k > 0$, indiquant qu'elle a déjà partagé la ressource disponible, alors nous éliminons toutes les actions de son espace d'actions qui proposent un $w_{x,k}$ supérieur à 0%. Sinon, nous éliminons les actions avec $w_{x,k} = 0$ comme s'est illustré dans l'équation (3.17).

$$A_k = \begin{cases} \{a_x \in A_k \mid w_{x,k} = 0\}, & \text{if } r_k > 0 \\ \{a_x \in A_k \mid w_{x,k} = \text{Per}_{\text{available}}\}, & \text{otherwise} \end{cases} \quad (3.17)$$

Il convient de mentionner que cette partie se répète à chaque fois que nous recevons une demande de partage. L'algorithme (5) illustre la fonction d'élimination des actions irréalisables.

Algorithm 5 Pseudo-code of Infeasible Action Elimination

Require:

- S : Vector of the actual state of the environment.
- A_k : The set of possible actions

```

1: for each  $a_x \in A_k$  do
2:   for  $i \leftarrow 1$  to  $K$  do
3:     if  $a_x[i] > (100 - S[i])$  then
4:       Delete  $a_x$  from the action space of state  $A_k$ 
5:     end if
6:     if the  $k$ -th slice is type 2 then
7:       if  $r_k > 0$  and  $a_x[i] > 0$  then
8:         Delete  $a_x$  from the action space of state  $A_k$ 
9:       else if  $r_k = 0$  and  $a_x[i] = 0$  then
10:        Delete  $a_x$  from the action space of state  $A_k$ 
11:       end if
12:     end if
13:   end for
14: end for

```

Algorithme 6 montre comment l'algorithme 3, l'algorithme 4, et l'algorithme 5 contribuent au processus d'entraînement de l'agent Deep Q-Learning.

Algorithm 6 Deep Q-Learning Training Algorithm

Require:

- Initialize the evaluation network Q with random weights $\theta_0 = \theta_{\text{Random}}$, then initialize the target network Q_{target} with weights $\theta^- = \theta_0$.
- Initialize the replay memory dataset D with a size of N .
- $\{A_1, \dots, A_k\}$ Initialize all the slices' action spaces through Algorithm 4

```

1: for episode  $\leftarrow 1$  to  $M$  do
2:   Initialize the index  $t = 0$ ;
3:   Reset environment and obtain initial state  $s_0$ ;
4:   if mod(episode,  $C$ ) == 0 then
5:     the agent assigns weights  $\theta$  of the evaluation network  $Q$  to the weights of the
       target network  $Q_{\text{target}}$  as  $\theta^- = \theta$ ;
6:   end if
7:   while  $t < t_{\text{max}}$  do
8:     Elimination of the unfeasible actions from the action space  $A_k$  for the actual state
        $S_n$ , resulting in  $A_k, s$  through Algorithm 5
9:      $\rho$  is initialized to a random number between 0 and 1;
10:    Based on greedy policy, calculate  $\epsilon(t)$ ;
11:    if  $\rho < \epsilon(t)$  then
12:      a random action;
13:    else
14:       $a_t = \max_{a \in A} Q(s_{t+1}, a, \theta)$ 
15:    end if
16:    Execute action and observe reward and state according to Algorithm 3
17:    The agent stores the episode experience  $e_t = \langle s_t, a_t, s_{t+1}, R(s_t, a_t) \rangle$  into  $D$ ;
18:    Sample random minibatch of transitions  $(s, a, s', r)$  from  $D$ ;
19:    if state is terminal then
20:       $Q_{\text{target}}(s_t, a_t) = R(s_t, a_t)$ ;
21:    else
22:       $Q_{\text{target}}(s_t, a_t) = R(s_t, a_t) + \gamma \max_{a'} Q_{\text{target}}(s_{t+1}, a')$ ;
23:    end if
24:    Calculate the loss function;
25:    The agent updates the weights  $\theta_t$  for the evaluation network using a gradient-
       based approach;
26:     $t = t + 1$ ;
27:  end while
28: end for

```

3.5.4 Scénario de partage de ressources

Pour illustrer toutes les étapes du mécanisme de partage de ressources proposé et montrer comment celui-ci est déclenché, un exemple d'un scénario spécifique est présenté ci-dessous et résumé dans la figure 3.7.

Considérons une zone urbaine où un environnement de découpage réseau V2X 5G est déployé avec plusieurs cellules, chacune prenant en charge les quatre tranches de réseau V2X. Chaque tranche dispose de ses ressources dédiées. Dans l'une des cellules, à un

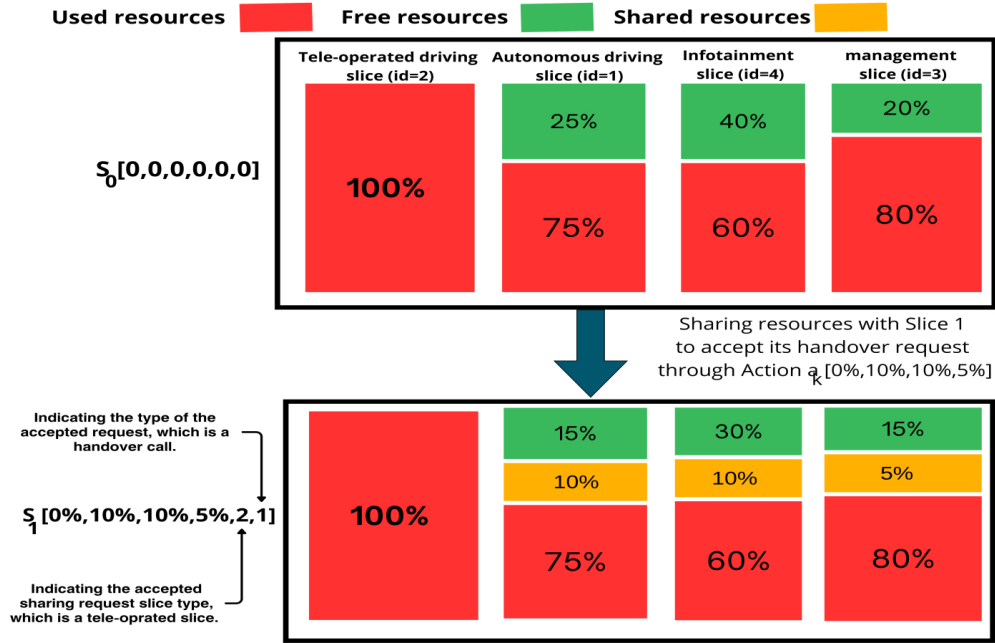


FIGURE 3.7 – Illustration d'un exemple scenario

certain taux d'arrivée des véhicules, la tranche télé-opérée devient surchargée, c'est-à-dire que ses ressources dédiées sont entièrement utilisées et qu'elle ne peut plus accepter de nouvelles demandes de connexion. En conséquence, sa politique de contrôle d'admission rejette les nouvelles demandes. Pendant ce temps, les trois autres tranches V2X restent sous-utilisées. Par exemple, dans notre scénario, la tranche 1 est sous-utilisée à hauteur de 25% de ses ressources dédiées, la tranche 3 à hauteur de 20%, et la tranche 4 à hauteur de 40%. Sans un mécanisme efficace de partage de ressources, le réseau fera face à une utilisation inefficace des ressources, car la tranche télé-opérée rejettera les demandes de transfert d'appels (handover) et les nouvelles demandes d'appels en raison de la surcharge. Dans cette situation, supposons qu'une nouvelle demande de handover soit reçue par la tranche télé-opérée, les étapes suivantes se dérouleront :

1. Le contrôle d'admission de la tranche rejette la demande en raison de la surcharge.
2. Le mécanisme de partage de ressources est déclenché en envoyant une demande de partage à l'agent DQL.
3. L'espace d'actions de la tranche télé-opérée est préparé en éliminant les actions non réalisables à l'aide de l'algorithme 5.
4. Notre agent DQL entraîné sélectionne une décision de partage de ressources appropriée (par exemple, $a_k = [0\%, 10\%, 10\%, 5\%]$) dans l'espace d'actions de la tranche télé-opérée. Cela signifie que nous acceptons la demande de transfert en empruntant 10% des ressources dédiées de la tranche 2, 10% des ressources de la tranche 3 et 5% des ressources de la tranche 4.
5. Puisque notre environnement était initialement dans l'état S_0 , en acceptant la demande de transfert grâce au processus de partage, l'environnement passe à l'état

suivant $S_1 = [0\%, 10\%, 10\%, 5\%, 2, 1]$. La valeur "2" indique que la demande de partage acceptée provient de la tranche télé-opérée, et la valeur "1" indique qu'il s'agit d'une demande de handover.

3.6 Évaluation des performances

Dans cette section, nous analysons les performances de notre contribution et les comparons à celles d'autres solutions en fonction de plusieurs paramètres.

3.6.1 Environnement du travail

Nous évaluons les performances de notre proposition en implémentant l'environnement de découpage V2X sur le simulateur OMNET++ 6.0 [115] en utilisant Simu5G [116] pour les communications réseau 5G et le framework SUMO [117] pour la mobilité des véhicules sur les routes. La taille de la carte de la zone de simulation est d'environ $3 \times 2,5 km^2$, couverte par trois gNodeB 5G. La portée de transmission radio de chaque gNodeB est de 1 km, et la capacité de ressources dédiée aux quatre tranches V2X dans chaque cellule est de 150 Mbps. Nous utilisons la bibliothèque open source Keras pour construire notre modèle d'apprentissage profond par renforcement Q-Learning. Le simulateur OMNET++ (C++) et la bibliothèque Keras (Python) sont intégrés via des fichiers textes pour créer un mécanisme de communication entre eux. Le tableau 3.1 montre les valeurs des paramètres utilisées dans la simulation.

TABLE 3.1 – Paramètres du réseau

Paramètres	Valeurs		
Slice k	Slice capacity B_k	Required data rate	Priority α_k
Autonomous driving slice	0.25*B	10 Mbps	$\alpha_1 = 0.5$
Tele-operated driving slice	0.4*B	20 Mbps	$\alpha_2 = 1$
Vehicle remote diagnostics and management slice	0.1*B	1 Mbps	///
Vehicular infotainment slice	0.25*B	5 Mbps	///
R_{max}	10		
C	$-1 \times R_{max}$		
Percentage of values the slice can offer	0%, 5%, 10%, 15%,, 40%		
p_h	1		
p_n	0.6		

3.6.2 Description du scénario

Le taux d'arrivée des véhicules sera varié lors des simulations, et le temps passé par chaque véhicule dans chaque cellule dépend de sa vitesse, qui sera en moyenne de 40 km/h . Il est important de noter qu'un véhicule peut avoir plusieurs flux de connexion admis sur différentes tranches en même temps. Étant donné que notre travail se concentre sur l'environnement de découpage V2X, les tranches de conduite télé-opérée et de conduite autonome auront la priorité la plus élevée parmi les tranches V2X. Ainsi, le partage des ressources n'est déclenché que lorsque ces deux tranches sont surchargées. De plus, la terminaison forcée d'un appel en cours est plus gênante qu'un appel bloqué au début. Par conséquent, la priorité est donnée aux demandes de handover par rapport aux nouvelles demandes d'appels.

Le nombre d'agents DQL sera égal au nombre de cellules, ce qui signifie que chaque cellule dispose de son propre agent dédié, responsable de prendre les décisions de partage. Il est important de noter que nous utilisons des agents DQL déjà entraînés pendant la simulation, ce qui signifie qu'au cours de la phase d'entraînement, l'agent a déjà exploré l'environnement pour apprendre à prendre les actions les plus rentables. Le processus d'apprentissage a été arrêté lorsque l'agent a commencé à fournir une récompense relativement constante. Des détails supplémentaires sur les paramètres de l'agent sont fournis dans le tableau 3.2 .

TABLE 3.2 – Paramètres de formation des agents

Paramètres	Valeurs
The policy	DNNs
Activation function for all layers	Rectified Linear Unit (ReLU)
Batch size	128 samples
Target Net. update period	30
Discount factor γ	0.8
Initial exploration rate	0.9
Final exploration rate	0.01

3.6.3 Approches d'analyse comparative

Pour évaluer les performances de la contribution proposée, nous comparerons notre modèle de Deep Q-learning aux autres solutions existantes :

- **Politique de découpage strict** : Cette stratégie n'autorise aucun partage avec les tranches surchargées, ce qui signifie une isolation totale des tranches. Toutes les demandes de partage seront rejetées dans ce mécanisme de partage de ressources.
- La solution de **négociation de Nash** pour le jeu I, désignée par *NBS1*, est proposée dans [51]. Cette solution est basée sur un jeu de négociation, qui est une forme spéciale de jeux coopératifs. Les tranches peuvent participer à ce jeu de négociation en offrant des ressources aux tranches surchargées tout en respectant certaines conditions.

- La politique **SoftSLICE**, où les ressources libres sont allouées des tranches de réseau sous-utilisées aux tranches surchargées, est basée sur une méthode de partage prédéfinie, proposée dans [59].

Il est à noter que la solution *NBS1* et la politique *SoftSLICE* acceptent les demandes de partage sur une base de premier arrivé, premier servi, jusqu'à ce que leurs seuils soient atteints.

3.6.4 Métriques

Pour évaluer notre proposition, nous avons pris en compte les paramètres suivants :

1) Taux d'utilisation des ressources : Cette métrique nous permet d'évaluer l'efficacité du schéma de gestion des ressources pour éviter une utilisation inefficace des ressources. Il est important de mentionner que la métrique d'utilisation des ressources est obtenue comme la valeur moyenne de l'utilisation des ressources de toutes les tranches du réseau et elle est évaluée selon l'équation (3.18).

$$\text{Resource}_{\text{utilisation}} = \frac{\sum_{c=1}^C \sum_{k=1}^K U_{c,k}}{\sum_{c=1}^C \sum_{k=1}^K B_{c,k}} \quad (3.18)$$

où :

- " C " : est le nombre total de cellules.
- " K " : est le nombre total de tranches dans chaque cellule.
- " $U_{c,k}$ " : est le nombre de ressources utilisées par la tranche " k " dans la cellule " c ".
- " $B_{c,k}$ " : est la capacité totale de ressources dédiées à la tranche " k " dans la cellule " c ".

L'objectif de notre contribution est de réduire les échecs des nouveaux appels et des demandes de transfert dans les tranches surchargées tout en maintenant une isolation élevée entre les tranches. Pour ce faire, et afin d'illustrer l'efficacité de notre solution par rapport aux autres solutions proposées en termes d'isolation, nous calculons, dans chaque cellule, la probabilité de blocage des nouveaux appels ainsi que la probabilité d'abandon des appels en transfert, séparément pour les tranches surchargées et les tranches prêteuses. Le résultat final est obtenu en prenant la moyenne de toutes les cellules du réseau.

Dans le scénario étudié, nous avons considéré que les tranches surchargées sont celles qui ne sont pas en mesure de répondre à de nouveaux appels ou à des demandes de transfert avec leurs ressources dédiées, et qui ont la permission de demander un partage de ressources aux autres tranches de la cellule, appelées tranches prêteuses.

2) Probabilité d'abandon du transfert dans une tranche : C'est une mesure qui évalue la fréquence à laquelle les appels en cours de transfert sont interrompus ou abandonnés lorsqu'ils passent d'une tranche d'une cellule au même type de tranche d'une autre cellule avant leur finalisation. Cette probabilité est cruciale pour évaluer la performance du système de gestion des ressources dans des scénarios de mobilité élevée. Cette métrique est évaluée selon l'équation (3.19).

$$\text{probabilité}_{\text{abandon}} = \frac{N_{\text{abandons}}}{N_{\text{transferts}}} \quad (3.19)$$

où :

- N_{abandons} : est le nombre d'appels en cours de transfert qui ont été interrompus ou abandonnés dans une tranche.
- $N_{\text{transferts}}$: est le nombre total d'appels en cours de transfert reçus dans une tranche.

Probabilité d'abandon du transfert dans les tranches surchargées :

La probabilité d'abandon du transfert dans les tranches surchargées est calculée comme c'est exprimée par l'équation (3.20).

$$\text{Probabilité}_{\text{abandon, surchargées}} = \frac{\sum_{c=1}^C \sum_{k=1}^{K_{\text{surchargées}}} \text{Probabilité}_{\text{abandon},k}}{\sum_{c=1}^C K_{\text{surchargée},c}} \quad (3.20)$$

où :

- " C " : est le nombre de cellules dans le système.
- " $K_{\text{surchargée},c}$ " : est le nombre de tranches surchargées dans la cellule " c ".
- " $\text{Probabilité}_{\text{abandon},k}$ " : est la probabilité d'abandon du transfert dans la tranche " k ".

Probabilité d'abandon du transfert dans les tranches prêteuses :

La probabilité d'abandon du transfert dans les tranches prêteuses est évaluée selon l'équation (3.21).

$$\text{Probabilité}_{\text{abandon, prêteuses}} = \frac{\sum_{c=1}^C \sum_{k=1}^{K_{\text{prêteuses},c}} \text{Probabilité}_{\text{abandon},k}}{\sum_{c=1}^C K_{\text{prêteuses},c}} \quad (3.21)$$

où :

- " C " : est le nombre de cellules dans le système.
- " $K_{\text{prêteuses},c}$ " : est le nombre de tranches prêteuses dans la cellule " c ".
- " $\text{Probabilité}_{\text{abandon},k}$ " : est la probabilité d'abandon du transfert dans la tranche " k ".

3) Probabilité de blocage des nouveaux appels dans une tranche :

C'est une mesure de la fréquence à laquelle les nouveaux appels dans une tranche ne peuvent pas être acceptés en raison de l'absence de ressources suffisantes dans cette tranche. Cette métrique est évaluée selon l'équation (3.22).

$$\text{probabilité}_{\text{blocage}} = \frac{N_{\text{blocage}}}{N_{\text{nouveaux appels}}} \quad (3.22)$$

Où :

- " N_{abandons} " : est le nombre de nouveaux appels qui ont été bloqués ou rejetés dans une tranche.
- " $N_{\text{nouveaux appels}}$ " : est le nombre total de nouveaux appels arrivés dans la tranche.

Probabilité de blocage des nouveaux appels dans les tranches surchargées :

La probabilité de blocage des nouveaux appels dans les tranches surchargées est exprimée selon l'équation (3.23).

$$\text{Probabilité}_{\text{blocage, surchargées}} = \frac{\sum_{c=1}^C \sum_{k=1}^{K_{\text{surchargées}}} \text{Probabilité}_{\text{blocage},k}}{\sum_{c=1}^C K_{\text{surchargée},c}} \quad (3.23)$$

où :

- " C " : est le nombre de cellules dans le système.
- " $K_{\text{surchargée},c}$ " : est le nombre de tranches surchargées dans la cellule " c ".
- " $\text{Probabilité}_{\text{blocage},k}$ " : est la probabilité de blocage des nouveaux appels dans la tranche " k ".

Probabilité de blocage des nouveaux appels dans les tranches prêteuses :

La probabilité de blocage des nouveaux appels dans les tranches prêteuses est calculée à l'aide de l'équation (3.24).

$$\text{Probabilité}_{\text{blocage, prêteuses}} = \frac{\sum_{c=1}^C \sum_{k=1}^{K_{\text{prêteuses}}} \text{Probabilité}_{\text{blocage},k}}{\sum_{c=1}^C K_{\text{prêteuses},c}} \quad (3.24)$$

où :

- " C " : est le nombre de cellules dans le système.
- " $K_{\text{prêteuses},c}$ " : est le nombre de tranches prêteuses dans la cellule " c ".
- " $\text{Probabilité}_{\text{blocage},k}$ " : est la probabilité de blocage des nouveaux appels dans la tranche " k ".

3.6.5 Taux d'utilisation des ressources

Nous commençons l'évaluation des performances du réseau avec un indicateur clé : l'utilisation des ressources. Cette métrique nous permet d'évaluer l'efficacité du mécanisme de gestion des ressources dans la prévention de l'utilisation inefficace des ressources. Il est important de noter que la métrique d'utilisation des ressources est obtenue en calculant la valeur moyenne de l'utilisation des ressources pour toutes les tranches du réseau.

La figure 3.8 montre que notre solution de partage basée sur DQL atteint un taux d'utilisation des ressources de 98,05 % avec un faible taux d'arrivée de véhicules de 45 véhicules par minute, augmentant jusqu'à près de 99 % à des taux d'arrivée plus élevés. Il est clairement observé qu'à mesure que le taux d'arrivée de véhicules augmente, le taux d'utilisation des ressources du système s'améliore progressivement. En comparaison, la méthode de découpage strict n'atteint qu'environ 91,45 % à des taux d'arrivée faibles et 91,89 % à des taux plus élevés, démontrant une amélioration allant jusqu'à 6,96 % avec notre méthode. Cela souligne l'efficacité de notre solution dans l'optimisation de l'utilisation des ressources via le processus de partage. Le schéma de découpage strict fournit la plus faible utilisation des ressources en raison de sa stratégie d'isolation complète, qui empêche le partage des ressources avec les demandes qui ont été rejetées ou bloquées par leurs tranches.

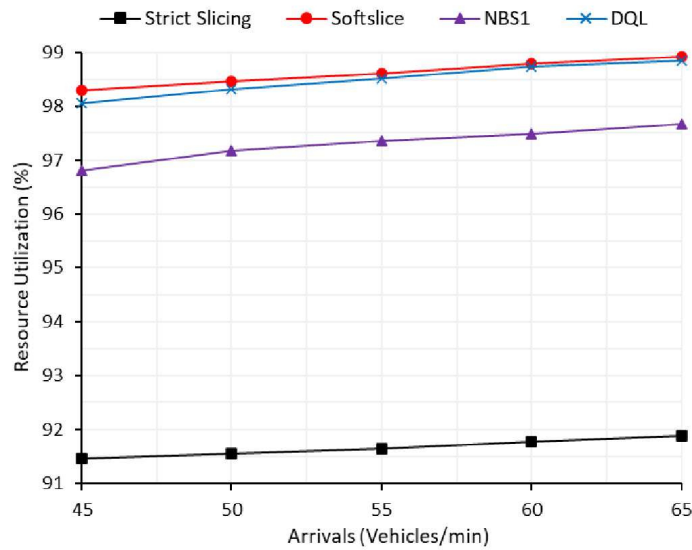


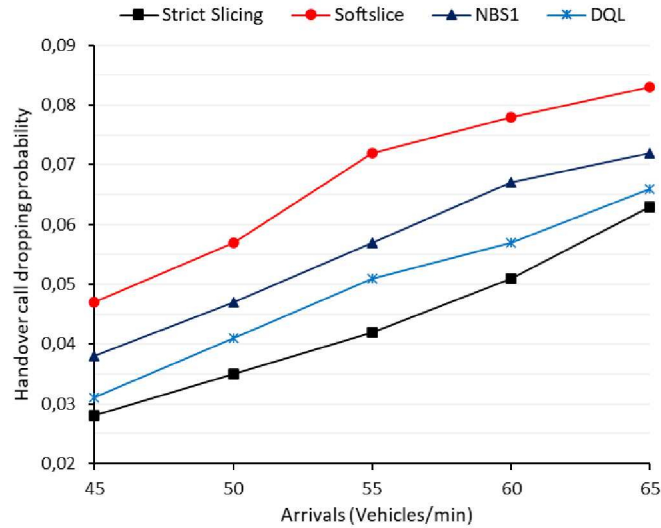
FIGURE 3.8 – Taux d'utilisation des ressources

En outre, les résultats montrent que notre solution proposée est plus efficace que la solution NBS1 pour atteindre un taux d'utilisation plus élevée des ressources. Cela s'explique par le fait que notre mécanisme DQL permet un échange de ressources entre les tranches, aussi bien pour les nouveaux appels que pour les demandes de transfert, lorsque leurs tranches sont surchargées. En revanche, la solution NBS1 ne permet le partage de ressources que pour les demandes de transfert. En comparant notre solution DQL à SoftSlice, la solution SoftSlice atteint des niveaux élevés d'utilisation des ressources à un faible taux d'arrivée car elle ne tient pas compte de l'isolation entre les tranches. SoftSlice ne fait que vérifier les ressources disponibles dans d'autres tranches de la même cellule au moment où la demande de partage est reçue, contrairement à notre approche, qui prend

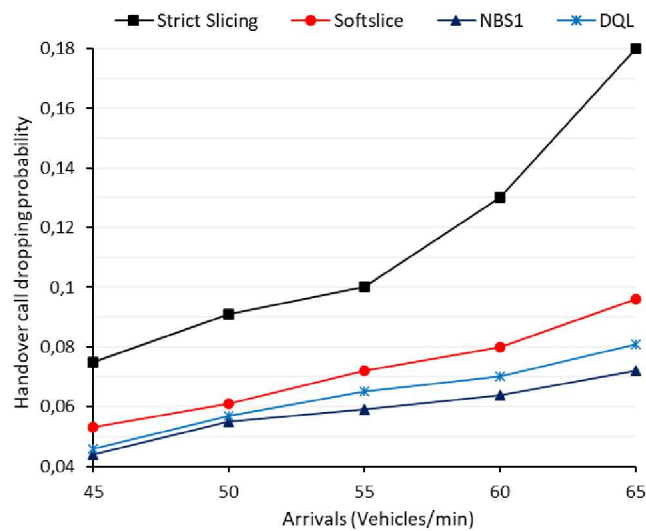
en compte l'isolation des tranches. En outre, à mesure que le taux d'arrivée augmente, la performance de notre méthode proposée se rapproche de celle du schéma SoftSlice. Cela s'explique par le fait qu'un taux d'arrivée plus élevé de véhicules entraîne plus de demandes de partage, ce qui permet à notre méthode d'accepter davantage de demandes sans dégrader la performance des tranches prêteuses.

3.6.6 Probabilité d'abandon du transfert d'appel

L'un des principaux objectifs de notre solution est de réduire l'échec des demandes de transfert d'appel dans le réseau tout en maintenant une forte isolation entre les tranches. Ainsi, afin de valider l'efficacité de notre approche proposée, nous évaluons séparément la probabilité d'abandon des appels en transfert dans les tranches surchargées et dans les tranches prêteuses. Pour cela, nous faisons varier le taux d'arrivée des véhicules.



(a) Probabilité d'abandon du transfert dans les tranches prêteuses



(b) Probabilité d'abandon du transfert dans les tranches surchargées

FIGURE 3.9 – Probabilité d'abandon du transfert

Les figures 3.9a et 3.9b montrent que la méthode de découpage strict entraîne une probabilité élevée de rejet des appels de transfert (handover) dans les tranches surchargées. Cela s'explique par le fait qu'avec un grand nombre de demandes de transfert, l'allocation statique peut ne pas être en mesure de répondre à toutes les demandes en raison de la quantité fixe de ressources allouées à chaque tranche. Cependant, on peut également en déduire que cette absence de partage des ressources garde une probabilité plus faible de rejet des appels de transfert dans les tranches prêtes.

Le schéma SoftSlice offre de meilleures performances que le découpage strict en termes de réduction de la probabilité de rejet des appels de transfert dans les tranches surchargées, car il permet le partage des ressources entre les tranches, même s'il ne donne pas la priorité pour les demandes de transfert. Toutefois, cet aspect de partage des ressources se traduit par un taux de rejet des demandes de transfert plus élevé par rapport au découpage statique, au NBS1 et à notre approche DQL. En revanche, notre solution proposée réduit la probabilité de rejet des appels de transfert dans les tranches surchargées tout en offrant la probabilité la plus faible de rejet d'appels dans les tranches prêtes, par rapport aux autres schémas permettant le partage des ressources. Cela garantit un niveau élevé d'isolation entre les tranches. Cette amélioration est due au fait que les agents d'apprentissage profond (agents DQL) ont appris, grâce aux pénalités imposées lors du processus d'entraînement, à éviter les actions de partage susceptibles qui entraîneraient une forte probabilité de rejet des transferts dans les tranches prêtes. De plus, en comparant SoftSlice et l'approche DQL proposée, cette dernière présente un taux de rejet des appels de transfert inférieur, car nous accordons une légère priorité aux appels de transfert. En effet, nous pouvons déduire que NBS1 obtient la probabilité de rejet la plus faible, ce qui s'explique par le fait que ce schéma ne déclenche le partage des ressources des tranches prêtes que pour les demandes de transfert.

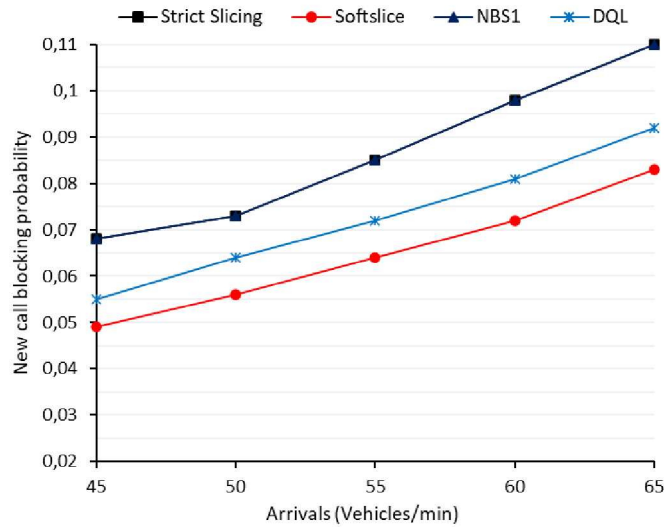
3.6.7 Probabilité de blocage des nouveaux appels

Un autre indicateur important pour évaluer notre méthode proposée est la probabilité de blocage des nouveaux appels dans les tranches surchargées ainsi que dans les tranches prêtes. Les figures 3.10a et 3.10b montrent que la solution SoftSlice présente la plus faible probabilité de blocage des nouveaux appels dans les tranches surchargées par rapport à toutes les autres solutions. Cela s'explique par le fait que notre solution, via l'agent DQL, priorise les demandes de transfert par rapport aux nouveaux appels, tandis que la solution NBS1 n'autorise pas le partage pour les nouveaux appels dans les tranches surchargées.

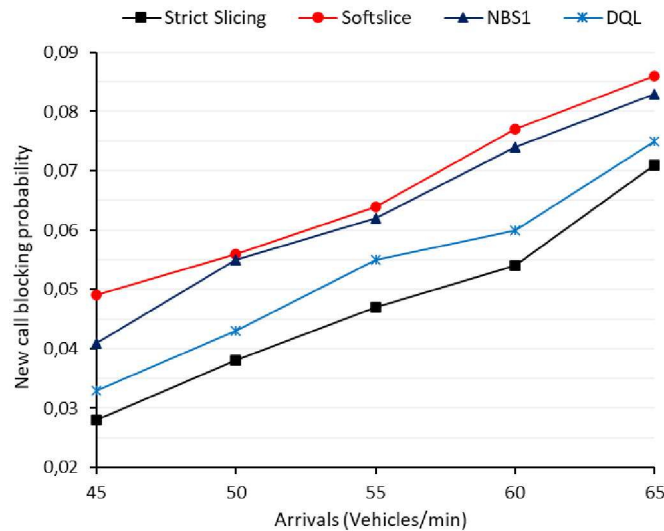
De plus, bien que notre approche améliore considérablement la performance du réseau en réduisant la probabilité de blocage des nouveaux appels dans les tranches surchargées, on observe une légère dégradation de cette probabilité dans les tranches prêtes. Contrairement à la solution SoftSlice qui affecte énormément l'isolation entre les tranches en augmentant fortement la probabilité de blocage des nouveaux appels dans les tranches prêtes.

En ce qui concerne la solution de découpage strict et la solution NBS1, on constate que toutes les deux entraînent la plus haute probabilité de blocage des nouveaux appels dans les tranches surchargées. En outre, la solution NBS1 provoque également une forte dégradation de la probabilité de blocage des nouveaux appels dans les tranches prêtes.

Cela peut s'expliquer par les ressources partagées avec les demandes de transfert provenant des tranches surchargées.



(a) Probabilité de blocage des nouveaux appels dans les tranches surchargées



(b) Probabilité de blocage des nouveaux appels dans les tranches prêtes

FIGURE 3.10 – Probabilité de blocage des nouveaux appels

3.6.8 Récompenses obtenues par différentes stratégies

La récompense est calculée comme la valeur moyenne obtenue au niveau de toutes les cellules du réseau. En d'autres termes, nous mesurons la récompense de chaque cellule, puis nous calculons le résultat final en faisant la moyenne des récompenses de toutes les cellules du réseau.

La figure 3.11 illustre les récompenses moyennes obtenues avec différentes stratégies. On peut observer que notre approche DQL surpasse les autres stratégies en termes de

récompense moyenne, grâce à sa capacité à optimiser le processus de prise de décision pour le partage des ressources et à maximiser les récompenses à long terme. Contrairement aux autres stratégies, qui manquent de prévoyance quant aux résultats futurs. Pour un taux d'arrivée faible de véhicules, parmi les stratégies alternatives, la politique SoftSlice obtient la meilleure récompense car elle permet le partage des ressources pour les demandes de transfert et les nouveaux appels. En revanche, la solution NBS1 permet uniquement le partage pour les demandes de "handover". Cependant, à mesure que le taux d'arrivée augmente, la politique SoftSlice devient sous-optimale du point de vue des récompenses à long terme, ce qui entraîne des récompenses moyennes inférieures à celles de la solution NBS1, conduisant finalement à une récompense négative. En effet, la plupart des demandes de partage acceptées dégradent les performances des tranches prêteuses, générant ainsi des récompenses négatives. D'un autre côté, la politique de découpage stricte donne toujours une récompense de 0, car elle ne permet aucun partage de ressources, ce qui l'empêche de gagner des récompenses positives ou négatives. De plus, lorsque l'on compare les différentes stratégies, nous pouvons constater qu'à mesure que le taux d'arrivée augmente, toutes les stratégies conduisent à des récompenses moyennes inférieures. Cela est dû au fait que moins de ressources sont disponibles pour le partage.

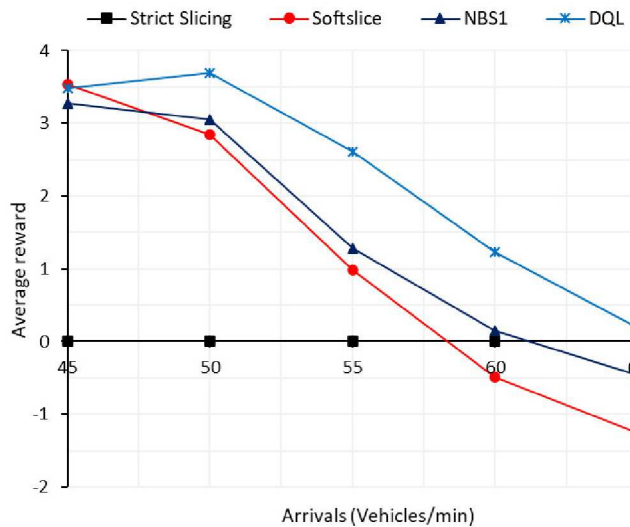


FIGURE 3.11 – Récompenses obtenues avec différentes approches

3.6.9 Impact des paramètres de priorités

Une caractéristique pertinente de l'approche d'apprentissage profond par renforcement proposée est sa flexibilité. Elle peut être configurée pour adopter différents comportements en ajustant correctement les valeurs des paramètres de priorité dans la fonction de récompense.

Dans cette évaluation, nous commençons par analyser l'impact du paramètre de priorité des types de demandes sur le pourcentage de nouveaux appels et de demandes de transfert parmi l'ensemble des demandes de partage acceptées, en fixant la priorité des demandes de transfert à 1 et en faisant varier la priorité des nouveaux appels entre 0 et 1. Ensuite, nous évaluons l'impact du paramètre de priorité des tranches sur le pourcentage

de demandes de partage acceptées provenant des deux tranches parmi toutes les demandes de partage acceptées. La tranche de conduite télé-opérée, dont la priorité sera fixée à 1, et la tranche de conduite autonome, dont la priorité variera entre 0 et 1. Le taux d'arrivée des véhicules est fixé à 55 véhicules/min.

La figure 3.12 montre qu'avec une priorité $p_n = 0$, les agents DQL rejettent toutes les demandes de partage pour les nouveaux appels. Cela s'explique par le fait qu'aucune récompense positive n'est générée en acceptant ces nouveaux appels. En conséquence, toutes les demandes de partage acceptées concernent les demandes de transfert. Comme on peut le voir sur cette figure, l'augmentation de p_n entraîne une augmentation du pourcentage de nouveaux appels acceptés dans le processus de partage, puisqu'une récompense positive est générée. Par exemple, pour $p_n = 0,6$, notre approche permet d'accepter 34,7% des nouveaux appels parmi toutes les demandes de partage acceptées. De plus, dans le cas où les priorités sont égales ($p_n = p_h = 1$), on observe que le pourcentage de nouveaux appels acceptés augmente pour représenter 56,2%, tandis que les demandes de transfert sont réduites à 43,8%.

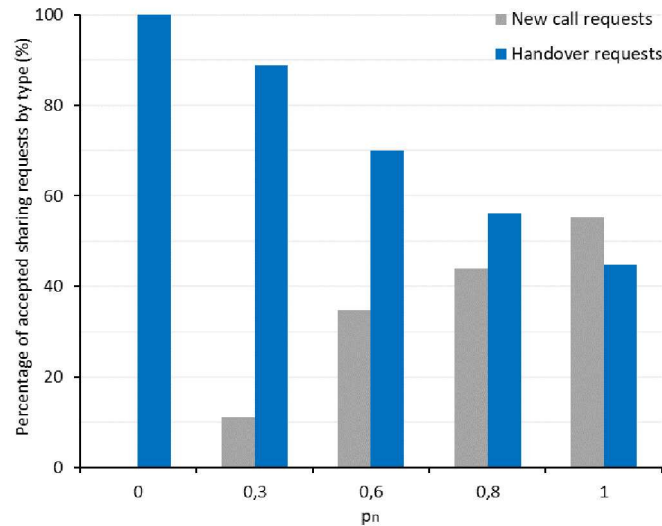


FIGURE 3.12 – Impact du paramètre “type priority”

La figure 3.13 montre qu'en faisant simplement varier le paramètre de priorité α_1 , la solution proposée permet d'obtenir un équilibre différent entre le nombre de demandes de partage acceptées provenant de la tranche de conduite télé-opérée et celles provenant de la tranche de conduite autonome. Cependant, avec $\alpha_1 = 0,5$, notre approche aboutit à ce que 42,1% de toutes les demandes de partage acceptées soient des connexions de la tranche de conduite télé-opérée. De plus, en augmentant le paramètre de priorité à $\alpha_1 = 1$, le pourcentage de demandes de partage acceptées provenant de la tranche de conduite télé-opérée diminue à 12,85%. Cela s'explique par le fait qu'une connexion de la tranche de conduite télé-opérée nécessite plus de ressources qu'une connexion de la tranche de conduite autonome. Cela confirme que la solution DQL proposée offre une grande flexibilité dans la gestion des ressources grâce au processus de partage entre les différents types de tranches et de demandes.

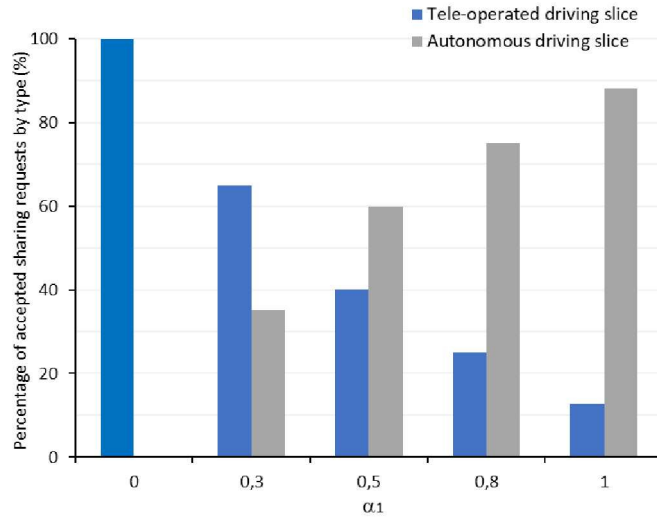


FIGURE 3.13 – Impact du paramètre “slice priority”

3.6.10 Évaluation de l’agent DQL durant l’apprentissage

Dans cette étude, nous illustrons les performances de convergence de notre processus d’entraînement.

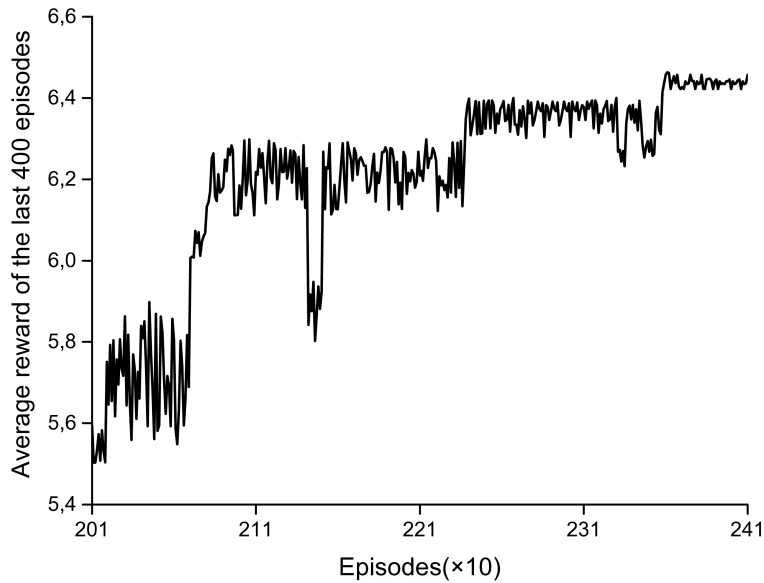


FIGURE 3.14 – Entraînement de l’agent DQL

La figure 3.14 montre la convergence de l’agent DQL en fonction des périodes d’apprentissage. Elle présente les récompenses moyennes par période sur les 400 dernières périodes du processus d’apprentissage. Il est clair que les récompenses moyennes s’améliorent à mesure que le nombre de périodes d’entraînement augmente, ce qui démontre l’efficacité de l’algorithme proposé. On observe que l’agent DQL converge à partir de la période numéro 2360. Il est important de noter que la convergence de l’algorithme n’est pas totalement fluide et comporte quelques fluctuations, principalement dues à la forte

mobilité dans l'environnement véhiculaire. Sur la base des récompenses moyennes obtenues pendant l'entraînement, nous avons conclu que, pour garantir la convergence, l'agent doit être entraîné hors ligne pendant au moins 2400 périodes.

3.7 Conclusion

Dans ce chapitre, nous avons présenté notre solution pour la gestion des ressources dans l'environnement de découpage réseau V2X, en particulier un mécanisme de partage des ressources entre les tranches V2X. Nous avons modélisé le problème en utilisant un modèle de processus de décision markovien, et pour le résoudre, nous avons appliqué l'algorithme de deep Q-learning. L'objectif de l'agent est de prendre des décisions concernant le partage des ressources, soit en acceptant, soit en rejetant les demandes de partage. En cas d'acceptation, l'agent doit déterminer la contribution de chaque tranche afin de répondre aux besoins des demandes de connexion provenant des tranches surchargées. Avant le début de l'entraînement de l'agent, les actions doivent être classées dans différents espaces d'actions spécifiques aux tranches, qui répondent aux exigences en matière de ressources de chaque tranche. Ces espaces seront actualisés à chaque réception de demandes de partage, en fonction de la disponibilité des ressources des tranches prêteuses. L'agent DQL sélectionne des actions à partir de ces espaces et les applique pour s'entraîner à prendre des décisions qui évitent toute dégradation dans les tranches prêteuses. Cela signifie que toute dégradation est pénalisée pendant le processus d'entraînement.

De plus, l'analyse des performances démontre l'efficacité du mécanisme proposé en termes du taux d'utilisation des ressources, de la probabilité de blocage des nouveaux appels et la probabilité d'abandon des demandes de transfert, tout en garantissant une isolation élevée. À cela s'ajoute la flexibilité de l'agent DQL dans l'allocation des ressources, grâce aux paramètres de priorités des tranches et des types de demandes intégrés dans la fonction de récompense.

Chapitre 4

Mécanisme de dissemination des informations du trafic véhiculaire dans les VANETs

Chapitre 4

Mécanisme de dissémination des informations du trafic véhiculaire dans les VANETs

4.1 Introduction

Les réseaux ad hoc véhiculaires (VANETs) diffèrent des réseaux ad hoc mobiles "Mobile Adhoc NETworks (MANETs)" en raison de leurs caractéristiques distinctes, notamment leur mobilité restreinte, car les véhicules sont limités à des voies de circulation prédéfinies. En revanche, les noeuds dans les MANETs ne sont pas soumis à de telles limitations et peuvent se déplacer librement. Par conséquent, il est essentiel de développer des algorithmes de routage adaptés aux VANETs qui tiennent compte de ces attributs uniques. L'une des approches les plus adaptées à cet environnement est l'utilisation de protocoles de routage basés sur la position, qui sont très efficaces pour gérer les changements dynamiques de la topologie VANET.

Dans ce chapitre, nous présentons un protocole de routage basé sur une approche de transmission gloutonne, appelé Efficient Routing Greedy Forwarding (ERGF) [118]. Ce protocole de routage basé sur la position utilise des informations actualisées sur le trafic véhiculaire dans le processus de routage, ce qui lui permet de s'adapter efficacement aux changements dynamiques de la topologie des réseaux VANETs, aboutissant à une performance de routage améliorée. Plus précisément, notre contribution améliore certaines des limitations du protocole "Partial backwards routing protocol (PBRP)" [119].

Des diverses simulations du protocole ERGF en utilisant le simulateur OMNET++ ont été réalisées pour évaluer ses performances ainsi que sa résistance aux changements fréquents de topologie. Les résultats de l'évaluation ont montré que le protocole proposé offre des performances supérieures par rapport aux autres protocoles de routage existants.

4.2 Contexte

Le mécanisme de routage joue un rôle important dans les réseaux VANETs, car tous les services pris en charge, qu'ils soient monocast ou multicast, s'appuient sur des communications multi-sauts pour la transmission des données. La gestion du processus de routage

doit être effectuée avec un contrôle et une consommation de bande passante minimales. Les protocoles de routage basés sur la position exploitent des informations géographiques pour effectuer le routage. Le véhicule source utilise sa position géographique, plutôt qu'une adresse réseau, pour transmettre les données contenues dans les paquets. Chaque nœud est généralement équipé d'un module GPS permettant de localiser sa position en temps réel, et les paquets sont transmis vers leur destination à l'aide d'un mécanisme de routage multi-saut. Par rapport aux autres protocoles, les protocoles de routage géographique se distinguent par leur grande évolutivité et un temps de communication minimal, garantissant ainsi que le véhicule n'a pas encore changé de position. Ces caractéristiques les rendent plus stables et particulièrement adaptés aux VANETs.

Dans ce type de routage, chaque véhicule désigne un voisin comme saut-suivant selon une stratégie prédéfinie. Chaque véhicule dispose d'une table de voisinage contenant des informations sur tous ses voisins. De plus, un mécanisme d'échange de paquets "Hello" est utilisé entre les véhicules pour mettre à jour leurs tables de voisinage. Les véhicules étant limités à la circulation sur les routes, l'une des meilleures façons d'adapter ces types de protocoles à l'environnement véhiculaire est de prendre en compte les segments des routes présentant une excellente connectivité. Cela peut être réalisé en prenant en compte à la fois l'état du segment et la position des véhicules. De plus, l'état du réseau et la charge de trafic de données par segment doivent être pris en compte lors de la prise de décisions de routage afin d'éviter de sélectionner des itinéraires encombrés pour la transmission des paquets de données. À cette fin, le routage à travers un segment doit prendre en compte les informations sur le trafic et l'état du réseau, notamment la densité, la qualité des connexions entre véhicules, la répartition des véhicules et la charge de communication réseau par segment.

4.3 Travaux connexes

Plusieurs protocoles de routage ont été conçus pour s'adapter aux VANETs, voici quelques travaux notables sur les protocoles de routage pour les réseaux VANETs.

Le travail [120] présente un protocole de routage basé sur l'état des liens appelé "Link State Aware Geographic Opportunistic (LSGO)". Dans ce protocole, chaque véhicule calcule le nombre attendu de transmissions "Expected transmission count (ETX)" à l'aide de paquets périodiques. LSGO donne la priorité aux nœuds en fonction de la valeur ETX et de la distance jusqu'à la destination, le véhicule ayant la priorité la plus élevée transmettant le paquet directement à la destination. Pendant ce temps, le nœud ayant la priorité la plus basse attend un délai prédéfini. Si le délai expire et que le véhicule ayant la priorité la plus élevée n'a pas transmis le paquet, le véhicule ayant la priorité la plus basse est autorisé à l'envoyer. Bien que ce protocole soit considéré comme un protocole fiable en raison de cette stratégie de transmission, il entraîne un grand overhead de routage.

Un protocole de routage partiel vers l'arrière "PBRP" a été proposé dans [119], qui comprend trois stratégies : la dissémination des informations sur le trafic routier, l'algorithme de transfert partiel et la stratégie de récupération vers l'arrière. Ces stratégies fonctionnent en tandem pour fournir en continu des informations sur le trafic des véhicules, ce qui permet à l'algorithme de routage de résister efficacement aux changements dynamiques de la topologie du réseau.

Dans [121], une approche de routage a été proposée, appelée réseau backbone basé

sur la sélection de relais hybride pour la dissémination de données selon un schéma de routage multi-sauts dans les VANETs (BACKNET-HSR). Dans ce protocole, les règles de sélection de relais sont modifiées selon un nouveau modèle de propagation de messages pour éviter une distribution de relais chaotique et améliorer la couverture des informations. L'objectif est de créer une synergie entre les candidats relais potentiels pour établir un réseau fédérateur (backbone) à la volée qui fonctionne parallèlement à la topologie du réseau routier de manière distribuée. Cependant, la mise en œuvre du modèle de propagation de messages entre les véhicules peut affecter négativement la réception réussie des messages entre les véhicules.

Le protocole proposé dans [122] est un protocole de routage géographique basé sur les intersections pour les VANETs en milieu urbain, appelé "Greedy Traffic Aware Routing (GyTAR)". Dans GyTAR, la stratégie de sélection des intersections est déterminée par la densité de la route (le nombre de véhicules sur la route) et la distance curvemétrique jusqu'à la destination. GyTAR utilise des paquets de densité de cellules pour sélectionner les intersections potentielles. Cependant, le lancement de ces paquets n'est pas régulier dans le temps car il dépend du moment où un véhicule est sur le point de quitter sa route actuelle. Une version améliorée de GyTAR, appelée "Enhanced Greedy Traffic Aware Routing (EGyTAR)", a été introduite dans [123]. Dans cette version améliorée, la stratégie de sélection des intersections est basée sur le nombre de véhicules se déplaçant vers les intersections candidates.

Dans [124], le protocole de routage proposé pour les VANETs est basé sur un modèle d'apprentissage par renforcement profond (DRL). Dans ce modèle, l'unité routière "RSUs" suit les informations sur le trafic routier et le DRL prédit le mouvement du véhicule pour établir un chemin de routage optimal pour la transmission de paquets avec une capacité de transmission améliorée. Cependant, le modèle ne prend pas en compte la densité routière dans la prédiction, ce qui a entraîné une légère diminution du taux de livraison des paquets.

Dans [125], les auteurs ont proposé un protocole basé sur le comportement des liens dynamiques entre couches qui prend en compte des facteurs tels que le taux de perte, les interférences physiques et les frais généraux de contrôle d'accès au support. En utilisant ce modèle basé sur les liens, ils ont développé un protocole de routage géographique appelé GeoCAP pour les VANETs dans les milieux urbains. Ce protocole se concentre sur la sélection des liens les plus fiables et sur la réaction rapide aux changements de charge des nœuds et des conditions du réseau, ce qui conduit à des performances améliorées en termes de taux de livraison des paquets et de délai.

4.4 Contribution : Le protocole ERGF

Dans notre contribution, nous proposons un protocole de routage efficace pour les VANETs dans les environnements urbains, prenant en compte les états des routes et les intersections qui les relient, ce protocole est appelé ERGF. En nous appuyant sur les informations sur le trafic routier, nous avons amélioré le mécanisme de dissémination des données sur le trafic routier présenté dans le protocole PBRP [119] où ce mécanisme est chargé de fournir des informations actualisées sur la topologie dans les réseaux VANETs. Le protocole proposé comprend deux algorithmes principaux : l'algorithme de dissémination des informations sur le trafic routier et l'algorithme de transfert des données basé sur une l'approche de transfert gloutonne (greedy forwarding). En implémentant ces amélio-

rations, le protocole proposé exploite les informations sur le trafic routier en temps réel, offrant une nouvelle perspective sur les changements dynamiques de la topologie VANET. Cela se traduit par des performances améliorées par rapport au protocole PBRP.

4.4.1 Dissemination d'informations sur le trafic routier

Notre contribution se concentre sur la stratégie de dissemination des informations sur le trafic routier entre les véhicules comme illustré dans la figure 4.1, en tenant compte de la nature dynamique des VANETs, où chaque route conserve des informations de trafic mises à jour et les partage avec les routes adjacentes. Pour faciliter cela, la communication entre les véhicules de notre environnement est réalisée à l'aide d'un paquet appelé "Road Traffic Information Packet (RTIP)". Ce paquet est spécifiquement conçu pour gérer et transmettre des données de trafic routier, y compris des mesures telles que la densité de la route (nombre de véhicules), la connectivité routière et le délai de traversée des paquets pour un segment de route donné. Le paquet RTIP est initié par un véhicule appelé "RTIP Initiator Node (RIN)", qui est le véhicule le plus proche d'une intersection. Ces véhicules RIN sont chargés de mettre à jour le paquet RTIP et de le transmettre à tous les autres intersections voisines.

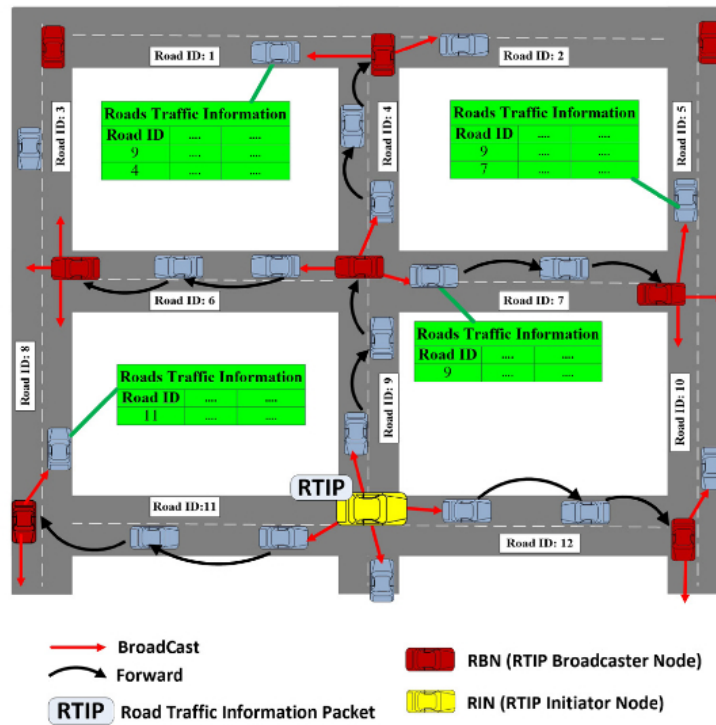


FIGURE 4.1 – Dissemination des informations sur le trafic routier

Dans notre approche, le véhicule RIN met à jour le paquet d'informations sur le trafic routier pour calculer la connectivité, plutôt que de le réinitialiser comme cela se fait dans le protocole PBRP. De plus, notre stratégie prend en compte la vitesse des véhicules en milieu urbain pour des mises à jour plus précises et plus efficaces. À chaque extrémité d'une route, un véhicule est désigné comme le nœud responsable de la dissemination du paquet RTIP, appelé "RTIP Broadcast Node (RBN)". Le RBN transmettra le paquet

RTIP vers toutes les intersections voisines, mais seulement si le délai d'attente n'a pas expiré. Le RTIP se compose de trois parties comme s'est illustré dans la figure 4.2 :

- **En-tête statique** : Initialisé une seule fois par le nœud initiateur RTIP (RIN), le paquet comprend deux champs : un champ time-stamp représentant le moment de l'initiation du paquet RTIP, et un champ Source-Junction ID comprenant l'ID de l'intersection courante du RIN.
- **En-tête dynamique** : Utilisé par le RIN et le nœud diffuseur du paquet RTIP (RBN), l'en-tête dynamique, contrairement à l'en-tête statique, peut être mis à jour à chaque processus de dissémination. Il contient les informations pertinentes concernant la route suivante qui sera traitée. Il se compose de : un champ "From-JunctionID" représentant l'ID de jonction actuelle du (RIN ou RBN), un champ "ToJunctionID" désignant l'ID de la jonction de destination, et un champ Time-stamp représentant le temps de début du processus de dissémination du paquet RTIP.
- **Informations sur les routes** : Cette partie contient les entrées d'informations sur les routes. Chaque entrée est générée par un RBN (ou par le RIN lors de l'initialisation du RTIP). De plus, toutes les entrées incluent les champs suivants : un champ "Road-ID" qui est l'identifiant de la route actuelle, un champ "Road-Connectivity" indiquant la valeur de la connectivité routière (qui doit être mise à jour pour tous les chefs de groupe "Group Leader (GL)" situés sur la route actuelle), un champ "Road-Density" désignant la densité de la route (le nombre de véhicules), un champ "Delay" représentant le temps nécessaire pour traverser la route actuelle, et un champ Time-stamp indiquant l'heure de création de l'entrée de la table d'informations sur la route.

Il est important de mentionner que le paquet RTIP a une durée de vie de 0,8 seconde, une période brève mais suffisante avant que des changements significatifs dans la connectivité d'une route ne soient constatés. L'objectif principal de ce travail est de fournir au schéma de routage des informations actualisées sur le trafic de chaque route et de les partager avec les routes suivantes. Pour ce faire, la période de transmission des paquets RTIP est ajustée en fonction de la date de leur réception, garantissant qu'au cours de cette période, la connectivité de la route reste relativement stable.

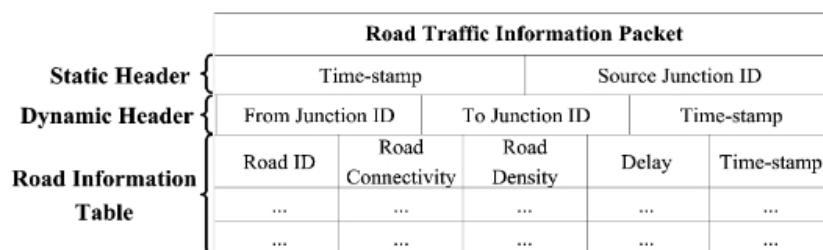


FIGURE 4.2 – Paquet d'information sur le trafic routier

Notre protocole divise chaque route en plusieurs cellules de groupes de véhicules, comme illustré dans la figure 4.3 . Dans chaque cellule, le véhicule le plus proche du

centre est sélectionné comme Chef de Groupe (GL). Lorsque le paquet RTIP passe par tous les GLs sur chaque route, chaque GL ajoute les informations de son groupe (telles que la densité, la connectivité, le délai, etc.) au paquet RTIP. Afin de remédier au retard accumulé lors de la dissémination des paquets RTIP dans le protocole PBRP, nous avons proposé dans le protocole ERGF que la période de dissémination des paquets RTIP soit mise à jour, plutôt que d'initialiser cette période et d'attendre la dissémination du prochain paquet RTIP après 0,8 seconde.

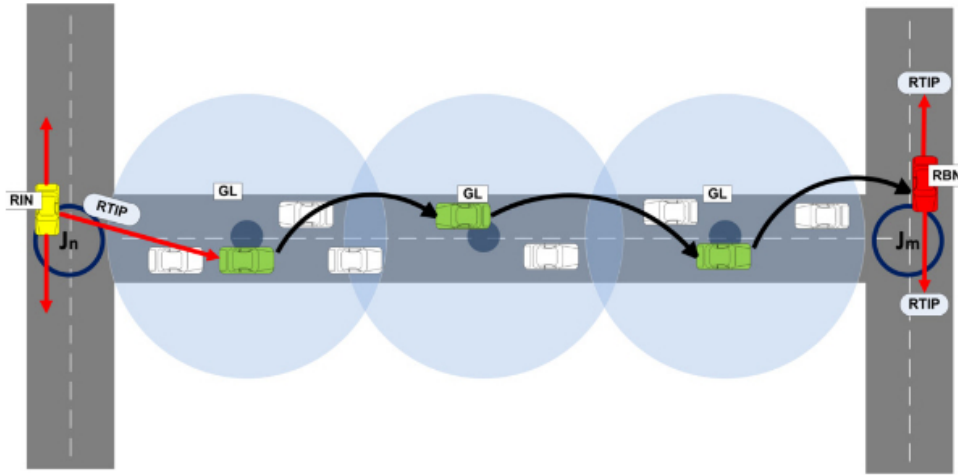


FIGURE 4.3 – Schéma de transmission des informations sur le trafic routier

4.4.2 Création du paquet RTIP

Afin de garantir un partage continu et régulier des informations sur le trafic routier, les paquets RTIPs sont périodiquement réinitialisés selon l'algorithme "Configuration de l'information périodique sur le trafic routier", comme illustré dans l'algorithme 1. Concrètement, lorsqu'un véhicule traverse une intersection et qu'il est le plus proche du centre de l'intersection actuelle (c'est-à-dire le RIN) avec un temporisateur expiré, il crée un nouveau paquet RTIP et initialise son en-tête statique. Ensuite, le nœud RIN diffuse ce paquet RTIP à tous les GLs des cellules voisines et réinitialise son temporisateur à 0. Il est important de noter que l'en-tête statique du paquet RTIP n'est modifié qu'à sa création.

4.4.3 Dissémination du paquet RTIP

Le paquet RTIP nouvellement créé passera par tous les GLs des cellules de chaque route afin de collecter les informations de trafic sur les routes qu'il traverse jusqu'à atteindre la fin de chaque segment. À chaque véhicule qu'il croise, la durée de vie du RTIP, qui est de 0,8 seconde, est vérifiée ; si cette durée est expirée, le RTIP est immédiatement supprimé. De plus, si le véhicule est un GL dans sa cellule, les champs de la partie "Road Information Table (RIT)" du paquet sont mis à jour en y ajoutant les informations de la cellule de ce GL. Si le GL actuel se trouve à la fin d'une route, le paquet doit être transmis au nœud RBN (RTIP Broadcast Node), où l'en-tête dynamique et la RIT du RTIP seront

Algorithm 1 Setup of Periodic Road Traffic Information

Require:

- V : Vehicle
- $RTIP_p$: RTIP packet
- $Cell_n$: Neighbor cell
- $Current_j$: Current junction of V
- $First_j$: First junction of $Cell_n$
- $Last_j$: Last junction of $Cell_n$
- SH : Static header of $RTIP_p$
- DH : Dynamic header of $RTIP_p$
- RIN : Road Traffic Information Initiator Node

```
1: if  $V$  is closer to the center of a junction then
2:   if  $V$  receives  $RIN\_Setup\_msg$  then
3:      $timer = 0$ 
4:     Start timer( $V$ )
5:   end if
6:   if  $V$  is  $RIN$  and  $timer > 0.8$  sec then
7:      $RTIP_p = \text{new RTIP packet}$ 
8:      $RTIP_p.SH.time\_stamp = currentTime()$ 
9:      $RTIP_p.DH.time\_stamp = currentTime()$ 
10:     $RTIP_p.SH.source\_junction = current_j$ 
11:    for all neighboring cells  $Cell_n$  do
12:       $RTIP_p.DH.from\_junction = Cell_n.First_j$ 
13:       $RTIP_p.DH.to\_junction = Cell_n.Last_j$ 
14:      Forward  $RTIP_p$  to  $GL$  of  $Cell_n$ 
15:    end for
16:    Broadcast  $RIN\_Setup\_msg$  around  $Current_j$ 
17:  end if
18: end if
```

mis à jour. Le paquet RTIP sera ensuite diffusé à tous les véhicules présents au niveau de l'intersection actuelle et aux GLs des cellules voisines. Les véhicules désignés comme RBN et RIN sont les seuls responsables de la mise à jour de l'en-tête dynamique du RTIP et de l'ajout d'entrées à la partie RIT. Il est important de mentionner que chaque véhicule met à jour son RIT lorsqu'il reçoit un paquet RTIP provenant du nœud RBN. L'algorithme 2 illustre le processus de dissémination des paquets RTIP et de mise à jour de la table RTI, ainsi que de la mise à jour des paquets RTIP.

4.4.4 Améliorations proposées du mécanisme de dissémination sur le trafic routier

Dans notre contribution, nous visons à proposer un mécanisme amélioré de dissémination des informations sur le trafic routier par rapport à celui présenté dans le protocole PBRP. Ce protocole s'appuie sur les informations de trafic, intégrées dans le paquet RTIP, pour prendre des décisions de routage. Ces paquets sont périodiquement initiés dans le

Algorithm 2 Road Traffic Information Forwarding Algorithm

Require:

- V : Vehicle
- $RTILF$: Road Traffic Information lifetime
- GL : Group Leader

```

1:  $V$  receives  $RTIP_p$ 
2: if ( $CurrentTime - RTIP_p.SH.time\_stamp > RTILF$ ) then
3:   Update RTI table
4:   return
5: end if
6:  $To_j = RTIP_p.DH.To\_junction$ 
7:  $From_j = RTIP_p.DH.To\_junction$ 
8: if  $V$  is  $GL$  of  $CurrentCell$  then
9:   Update  $RTIP_p$ 
10:  if  $CurrentCell$  is last cell then
11:    if  $V$  is around  $To_j$  then
12:      Forward  $RTIP_p$  to  $RBN$ 
13:    else
14:      Forward  $RTIP_p$  to closest neighbor to  $RBN$ 
15:    end if
16:    Forward  $RTIP_p$  to  $GL$  of  $Cell_n$ 
17:  else
18:    Forward  $RTIP_p$  to  $GL$  of next cell
19:  end if
20: else
21:   Forward  $RTIP_p$  to  $GL$  of  $CurrentCell$ 
22: end if
23: if  $V$  is  $RBN$  of  $To_j$  then
24:    $RoadID = getRoadID(From_j, To_j)$ 
25:    $Delay = CurrentTime - RTIP_p.DH.time\_stamp$ 
26:    $RTIP_p.RoadTable[RoadID].Delay = Delay$ 
27:   Broadcast  $RTIP_p$  around  $To_j$ 
28:   for all neighboring cells  $Cell_n$  do
29:      $RTIP_p.DH.From\_junction = Cell_n.First_j$ 
30:      $RTIP_p.DH.To\_junction = Cell_n.Last_j$ 
31:     Forward  $RTIP_p$  to  $GL$  of  $Cell_n$ 
32:   end for
33: end if

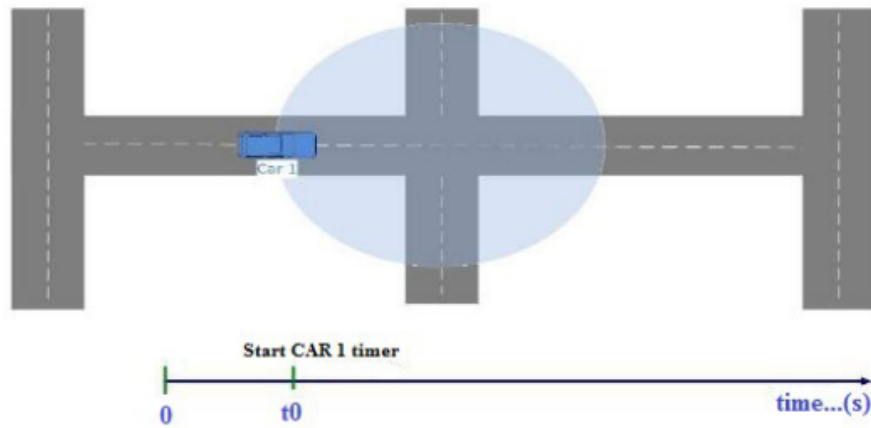
```

réseau par l'algorithme "Setup Periodic Road Traffic Information". Lors de nos tests, nous avons identifié certaines situations spécifiques où les performances de cet algorithme ont montré des limitations, avec un comportement parfois sous-optimal. Cela nous a encouragé à développer des solutions améliorées pour garantir un fonctionnement optimal dans ces contextes.

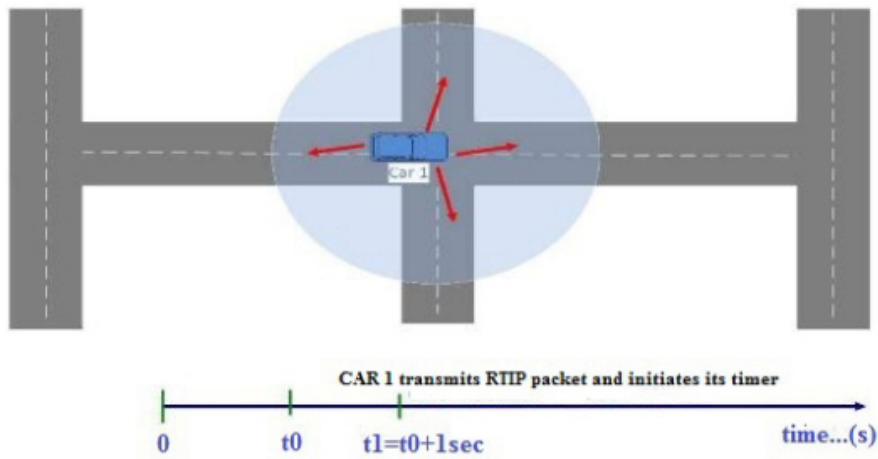
a) **Première limitation dans le protocole PBRP**

L'initiation du paquet RTIP aux intersections joue un rôle essentiel dans la diffusion des informations sur le trafic. Toute absence de ces informations ou de communication peut entraîner des échecs et des retards dans la transmission des données sur les routes connectées à ces intersections. Nous abordons ici la première limitation où l'algorithme "Setup Periodic Road Traffic Information" présente des lacunes. Comme illustré dans les figures suivantes, le scénario se déroule comme suit :

1. Au moment t_0 , un véhicule, appelé CAR_1 , entre dans une intersection et initialise son chronomètre comme montre la figure 4.4a. Une seconde plus tard, à t_1 , CAR_1 diffuse un paquet RTIP vers toutes les routes connectées à cette intersection comme s'est illustré dans la figure 4.4b.



(a) Premier scénario : L'arrivée du premier véhicule à l'intersection

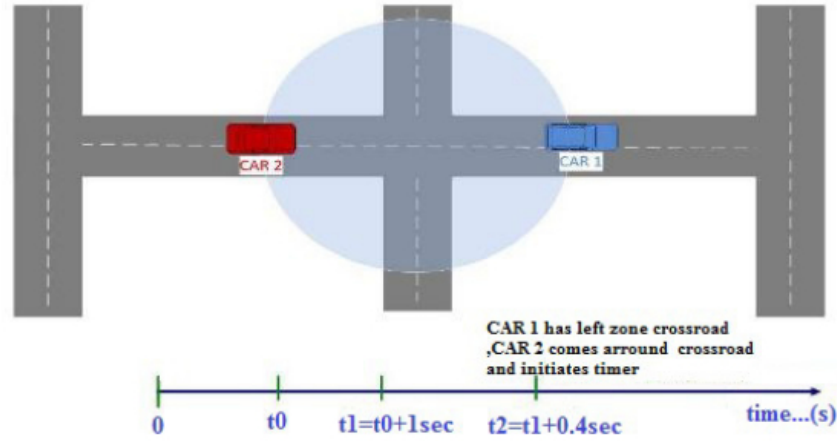


(b) Premier scénario : Diffusion du premier paquet RTIP

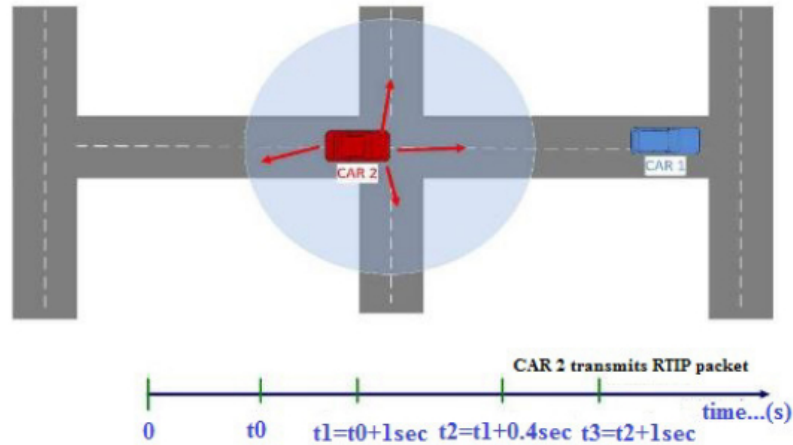
FIGURE 4.4 – Premier scénario : Partie 1

2. À t_2 , CAR_1 quitte la zone de l'intersection et un autre véhicule CAR_2 , entre en suivant le même processus que CAR_1 (voir Figure 4.5a). Ces événements se produisent 0,4 seconde après la dernière diffusion effectuée par CAR_1 ($t_2 = t_1 + 0.4 \text{ second}$).

Une seconde après l'arrivée de CAR_2 ($t_3 = t_2 + 1 \text{ second}$), CAR_2 initie un nouveau paquet RTIP pour toutes les routes connectées à l'intersection, réinitialise son chronomètre à zéro et informe les autres véhicules dans la même zone de réinitialiser également leurs chronomètres à zéro comme montre la figure 4.5b).



(a) Premier scénario : Arrivée du deuxième véhicule et la sortie du premier



(b) Premier scénario : Diffusion du deuxième RTIP par le deuxième véhicule

FIGURE 4.5 – Premier scénario : Partie 2

Comme décrit dans ce scénario, nous observons que lorsque CAR_1 initie le paquet RTIP, réinitialise son minuteur et demande à tous les autres véhicules présents dans l'intersection de faire de même, CAR_2 n'était pas présent dans la cellule à ce moment-là. Par conséquent, il n'a pas reçu cette information pour commencer le comptage à partir du moment du lancement du dernier paquet RTIP. Cela entraînera un retard dans le lancement du nouveau paquet RTIP, car CAR_2 attendra une seconde après être entré dans l'intersection avant de l'initier, sachant que son arrivée à cette intersection a eu lieu 0,4 seconde après le lancement du dernier paquet RTIP. En appliquant l'algorithme de diffusion du paquet RTIP tel qu'il est dans le protocole PBRP dans cette situation, un

retard de 0,4 seconde apparaîtra pour le lancement du paquet RTIP. Avec la répétition de ce scénario plusieurs fois, ce retard deviendra considérable et réduira l'efficacité du schéma de routage, car aucune information de trafic frais ne circulera. La figure 4.6 illustre le scénario actuel de diffusion des paquets RTIP en utilisant la version initiale de l'algorithme dans le protocole PBRP, ainsi que le scénario idéal dans lequel la diffusion devrait avoir lieu selon notre protocole ERGF.

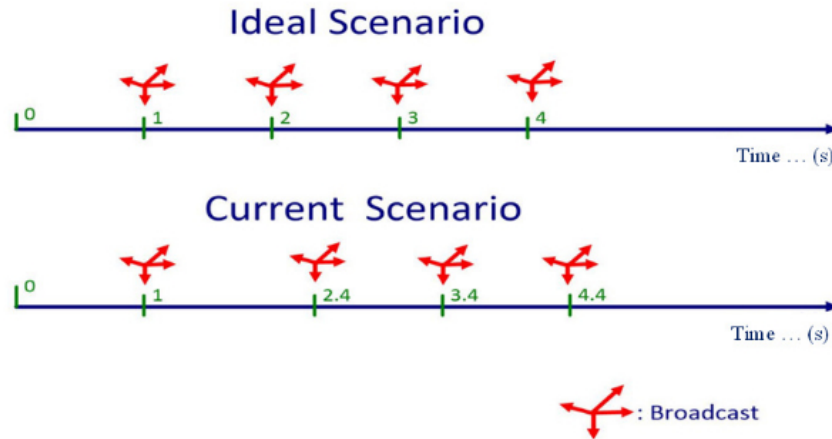


FIGURE 4.6 – Limitation 1 : Scénario idéal vs scénario actuel

b) Première amélioration

Pour gérer des situations comme celle-ci, nous avons proposé un mécanisme qui permet de se rapprocher du scénario idéal dans ce cas précis. Nous introduisons une solution basée sur le fait que lorsqu'un véhicule quitte une zone d'intersection, il doit communiquer la valeur de son minuteur à tous les véhicules présents dans cette zone. Cela permettra aux véhicules récepteurs de commencer leur comptage à partir de la valeur du minuteur reçue, si celle-ci est supérieure à leur propre valeur de minuteur. Le véhicule entrant dans la zone de l'intersection saura alors quand le dernier paquet RTIP a été initié et pourra lancer le suivant exactement 1 seconde après le précédent. Cette solution est détaillée dans l'algorithme 3.

c) Deuxième limitation présentée dans le protocole PBRP

Étant donné que l'environnement véhiculaire est très dynamique et dépend de plusieurs critères, tels que la période de transition (jour ou nuit), il peut arriver que certaines intersections restent vides pendant des périodes supérieures à 1 seconde. Nous illustrons ce scénario dans les figures 4.7a et 4.7b.

Nous considérons que les instants t_0 et t_1 se produisent comme dans le premier scénario décrit. Dans ce cas, 5 secondes après l'initiation du paquet RTIP par CAR_1 , ce véhicule a quitté la zone. Un autre véhicule, CAR_2 , est alors entré dans cette intersection et a lancé son propre minuteur. Bien qu'il n'y ait eu aucune information pendant 5 secondes sur les 4 routes de cette intersection, CAR_2 doit attendre 1 seconde avant d'initier le paquet RTIP, puis réinitialiser son minuteur à 0, et ainsi de suite. Cela aggrave la situation en ajoutant, aux 5 secondes d'absence d'informations, une autre seconde, ce qui entraîne un retard

Algorithm 3 First Improvement**Require:**

- V : Vehicle
- $Timer_V$: the current timer of vehicle V
- $RTIP_p$: Road traffic information packet
- $Last_RTIP_Timer$: the timer of the last vehicle leaving the junction
- $Cell_n$: neighbor cell
- $Current_j$: the current junction of V
- $First_j$: the first junction of $Cell_n$
- $Last_j$: the last junction of $Cell_n$
- SH : static header of $RTIP_p$
- DH : dynamic header of $RTIP_p$

```

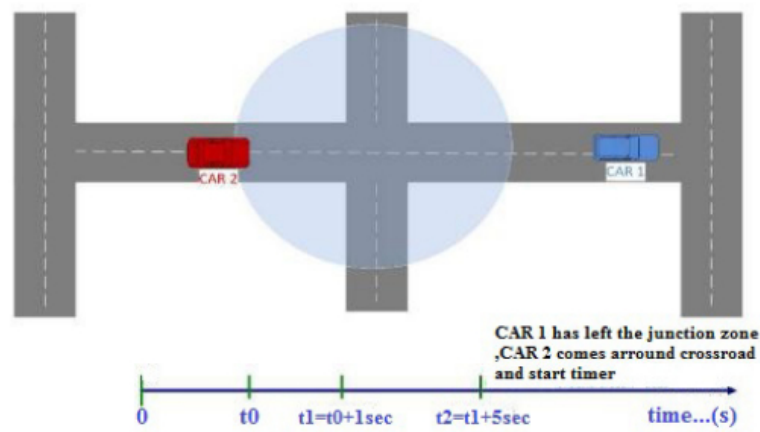
1: while True do
2:   if  $V$  receives  $Last\_RTIP\_Timer$  then
3:     if  $Last\_RTIP\_Timer > Timer_V$  then
4:        $Timer_V = Last\_RTIP\_Timer$ 
5:     end if
6:   end if
7:   if  $Timer_V \geq 1$  then
8:      $RTIP_p = \text{new } RTIP\_packet()$ 
9:      $RTIP_p.SH.time\_stamp = CurrentTime$ 
10:     $RTIP_p.DH.time\_stamp = CurrentTime$ 
11:     $RTIP_p.SH.source\_junction = Current_j$ 
12:    for all neighboring cells  $Cell_n$  do
13:       $RTIP_p.DH.from\_junction = Cell_n.First_j$ 
14:       $RTIP_p.DH.to\_junction = Cell_n.Last_j$ 
15:      Forward  $RTIP_p$  to  $GL$  of  $Cell_n$ 
16:    end for
17:    Broadcast  $Setup\_msg$  around  $Current_j$ 
18:     $Timer_V = 0$ 
19:   end if
20: end while

```

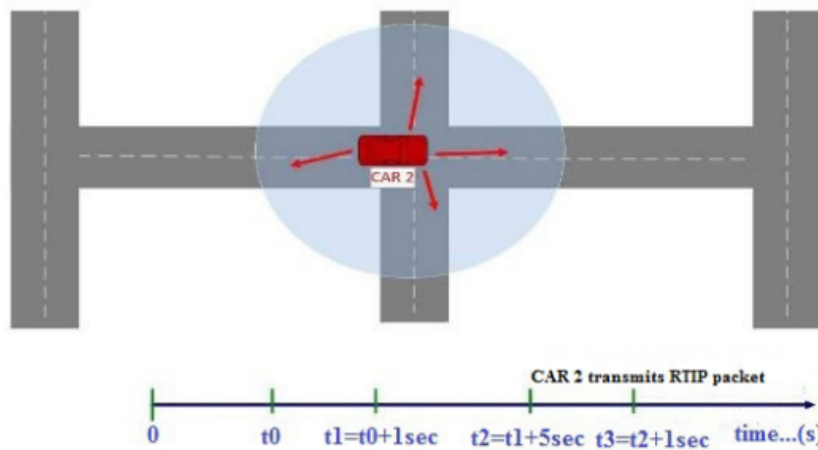
supplémentaire dans la diffusion des informations sur le trafic. La figure 4.8 représente le scénario actuel de diffusion des paquets RTIP avec l'algorithme d'origine du protocole PBRP, ainsi que le scénario idéal de diffusion tel qu'envisagé dans le protocole ERGF.

d) Deuxième amélioration

Pour gérer ce scénario et éviter les retards considérables et les périodes de vide dans l'initiation des paquets RTIP observées dans le protocole PBRP, nous proposons que tout véhicule venant de cellules situées au début ou à la fin d'une route, près de la zone d'intersection, initie immédiatement un paquet RTIP dès qu'il entre dans la zone de l'intersection. L'algorithme 4 résume la solution proposée dans le protocole ERGF pour surmonter cette deuxième limitation présentée dans le protocole PBRP.



(a) Deuxième scénario : Arrivée du deuxième véhicule



(b) Deuxième scénario : Diffusion de paquet RTIP deuxième véhicule

FIGURE 4.7 – Deuxième scénario

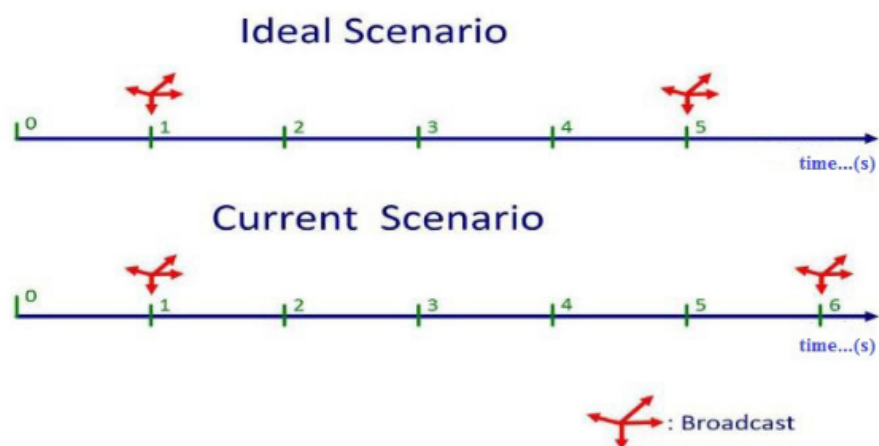


FIGURE 4.8 – Limitation 2 : Scénario idéal vs scénario actuel

Algorithm 4 Second Improvement**Require:**

- V : Vehicle
- $Timer_V$: the current timer of vehicle V
- $RTIP_p$: Road traffic information packet
- $Last_RTIP_Timer$: the timer of the last vehicle leaving the junction
- $Cell_n$: neighbor cell
- $Current_j$: the current junction of V
- $First_j$: the first junction of $Cell_n$
- $Last_j$: the last junction of $Cell_n$
- SH : static header of $RTIP_p$
- DH : dynamic header of $RTIP_p$

```

1: for Each time a vehicle changes cell do
2:   if  $CurrentCell == FirstCell \vee CurrentCell == LastCell$  then
3:      $RTIP_p = \text{new } RTIP\_packet()$ 
4:      $RTIP_p.SH.time\_stamp = CurrentTime$ 
5:      $RTIP_p.DH.time\_stamp = CurrentTime$ 
6:      $RTIP_p.SH.source\_junction = Current_j$ 
7:     for all neighboring cells  $Cell_n$  do
8:        $RTIP_p.DH.from\_junction = Cell_n.First_j$ 
9:        $RTIP_p.DH.to\_junction = Cell_n.Last_j$ 
10:      Forward  $RTIP_p$  to  $GL$  of  $Cell_n$ 
11:     end for
12:     Broadcast  $Setup\_msg$  around  $Current_j$ 
13:      $Timer_V = 0$ 
14:   end if
15: end for

```

4.5 Evaluation des performances

Dans cette section, nous évaluons les performances de nos améliorations proposées et les comparons à celles des protocoles EGYTAR et PBRP en termes de taux de livraison des paquets "Packet Delivery Ratio (PDR)", de délai de bout en bout et de l'overhead. Le taux de livraison des paquets représente le rapport entre le nombre total de paquets reçus et le nombre total de paquets générés par un véhicule source. L'overhead de routage est mesurée en fonction du nombre de paquets de contrôle, y compris les paquets RTIPs, pour le protocole PBRP et notre protocole proposé.

4.5.1 Environnement de simulations

Nous avons implémenté le protocole proposé à l'aide du simulateur OMNET++ et du framework INET-3.4.0, en utilisant le modèle de mobilité Manhattan Grid, qui décrit le mouvement des véhicules dans une grille de $1500 \times 1500 m^2$ composée de 16 routes droites. Chaque scénario a été réalisée 10 fois, et les résultats des performances sont calculés en se basant sur la moyenne de ces simulations. Les paramètres des simulations sont présentés dans le tableau 4.1.

TABLE 4.1 – Paramètres de simulation

Paramètres	Valeurs
Simulation time	100 sec
Beacon interval	0.5 sec
RTIP lifetime	0.8 sec
MAC protocol	802.11p
Network area	$1500 \times 1500 \text{ m}^2$
Transmission range	250 meters
Mobility model	Manhattan Mobility
Vehicle speed mean	50 Km/h
Channel capacity	6 Mbps
Number of junctions	16

4.5.2 Résultats de simulation

Dans cette section, nous présentons une comparaison des performances de nos améliorations avec d'autres protocoles tels que EGyTAR et PBRP.

a) Taux de livraison de paquets

La figure 4.9 montre que plus la densité de véhicules augmente, plus le nombre de paquets livrés à leurs véhicules de destination s'accroît. Cela s'explique par le fait que la connectivité des routes s'améliore avec une densité accrue. De plus, on peut observer que notre solution offre de meilleures performances que le protocole PBRP, grâce à notre mécanisme qui traite les situations où l'algorithme d'initiation du paquet RTIP ne se comporte pas de manière idéale. En effet, le protocole PBRP présente un ratio de livraison de paquets plus faible, principalement en raison du retard dans l'initiation des paquets RTIP. Ce retard affecte le processus de prise de décision de routage, car les informations sur le trafic routier contenues dans les paquets RTIP sont utilisées pour calculer le chemin des paquets de données. En revanche, notre solution proposée résout ce problème en évitant ces retards dans l'initiation des paquets RTIP. Enfin, en comparant notre protocole au protocole EGyTAR, on constate que le notre maintient un ratio de livraison de paquets supérieur. Cela s'explique par le fait que le protocole ERGF proposé utilise une vision avancée et des informations sur le trafic routier continuellement actualisées et précises, ce qui lui permet d'éviter les zones vides où le réseau est déconnecté.

b) Délai de bout en bout

La figure 4.10 démontre que notre protocole offre un délai de bout en bout inférieur à celui des autres protocoles EGyTAR et PBRP. En effet, en comparant notre solution au protocole PBRP, notre gestion efficace des situations où l'algorithme d'initiation des paquets RTIP ne fonctionne pas idéalement nous a permis de réduire les délais générés. En revanche, le protocole PBRP présente un délai plus élevé en raison de l'accumulation des retards lors du traitement des paquets RTIP par les différents véhicules entrant dans la zone d'intersection. Ce retard dans l'initiation des paquets RTIP impacte le processus de prise de décision pour le routage. Cependant, grâce aux solutions que nous avons proposées

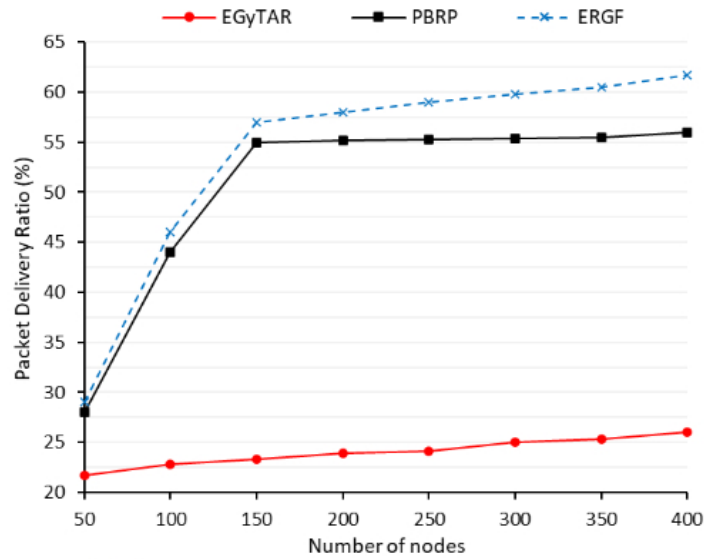


FIGURE 4.9 – Taux de livraison des paquets dans les protocoles EGyTAR, PBRP, ERGF

pour remédier aux retards causés par une gestion inefficace des paquets RTIP, le protocole ERGF a permis de réduire de manière significative ces délais pour le routage des paquets RTIP entre les véhicules. En outre, notre protocole propose un délai plus faible que le protocole EGyTAR, car ce dernier utilise la stratégie "Carry and Forward", qui consiste à conserver et transporter les paquets, ce qui augmente les délais. En revanche, dans le mode de récupération de notre protocole, nous appliquons l'algorithme de routage "Backward Forwarding", ce qui permet de réduire les délais de transmission.

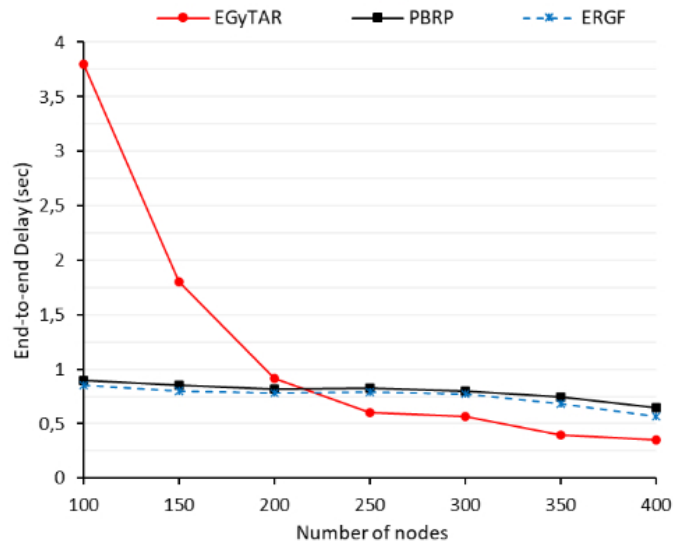


FIGURE 4.10 – Délai de bout en bout dans les protocoles EGyTAR, PBRP, ERGF

c) Overhead du routage

La métrique d'overhead de routage représente le nombre de paquets de contrôle, représentés par les paquets RTIP dans les deux protocoles. Nous analyserons l'initiation des paquets RTIP dans les deux protocoles (ERGF et PBRP), ainsi que l'impact de cette initiation sur la diffusion des paquets RTIP par les nœuds RBN vers tous les véhicules voisins.

1- Évaluation des paquets d'initiation par le nœud RIN

La figure 4.11 montre le nombre d'initiations et de lancements des paquets RTIP dans le réseau pendant la simulation par le véhicule RIN. Nous pouvons conclure qu'avec le protocole ERGF, la stratégie de diffusion conduit à un nombre plus élevé d'initiations de paquets RTIP par rapport au protocole PBRP. Cela peut s'expliquer comme suit : en raison de la fréquence d'apparition de certaines situations aux intersections, l'algorithme d'initiation des paquets RTIP dans le protocole PBRP ne parvient pas à lancer les paquets RTIP aux moments idéaux attendus, ce qui crée un retard et réduit le nombre d'initiations de paquets RTIP. Bien que notre solution entraîne une augmentation de l'overhead de routage, nous pouvons considérer que cette quantité de paquets de contrôle est nécessaire pour garder l'efficacité des décisions du processus de routage, car ce mécanisme de routage nécessite les informations sur le trafic véhiculaire contenues dans les paquets RTIP pour transmettre davantage de paquets.

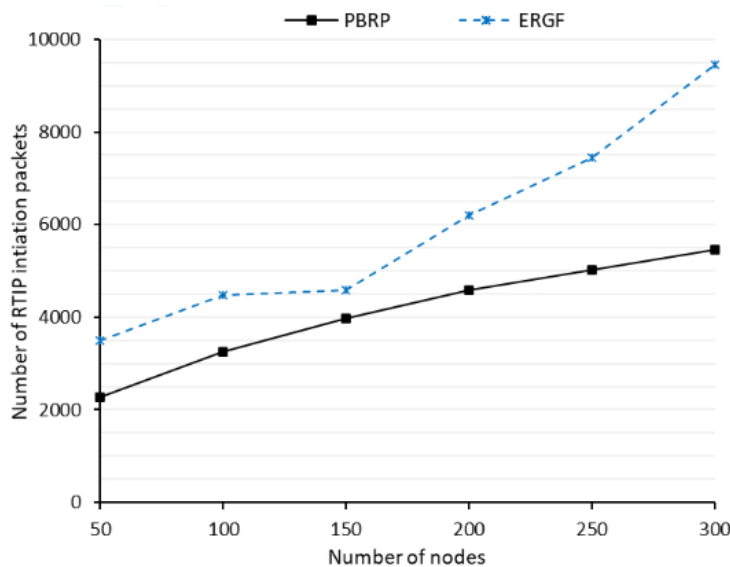


FIGURE 4.11 – Nombre de paquets RTIP dans les protocoles PBRP et ERGF

2- Diffusion des paquets RTIP par le nœud RBN

Nous étudions l'impact de nos améliorations de l'algorithme d'initiation sur la diffusion des paquets RTIP par les véhicules RBN, ce qui montre que l'information a pu atteindre un grand nombre de nœuds, puisque les nœuds RBN sont responsables de la diffusion des paquets RTIP à la fin de chaque route. Comme montré dans la figure 4.12, l'amélioration de l'algorithme d'initiation a conduit à un plus grand nombre de paquets RTIP diffusés

par rapport au protocole PBRP. Cela confirme que davantage de paquets RTIP ont atteint la fin des routes, ce qui a permis à un plus grand nombre de nœuds RBN de recevoir les paquets RTIP et de les diffuser aux véhicules voisins. Cette augmentation de la diffusion des paquets RTIP est une confirmation que l'information a pu atteindre un grand nombre de nœuds, ce qui aide ces nœuds dans le processus de prise de décision de routage.

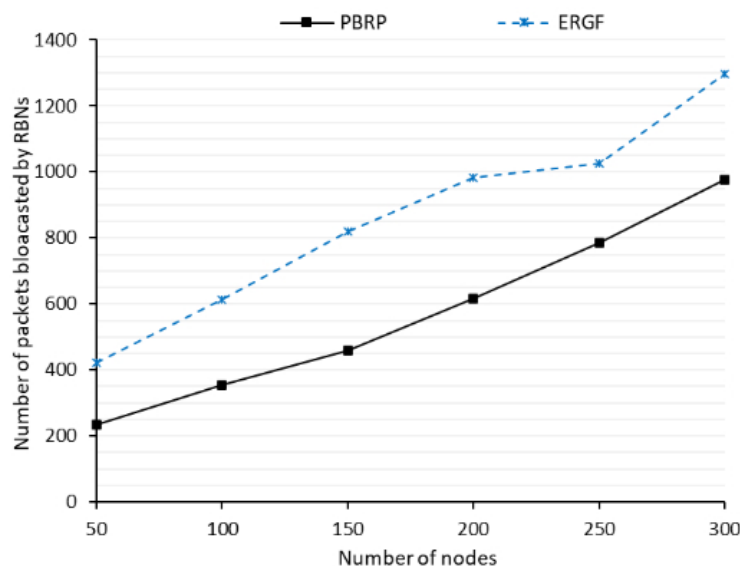


FIGURE 4.12 – Nombres de paquets RTIP diffusés par les RBNs dans les protocoles PBRP et ERGF

4.6 Conclusion

Dans ce chapitre, nous avons proposé des améliorations au mécanisme de diffusion des informations sur le trafic véhiculaire en utilisant une version optimisée de l'algorithme "Periodic Road Traffic Information Initiation" dans notre protocole ERGF, comparé à sa version initiale dans le protocole PBRP. Cela a donné à notre protocole ERGF la capacité de fournir des informations permanentes et avancées sur le trafic véhiculaire incluses dans les paquets RTIP.

Nous avons implémenté et évalué les améliorations proposées à l'aide de simulations réalisées avec le simulateur OMNET++. Les résultats obtenus démontrent que notre contribution améliore l'exploitation des informations sur le trafic véhiculaire de manière efficace, offrant des performances significatives en termes de taux de livraison des paquets et de délai de bout en bout, par rapport aux protocoles GyTAR/EGyTAR. Cela est dû à l'utilisation d'une information continue sur le trafic routier grâce au mécanisme du diffusion.

Conclusions et Perspectives

Conclusions et Perspectives

Le paradigme de découpage réseau est considéré comme une solution permettant d'assurer l'efficacité de la coexistence de multiples services avec des exigences de QoS différentes au sein de la même infrastructure 5G, où chaque tranche du réseau est dédiée à satisfaire un type spécifique de service. La diversité des services véhiculaires rend la mise en place du découpage réseau dans la 5G une nécessité, afin de répondre aux différentes exigences en matière de QoS pour chaque service. Cependant, cette coexistence au sein d'une même infrastructure pose de nombreux défis, notamment en ce qui concerne l'allocation des ressources entre ces différents services, en particulier dans la partie du réseau d'accès radio (RAN), où les ressources telles que la bande passante sont limitées. Afin de garantir une meilleure utilisation de la bande passante tout en répondant aux exigences des services véhiculaires, il est judicieux que la gestion des ressources soit efficace. Dans le cadre de cette thèse, nous avons proposé une solution axée sur la gestion des ressources dans un environnement 5G V2X avec le découpage réseau.

L'approche de réservation complète, appelée "strict slicing", est l'une des approches de référence visant à fournir une isolation totale. Cependant, dans un environnement de véhicules hautement dynamique, cette isolation totale entraîne un manque de flexibilité, ce qui rend cette approche inadaptée à être appliquée telle quelle. Dans notre contribution, Nous avons développé un mécanisme de partage des ressources entre les tranches V2X, qui sera déclenché en cas de surcharge d'une tranche. Pour ce faire, nous avons modélisé notre problème en utilisant un processus de décision markovien (MDP) et appliqué l'apprentissage par renforcement profond "Deep reinforcement Learning (DQL)" pour déterminer les décisions de partage de ressources appropriées. Notre solution permet d'ajouter de la flexibilité à l'approche strict slicing en permettant de trouver des ressources dans les tranches sous-utilisées pour accepter les demandes de connexion provenant des tranches surchargées, tout en maintenant la personnalisation de la politique d'allocation des ressources au sein de chaque tranche. Elle prend également en compte l'aspect de l'isolation entre les tranches, où les décisions de partage qui entraînent une dégradation des performances dans les tranches prêteuses sont considérées comme de mauvaises décisions. De plus, nos évaluations de performances ont démontré que le mécanisme de partage des ressources proposé est à la fois efficace dans l'optimisation de l'utilisation des ressources, tout en protégeant l'aspect d'isolation entre les tranches dans un environnement dynamique tel que l'environnement véhiculaire.

Dans un autre contexte, nous avons traité une problématique liée au routage dans l'environnement VANETs, en proposant des améliorations au mécanisme chargé de la dissémination des informations sur le trafic routier et l'état du réseau. Ces améliorations visent à fournir une vision précise et actualisée des conditions du trafic routier, permettant ainsi une prise de décision plus efficace dans le processus de routage.

En tant que perspectives :

- Nous aspirons à améliorer davantage le partage des ressources en proposant un mécanisme de gestion des ressources distribué basé sur l'apprentissage par renforcement profond multi-agent (DRL). Ce mécanisme optimisera le partage à deux niveaux : entre les tranches au sein de la même cellule et entre les tranches de cellules adjacentes.
- Nous proposons d'améliorer le processus d'entraînement en proposant un modèle d'apprentissage automatique, tel que les réseaux Long Short-Term Memory (LSTM), qui capture les caractéristiques temporelles d'un scénario réel de mobilité des véhicules et leur comportement de connexion dans chaque tranche sur une période de temps définie, afin de les classifier. Cette approche permettra de former l'agent d'apprentissage par renforcement profond (DQL) de manière adaptée, en fonction du scénario spécifique.
- Nous envisageons également d'explorer l'utilisation de modèles d'apprentissage par renforcement dans la gestion des différentes ressources, telles que les ressources de calcul dans les environnements de edge computing, de cloud computing et dans le réseau central pour les tranches de réseau véhiculaires.

Références Bibliographiques

Bibliographie

- [1] The 5G Infrastructure Public Private Partnership, “5G Automotive Vision,” *5G-PPP Initiative*, pp. 1–67, 2015.
- [2] S. A. Abdel Hakeem, A. A. Hady, and H. W. Kim, “5G-V2X : standardization, architecture, use cases, network-slicing, and edge-computing,” *Wireless Networks*, vol. 26, no. 8, pp. 6015–6041, 2020.
- [3] M. Agiwal, A. Roy, and N. Saxena, “Next Generation 5G Wireless Networks : A Comprehensive Survey,” vol. 18, no. 3, pp. 1617–1655, 2016.
- [4] A. Osseiran, J. F. Monserrat, and P. Marsch, *5G mobile and wireless communications technology*. 2016.
- [5] T. S. Rappaport, S. Sun, R. Mayzus, H. Zhao, Y. Azar, K. Wang, G. N. Wong, J. K. Schulz, M. Samimi, and F. Gutierrez, “Millimeter wave mobile communications for 5G cellular : It will work!,” *IEEE Access*, vol. 1, pp. 335–349, 2013.
- [6] T. Doukoglou, V. Gezerlis, K. Trichias, N. Kostopoulos, N. Vrakas, M. Bougioukos, and R. Legouable, “Vertical Industries Requirements Analysis Targeted KPIs for Advanced 5G Trials,” *2019 European Conference on Networks and Communications, EuCNC 2019*, pp. 95–100, 2019.
- [7] Thales, “Introducing 5G technology and networks (definition, use cases and rollout),” 2019.
- [8] A. B. Lindquist, R. R. Seeber, and L. W. Cdmeau, “A Time-Sharing System Using an Associative Memory,” *Proceedings of the IEEE*, vol. 54, no. 12, pp. 1774–1779, 1966.
- [9] A. A. Barakabitze, A. Ahmad, R. Mijumbi, and A. Hines, “5G network slicing using SDN and NFV : A survey of taxonomy, architectures and future challenges,” *Computer Networks*, vol. 167, p. 106984, 2020.
- [10] D. Kreutz, F. M. Ramos, P. E. Verissimo, C. E. Rothenberg, S. Azodolmolky, and S. Uhlig, “Software-defined networking : A comprehensive survey,” *Proceedings of the IEEE*, vol. 103, no. 1, pp. 14–76, 2015.
- [11] H. Zhang, N. Liu, X. Chu, K. Long, A. H. Aghvami, and V. C. Leung, “Network Slicing Based 5G and Future Mobile Networks : Mobility, Resource Management, and Challenges,” *IEEE Communications Magazine*, vol. 55, no. 8, pp. 138–145, 2017.

- [12] Next Generation Mobile Networks Alliance 5G Initiative, “5G White Paper,” *A Deliverable by the NGMN Alliance*, p. 124, 2015.
- [13] ITU-T Study Group 13, “ITU-T Rec. Y.3011 (01/2012) Framework of network virtualization for future networks,” p. 28, 2012.
- [14] A. Neal, “Feasibility Study on New Services and Markets Technology Enablers for Massive Internet of Things (28 pages),” tech. rep.
- [15] TSGS, “TS 128 533 - V16.4.0 - 5G ; Management and orchestration ; Architecture framework (3GPP TS 28.533 version 16.4.0 Release 16),” vol. 0, pp. 0–31, 2020.
- [16] C. De Alwis, P. Porambage, K. Dev, T. R. Gadekallu, and M. Liyanage, “A Survey on Network Slicing Security : Attacks, Challenges, Solutions and Research Directions,” *IEEE Communications Surveys and Tutorials*, vol. 26, no. 1, pp. 534–570, 2024.
- [17] A. A. Barakabitze, L. Sun, I.-H. Mkwawa, and E. Ifeakor, “A Novel QoE-Aware SDN-enabled, NFV-based Management Architecture for Future Multimedia Applications on 5G Systems,” no. August, 2019.
- [18] Open Networking Foundation, “Network Functions Virtualisation : An Introduction, Benefits, Enablers, Challenges & Call for Action,” *SDN and OpenFlow World Congress (white paper)*, vol. 1, no. 1, pp. 1–16, 2012.
- [19] A. Doria, J. Hadi Salim, Z. R. Haas, E. H. IBM Khosravi, E. W. Intel Wang, and E. L. Dong, “Internet Engineering Task Force (IETF) Forwarding and Control Element Separation (ForCES) Protocol Specification,” *Internet Engineering Task Force (IETF)*, pp. 1–124, 2010.
- [20] A. Lara, A. Kolasani, and B. Ramamurthy, “Network innovation using open flow : A survey,” *IEEE Communications Surveys and Tutorials*, vol. 16, no. 1, pp. 493–512, 2014.
- [21] O. Members, “Applying SDN Architecture to 5G Slicing, Issue 1 (19 pages),” tech. rep., Open Networking Foundation (ONF), 2016.
- [22] J. Ordonez-Lucena, P. Ameigeiras, D. Lopez, J. J. Ramos-Munoz, J. Lorca, and J. Folgueira, “Network Slicing for 5G with SDN/NFV : Concepts, Architectures, and Challenges,” *IEEE Communications Magazine*, vol. 55, no. 5, pp. 80–87, 2017.
- [23] I. Farris, T. Taleb, Y. Khettab, and J. Song, “A survey on emerging SDN and NFV security mechanisms for IoT systems,” *IEEE Communications Surveys and Tutorials*, vol. 21, no. 1, pp. 812–837, 2019.
- [24] R. Mijumbi, J. Serrat, J. L. Gorricho, N. Bouten, F. De Turck, and R. Boutaba, “Network function virtualization : State-of-the-art and research challenges,” *IEEE Communications Surveys and Tutorials*, vol. 18, no. 1, pp. 236–262, 2016.
- [25] S. M. A. Kazmi, L. U. Khan, N. H. Tran, and C. S. Hong, *Network Slicing for 5G and Beyond Networks*. Springer Cham, 2019.

-
- [26] S. Bakri, *Towards enforcing network slicing in 5G networks*. PhD thesis, Doctoral School of Informatics, Telecommunications and Electronics of Paris EURECOM, 2021.
 - [27] Z. Sanaei, S. Abolfazli, A. Gani, and R. Buyya, "Heterogeneity in mobile cloud computing : Taxonomy and open challenges," *IEEE Communications Surveys and Tutorials*, vol. 16, no. 1, pp. 369–392, 2014.
 - [28] G. Members, "Technical Specification Group Services and System Aspects; Study on Enhancement of 3GPP Support for 5G V2X Services (Release 16)," tech. rep., 3GPP support office address, 2018.
 - [29] X. Foukas, G. Patounas, A. Elmokashfi, and M. K. Marina, "Network Slicing in 5G : Survey and Challenges," *IEEE Communications Magazine*, vol. 55, no. 5, pp. 94–100, 2017.
 - [30] F. Z. Yousaf, P. Loureiro, F. Zdarsky, T. Taleb, and M. Liebsch, "Cost analysis of initial deployment strategies for virtualized mobile core network functions," *IEEE Communications Magazine*, vol. 53, no. 12, pp. 60–66, 2015.
 - [31] H. Hawilo, A. Shami, M. Mirahmadi, and R. Asal, "NFV : State of the art, challenges, and implementation in next generation mobile networks (vEPC)," *IEEE Network*, vol. 28, no. 6, pp. 18–26, 2014.
 - [32] M. Richart, J. Baliosian, J. Serrat, and J. L. Gorricho, "Resource Slicing in Virtual Wireless Networks : A Survey," *IEEE Transactions on Network and Service Management*, vol. 13, no. 3, pp. 462–476, 2016.
 - [33] T. Specification and G. Services, "3rd Generation Partnership Project ; Security aspect for LTE support of Vehicle-to-Everything (V2X)," tech. rep., European Telecommunications Standards Institute, 2017.
 - [34] Y. B. Lin, C. C. Tseng, and M. H. Wang, "Effects of transport network slicing on 5g applications," *Future Internet*, vol. 13, no. 3, 2021.
 - [35] D. Systems, "ITU-T," 2018.
 - [36] Z. Shu and T. Taleb, "A Novel QoS Framework for Network Slicing in 5G and beyond Networks Based on SDN and NFV," *IEEE Network*, vol. 34, no. 3, pp. 256–263, 2020.
 - [37] B. Cornaglia, G. Young, and A. Marchetta, "Fixed Access Network Sharing," *Optical Fiber Technology*, vol. 26, pp. 2–11, 2015.
 - [38] I. F. Akyildiz, P. Wang, and S. C. Lin, "SoftAir : A software defined networking architecture for 5G wireless systems," *Computer Networks*, vol. 85, pp. 1–18, 2015.
 - [39] C. Mobile, "C-RAN : the road towards green RAN," *White Paper, ver 3.0*, vol. 0, pp. 15–16, 2015.

- [40] I. Afolabi, T. Taleb, K. Samdanis, A. Ksentini, and H. Flinck, “Network slicing and softwarization : A survey on principles, enabling technologies, and solutions,” *IEEE Communications Surveys and Tutorials*, vol. 20, no. 3, pp. 2429–2453, 2018.
- [41] W. Guan, X. Wen, L. Wang, Z. Lu, and Y. Shen, “A service-oriented deployment policy of end-to-end network slicing based on complex network theory,” *IEEE Access*, vol. 6, no. June, pp. 19691–19701, 2018.
- [42] J. Mei, X. Wang, and K. Zheng, “Intelligent Network Slicing for V2X Services Toward 5G,” *IEEE Network*, vol. 33, no. 6, pp. 196–204, 2019.
- [43] A. Alalewi, I. Dayoub, and S. Cherkaoui, “On 5G-V2X Use Cases and Enabling Technologies : A Comprehensive Survey,” *IEEE Access*, vol. 9, pp. 107710–107737, 2021.
- [44] J. T. P. G. M. F. W. S. E. N. B. B. Antonio Eduardo Fernandez, Alain Servel, “LTE ; Service requirements for V2X services (3GPP TS 22.185 version 14.3.0 Release 14) (16 pages),” tech. rep., Fifth Generation Communication Automotive Research and innovation, 2017.
- [45] E. Members, “5G ; Vehicle-to-everything (V2X) ; Media handling and interaction (3GPP TR 26.985 version 17.0.0 Release 17) (29 pages),” tech. rep., European Telecommunications Standards Institute (ETSI), 2022.
- [46] G. Members, “Technical Specification Group Radio Access Network ; NR ; Study on NR Vehicle-to-Everything (V2X) (Release 16),” tech. rep., Alliance for Telecommunications Industry Solutions, 2019.
- [47] H. Ma, S. Li, E. Zhang, Z. Lv, J. Hu, and X. Wei, “Cooperative autonomous driving oriented MEC-Aided 5G-V2X : Prototype system design, field tests and AI-Based optimization tools,” *IEEE Access*, vol. 8, pp. 54288–54302, 2020.
- [48] Z. Ying, M. Ma, and L. Yi, “Bavpm : Practical autonomous vehicle platoon management supported by blockchain technique,” in *Proceedings of the 4th International Conference on Intelligent Transportation Engineering (ICITE)*, (Singapore, Singapore), pp. 256–260, IEEE, 2019.
- [49] C. Campolo, A. Molinaro, A. Iera, R. R. Fontes, and C. E. Rothenberg, “Towards 5G Network Slicing for the V2X Ecosystem,” in *Proceedings of the 4th International Conference on Network Softwarization and Workshops (NetSoft)*, (Montreal, Canada), pp. 400–405, IEEE, 2018.
- [50] C. Campolo, A. Molinaro, A. Iera, and F. Menichella, “5G Network Slicing for Vehicle-to-Everything Services,” *IEEE Wireless Communications*, vol. 24, no. 6, pp. 38–45, 2017.
- [51] N. Mouawad, R. Naja, and S. Tohme, “Inter-slice handover management in a V2X slicing environment using bargaining games,” *Wireless Networks*, vol. 26, pp. 3883–3903, 2020.

-
- [52] C. Campolo, R. D. R. Fontes, A. Molinaro, C. E. Rothenberg, and A. Iera, "Slicing on the Road : Enabling the Automotive Vertical through 5G Network Softwarization," *Sensors*, vol. 18, no. 12, p. 4435, 2018.
 - [53] P. Rost, C. Mannweiler, D. S. Michalopoulos, C. Sartori, V. Sciancalepore, N. Sastry, O. Holland, S. Tayade, B. Han, D. Bega, D. Aziz, and H. Bakker, "Network Slicing to Enable Scalability and Flexibility in 5G Mobile Networks," *IEEE Communications Magazine*, vol. 55, no. 5, pp. 72–79, 2017.
 - [54] X. Li, R. Ni, J. Chen, Y. Lyu, Z. Rong, and R. Du, "End-to-End Network Slicing in Radio Access Network, Transport Network and Core Network Domains," *IEEE Access*, vol. 8, pp. 29525–29537, 2020.
 - [55] K. Xiong, S. Leng, C. Huang, C. Yuen, and Y. L. Guan, "Intelligent Task Offloading for Heterogeneous V2X Communications," *IEEE Transactions on Intelligent Transportation Systems*, vol. 22, no. 4, pp. 2226–2238, 2021.
 - [56] T. Flynn, G. Grispos, W. B. Glisson, and W. Mahoney, "Knock! Knock! Who is there? Investigating data leakage from a medical internet of things hijacking attack," in *Proceedings of the 53rd Hawaii International Conference on System Sciences*, (Hawaii, USA), pp. 6456–6465, 2020.
 - [57] M. R. Raza, C. Natalino, P. Öhlen, L. Wosinska, and P. Monti, "Reinforcement Learning for Slicing in a 5G Flexible RAN," *Journal of Lightwave Technology*, vol. 37, no. 20, pp. 5161–5169, 2019.
 - [58] H. Zhang and V. W. S. Wong, "A Two-Timescale Approach for Network Slicing in C-RAN," *IEEE Transactions on Vehicular Technology*, vol. 69, no. 6, pp. 6656–6669, 2020.
 - [59] A. A. Gebremariam, M. Chowdhury, M. Usman, A. Goldsmith, and F. Granelli, "SoftSLICE : Policy-Based Dynamic Spectrum Slicing in 5G Cellular Networks," in *Proceedings of International Conference on Communications (ICC)*, (Kansas City, MO, USA), pp. 1–6, IEEE, 2018.
 - [60] M. Maule, P.-V. Mekikis, K. Ramantas, J. Vardakas, and C. Verikoukis, "Dynamic partitioning of radio resources based on 5G RAN Slicing," in *Proceedings of Global Communications Conference*, (Taipei, Taiwan), pp. 1–6, IEEE, 2020.
 - [61] H. Khan, S. Samarakoon, and M. Bennis, "Enhancing Video Streaming in Vehicular Networks via Resource Slicing," *IEEE Transactions on Vehicular Technology*, vol. 69, no. 4, pp. 3513–3522, 2020.
 - [62] T. Shuminoski and T. Janevski, "Lyapunov optimization framework for 5g mobile nodes with multi-homing," *IEEE Communications Letters*, vol. 20, no. 5, pp. 1026–1029, 2016.
 - [63] F. Mehmeti and T. F. L. Porta, "Admission Control for Consistent Users in Next Generation Cellular Networks," in *Proceedings of International Conference on Communications (ICC'2019)*, (Shanghai, China), pp. 1–7, IEEE, 2019.

- [64] T. Guo and R. Arnott, "Active LTE RAN Sharing with Partial Resource Reservation," in *Proceedings of the 78th Vehicular Technology Conference (VTC Fall)*, (Las Vegas, NV, USA), pp. 1–5, IEEE, 2013.
- [65] M. Kalil, A. Shami, and Y. Ye, "Wireless resources virtualization in LTE systems," in *Proceedings of International Conference on Computer Communications Workshops (INFOCOM WKSHPS)*, (Toronto, ON, Canada), pp. 363–368, IEEE, 2014.
- [66] R. Kokku, R. Mahindra, H. Zhang, and S. Rangarajan, "CellSlice : Cellular wireless resource slicing for active RAN sharing," in *Proceedings of the 5th International Conference on Communication Systems and Networks (COMSNETS)*, (Bangalore, India), pp. 1–10, IEEE, 2013.
- [67] M. I. Kamel, L. B. Le, and A. Girard, "LTE Wireless Network Virtualization : Dynamic Slicing via Flexible Scheduling," in *Proceedings of the 80th Vehicular Technology Conference (VTC2014-Fall)*, (Vancouver, BC, Canada), pp. 1–5, IEEE, 2014.
- [68] P. Caballero, A. Banchs, G. de Veciana, and X. Costa-Pérez, "Multi-Tenant Radio Access Network Slicing : Statistical Multiplexing of Spatial Loads," *IEEE/ACM Transactions on Networking*, vol. 25, no. 5, pp. 3044–3058, 2017.
- [69] J. Perez-Romero and O. Sallent, "Optimization of Multitenant Radio Admission Control through a Semi-Markov Decision Process," *IEEE Transactions on Vehicular Technology*, vol. 69, no. 1, pp. 862–875, 2020.
- [70] D. Nojima, Y. Katsumata, T. Shimojo, Y. Morihiro, T. Asai, A. Yamada, and S. Iwashina, "Resource Isolation in RAN Part while Utilizing Ordinary Scheduling Algorithm for Network Slicing," *IEEE Vehicular Technology Conference*, vol. 2018-June, pp. 1–5, 2018.
- [71] Y. Cai, Q. Zhang, and Z. Feng, "QoS-Guaranteed Radio Resource Scheduling in 5G V2X Heterogeneous Systems," in *Proceedings of IEEE Globecom Workshops (GC Wkshps)*, (Waikoloa, HI, USA), pp. 1–6, IEEE, 2019.
- [72] A. Ksentini and N. Nikaein, "Toward Enforcing Network Slicing on RAN : Flexibility and Resources Abstraction," *IEEE Communications Magazine*, vol. 55, no. 6, pp. 102–108, 2017.
- [73] K. Zhu and E. Hossain, "Virtualization of 5G Cellular Networks as a Hierarchical Combinatorial Auction," *IEEE Transactions on Mobile Computing*, vol. 15, no. 10, pp. 2640–2654, 2016.
- [74] Y. Gadallah, M. H. Ahmed, and E. Elalamy, "Dynamic LTE resource reservation for critical M2M deployments," *Pervasive and Mobile Computing*, vol. 40, pp. 541–555, 2017.
- [75] R. Kokku, R. Mahindra, H. Zhang, and S. Rangarajan, "NVS : A virtualization substrate for WiMAX networks," *Proceedings of the 16th annual international conference on Mobile computing and networking (MobiCom '10)*, pp. 233–244, 2010.

-
- [76] B. Liu and H. Tian, “A Bankruptcy Game-Based Resource Allocation Approach among Virtual Mobile Operators,” *IEEE Communications Letters*, vol. 17, no. 7, pp. 1420–1423, 2013.
 - [77] Y. Jia, H. Tian, S. Fan, P. Zhao, and K. Zhao, “Bankruptcy game based resource allocation algorithm for 5G Cloud-RAN slicing,” in *Proceedings of IEEE Wireless Communications and Networking Conference (WCNC)*, (Barcelona, Spain), pp. 1–6, IEEE, 2018.
 - [78] J. Hu, Z. Zheng, B. Di, and L. Song, “Tri-Level Stackelberg Game for Resource Allocation in Radio Access Network Slicing,” in *Proceedings of IEEE Global Communications Conference (GLOBECOM)*, (Abu Dhabi, United Arab Emirates), pp. 1–6, IEEE, 2018.
 - [79] P. Caballero, A. Banchs, G. De Veciana, X. Costa-Perez, and A. Azcorra, “Network slicing for guaranteed rate services : Admission control and resource allocation games,” *IEEE Transactions on Wireless Communications*, vol. 17, no. 10, pp. 6419–6432, 2018.
 - [80] M. Morcos, T. Chahed, L. Chen, J. Elias, and F. Martignon, “A two-level auction for resource allocation in multi-tenant c-ran,” *Computer Networks*, vol. 135, pp. 240–252, 2018.
 - [81] M. Vincenzi, E. Lopez-Aguilera, and E. Garcia-Villegas, “Timely Admission Control for Network Slicing in 5G with Machine Learning,” *IEEE Access*, vol. 9, pp. 127595–127610, 2021.
 - [82] J.-B. Monteil, J. Hribar, P. Barnard, Y. Li, and L. A. DaSilva, “Resource Reservation within Sliced 5G Networks : A Cost-Reduction Strategy for Service Providers,” in *2020 IEEE International Conference on Communications Workshops (ICC Workshops)*, (Dublin, Ireland), pp. 1–6, IEEE, 2020.
 - [83] A. Thantharate, R. Paropkari, V. Walunj, and C. Beard, “DeepSlice : A Deep Learning Approach towards an Efficient and Reliable Network Slicing in 5G Networks,” in *Proceedings of the 10th Annual Ubiquitous Computing, Electronics & Mobile Communication Conference (UEMCON)*, (New York, USA), pp. 0762–0767, IEEE, 2019.
 - [84] S. Khan, S. Khan, Y. Ali, M. Khalid, Z. Ullah, and S. Mumtaz, “Highly Accurate and Reliable Wireless Network Slicing in 5th Generation Networks : A Hybrid Deep Learning Approach,” *Journal of Network and Systems Management*, vol. 30, no. 2, pp. 1–22, 2022.
 - [85] J. A. H. Sánchez, K. Casilimas, and O. M. C. Rendon, “Deep Reinforcement Learning for Resource Management on Network Slicing : A Survey,” *Sensors*, vol. 22, no. 8, p. 35459015, 2022.
 - [86] H. D. R. Albonda and J. Pérez-Romero, “An Efficient RAN Slicing Strategy for a Heterogeneous Network With eMBB and V2X Services,” *IEEE Access*, vol. 7, pp. 44771–44782, 2019.

- [87] R. Li, Z. Zhao, Q. Sun, C.-L. I, C. Yang, X. Chen, M. Zhao, and H. Zhang, “Deep Reinforcement Learning for Resource Management in Network Slicing,” *IEEE Access*, vol. 6, pp. 74429–74441, 2018.
- [88] S. Bakri, B. Brik, and A. Ksentini, “On using reinforcement learning for network slice admission control in 5G : Offline vs. online,” *International Journal of Communication Systems*, vol. 34, no. 7, p. e4757, 2021.
- [89] A. Aijaz, “Hap-SliceR : A Radio Resource Slicing Framework for 5G Networks With Haptic Communications,” *IEEE Systems Journal*, vol. 12, no. 3, pp. 2285–2296, 2018.
- [90] B. Khodapanah, A. Awada, I. Viering, A. N. Barreto, M. Simsek, and G. Fettweis, “Slice Management in Radio Access Network via Deep Reinforcement Learning,” in *Proceedings of the 91st Vehicular Technology Conference (VTC2020-Spring)*, (Antwerp, Belgium), pp. 1–6, IEEE, 2020.
- [91] T. Li, X. Zhu, and X. Liu, “An End-to-End Network Slicing Algorithm Based on Deep Q-Learning for 5G Network,” *IEEE Access*, vol. 8, pp. 122229–122240, 2020.
- [92] M. Hao, D. Ye, S. Wang, B. Tan, and R. Yu, “URLLC Resource Slicing and Scheduling in 5G Vehicular Edge Computing,” *IEEE Vehicular Technology Conference*, vol. 2021-April, pp. 0–4, 2021.
- [93] W. Wu, N. Chen, C. Zhou, M. Li, X. Shen, W. Zhuang, and X. Li, “Dynamic RAN Slicing for Service-Oriented Vehicular Networks via Constrained Learning,” *IEEE Journal on Selected Areas in Communications*, vol. 39, no. 7, pp. 2076–2089, 2021.
- [94] D. Yan, B. K. Ng, W. Ke, and C.-T. Lam, “Deep Reinforcement Learning Based Resource Allocation for Network Slicing With Massive MIMO,” *IEEE Access*, vol. 11, pp. 75899–75911, 2023.
- [95] M. A. Tairq, M. M. Saad, M. T. R. Khan, J. Seo, and D. Kim, “DRL-based Resource Management in Network Slicing for Vehicular Applications,” *ICT Express*, vol. 9, no. 6, pp. 1116–1121, 2023.
- [96] X. Chang, T. Ji, R. Zhu, Z. Wu, C. Li, and Y. Jiang, “Toward an efficient and dynamic allocation of radio access network slicing resources for 5g era,” *IEEE Access*, vol. 11, pp. 95037–95050, 2023.
- [97] B. Hazarika, P. Saikia, K. Singh, and C.-P. Li, “Enhancing Vehicular Networks With Hierarchical O-RAN Slicing and Federated DRL,” *IEEE Transactions on Green Communications and Networking*, vol. 8, no. 3, pp. 1099–1117, 2024.
- [98] S. A. AlQahtani, “An efficient resource allocation to improve QoS of 5G slicing networks using general processor sharing-based scheduling algorithm,” *International Journal of Communication Systems*, vol. 33, no. 4, p. e4250, 2019.
- [99] Y. Chen, “An adaptive heuristic algorithm to solve the network slicing resource management problem,” *International Journal of Communication Systems*, vol. 36, no. 8, p. e5463, 2023.

- [100] H. Khan, P. Luoto, M. Bennis, and M. Latva-aho, "On the Application of Network Slicing for 5G-V2X," in *Proceedings of the 24th European Wireless Conference*, (Catania, Italy), pp. 1–6, IEEE, 2018.
- [101] B. Han, J. Lianghai, and H. D. Schotten, "Slice as an Evolutionary Service : Genetic Optimization for Inter-Slice Resource Management in 5G Networks," *IEEE Access*, vol. 6, pp. 33137–33147, 2018.
- [102] X. Yang, Y. Liu, I. C. Wong, Y. Wang, and L. Cuthbert, "Genetic Algorithm for Inter-Slice Resource Management in 5G Network with Isolation," in *Proceedings of International Conference on Software, Telecommunications and Computer Networks (SoftCOM)*, (Split, Croatia), IEEE, 2020.
- [103] D. Marabissi and R. Fantacci, "Highly Flexible RAN Slicing Approach to Manage Isolation, Priority, Efficiency," *IEEE Access*, vol. 7, pp. 97130–97142, 2019.
- [104] R. Adam, "Introduction to the Karush-Kuhn-Tucker (KKT) Conditions," tech. rep., Illinois Institute of Technology, Department of Applied Mathematics, 2018.
- [105] J. Kwak, J. Moon, H.-W. Lee, and L. B. Le, "Dynamic network slicing and resource allocation for heterogeneous wireless services," in *Proceedings of the 28th Annual International Symposium on Personal, Indoor, and Mobile Radio Communications (PIMRC)*, (Montreal, QC, Canada), pp. 1–5, IEEE, 2017.
- [106] N. Yarkina, Y. Gaidamaka, L. M. Correia, and K. Samouylov, "An Analytical Model for 5G Network Resource Sharing with Flexible SLA-Oriented Slice Isolation," *Mathematics*, vol. 8, no. 7, p. 1177, 2020.
- [107] F. Mehmeti and T. F. La Porta, "Optimizing 5G Performance by Reallocating Unused Resources," in *2019 28th International Conference on Computer Communication and Networks (ICCCN)*, (Valencia, Spain), pp. 1–9, IEEE, 2019.
- [108] F. Mehmeti and T. F. La Porta, "Optimizing 5G Performance by Reallocating Unused Resources," in *Proceedings of the 28th International Conference on Computer Communication and Networks (ICCCN)*, (Valencia, Spain), pp. 1–9, IEEE, 2019.
- [109] B. Han, J. Lianghai, and H. D. Schotten, "Slice as an Evolutionary Service : Genetic Optimization for Inter-Slice Resource Management in 5G Networks," *IEEE Access*, vol. 6, no. c, pp. 33137–33147, 2018.
- [110] X. Yang, Y. Liu, I. C. Wong, Y. Wang, and L. Cuthbert, "Genetic Algorithm for Inter-Slice Resource Management in 5G Network with Isolation," in *Proceedings of International Conference on Software, Telecommunications and Computer Networks (SoftCOM)*, (Split, Croatia), pp. 1–6, IEEE, 2020.
- [111] A. N. Saidi and M. Lehsaini, "A Reinforcement Q-Learning-based Resource Sharing Mechanism for V2X slicing Networks," in *Proceedings of the International Conference on Emerging Intelligent Systems for Sustainable Development (ICEIS 2024)*, (Aflou, Algeria), pp. 298–312, Atlantis, 2024.

- [112] A. N. Saidi and M. Lehsaini, “A Deep Q-Learning Approach for an Efficient Resource Management in Vehicle-to-Everything Slicing Environment,” *International Journal of Communication Systems*, vol. 38, no. 4, p. e6137, 2025.
- [113] M. van Otterlo and M. Wiering, *Reinforcement Learning and Markov Decision Processes*, pp. 3–42. Berlin, Heidelberg : Springer Berlin Heidelberg, 2012.
- [114] V. Mnih, K. Kavukcuoglu, D. Silver, A. A. Rusu, J. Veness, M. G. Bellemare, A. Graves, M. Riedmiller, A. K. Fidjeland, G. Ostrovski, S. Petersen, C. Beattie, A. Sadik, I. Antonoglou, H. King, D. Kumaran, D. W. S. Legg, and D. Hassabis, “Human-level control through deep reinforcement learning,” *Nature*, vol. 518, pp. 529–533, 2015.
- [115] G. Nardini, D. Sabella, G. Stea, P. Thakkar, and A. Viridis, “SiMu5G—An OM-NeT++ library for end-to-end performance evaluation of 5G networks,” *IEEE Access*, vol. 8, pp. 181176–181191, 2020.
- [116] G. Nardini, G. Stea, A. Viridis, and D. Sabella, “Simu5G : A System-level Simulator for 5G Networks,” in *Proceedings of the 10th International Conference on Simulation and Modeling Methodologies, Technologies and Applications (SIMULTECH 2020)*, (Online Streaming), pp. 68–80, Springer, 2020.
- [117] M. Behrisch, L. Bieker, J. Erdmann, and D. Krajzewicz, “SUMO – Simulation of Urban MObility : An Overview,” *Proceedings of the Third International Conference on Advances in System Simulation (SIMUL 2011)*, (Barcelona, Spain), pp. 63–68, ARIA, 2011.
- [118] M. Lehsaini, A. N. Saidi, T. Nebbou, and P. Lorenz, “Efficient Routing Protocol Using Fresh Vehicular Traffic Information for VANETs,” *Concurrency and Computation : Practice and Experience*, vol. 37, no. 9-11, p. e70081, 2025.
- [119] T. Nebbou, M. Lehsaini, and H. Fouchal, “Partial backwards routing protocol for VANETs,” *Vehicular Communications*, vol. 18, pp. 1–17, 2019.
- [120] X. Cai, Y. He, C. Zhao, L. Zhu, and C. Li, “LSGO : Link State aware Geographic Opportunistic routing protocol for VANETs,” *Eurasip Journal on Wireless Communications and Networking*, vol. 2014, no. 1, pp. 1–10, 2014.
- [121] A. Hajjej, M. Ayaida, S. Najeh, N. Messai, and L. Najjar, “Robust backbone network based on hybrid selection of relays for multi-hop data dissemination in VANETs,” *Vehicular Communications*, vol. 44, pp. 1–58, 2023.
- [122] M. Jerbi, S. M. Senouci, T. Rasheed, and Y. Ghamri-Doudane, “Towards efficient geographic routing in urban vehicular networks,” *IEEE Transactions on Vehicular Technology*, vol. 58, no. 9, pp. 5048–5059, 2009.
- [123] S. Bilal, S. Madani, and I. Khan, “Enhanced Junction Selection Mechanism for Routing Protocol in VANETs,” *International Arab Journal of Information Technology*, vol. 8, no. 4, pp. 422–429, 2011.

- [124] M. Saravanan and P. Ganeshkumar, “Routing using reinforcement learning in vehicular ad hoc networks,” *Computational Intelligence*, vol. 36, no. 2, pp. 682–697, 2020.
- [125] O. Alzamzami and I. Mahgoub, “Geographic routing enhancement for urban VANETs using link dynamic behavior : A cross layer approach,” *Vehicular Communications*, vol. 31, p. 100354, 2021.

Résumé

Avec l'évolution et l'émergence de nouveaux services visant à améliorer l'expérience des passagers de véhicules, la communication entre diverses entités telles que les véhicules, les unités routières, les systèmes de gestion du trafic et les applications cloud a considérablement augmenté le volume d'informations échangées. Cela a entraîné une plus grande complexité dans la gestion des diverses exigences de qualité de service, notamment dans la partie communication, au sein des réseaux Vehicular-to-Everything (V2X). En conséquence, la technologie 5G est considérée comme une solution indispensable pour prendre en charge efficacement les communications V2X grâce au paradigme de découpage de réseau. Le découpage du réseau permet à la technologie 5G de gérer divers services avec des exigences variées au sein d'une infrastructure partagée, en créant des réseaux virtuels personnalisés isolés, appelés tranches. Cela a conduit les chercheurs à proposer l'intégration du découpage réseau dans les VANETs afin de tirer parti de ses avantages. Cependant, ce paradigme nécessite une certaine adaptation pour s'intégrer efficacement avec les VANETs, qui présentent des spécificités uniques, telles que la mobilité élevée. Parmi les domaines nécessitant une adaptation, la gestion des ressources entre les tranches véhiculaires est particulièrement cruciale, notamment dans la partie radio où la bande passante est une ressource limitée. Dans le cadre de cette thèse, l'objectif est de relever les défis liés à la gestion des ressources dans l'environnement 5G V2X avec le paradigme du découpage réseau. Pour réaliser cet objectif, nous avons proposé un mécanisme de partage des ressources entre les tranches véhiculaires, basé sur l'apprentissage par renforcement profond. Les résultats de simulations de cette contribution ont été comparés à d'autres travaux pertinents en matière de performances et ils ont illustré une certaine supériorité par rapport à ces travaux. Une autre contribution abordée dans le cadre de cette thèse concerne l'amélioration du mécanisme de routage dans les VANETs afin de permettre une dissémination plus efficace et fiable des informations routières.

Mot clés : V2X, 5G, Découpage réseau, Gestion des ressources, Apprentissage par renforcement, VANETs, Protocoles de routage.

Abstract

With the evolution and emergence of new services aimed at improving the experience of vehicle passengers, communication between various entities such as vehicles, roadside units, traffic management systems and cloud applications has considerably increased the volume of information exchanged. This has led to greater complexity in managing the various quality of service requirements, particularly on the communication side, within Vehicular-to-Everything (V2X) networks. As a result, 5G technology is seen as an indispensable solution to efficiently support V2X communications through the network slicing paradigm. Network slicing enables 5G technology to manage various services with varying requirements within a shared infrastructure, by creating isolated custom virtual networks, called slices. This has led researchers to propose the integration of network slicing into VANETs to take advantage of its benefits. However, this paradigm requires some adaptation to integrate effectively with VANETs, which have unique features such as high mobility. Among the areas requiring adaptation, resource management between vehicular slices is particularly crucial, especially in the radio part where bandwidth is a limited resource. In this thesis, the goal is to address the challenges related to resource management in the 5G V2X environment with the network slicing paradigm. To achieve this goal, we proposed a resource sharing mechanism between vehicular slices based on deep reinforcement learning. The simulation results of this contribution have been compared with other relevant works in terms of performance and have shown a certain superiority over these works. Another contribution addressed in this thesis concerns the improvement of the routing mechanism in VANETs in order to allow a more efficient and reliable dissemination of traffic information.

Keywords: V2X, 5G, Network slicing, Resource management, Reinforcement learning, VANETs, Routing protocols.

ملخص

مع تطور وظهور خدمات جديدة تهدف إلى تحسين تجربة ركاب المركبات، أدى التواصل بين مختلف الكيانات مثل المركبات ووحدات الطرق وأنظمة إدارة المرور وتطبيقات السحابة إلى زيادة كبيرة في حجم المعلومات المتبادلة. وقد أدى هذا إلى مزيد من التعقيد في إدارة متطلبات جودة الخدمة المختلفة، وخاصة في جزء الاتصالات، ضمن شبكات المركبات إلى كل شيء (V2X). لذلك، تعتبر تقنية الجيل الخامس 5G حلاً لا غنى عنه لدعم اتصالات V2X بكفاءة من خلال نموذج تقطيع الشبكة. يتيح تقطيع الشبكة لتقنية الجيل الخامس إدارة خدمات متنوعة ذات متطلبات مختلفة ضمن بنية تحتية مشتركة من خلال إنشاء شبكات افتراضية معزولة ومخصصة، تسمى الشرائح. وقد دفع هذا الباحثين إلى اقتراح دمج تقسيم الشبكة في شبكات VANETs من أجل الاستفادة من فوائدها. ومع ذلك، يتطلب هذا النموذج بعض التكيف ليتكامل بشكل فعال مع شبكات VANETs، التي تتمتع بخصائص فريدة، مثل القدرة العالية على الحركة. ومن بين المجالات التي تتطلب التكيف، تعد إدارة الموارد بين شرائح المركبات أمراً بالغ الأهمية، وخاصة في الجزء اللاسلكي حيث يكون النطاق الترددي مورداً محدوداً. في هذه الأطروحة، الهدف هو معالجة تحديات إدارة الموارد في بيئة 5G V2X مع نموذج تقطيع الشبكة. ولتحقيق هذا الهدف، اقترحنا آلية لتقاسم الموارد بين شرائح المركبات، استناداً إلى التعلم التعزيزي العميق. تمت مقارنة نتائج المحاكاة لهذه المساهمة مع أعمال الأداء الأخرى ذات الصلة وأظهرت بعض التفوق على هذه الأعمال. تتضمن مساهمة أخرى تناولتها هذه الأطروحة تحسين آلية التوجيه في شبكات الطرق السريعة من أجل السماح بنشر معلومات الطرق بشكل أكثر كفاءة وموثوقية.

الكلمات المفتاحية : V2X، 5G، تقطيع الشبكة، إدارة الموارد، التعلم التعزيزي، شبكات VANETs، بروتوكولات التوجيه.