



D4.3_Final results on the management of radio resources in
subnetworks_v1.0

Dissemination Level: PU



Project: 101095738 – 6G-SHINE-HORIZON-JU-SNS-2022

Project no.:	101095738		
Project full title:	6G Short range extreme communication IN Entities		
Project Acronym:	6G-SHINE		
Project start date:	01/03/2023	Duration	30 months

D4.3 – FINAL RESULTS ON THE MANAGEMENT OF RADIO RESOURCES IN SUBNETWORKS IN THE PRESENCE OF LEGITIMATE AND MALICIOUS INTERFERERS

Due date	30/06/2025	Delivery date	30/06/2025
Work package	WP4		
Responsible Author(s)	Saeed Hakimi (AAU)		
Contributor(s)	Gilberto Berardinelli (AAU), Pedro Maia de Sant Ana (BOSCH), Fotis Foukalas (COGN), Christos Tsakos (COGN), Yasser Mestrah (IDE), Filipe Conceicao (IDE), Renato Barbosa Abreu (Nokia)		
Version	V1.0		
Reviewer(s)	Anders Berggren (Sony), Spilios Giannoulis (IMEC)		
Dissemination level	Public		



Horizon Europe Grant Agreement No. 101095738. Views and opinions expressed are however those of the author(s) only and do not necessarily reflect those of the European Union or SNS JU. Neither the European Union nor the granting authority can be held responsible for them.

VERSION AND AMENDMENT HISTORY

Version	Date (MM/DD/YYYY)	Created/Amended by	Changes
0.1	17/01/2025	Saeed Hakimi	ToC finalization
0.2	03/03/2025	Saeed Hakimi	Initial version
0.3	01/04/2025	Saeed Hakimi	Abbreviation list completed
0.4	01/05/2025	Saeed Hakimi	Editor revision and feedback
0.5	05/05/2025	Saeed Hakimi	Editor review completed
0.6	12/05/2025	Saeed Hakimi	Version ready for internal review
0.7	06/06/2025	Anders Berggren, Spilios Giannoulis	Feedback from the internal reviewers
0.8	10/06/2025	Saeed Hakimi	Updates based on reviewers' comments
0.9	23/06/2025	Berit H. Christensen	Final proofreading and layout check
1.0	30/06/2025	Saeed Hakimi	Submitted version

TABLE OF CONTENTS

FIGURES	5
TABLES	7
ABBREVIATIONS	8
EXECUTIVE SUMMARY	10
1 INTRODUCTION	11
1.1 FOUNDATIONS AND MOTIVATION FOR ADVANCED RRM IN 6G SUBNETWORKS.....	11
1.2 OVERVIEW OF FINALIZED RRM CONTRIBUTIONS.....	11
2 ADDRESSING 6G-SHINE USE CASES AND TARGETS	15
3 RRM FOR IN-X SUBNETWORKS	20
3.1.1 Spatio-Temporal Attention-Based model for RRM in outdated CSI.....	20
3.1.2 Resilient DNN for Joint Sub-Band and Power Control in Mobile In-FS.....	29
3.1.3 Simulation Results and Analysis for proposed Centralized RRM	34
3.2 DISTRIBUTED RRM FOR IN-X SUBNETWORKS	42
3.2.1 System Model	43
3.2.2 Simulation Setup.....	46
3.2.3 Results and Analysis	49
3.3 SUMMARY	50
4 GOAL ORIENTED RRM	52
4.1 VELOCITY CONTROL FOR INTER-SUBNETWORK INTERFERENCE MITIGATION IN MOBILE SUBNETWORKS.....	52
4.2 SYSTEM MODEL	53
4.2.1 Reliability Constraints (BLER Model)	53
4.3 ROBOT MOBILITY AND SPEED CONTROL PROBLEM	54
4.3.1 Robot Motion Model.....	54
4.3.2 Optimization problem formulation.....	54
4.4 REINFORCEMENT LEARNING APPROACH.....	55
4.4.1 PPO-Based Speed Control.....	55
4.5 EVALUATION AND RESULTS.....	55
4.5.1 Implementation and training.....	55
4.5.2 Simulation Setup.....	56
4.5.3 Performance Evaluation and Discussion	57
4.6 SUMMARY	59
5 ENABLERS FOR RRM IN SUBNETWORKS	60
5.1 HC-HC AND SNE-HC SIDELINK COMMUNICATION IN A SHARED BAND.....	61
5.2 CONSIDERATIONS FOR SUPPORTING SEMI-STATIC CHANNEL ACCESS IN UNLICENSED BANDS	66
5.3 OPPORTUNISTIC USAGE OF LICENSED AVAILABLE RESOURCES FOR SUBNETWORKS.....	71
5.4 SUMMARY	75
6 DETECTION AND MITIGATION MECHANISM OF EXTERNAL INTERFERENCE	77
6.1 ROBUST RADIO RESOURCE MANAGEMENT FOR IN-FACTORY SUBNETWORKS UNDER EXTERNAL INTERFERENCE	78
6.1.1 External Interference Model and Problem Formulation	79
6.1.2 Gradient Descent-based Resource Allocation Algorithm	80
6.1.3 Modified SISA-SIPA Algorithm	81
6.1.4 Simulation results and analysis for RRM in the presence of external interference	83
6.2 PERFORMANCE EVALUATION FRAMEWORK FOR EFFICIENT RECEIVER ADAPTATION OVER SUBNETWORKS.....	87

6.2.1	Receiver Design approaches.....	88
6.2.2	Approximation functions.....	88
6.2.3	Parameter estimation	90
6.2.4	In subnetwork receiver approximation	92
6.3	SUMMARY	94
7	CONCLUSIONS.....	95
	REFERENCES	97

FIGURES

Figure 1: Reference deployment architecture for the methods studied in this deliverable.....	12
Figure 2: Illustration of the Subnetwork Co-existence in Factory Hall Use Case.....	16
Figure 3: Illustration of immersive education showing some potential hierarchical subnetworks	18
Figure 4: Illustration of indoor interactive gaming within a subnetwork	19
Figure 5: System model illustrating the deployment of in-factory subnetworks	22
Figure 6: Illustration of the CSI update process, highlighting the delay between the last CSI measurement, decision-making, and the subsequent reconfiguration interval.	23
Figure 7: Overall framework of the proposed solution, utilizing delayed CSI as input to generate predicted CSI as output.....	25
Figure 8: Structure of the dual attention-based LSTM encoder-decoder, integrating spatial and temporal attention mechanisms for enhanced channel prediction.	28
Figure 9: Structure of the proposed DNN-based RRM framework, showing interconnected modules for sub-band allocation and power control, supported by preprocessing and loss optimization components.	31
Figure 10: Evolution of the probability of minimum SE violation and average SE across all subnetworks as a function of the RRM training epochs, illustrating the model's convergence behaviour and its ability to balance QoS constraints with system efficiency.	37
Figure 11: CDF of the binarization error for the RRM model, highlighting the effectiveness of the proposed framework in generating binary outputs for sub-band allocation with minimal deviation....	38
Figure 12: CDF of the average SE across all subnetworks, comparing the proposed DNN-based RRM and the benchmark under varying delay conditions	39
Figure 13: CDF of the individual SE across all subnetworks, comparing the proposed DNN-based RRM and the benchmark under varying delay conditions.	39
Figure 14: Training and validation loss trends for the dual attention LSTM and standard LSTM predictors.	40
Figure 15: CDF of the average SE across all subnetworks for DNN-based RRM with a $\tau = 4$ -sample delay, comparing sample-and-hold, Attention-LSTM, and LSTM predictors.	41
Figure 16: CDF of the SE for all subnetworks for DNN-based RRM with a $\tau = 4$ -sample delay, comparing sample-and-hold, Attention-LSTM, and LSTM predictors.....	41
Figure 17: AI/ML RRM through power control.....	43
Figure 18: MPNN framework's phases.....	45
Figure 19: Distributed power control through MPNN for interference management in 6G subnetworks.	46
Figure 20: Overall rate during model training for train and test data	48
Figure 21: The block diagram of the implementation setup shows how the processes exchange data between them.	49
Figure 22: Sum-Rate of the considered 6G subnetwork over SINR1 (pair-1) for (A) one, (B) two, and (C) three interfering pairs with different SINR, such as 5, 10 and 20 dB considering the MPNN (GNN) and the EPA solutions.	50

Figure 23: Each mobile robot carries a subnetwork to facilitate the wireless communication between devices in the robot. There is strong mutual interference between the subnetworks due to the proximity of the robots on the factory floor.....	52
Figure 24: Mobility model of the robot. The state consists of the lateral error e , the yaw angular error θ , their derivatives and the velocity error v	54
Figure 25: Cumulative reward per episode during the PPO training phase of the communication-aware dynamic speed control policy.....	56
Figure 26: The average arrival time for a travel distance of 25m vs the probability.....	58
Figure 27: The Complementary Cumulative Distribution Function (CCDF) of the achieved block error rate.	59
Figure 28: Example of subnetwork deployment where SNE-HC (blue) and HC-HC (orange) communication may coexist.....	63
Figure 29: UE satisfaction ratio for different number of subnetworks in unlicensed band with Type1/Type2 channel access.....	66
Figure 30: Example of timing of the channel access mechanism for FBE [41]	67
Figure 31: Modified sidelink slot configuration considered for semi-static channel access.....	68
Figure 32: UE satisfaction ratio for different number of subnetworks in unlicensed band with semi-static channel access.	68
Figure 33: IBE-aware Inter-UE coordination scheme (Red steps highlight impact on existing IUC procedure).....	70
Figure 34: UE satisfaction ratio for different number of subnetworks using IBE aware resource coordination.	71
Figure 35: Example of self-interference which can lead to sensitivity degradation.....	72
Figure 36: Example of signalling where the HC device acting as access point of a subnetwork obtains radio resource from NW.	73
Figure 37: UE satisfaction ratio for different number of subnetworks in licensed band.....	74
Figure 38: Outage probability of individual links across all subnetworks, under three scenarios: No Interference, Interference-Unaware, and Interference-Aware.	85
Figure 39: CDF of the average SE across all subnetworks, under two scenarios: Interference-Unaware and Interference-Aware.....	85
Figure 40: Outage probability of individual links across all subnetworks, under varying levels of external interference power.....	86
Figure 41: Outage probability of individual links across all subnetworks, under scenarios with no interference and with interference activated on 1, 2, and 3 sub-bands.	86
Figure 42 CDF of transmit powers across all subnetworks, under varying levels of external interference power.	87
Figure 43: Comparison of the optimal LLR shape with different approximations.....	89
Figure 44: Optimal region and BER as a function of a and b	91
Figure 45: Example procedure for reporting the best function approximating LLR and its validity window	93

TABLES

Table 1: Simulation parameters for centralized RRM.....	35
Table 2: Neural Network’s hyper-parameters.....	47
Table 3: Simulation Assumption.....	56
Table 4: Summary of evaluation assumptions.	64
Table 5: Simulation Parameters for RRM under external interference	83
Table 6: Execution Time per Sample for Different Batch Sizes and Algorithms.....	87

ABBREVIATIONS

Acronym	Description
3GPP	3rd Generation Partnership Project
AGC	Automatic Gain Control
AGV / AGVs	Automated Guided Vehicle(s)
AP	Access Point
BLER	Block Error Rate
CC	Centralized Controller
CCA	Clear Channel Assessment
CGC	Centralized Graph Colouring
CSI	Channel State Information
D2.2	Deliverable 2.2: Refined Definition of Scenarios, Use Cases, and Service Requirements for in-X Subnetworks
DNN	Deep Neural Network
ECDF	Empirical Cumulative Distribution Function
eMBB	Enhanced Mobile Broadband
EN	Entities
FBE	Frame-Based Equipment
FFP	Fixed Frame Period
FNN	Fully Connected Neural Network
FR1	Frequency Range 1 (sub-6 GHz)
GDRA	Gradient-Descent with Random Access
GNN	Graph Neural Network
HC / HCs	High-Capability Element(s)
IBE	Intra-Band Emissions
InF-S	In-Factory Subnetworks
ISR	Interference-to-Signal Ratio
IUC	Inter-UE Coordination
KPI / KPIs	Key Performance Indicator(s)
KVI / KVIs	Key Value Indicator(s)
LBE	Load-Based Equipment
LBT	Listen-Before-Talk
LC	Low Capabilities
LL	Log-Likelihood
LLR	Log-Likelihood Ratio
MCS	Modulation and Coding Scheme
ML	Machine Learning
MDP	Markov Decision Process
OCB	Occupied Bandwidth

PSCCH	Physical Sidelink Control Channel
PSD	Power Spectral Density
QoS	Quality of Service
RA	Random Allocation
RB	Resource Block
RL	Reinforcement Learning
RRC	Radio Resource Control
RRM	Radio Resource Management
RSRP	Reference Signal Received Power
SCS	Subcarrier Spacing
SE	Spectral Efficiency
SINR	Signal-to-Interference-plus-Noise Ratio
SISA	Sequential Iterative Sub-band Allocation
SN	Subnetwork
SNE	Subnetwork Element
STA	Spatio-Temporal Attention
SoA	State-of-the-Art
UE	User Equipment
URLLC	Ultra-Reliable and Low-Latency Communication
XR	Extended Reality

EXECUTIVE SUMMARY

The 6G-SHINE project develops innovative solutions for managing radio resources in dense and dynamic subnetwork environments, where multiple mobile or static subnetworks must operate reliably, autonomously, and with minimal interference. This deliverable, D4.3, presents the final technical achievements in the domain of radio resource management (RRM) for in-X subnetworks, addressing critical challenges such as scalability, interference management, latency assurance, and operational robustness under realistic deployment conditions.

Building upon the preliminary investigations of D4.1, this document advances both centralized and distributed RRM strategies, tackles external interference challenges, and introduces enabling technologies to strengthen intra- and inter-subnetwork communication.

Specifically, D4.3 presents:

- Centralized RRM techniques based on spatio-temporal attention-based channel prediction and resilient deep neural network resource allocation, enabling proactive adaptation to channel state information (CSI) delays while balancing spectral efficiency and fairness.
- Distributed and hybrid RRM solutions that rely on graph-based neural networks and decentralized coordination mechanisms, allowing subnetworks to autonomously optimize their resource usage even under limited or intermittent parent network connectivity.
- Goal-Oriented RRM approaches, where optimization targets shift from traditional communication metrics (e.g., SINR, throughput) to application-specific objectives such as minimizing robot mission time or ensuring control stability. Reinforcement learning methods, such as Proximal Policy Optimization (PPO), are deployed to dynamically adjust mobility patterns and resource usage based on observed network states.
- Spectrum access strategies for operation in licensed, unlicensed, or shared bands, including semi-static access schemes, sidelink-based resource coordination, and mechanisms to mitigate in-band emissions for increased spectral coexistence.
- External interference detection and mitigation frameworks, supporting robustness against natural, unintentional, or malicious interference sources through interference-aware resource allocation, jammer detection methods, and robust receiver designs.
- Adaptive receiver design techniques, leveraging likelihood ratio (LLR) approximations and unsupervised learning to maintain reliable communication even under impulsive or non-Gaussian interference conditions.

The developed solutions are systematically mapped to the 6G-SHINE project's defined use cases and technical objectives, such as latency, reliability, scalability, and spectral efficiency. They are designed to operate flexibly across a variety of deployment scenarios - centralized or decentralized - and spectrum regimes.

Overall, this deliverable represents a significant step toward enabling ultra-reliable, scalable, and efficient operation of 6G subnetworks in future wireless ecosystems.

1 INTRODUCTION

1.1 Foundations and Motivation for Advanced RRM in 6G Subnetworks

This document presents the final findings on radio resource management (RRM) strategies developed within the 6G-SHINE project, focusing on the highly dynamic and interference-sensitive environments of in-X subnetworks. The aim is to optimize performance under two constraints: (1) meeting the stringent requirements of latency-critical applications, and (2) coping with both internal (legitimate) and external (uncontrolled or malicious) interference sources.

The allocation of radio parameters such as power levels, spectrum bands, time slots, and modulation schemes becomes a highly complex task in these environments. This complexity is exacerbated by physical constraints such as signal blockage, rapidly changing channel conditions due to mobility, and the lack of coordination among densely coexisting subnetworks. The extreme connection density expected in 6G, projected to be approximately 10 times greater than in 5G deployments [1], introduces unprecedented challenges in sustaining communication quality, especially when multiple autonomous entities operate within the same physical and spectral space.

A key characteristic of future 6G subnetworks is their ability to function in both centralized and decentralized modes. Centralized RRM strategies can take advantage of a global view of the network to enable predictive and harmonized resource allocation. However, such strategies are limited by practical constraints such as CSI update delays and signalling overhead - issues that become particularly pronounced in rapidly evolving industrial and mission-critical scenarios.

To ensure robust and scalable operation, decentralized approaches, where each subnetwork node adapts its behaviour based on local observations, are highly valuable. This adaptability is essential in use cases such as factory automation, autonomous vehicle swarms, and immersive consumer experiences, where responsiveness and resilience cannot rely solely on centralized infrastructure.

Beyond managing interference among legitimate users, the critical issues of external interference -such as electromagnetic noise, cross-technology protocol collisions, and deliberate jamming - must also be effectively addressed. These disruptive factors can severely impact service quality and must be detected, characterized, and mitigated in real time to ensure continuous and reliable communication.

The results presented in this deliverable provide a cohesive set of technical advancements that support the efficiency, scalability, and adaptability targets of 6G subnetworks, across both licensed and unlicensed spectrum regimes. They make a substantial step forward in enabling autonomous, resilient, and intelligent RRM under the realities of future wireless environments.

1.2 Overview of Finalized RRM Contributions

Focusing on the challenges posed by dense, dynamic, and interference-prone in-X subnetwork environments, the finalized RRM solutions developed within the 6G-SHINE project are designed to ensure service continuity and quality under real-world constraints such as mobility-induced channel variability, unpredictable external interference, limited signalling capacity, and stringent latency/reliability requirements typical of mission-critical and immersive applications.

The work builds upon the preliminary findings in Deliverable D4.1 [2] and consolidates the methods into a comprehensive set of strategies addressing both centralized and decentralized RRM needs. The contributions span a spectrum of enablers - from proactive channel prediction to goal-oriented control optimization, from distributed learning to robust receiver adaptation - each designed to function under distinct deployment and coordination models.

The reference deployment architecture, introduced in D4.1 and shown in Figure 1, remains foundational to the structure of the developed RRM solutions. It illustrates the hierarchical and modular nature of in-X subnetwork deployments within the 6G-SHINE framework. Each entity (EN), such as a robot, vehicle, or production module, hosts one or more subnetworks (SN), composed of subnetwork elements (SNEs) coordinated by a high-capability controller (HC).

At the core of the deployment is the 6G base station (6G BS), functioning as a parent network with high processing capabilities. The 6G BS integrates compute nodes and RRM modules capable of coordinating multiple subnetworks under its coverage. When full connectivity is available, centralized RRM decisions can be made based on a global network view. However, in cases where SNs are disconnected from the 6G BS or low-latency coordination is infeasible, HCs are empowered to perform autonomous RRM, relying on local observations and goal-oriented decision-making. This dual-mode operation (combining centralized RRM via the 6G BS and decentralized RRM via local HCs) can be enabled by NR sidelink evolution and further enhanced by mechanisms for interference mitigation, including those targeting malicious jamming and cross-technology collisions.

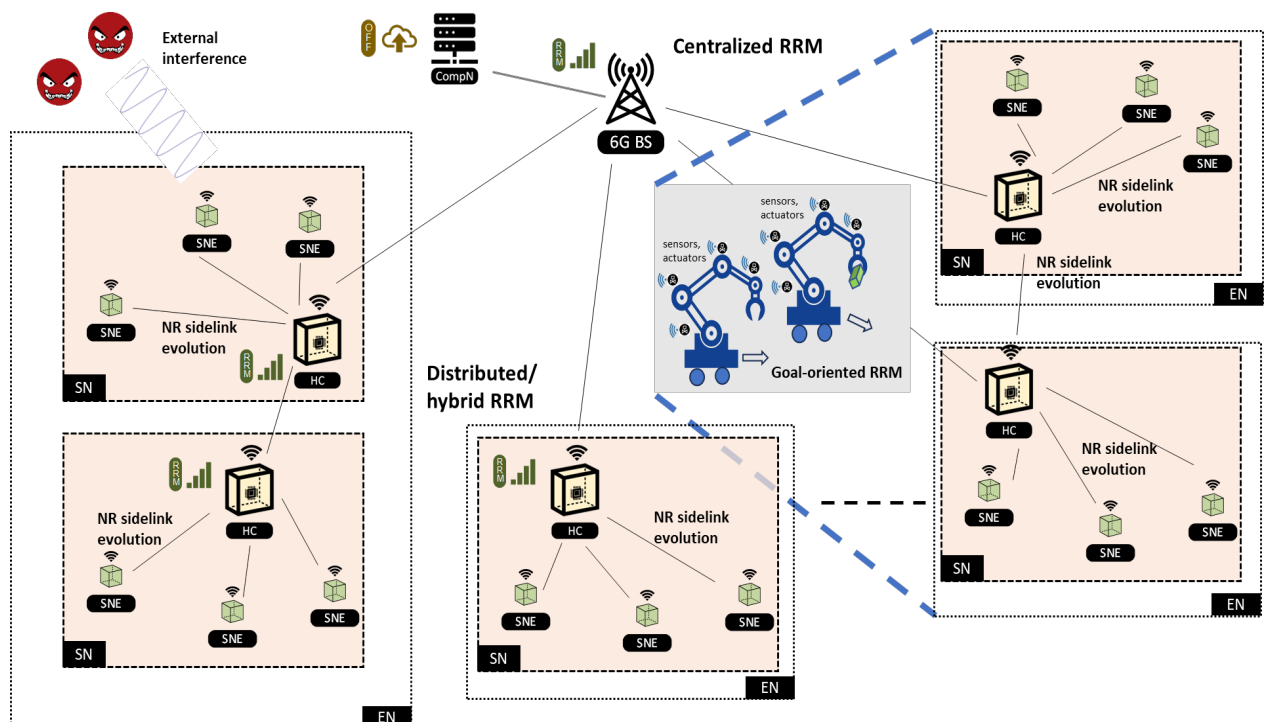


Figure 1: Reference deployment architecture for the methods studied in this deliverable.

This document is organized to highlight not only the technical solutions themselves but also the logical dependencies between them and the broader objectives of the 6G-SHINE project. The content is structured across six chapters: an introduction, a chapter mapping contribution to use cases and targets, and four chapters detailing the finalized RRM contributions, each addressing a key dimension of RRM in highly dynamic subnetwork environments. The solutions reflect a mix of centralized, distributed, and hybrid approaches that consider system-level limitations such as CSI aging, signalling delays, lack of central coordination, and interference from legitimate or malicious sources.

- Chapter 2 maps each contribution to the use cases defined in Deliverable D2.2 and explicitly connects the proposed methods to the technical targets stated in the project proposal. These targets include reliability, latency, scalability, and spectral efficiency. The use cases span industrial, and consumer categories, with examples like *Subnetwork Swarms in Factory Halls* and *Indoor Immersive Education* illustrating the link between proposed solutions and real-world deployment scenarios.
- Chapter 3 introduces the core radio RRM algorithms for joint sub-band and power control, covering both centralized and distributed paradigms. The centralized methods leverage spatio-temporal attention-based CSI prediction to proactively counteract CSI aging and optimize resource allocation using resilient deep neural networks. In parallel, the chapter presents distributed strategies based on Graph Neural Networks (GNNs) and over-the-air aggregation, enabling each subnetwork to optimize its transmit power locally with minimal coordination overhead. These approaches are designed to support scalability and low-latency decision-making in highly dynamic and dense deployments, where centralized control may be infeasible. All methods are evaluated under realistic assumptions regarding CSI delays, device density, and network dynamics, providing insights into their performance in both industrial and consumer scenarios.
- Chapter 4 explores goal-oriented RRM, where optimization is based not only on communication-centric KPIs like throughput or SINR but also on application-specific performance metrics such as mission time or control error. The chapter introduces reinforcement learning (RL)-based methods, particularly using Proximal Policy Optimization (PPO), to jointly adapt robot mobility patterns and RRM parameters in response to network feedback. These co-design methods are crucial for mission-critical applications where network and control performance are interdependent.
- Chapter 5 turns attention to spectrum access and coexistence in unlicensed or mixed-band environments. It presents methods such as semi-static channel access in sidelink, IBE-aware resource coordination, and licensed-assisted operation to improve spectrum usage and reduce latency in dense environments. These mechanisms enable 10× improvements in XR capacity and support higher subnetwork densities compared to baseline sidelink access and resource allocation schemes.
- Chapter 6 focuses on the detection and mitigation of external interference, including uncoordinated cross-technology interferers and malicious jammers. A Gradient Descent-based Resource Allocation (GDRA) algorithm is introduced for joint power control and sub-band selection under probabilistic interference models. Furthermore, this chapter presents a novel

receiver-side adaptation framework, including LLR approximation and adaptive demapper selection, to maintain reliable decoding in impulsive or jammer-affected noise conditions.

A key characteristic of these solutions is their adaptability. Depending on the deployment scenario, RRM decisions may be shaped by:

- the availability and freshness of CSI (e.g., full, delayed, or partial),
- the spectrum regime (licensed, unlicensed, or hybrid),
- the coordination mode (centralized, distributed, or hybrid),
- the interference source (legitimate, cross-technology, or malicious),
- and the application's QoS constraints (e.g., reliability, latency, scalability).

To ensure feasibility in practical deployment scenarios, many of the proposed methods are designed under realistic signalling assumptions, including limited feedback, sidelink-based coordination, and lightweight over-the-air communication protocols. Each solution is benchmarked against relevant state-of-the-art baselines and analysed in simulation environments that mirror industrial and consumer settings with high device density and mobility.

2 ADDRESSING 6G-SHINE USE CASES AND TARGETS

This chapter presents a detailed mapping between the use cases defined in Deliverable D2.2: *“Refined definition of scenarios use cases and service requirements for in-X subnetworks”* [3] and the technical advancements introduced in the current deliverable. It also highlights how these contributions address the objectives and targets of the 6G-SHINE project, with particular emphasis on objective 5: *“Develop cost-effective centralized, distributed, or hybrid radio resource management techniques (considering both legitimate and malicious interferers) in hyper-dense dynamic subnetwork deployments.”*

Deliverable D2.2 outlines several in-X subnetwork use cases across different categories. Among these, the most relevant to this deliverable is the "Subnetwork Swarms: Subnetwork Co-existence in Factory Hall" use case from the industrial subnetwork category (Figure 2). This use case is characterized by multiple mobile subnetworks, installed in robots, operating in proximity, leading to severe inter-subnetwork interference.

Key challenges include:

- Real-time adaptation of radio resource management to mobility-induced CSI variations;
- Guaranteed QoS for latency- and reliability-sensitive operations;
- Distributed interference coordination with minimal signalling between SNEs and HC, and among HCs.

To address mobility-induced CSI delay, one of the primary challenges in this scenario, a centralized Spatio-Temporal Attention-Based Channel Prediction method has been developed. This approach mitigates the adverse effects of outdated CSI, which otherwise leads to inefficient sub-band allocation, suboptimal power control, and poor interference management, ultimately degrading spectral efficiency and QoS. We consider a scenario where a central controller can perform decisions for the industrial mobile subnetworks in the swarm. The proposed mechanism leverages dual attention across space and time to accurately forecast future CSI. This capability allows the RRM to proactively adapt to changing channel conditions, which is critical for time-sensitive tasks such as robot coordination and AGV routing in factory halls.

In parallel, a Resilient DNN-based joint sub-band and power allocation scheme has been designed to manage radio resources while maintaining fairness and QoS.

Together, these two components form a centralized RRM solutions that directly supports objective 5 by providing scalable and intelligent RRM strategies that strike a balance between spectral efficiency and fairness in dense, dynamic environments. They also align with broader project objectives related to robustness and adaptability to environmental changes. As discussed in greater detail in deliverable, the proposed solution achieves a minimum spectral efficiency (SE) that is 53% higher than the SoA without a predictor, and 94% higher than the SoA with a predictor in scenarios with 4-sample delayed CSI. These results indicate that the proposed solution is approaching the associated target ($\sim 2\times$ improvement over SoA) and substantially enhances the reliability and effectiveness of RRM in mobile subnetwork scenarios. These results are obtained at a density of 25,000 subnetworks per km^2 , aligning with the scalability

target of achieving approximately 10× higher cell density compared to typical 5G ultra-dense deployments (~2,500 cells per km² [4],[5])

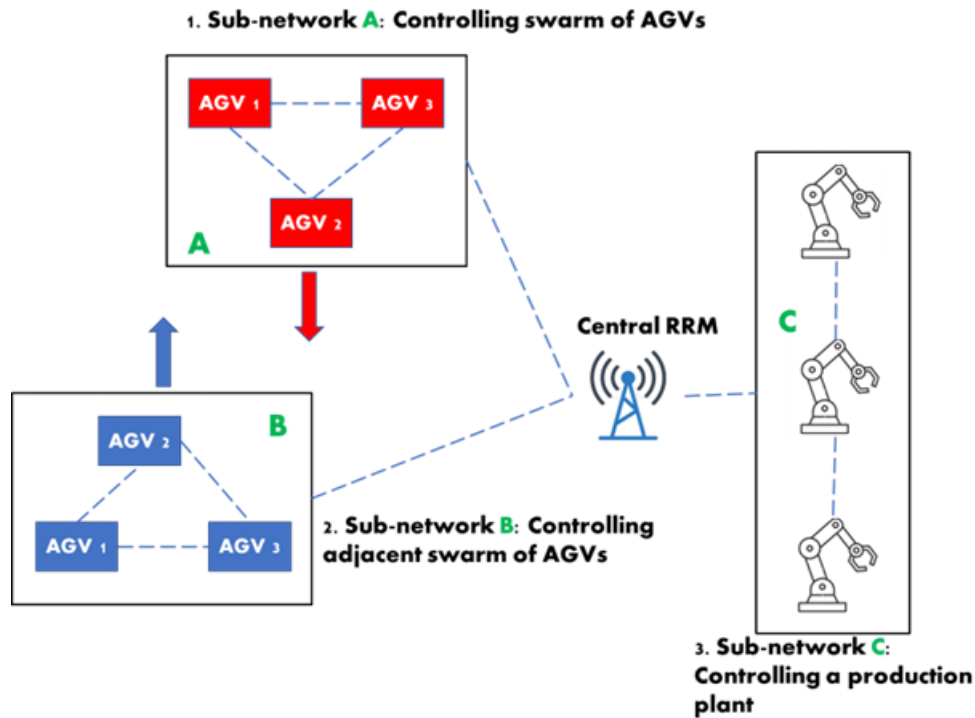


Figure 2: Illustration of the Subnetwork Co-existence in Factory Hall Use Case

Furthermore, to tackle inter-subnetwork interference in case no central controller is available, a key challenge in the "Subnetwork Swarms" use case, a distributed RRM scheme based on GNNs and over-the-air aggregation has been introduced. This approach enables each subnetwork to optimize its transmit power independently while preserving overall spectral efficiency and minimizing coordination overhead. The strategy is scalable, low-overhead, and compliant with 3GPP protocols, making it suitable for industrial deployments. Experimental evaluation through a 3GPP compliant platform shows that the proposed framework achieves a 7% improvement in spectral efficiency compared to Equal Power Allocation, with gains reaching 13.16% under heterogenous channel conditions.

Subnetwork swarms are also typically involved in a common mission, to be completed in a minimum time. In other terms, each mobile robot has its own installed subnetwork for local control tasks, while they are all moving towards a common mission. Therefore, intelligent mobility adaptation is explored to mitigate communication degradation caused by robot dynamics and generated interference. A reinforcement learning-based speed control algorithm adjusts robot mobility patterns in response to SINR feedback. This mechanism maintains URLLC performance with modest latency increases at the mobile robot mission time, achieving a 20% higher probability of meeting the same block error rate

(BLER) compared to state-of-the-art methods. This co-design of mobility and RRM represents a significant advancement toward adaptive, reliable and robust radio management in industrial subnetworks.

The deliverable also addresses consumer-centric use cases from D2.2, specifically, indoor immersive education and interactive gaming, illustrated in Figure 3 and Figure 4. These represent high-density, low-latency environments with multiple tightly localized devices (e.g., VR headsets, sensors, consoles) requiring consistent, real-time communication. Such deployments pose significant challenges for RRM, particularly in unlicensed or mixed-spectrum bands, where interference, contention, and unpredictable latency must be tightly controlled.

Although the evaluation focuses on consumer use cases, the underlying mechanisms are broadly applicable to other use cases. The focus on consumer use cases is motivated by three key factors:

1. Sidelink as the foundation for subnetworks: Our design builds on 3GPP sidelink, and recent RAN1 work highlights growing industry interest in expanding NR sidelink to commercial use cases - making consumer scenarios a natural reference point.
2. Practical evaluation across spectrum regimes: Consumer deployments are most likely to operate in unlicensed bands, which are cost-free, globally harmonized, and expanding (e.g., up to 1.2 GHz in the 6 GHz band). These attributes allow for meaningful performance comparisons across licensed and unlicensed settings.
3. Relevance to 6G standardization: As emphasized in the approved 6G work item on “Study on 6G Scenarios and Requirements”, consumer broadband services are expected to guide core radio design decisions, with additional adaptability to vertical needs. Thus, technical enablers validated in consumer settings are highly likely to shape future 6G standards.

Building on the 3GPP sidelink framework introduced as a baseline in Deliverable D4.1, this deliverable identifies key limitations of existing solutions that hinder efficient and scalable subnetwork operation in dense deployment scenarios. These limitations pose challenges to achieving reliable, low-latency, and high-data rate communication required for emerging consumer and industrial applications. The main issues are:

- In-band emissions (IBE) that degrade reception quality between adjacent resource blocks due to front-end imperfections or poor isolation. In dense deployments where multiple subnetworks operate in close proximity and frequently reuse neighbouring frequency resources, IBE can significantly degrade the signal-to-interference-plus-noise ratio (SINR), especially for low-power devices. This not only limits achievable data rates but also undermines the reliability of latency-sensitive links such as those used for XR control or sensor data exchange.
- Dynamic channel access mechanisms (e.g., LBT with random backoff) that introduce random and often excessive access delays under high traffic conditions. These mechanisms are mandated in unlicensed bands to ensure fair coexistence across different technologies. However, under high traffic loads and in dense subnetwork deployments, the randomized nature of LBT introduces unpredictable and often excessive delays. These delays are especially problematic for applications with strict latency budgets.

- Lack of coordination in shared spectrum, which limits scalability and reliability for intra- and inter-subnetwork communication. In many subnetwork deployments, particularly in unlicensed bands, there is no central authority to coordinate access among multiple subnetworks. This lack of coordination leads to unstructured and overlapping resource usage, increasing the probability of collisions, interference, and inefficient spectrum utilization. Intra-subnetwork coordination and inter-subnetwork coordination are both negatively affected, limiting the overall scalability and network reliability as the number of subnetworks grows.

To support these challenging scenarios and meet the project's scalability target of approximately 10× higher subnetwork density than typical 5G ultra-dense deployments (i.e., ~2,500 cells per km² [4][5]), the following RRM enhancements are introduced:

- IBE mitigation strategies, such as UE front-end enhancements and IBE-aware coordination, increasing the number of supported subnetworks by up to 40% and 19%, respectively, in dense unlicensed deployments.
- Semi-static channel access, adapted for sidelink operation, provides deterministic communication and enables 10× higher XR capacity than traditional dynamic access under strict latency constraints in shared bands.
- Opportunistic use of licensed spectrum, negotiated between a subnetwork HC and the parent network, eliminates LBT delays and reduces self-interference, supporting up to 67% more subnetworks compared to semi-static access in unlicensed bands.

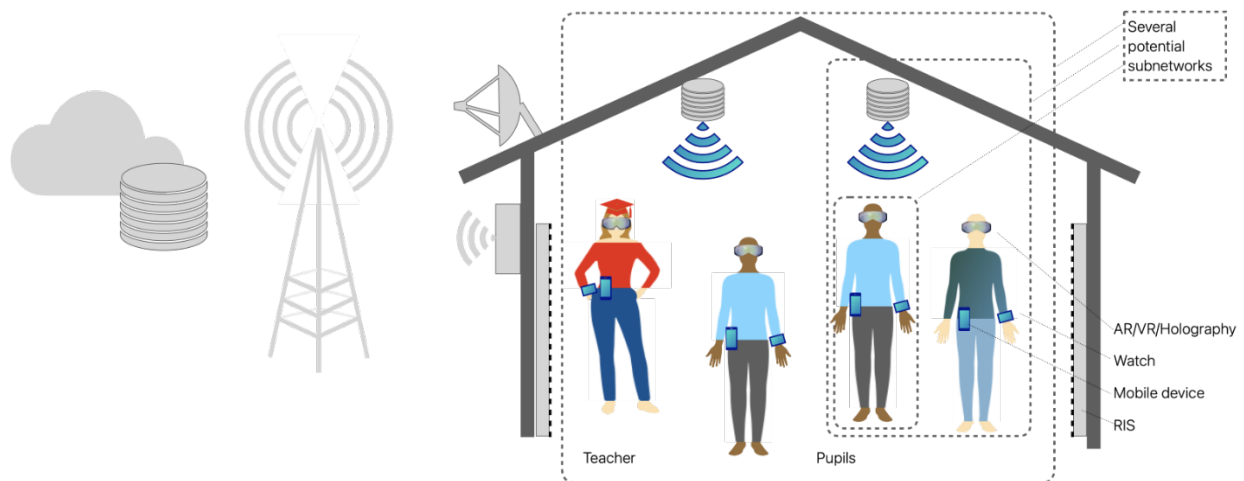


Figure 3: Illustration of immersive education showing some potential hierarchical subnetworks

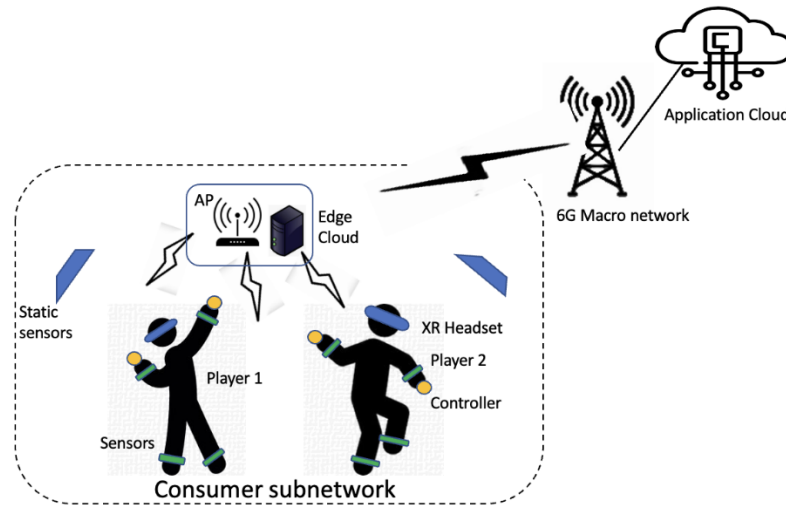


Figure 4: Illustration of indoor interactive gaming within a subnetwork

While inter-subnetwork interference is managed through coordinated RRM, external interference, from unmanaged industrial devices or malicious jammers, remains a critical challenge in realistic deployments. To address this, a robust RRM algorithm has been developed, formulated as a joint sub-band and power allocation problem and solved using the proposed Gradient-Descent with Random Access (GDRA) method [6].

This solution supports Objective 5 by enabling reliable subnetwork operation under unpredictable interference conditions in both industrial and consumer indoor environments. Specifically, with external interference set 20 dB below the subnetwork's maximum power and active on 50% of the subbands, the algorithm limits spectral efficiency loss to 9.7% at the median percentile, compared to 13.3% with SoA methods - meeting the project target of <10% loss.

Furthermore, the presence of interference from any source can have severe impact on key communication metrics such as data rate, delays, and overall error rates. The constrained nature of subnetwork nodes, especially in terms of computation, makes this problem even more complex. It is therefore important to address receiver design approaches that are suitable for the nature of constrained devices of a subnetwork. LLR approximations are presented in chapter 6, that enable constrained devices to perform better BER estimation, by means of methods of approximation of LLRs that present a wide range of complexities. Without these mechanisms, the integral of the FFT could be executed for polynomials of order equal or higher than 2, which is estimated as a $O(N^2)$ complexity. Moreover, once a function and its parameters are selected, it can be used to approximate the LLR during a validity window. This means the approximation only needs to be executed during some time instances, which results in no computational processing during that window.

3 RRM FOR IN-X SUBNETWORKS

In the previous deliverable [2], we provided a comprehensive overview of centralized, distributed, and hybrid RRM strategies tailored for densely deployed and highly mobile in-X subnetworks. These strategies addressed the diverse service requirements of industrial environments, such as Ultra-Reliable Low Latency Communications (URLLC) and enhanced Mobile Broadband (eMBB), through techniques ranging from heuristic sub-band allocation to deep learning-based resource optimization. Building upon that foundation, this chapter presents advanced technical contributions that further enhance RRM under realistic challenges, including outdated channel state information (CSI), high subnetwork density, and external interference. We focus on two complementary directions: (i) a centralized framework that leverages predictive CSI via a novel spatio-temporal attention-based model and resilient DNN-based joint sub-band and power control [7],[8], and (ii) a distributed AI-driven solution based on GNNs with over-the-air aggregation to enable low-overhead power coordination across independently operating subnetworks. Related research has also explored distributed reinforcement learning approaches for sub-band and power control in dense environments with stringent QoS demands, such as extended reality over in-body subnetworks [9]. These contributions aim to improve spectral efficiency, ensure QoS, and support the scalability and robustness required by 6G-enabled industrial systems.

3.1 Centralized RRM for in-X subnetworks

3.1.1 Spatio-Temporal Attention-Based model for RRM in outdated CSI

A fundamental challenge in RRM is the reliance on CSI, which often becomes outdated due to acquisition and processing delays. This outdated CSI can lead to suboptimal decisions in resource allocation, adversely affecting SE and QoS.

A major challenge is their reliance on frequent, periodic, and synchronous CSI updates, which are often impractical in scenarios with rapidly changing channel conditions and short coherence times. Current machine learning (ML)-based methods frequently overlook the impact of outdated CSI, leading to significant performance degradation in high-mobility environments.

This section introduces a novel Spatio-Temporal Attention-Based model designed to address these challenges. The model integrates attention mechanisms to accurately predict future CSI, leveraging both spatial and temporal correlations within the network. These predictions enable proactive and informed RRM decisions, mitigating the adverse effects of CSI delays and improving overall network performance. This activity has been carried out in collaboration with the SNS CENTRIC project.

3.1.1.1 System Model

The proposed system model consists of multiple in-factory subnetworks (InF-Ss) deployed within an industrial environment, each functioning as a localized wireless cell. Each subnetwork comprises an HC and multiple LCs or SNEs, facilitating wireless connectivity and coordination of industrial tasks. At the core of the system, a Centralized Resource Manager (CRM) is responsible for RRM and channel prediction, ensuring efficient orchestration of network resources across the factory. The deployment is

characterized by high mobility, dense subnetwork distribution, and localized communication, requiring robust interference mitigation and adaptive resource allocation strategies.

The system is deployed within a factory area of $L \times L$ square meters, designed to accommodate autonomous robots and industrial machinery. The layout supports seamless movement of SNEs while ensuring efficient communication between subnetworks. Each subnetwork consists of a central HC that serves as a communication hub, processing local data such as HC/SNE inputs and their controls. The HC facilitates real-time connectivity, allowing SNEs to efficiently coordinate their operations.

The communication coverage of each HC is defined by a circular transmission range with a radius R , ensuring that all associated SNEs remain within its connectivity zone. SNEs are positioned at distances ranging from $d_{\min}$ to R , ensuring compliance with minimum proximity constraints while maintaining reliable wireless links.

Subnetworks exhibit controlled mobility, following predefined trajectories that replicate the movement of autonomous mobile robots (AMRs) and automated guided vehicles (AGVs) transporting materials across the factory floor. The velocity of each subnetwork, represented as $v = \{v_1, \dots, v_N\}$, varies dynamically to reflect real-time industrial operations. Movement patterns may be adjusted based on environmental factors such as congestion, priority-based tasks, or safety constraints, ensuring adaptive and efficient navigation.

The network topology comprises N subnetworks, denoted as $\mathcal{N} = \{1, \dots, N\}$, operating independently while coexisting within a shared spectrum, which is divided into K sub-bands, $\mathcal{K} = \{1, \dots, K\}$. Given the high density of subnetworks, mutual interference presents a significant challenge. Consequently, effective RRM techniques, including dynamic sub-band allocation and power control, are essential to maintain spectral efficiency and mitigate interference.

The system's deployment, which is illustrated in Figure 5, aligns with real-world industrial scenarios where each autonomous robot functions as a self-contained subnetwork. Internal communication between LCs and SNEs is facilitated by the HC, ensuring uninterrupted task execution while navigating the factory environment. The CRM further enhances operational efficiency by managing network-wide resource allocation, ensuring seamless connectivity and minimizing performance degradation due to interference.

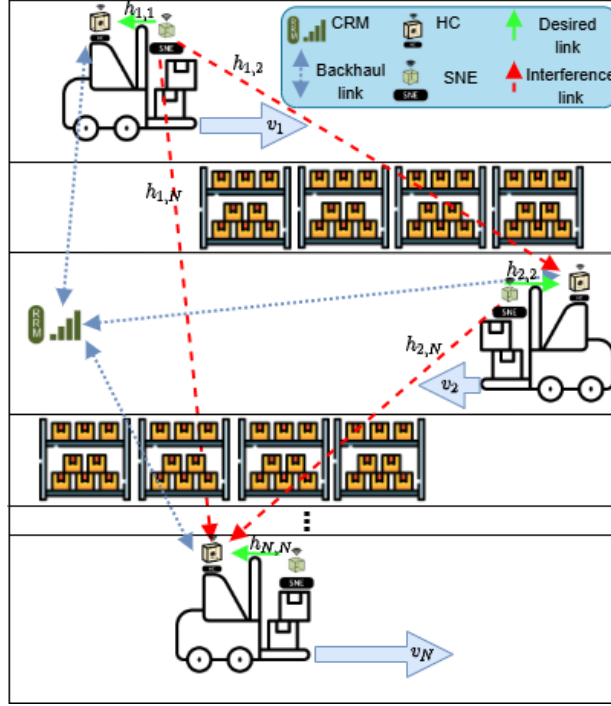


Figure 5: System model illustrating the deployment of in-factory subnetworks

3.1.1.2 Channel Model

The communication channel between LC/SNEs and HCs follows the 3rd Generation Partnership Project (3GPP) specifications for in-factory environments [10]. The channel gain for a link between SNE m and HC n at time t is given by:

$$h_t^{m,n} = |g_t^{m,n}|^2 \cdot \psi_t^{m,n} \cdot \Gamma_t^{m,n}$$

where:

- $\psi_t^{m,n}$ represents small-scale fading (Rayleigh fading).
- $\psi_t^{m,n}$ accounts for shadowing effects (modeled as a Gaussian random field) [11].
- $\Gamma_t^{m,n}$ denotes path loss, which depends on the carrier frequency f_c and the distance $d_{m,n}$ between nodes.

CSI is represented as a global matrix H_t , which contains all channel gains for each subnetwork pair. Specifically:

- $h_t^{n,n}$ represents the desired link (LC/SNEs to its associated HC).

- $h_t^{m,n}$ for $m \neq n$ corresponds to interfering links (between a LC/SNE in subnetwork m and an HC in subnetwork n).

Our focus is specifically on the FR3 frequency band, and the model is aligned with empirical measurements reported in deliverable D2.3 [12]. Detailed empirical characterization and validation of similar industrial propagation channels, including multi-frequency measurements, can also be found in this comprehensive study.

3.1.1.3 CSI Acquisition, Reporting, and Challenges of Outdated CSI

Efficient CSI acquisition and reporting are essential for adaptive RRM in industrial wireless networks. In each subnetwork, **entities (e.g., SNEs)** are configured to periodically transmits reference sequences, such as sounding reference signals (SRS), to enable neighbouring subnetworks to measure interference levels. The CSI is then reported to the CRM via dedicated backhaul channels, where it is used to optimize sub-band allocation and power control.

However, several limitations in the CSI acquisition and reporting process introduce latency and inaccuracies, leading to outdated CSI. In centralized architectures, the time required for data processing, backhaul transmission, and channel estimation causes a temporal discrepancy between when CSI is measured and when it is used for decision-making. The timing process is illustrated in Figure 6. In this diagram, Δt denotes the interval between consecutive transmissions of sounding reference signals, while τ quantifies the overall CSI feedback delay. This mismatch reduces the effectiveness of adaptive RRM strategies, resulting in suboptimal resource allocation and degraded SE.

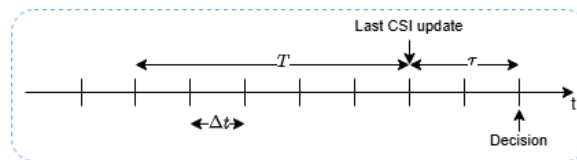


Figure 6: Illustration of the CSI update process, highlighting the delay between the last CSI measurement, decision-making, and the subsequent reconfiguration interval.

Challenges in Outdated CSI

Outdated CSI refers to channel state information that no longer accurately reflects the real-time wireless environment. Several factors contribute to this issue, making traditional real-time RRM strategies ineffective.

Latency in CSI Reporting

The process of CSI estimation, feedback transmission, and aggregation introduces delays at multiple levels, including:

Propagation and processing delays occur as the time taken for signals to travel between SNEs and HCs, followed by computational overhead at HCs, increases latency. Additionally, backhaul bandwidth constraints in centralized networks lead to queuing delays, slowing down CSI aggregation. Feedback overhead is another major factor, as reporting CSI for a large number of links in dense deployments results in excessive signalling congestion.

Rapid Channel Variability

High mobility of InF-Ss, such as autonomous mobile robots (AMRs) and automated guided vehicles (AGVs), causes fast-changing channel conditions, shortening the coherence time of the channel. The dynamic nature of industrial environments introduces constantly evolving interference conditions, which further distorts CSI accuracy. As a result, the CSI available at the CRM at decision-making time is often an outdated version of the actual channel state at that moment. The CSI available at the CRM at decision-making time $t + \tau$ is typically an outdated version of the actual channel state at time t . This temporal discrepancy, quantified by a delay factor τ , encompasses the entire acquisition, processing, and reporting delay chain.

Impact of Outdated CSI on RRM

The timing process of CSI acquisition and updates follows a cycle, where Δt represents the interval between consecutive CSI updates and τ , represents the total feedback delay, from acquisition to decision-making.

The mismatch between real-time channel conditions and delayed CSI feedback negatively affects multiple aspects of RRM. Sub-band allocation becomes inefficient, as frequency assignments based on outdated CSI fail to reflect current interference conditions, leading to higher interference and underutilized bandwidth. Similarly, power control decisions made using outdated CSI can cause power wastage, QoS violations, and inefficient energy use. Furthermore, interference management is affected, as the inability to accurately estimate interference levels results in poor coordination between subnetworks, leading to increased QoS degradation.

To mitigate these negative effects, the CRM maintains a buffer of past CSI samples. This enables temporal correlation analysis, improving prediction accuracy and allowing for better RRM decisions even under CSI delays [13].

Strategies to Overcome the Challenges of Outdated CSI

To compensate for outdated CSI and improve RRM efficiency, several advanced techniques can be implemented.

Predictive Channel Modelling

ML techniques can forecast future CSI based on historical patterns [14], [15]. Spatio-temporal deep learning frameworks, such as Long Short-Term Memory (LSTM)-based predictors, leverage both spatial correlations between subnetworks and temporal dependencies in CSI evolution [16].

Optimized CSI Feedback Mechanisms

Reducing feedback overhead while maintaining accuracy is essential for real-time RRM. This can be achieved through compressed CSI feedback, which transmits only the most relevant CSI components, and hybrid CSI estimation, which combines real-time CSI with historical data trends to improve estimation accuracy. Adaptive feedback scheduling, where feedback intervals dynamically adjust based on network conditions, can further optimize CSI reporting.

Leveraging Historical CSI for Enhanced Estimation

The CRM can store and analyse past CSI data to develop better resource allocation strategies. Advanced deep learning models, such as transformers and LSTMs, can process historical CSI data to make reliable predictions, reducing the dependency on real-time feedback.

Channel Prediction for CSI Estimation

Predicting future CSI is a time-series forecasting problem, requiring models that capture both short-term fluctuations and long-term trends. Traditional model-based approaches, such as autoregressive models and stochastic processes, struggle to handle the dynamic nature of industrial subnetworks [17].

In contrast, machine learning-based predictors offer a more adaptable and accurate solution. These models identify spatial correlations among subnetworks to understand interference relationships, model temporal dependencies in CSI evolution, and predict CSI over a delay horizon to compensate for reporting delays.

As depicted in Figure 7, the framework processes historical, delayed CSI data as input to determine optimal RRM strategies. By replacing outdated CSI with predictive CSI, RRM decisions become more accurate, spectral efficiency improves, and QoS adherence is maintained, ensuring reliable and high-performance wireless communication in 6G-enabled industrial environments.

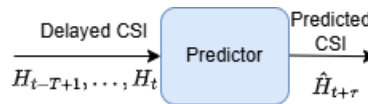


Figure 7: Overall framework of the proposed solution, utilizing delayed CSI as input to generate predicted CSI as output.

3.1.1.4 Dual Attention-Based Channel Prediction

Channel samples exhibit complex spatio-temporal correlations, particularly in dynamic industrial environments where the mobility of InF-Ss and fluctuating interference patterns cause rapid variations in the channel state. The dense deployment of subnetworks further introduces strong spatial dependencies, as neighbouring SNEs often experience similar interference and fading effects. These combined factors make channel prediction a challenging task, requiring models that can effectively capture both temporal dependencies and spatial correlations [18], [19].

Traditional time-series models, such as recurrent neural networks (RNNs), have been widely applied to channel prediction tasks. However, standard RNNs suffer from the vanishing gradient problem when handling long-term dependencies. LSTM networks mitigate this issue by introducing memory cells and gating mechanisms, enabling them to retain long-term dependencies. LSTMs have demonstrated superior performance in various applications, including language modelling, time-series forecasting, and industrial sensor analysis [20],[21],[22].

To improve channel prediction in dynamic environments, an LSTM-based encoder-decoder architecture is employed, enhanced with dual attention mechanisms. These mechanisms prioritize spatially and temporally significant features, allowing the model to capture critical patterns in channel evolution. The encoder extracts relevant features from historical CSI, while the decoder generates accurate future CSI predictions by leveraging attention-weighted hidden states.

LSTM-Based Encoder-Decoder Architecture

The core of the predictive model is an LSTM network structured as an encoder-decoder framework. Each CSI matrix, denoted as $\mathbf{H}_t$, represents the channel state at time t . Before processing, it is flattened into a one-dimensional vector:

$$\text{vec}(\mathbf{H}_t) = [h_t^{1,1}, \dots, h_t^{m,n}, \dots, h_t^{N,N}] \in \mathbb{R}^{N^2 \times 1}$$

This transformation ensures compatibility with the LSTM input format.

The LSTM unit processes the CSI vector $\mathbf{x}_t$ along with its previous hidden state $\mathbf{z}_{t-1}$ and memory cell state $\mathbf{m}_{t-1}$. These inputs determine which information is retained, updated, or discarded. The forget gate selectively removes irrelevant past information, while the update gate integrates new relevant information into the memory cell. The output gate regulates the final contribution of the memory cell to the hidden state. This structure enables the LSTM to capture temporal dependencies while avoiding vanishing gradients.

The LSTM encoder processes the historical CSI sequence $\{\mathbf{H}_{t-T+1}, \dots, \mathbf{H}_t\}$ and encodes it into a set of hidden states, which are then passed to the decoder. The decoder, in turn, predicts the future CSI based on these hidden states and previously generated outputs.

While LSTMs effectively model long-term dependencies, they treat all input features uniformly, which limits their ability to focus on critical time steps or spatially significant regions in the input sequence. To overcome this, dual attention mechanisms - spatial and temporal attention - are introduced within the encoder-decoder framework.

Spatial and Temporal Attention Mechanisms

Figure 8 illustrates the dual attention-based encoder-decoder structure, consisting of an encoder and a decoder. Attention mechanisms enhance the model by dynamically prioritizing important features within the input sequence. Spatial attention focuses on selecting the most relevant channel variables from each CSI matrix, while temporal attention identifies the most informative time steps in historical data.

The spatial attention mechanism assigns varying importance to different CSI components based on their contribution to channel prediction. The attention weights are computed using:

$$\hat{\alpha}_t^i = \mathbf{V}_i^{\text{SA}} \tanh(\mathbf{W}_i^{\text{SA}} \mathbf{s}_{t-1} + \mathbf{U}_i^{\text{SA}} \text{vec}(\mathbf{H}_t)^i + \mathbf{b}_i^{\text{SA}}), \quad 1 \leq i \leq N^2,$$

$$\alpha_t^i = \frac{e^{\hat{\alpha}_t^i}}{\sum_{j=1}^d e^{\hat{\alpha}_t^j}}, \quad 1 \leq i \leq N^2,$$

where $\mathbf{s}_{t-1}$ is the decoder's previous hidden state, and $\mathbf{V}_i^{\text{SA}}$, $\mathbf{W}_i^{\text{SA}}$, $\mathbf{U}_i^{\text{SA}}$, and $\mathbf{b}_i^{\text{SA}}$ are trainable parameters. The final spatial attention-weighted input is obtained as:

$$\text{vec}(\bar{\mathbf{H}}_t) = \alpha_t \odot \text{vec}(\mathbf{H}_t).$$

where $\alpha_t = [\alpha_t^1, \dots, \alpha_t^{N^2}]$ represents the spatial attention weights. An LSTM network is subsequently employed to process the attention-weighted inputs and extract latent features, denoted as $\{z_{t-T+l}, \dots, z_t\}$.

Once the encoder has extracted features from the spatially weighted input, the temporal attention mechanism prioritizes significant time steps for predicting the next CSI value. The temporal attention weights are computed as:

$$\hat{\beta}_t^l = \mathbf{V}_l^{\text{TA}} \tanh(\mathbf{W}_l^{\text{TA}} \mathbf{s}_{t-1} + \mathbf{U}_l^{\text{TA}} \mathbf{z}_{t-T+l} + \mathbf{b}_l^{\text{TA}}),$$

$$\beta_t^l = \frac{e^{\hat{\beta}_t^l}}{\sum_{q=1}^T e^{\hat{\beta}_t^q}}, \quad 1 \leq l \leq T,$$

where $\mathbf{z}_{t-T+l}$ represents the encoder's hidden states at previous time steps, $\mathbf{s}_{t-1}$ is the output of the $t-1$ -th unit of the LSTM decoder, and $\mathbf{V}_k^{\text{TA}}$, $\mathbf{W}_k^{\text{TA}}$, $\mathbf{U}_k^{\text{TA}}$, and $\mathbf{b}_k^{\text{TA}}$ are parameters to be learned. The normalized value β_t^l quantifies the relevance of the l -th hidden state to the current decoding step. The weighted sum of the encoder's hidden states forms the context vector:

$$\mathbf{c}_t = \sum_{l=1}^T \mathbf{z}_{t-T+l} \beta_t^l.$$

By integrating both spatial and temporal attention, the model ensures that only the most relevant features contribute to the prediction, significantly improving CSI forecasting accuracy in dynamic industrial environments.

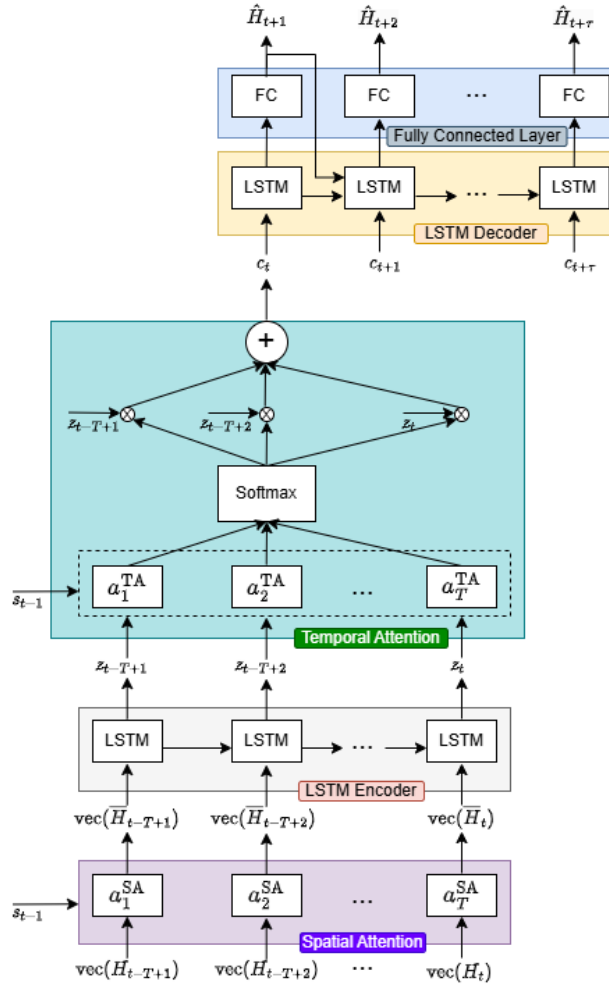


Figure 8: Structure of the dual attention-based LSTM encoder-decoder, integrating spatial and temporal attention mechanisms for enhanced channel prediction.

Prediction and Model Optimization

The decoder module generates the predicted CSI using the LSTM-based hidden state and the temporal attention-weighted context vector. The final prediction is computed as:

$$\hat{\mathbf{H}}_{t+1} = \mathbf{W}_f \mathbf{s}_t + \mathbf{b}_f,$$

where $\mathbf{W}_f$ and $\mathbf{b}_f$ are trainable parameters.

The prediction model is trained using the mean squared error (MSE) loss function:

$$\text{MSE} = E_{B_P} \{ |\mathbf{H} - \hat{\mathbf{H}}|^2 \},$$

where $\mathbf{H}$ and $\hat{\mathbf{H}}$ represent the actual and predicted CSI, respectively, and B_P is the batch size.

To ensure efficient optimization, the Adaptive Moment Estimation (ADAM) algorithm is used. ADAM combines momentum and RMSprop techniques, facilitating faster convergence and robust parameter updates, making it particularly suitable for complex models such as the dual attention-based encoder-decoder.

The proposed dual attention-based LSTM encoder-decoder effectively captures both spatial correlations and temporal dependencies in CSI evolution, enhancing channel prediction in dynamic factory environments. By prioritizing the most relevant features and time steps, the model mitigates the impact of outdated CSI and significantly improves the accuracy of resource management strategies in dense industrial networks.

3.1.2 Resilient DNN for Joint Sub-Band and Power Control in Mobile In-FS

In this section we introduce the problem formulation and then architecture of the proposed ML model and the associated learning strategy for efficient resource allocation. The DNN framework is specifically designed to address sub-band allocation and power control in an integrated manner while adhering to resource allocation constraints.

3.1.2.1 Problem Formulation

As illustrated in Figure 5, this study focuses on RRM for uplink transmissions, where SNEs communicate with their respective HCs within each subnetwork. The system assumes the available spectrum to be divided into sub-bands, denoted by $\mathcal{K} = \{1, \dots, K\}$, shared among all SNEs. Each subnetwork is assumed to serve a single SNE, with its transmission fully utilizing the assigned sub-band.

The primary objective is to develop a resource allocation strategy that jointly optimizes sub-band assignment and power control to maximize the average SE across all subnetworks, while ensuring that each subnetwork meets a minimum SE requirement, $SE_{\min}$. Sub-band allocation is denoted by a_n^k , and power levels are represented by p_n . These variables are adjusted based on the estimated CSI, $\widehat{H}_t$. The indicator $a_n^k \in \{0,1\}$ specifies whether subnetwork n transmits on sub-band k , ($a_n^k = 1$ for active transmission). Each subnetwork is constrained to use a single sub-band, expressed as $\sum_{k=1}^K a_n^k = 1$. The allocated sub-band is determined by converting the one-hot vector $a_n = [a_n^1, \dots, a_n^K]$ into a categorical value $a_n = \sum_{k=1}^K k \cdot a_n^k$. Transmit power p_n is modeled as a continuous variable, constrained by $P_{\max}$ offering greater flexibility compared to discrete power levels.

$$SE_n^k = \log_2 \left(1 + \frac{h_t^{n,n} a_n^k p_n}{\gamma_{m,n}^2 + \sum_{m \in \mathcal{N} \setminus \{n\}} h_t^{m,n} a_m^k p_m} \right),$$

where $h_t^{n,n}$ represents the desired channel link for subnetwork n , and $\gamma_{m,n}^2$ is the receiver noise power, calculated as:

$$\gamma_{m,n}^2 = 10^{\frac{-174 + NF + 10 \log_{10}(W_k)}{10}}$$

where W_k denotes the sub-band bandwidth (in Hz), and NF represents the receiver noise figure (in dB).

The joint optimization problem for sub-band allocation and power control is formulated as:

$$\begin{aligned}
 & \max_{\{a_n^k, p_n\}} \frac{1}{N} \sum_{n \in \mathcal{N}} \sum_{k \in \mathcal{K}} \text{SE}_n^k \\
 & \text{s.t. } \sum_{k \in \mathcal{K}} \text{SE}_n^k \geq \text{SE}_{\min}, \quad \forall n \in \mathcal{N} \\
 & a_n^k \in \{0,1\}, \quad \forall n \in \mathcal{N}, \forall k \in \mathcal{K} \\
 & \sum_{k \in \mathcal{K}} a_n^k = 1, \quad \forall n \in \mathcal{N} \\
 & 0 \leq p_n \leq P_{\max}, \quad \forall n \in \mathcal{N}
 \end{aligned}$$

This mixed-integer nonlinear programming (MINLP) problem is computationally challenging due to its non-convex nature. Conventional methods, such as branch-and-bound algorithms, dynamic programming, and convex relaxation techniques, can be employed to solve such problems; however, these approaches often suffer from prohibitive computational complexity, particularly in large-scale and dynamic scenarios. To address these challenges, deep learning techniques are employed to approximate the optimal mapping function, leveraging their universal approximation capabilities [23],[24].

3.1.2.2 Structure of the DNN Model

The detailed architecture of the developed deep neural network (DNN) is illustrated clearly in Figure 9. This architecture comprises two separate yet interconnected modules, each integrating M_U fundamental computational units designed explicitly for extracting and learning complex relationships from the input features and subsequently generating resource allocation decisions.

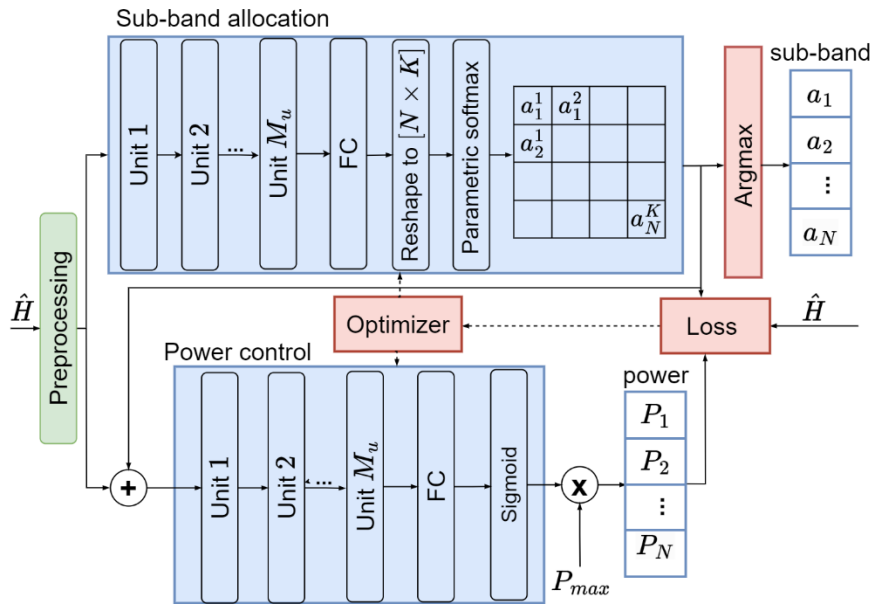


Figure 9: Structure of the proposed DNN-based RRM framework, showing interconnected modules for sub-band allocation and power control, supported by preprocessing and loss optimization components.

Each fundamental computational unit is structured around four essential layers. Initially, a fully connected (FC) layer captures high-dimensional feature representations from the input data, creating a rich feature set for subsequent processing. The next stage involves a batch normalization (BN) layer, strategically positioned to stabilize the intermediate outputs. This stabilization significantly accelerates the training convergence process by reducing internal covariate shifts and improving gradient flow.

Subsequently, the activation layer introduces non-linear transformations critical for the model's ability to capture intricate patterns within the input data. Depending on the position of the unit within the network, this layer leverages either an Exponential Linear Unit (ELU) or a Rectified Linear Unit (ReLU). Specifically, ELU is utilized within initial computational units to handle potential negative inputs effectively, thereby ensuring smoother gradient propagation and stable updates during early training phases. Conversely, deeper layers employ ReLU to focus exclusively on positive activations, simplifying computational complexity while preserving overall model performance and accuracy.

The final component of each unit is the dropout layer, a vital mechanism to combat overfitting. This layer randomly deactivates neurons during the training phase, thereby enhancing the model's generalization capability and ensuring robustness when applied to unseen data.

Each fundamental computational unit incorporates M_H hidden nodes, a configuration that allows the network to intricately model subtle and complex data relationships. Before processing, the input—consisting of an estimated channel gain matrix $\hat{\mathbf{H}}$ —undergoes a transformation to the decibel (dB) scale. This conversion ensures numerical stability and uniform scaling across inputs. Following this transformation, the matrix is normalized to achieve a zero mean and unit variance, thus preventing potential biases or skewed distributions in the data. The normalized channel gain matrix, initially structured as an $N \times N$ matrix, is flattened into a one-dimensional feature vector of length N^2 , aligning with the approach recommended in prior literature [23].

The overall DNN architecture is segmented into two modules that share a similar structural foundation but are individually tailored for specific resource allocation functions. The first module addresses sub-band allocation, while the second module focuses explicitly on power control, collectively ensuring a comprehensive and efficient handling of the resource allocation tasks.

Sub-Band Allocation Module

The sub-band allocation module outputs a total of NK values, initially structured into an $N \times K$ matrix. To meet the sub-band allocation constraints, each subnetwork's outputs are individually processed through a dedicated softmax function. This ensures each subnetwork's sub-band allocation probabilities collectively sum to unity, meeting the probabilistic constraints necessary for meaningful allocations.

During inference, the allocation decision for subnetwork n , represented as a_n , is finalized by selecting the sub-band with the maximum predicted probability: $a_n = \arg \max_k a_n^k$, thus adhering strictly to the discrete constraints articulated in problem formulation. However, during the training phase, outputs remain continuous probabilities a_n^b for all sub-bands. To effectively bridge the gap between continuous probability outputs during training and discrete allocation decisions during inference, a soft binarization approach with an adjustable sharpness parameter, δ , is introduced. By gradually tuning this parameter, the model progressively transitions from a continuous output regime during training towards a more discrete, decision-focused regime during inference, refining the predictive accuracy and robustness of the allocations over successive training iterations.

Power Control Module

The second module of the DNN architecture is dedicated to predicting optimal power levels for each subnetwork. After passing through the final fully connected layer, the outputs are further refined via a sigmoid activation function. This sigmoid transformation ensures that the predicted power levels are confined within a continuous range from 0 to 1, facilitating smooth and differentiable outputs suitable for gradient-based optimization. The resulting normalized predictions are then scaled by the maximum permissible transmission power P_{max} , directly ensuring compliance with the power control constraints detailed in problem formulation. The implementation of this sigmoid-based output mechanism significantly aids in the optimization process, enabling precise and efficient adjustments to the subnetwork power levels in response to varying channel conditions and operational requirements.

Loss Function and Training Methodology for the DNN Model

The developed DNN model employs an unsupervised learning strategy, which circumvents the necessity for labelled datasets that provide precomputed optimal solutions. Instead, the network directly targets the optimization of the resource allocation problem using a custom-designed loss function. This approach significantly simplifies the data preparation process, enhancing computational efficiency and making the model particularly suitable for real-time deployment scenarios. The model is trained offline, ensuring that computationally intensive optimization steps are completed beforehand, thereby substantially reducing inference complexity compared to traditional iterative optimization methods.

The formulated loss function strategically addresses dual optimization goals: maximizing the average SE across all subnetworks while rigorously enforcing compliance with QoS constraints. The formal expression for the loss function is given by: $L = -\frac{1}{N} \sum_{n \in \mathcal{N}} SE_n^w + \lambda \sigma(SE_{\min} - SE_n^w)$

where λ serves as a critical weighting parameter, carefully balancing the trade-off between achieving high overall SE performance and satisfying stringent QoS constraints.

In this context, SE_n^w denotes the weighted spectral efficiency for subnetwork n and is calculated as: $SE_n^w = \sum_{k \in \mathcal{K}} a_n^k SE_n^w$.

The loss function inherently pursues two intertwined objectives. The primary term targets the maximization of the weighted SE across all subnetworks, steering the model towards optimal resource utilization outcomes. Concurrently, the secondary term introduces a structured penalty to address QoS deviations effectively. Specifically, the penalty leverages a sigmoid function σ , offering a smooth, continuous, and differentiable penalty mechanism whenever a subnetwork's spectral efficiency drops below the predefined minimum threshold, $SE_{\min}$. The differentiability and smoothness of this penalty component are particularly advantageous for gradient-based training methods, as they enable stable gradient updates and focus optimization on substantial QoS violations without introducing abrupt shifts in training behavior.

The adjustable parameter λ provides the necessary flexibility to tailor model behavior to specific operational priorities and network conditions. Lower values of λ direct optimization efforts predominantly toward improving the overall network efficiency, whereas higher values strongly enforce compliance with QoS constraints. Such flexibility enables adaptive model performance across varying scenarios and requirements.

Addressing the challenge posed by the difference between continuous outputs during training and discrete decisions during inference, the training methodology incorporates a soft binarization approach. This method uses a parameterized softmax function described as follows:

$$\phi_\delta(x_n^k) = \frac{e^{x_n^k/\delta}}{\sum_{k=1}^K e^{x_n^k/\delta}}$$

where x_n^k represents the input logits to the softmax layer, and δ controls the sharpness of the probability distribution. Initially, δ is set to a large value, ensuring smoother and broader distributions, thus enhancing gradient stability and network exploration during early training stages. Subsequently, δ is progressively decreased using a systematic adaptive scaling schedule:

$$\delta^{(m)} = \delta^{(m-1)} \cdot \gamma, \text{ if } m = j \cdot I_{\text{update}}, j \in \mathbb{Z}^+$$

where γ denotes the scaling factor, I_{update} defines the interval at which updates occur, and m signifies the current training epoch. This adaptive reduction approach smoothly transitions the model's outputs from continuous probabilities to sharply defined discrete allocations, facilitating a seamless shift from broad exploration of resource allocation possibilities to precise exploitation of optimal solutions.

Si The experimental validation of the proposed scheme was conducted using high-performance cloud computing infrastructure, leveraging an AMD EPYC-Rome processor with 40 cores operating at 2.9 GHz, complemented by an NVIDIA A40 GPU and supported by 64 GB of RAM. This robust computational

environment ensured efficient training and evaluation, providing accurate and timely assessments of the proposed model's capabilities.

The channel prediction performance of the proposed method was benchmarked against a conventional LSTM model with identical configuration parameters. This direct comparison highlights the effectiveness and added value of integrating a dual attention mechanism within an encoder-decoder architecture, specifically designed to capture intricate spatial and temporal dependencies within the channel data. Such a mechanism significantly enhances prediction accuracy compared to standard recurrent neural network approaches.

For evaluating the RRM component, the proposed scheme was rigorously benchmarked against established SoA methods. Specifically, the comparative analysis included the Sequential Iterative Sub-band Allocation (SISA) algorithm paired with the Weighted Minimum Mean Square Error (WMMSE) technique. The SISA algorithm, as detailed in [25], employs an iterative, centralized approach explicitly aimed at minimizing interference during sub-band allocation. Subsequently, the WMMSE method, referenced in [26], is applied to optimize power allocation decisions based on the sub-band assignments provided by SISA. This combination represents a strong, optimization-driven baseline suitable for assessing the efficacy of our proposed scheme.

Moreover, the robustness and general applicability of the proposed RRM strategy were evaluated under highly dynamic and less structured scenarios by introducing a random allocation baseline. In this scenario, sub-bands were randomly assigned from the available K sub-bands for each subnetwork, and the transmission power levels were uniformly sampled within the allowed range $[0 P_{max}]$. This random assignment approach provided a practical lower-bound performance measure, facilitating a clearer understanding of the incremental improvements offered by the proposed model and baseline methods under challenging, interference-prone conditions.

Overall, the selected evaluation framework provides a comprehensive assessment, effectively demonstrating the proposed scheme's superior capabilities in both channel prediction accuracy and RRM efficiency. Results consistently indicated that the proposed methodology significantly outperformed traditional optimization methods and random strategies, particularly regarding maintaining high spectral efficiency and reliably meeting QoS requirements in dynamically evolving network conditions.

3.1.3 Simulation Results and Analysis for proposed Centralized RRM

3.1.3.1 Simulation Setup

The simulations were performed in a factory environment, modelled as a high-density deployment scenario within a confined area of $20 \times 20 \text{ m}$ area (i.e., 400 m^2). This corresponds to a density of 25,000 subnetworks per square kilometre, reflecting realistic industrial conditions. A total of 10 mobile subnetworks were simulated, each moving at randomly assigned velocities within the range of 0–10 m/s along parallel lanes. These lanes were equally spaced, representing controlled yet dynamic industrial movements. Each InF-S was modelled with a circular coverage radius of 1 m, ensuring localized

communication. Additionally, the minimum separation between HCs and SNEs within each subnetwork was constrained to 0.8 m, consistent with practical industrial deployments.

The wireless channel characteristics followed the detailed description in Section 3.1.1.2, reflecting typical dense clutter and low base station heights inherent to industrial factory scenarios (InF scenarios) as per standard [10]. Our simulation environment, particularly in the FR3 frequency band (around 10 GHz), aligns closely with empirical measurements conducted within the 6G-SHINE project's Work Package 2 [4]. Clutter elements were uniformly distributed with a size of 1 m and density of 70% across the simulated area. Shadowing effects were simulated with a standard deviation of 4 dB and a decorrelation distance of 5 m, ensuring accurate representation of spatially correlated signal impairments. The number of sub-bands was limited to $K = 3$, compelling the subnetworks to efficiently share these resources. The radio propagation parameters included a carrier frequency $f_c = 10$ GHz, channel bandwidth per sub-band $W_k = 40$ MHz, maximum transmit power $P_{max} = 0$ dB, and a noise figure (NF) of 5 dB. These parameters align closely with typical high-frequency industrial wireless network configurations. For consistency, the sounding reference signal period and the reconfiguration interval for RRM were both set to $\Delta t = 100$ ms. The buffer length for past CSI samples was fixed at $T = 5$ time steps, and the prediction horizon τ , representing the maximum allowable delay, was set to 4-time steps. The predictor employed a Long Short-Term Memory (LSTM)-based architecture consisting of $L_{LSTM} = 2$ hidden layers, each with $z_{LSTM} = 512$ neurons for both the encoder and decoder. The predictor training used a learning rate of $\alpha_P = 10^{-4}$, batch size of $B_P = 1024$, and was trained for $E_P = 500$ epochs to ensure robust predictive performance.

The RRM model was designed using $M_U = 4$ basic units, each containing $H_R = 512$ hidden nodes. The training parameters included a learning rate $\alpha_R = 10^{-5}$, dropout rate $r_d = 0.1$, and batch size $B_R = 1024$. The weighting parameter λ was set at 20, effectively balancing SE maximization and minimum QoS adherence. The minimum SE constraint SE_{min} was set at 4, emphasizing stringent QoS demands characteristic such as real-time industrial control, high-data-rate sensor aggregation, and machine-to-machine communication tasks [3]. The RRM model training spanned $E_R = 150$ epochs to ensure robust convergence and adaptability.

The dataset for evaluation was generated by reconstructing the environment 10,000 times, with each subnetwork moving for 10 s per simulation instance. A sliding window approach extracted samples from all HC-SNE pairs, yielding a training dataset with 900,000 samples and a test dataset comprising 100,000 samples.

The comprehensive list of simulation parameters is summarized clearly in Table 1, categorizing parameters related to system deployment, predictor model configuration, RRM model design, and dataset preparation to facilitate clarity and reproducibility.

Table 1: Simulation parameters for centralized RRM

Parameter	Value
System Deployment and Channel Model	
Factory area, $L \times L$	$20 \times 20 \text{ m}^2$

Number of subnetworks, N	10
Subnetwork radius, R	1 m
Minimum distance between HC and SNE	0.8 m
Clutter density	70%
Clutter size	1 m
Shadowing standard deviation, σ_s	4 dB
De-correlation distance, d_c	5 m
Number of sub-bands, K	3
Sub-band bandwidth, W_k	40 MHz
Carrier frequency, f_c	10 GHz
Maximum transmit power, P_{max}	0 dBm
Noise figure, NF	5 dB
Sounding reference signal period, Δt	100 ms
CSI buffer length, T	5
Prediction length (delay), τ	4
Predictor Model Hyperparameters	
Number of hidden layers, L_{LSTM}	2
Number of hidden neurons, z_{LSTM}	512
Learning rate, α_P	10^{-4}
Batch size, B_P	1024
Training epochs, E_P	500
RRM Model Hyperparameters	
Basic units, M_U	4
Hidden nodes per unit, H_R	512
Learning rate, α_R	10^{-5}
Dropout rate, r_d	0.1
Batch size, B_R	1024
Training epochs, E_R	150
Weighting parameter, λ	20
Minimum SE, SE_{min}	4
Initial Softmax tunable parameter, $\delta_{initial}$	1
Scaling factor, γ	0.9
Interval between updates, I_{update}	10
Dataset Parameters	
Reconstructed environments	10,000
Training samples	900,000
Testing samples	100,000
Simulation duration per reconstruction	10 s
Sliding window size	$T + \tau = 9$

3.1.3.2 Comparison of RRM Approaches

In this section, we evaluate the performance of the proposed DNN-based resource management model by comparing it against state-of-the-art (SoA) benchmark algorithms.

Initially, we consider an ideal scenario without delay, assuming instantaneous availability of CSI to the CRM. The primary objective here is to evaluate how effectively the proposed model maximizes SE while complying with the minimum SE constraints.

Figure 10 depicts the progression of two key metrics over the training epochs of the RRM model: the probability of minimum SE violations (interpreted as outage probability) and the average SE across all subnetworks. The figure clearly illustrates the model's convergence behaviour and its capability to balance individual QoS constraints against overall system efficiency. Initially, the probability of SE violations is high due to the model's limited initial optimization knowledge. However, as training progresses, a significant reduction in violations occurs, indicating improved capability in satisfying QoS constraints. Concurrently, the average SE consistently increases, reflecting enhanced resource utilization. The stabilization of both metrics toward the end of training demonstrates the robustness and efficacy of the designed loss function in achieving an optimal and fair resource allocation for Industrial systems.

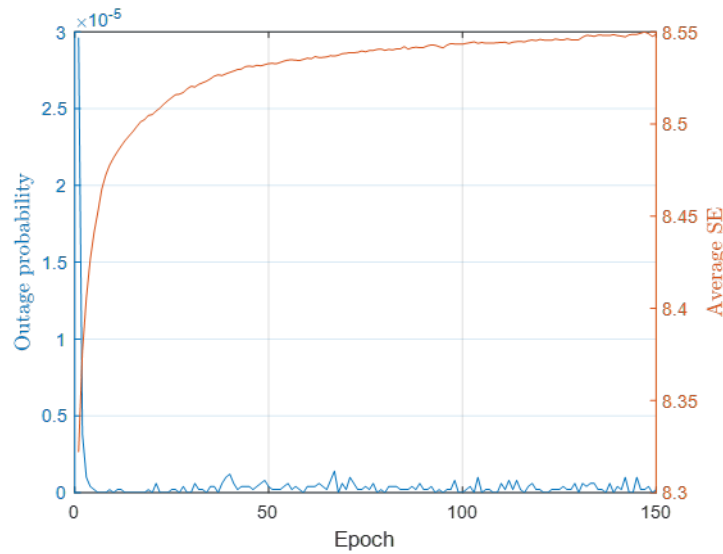


Figure 10: Evolution of the probability of minimum SE violation and average SE across all subnetworks as a function of the RRM training epochs, illustrating the model's convergence behaviour and its ability to balance QoS constraints with system efficiency.

To assess the effectiveness of the model in producing accurate binary decisions post-training, Figure 11 presents the cumulative distribution function (CDF) of the binarization error, defined mathematically as $E[|a_n - \text{round}(a_n)|]$, for each InF-S. The CDF is plotted on a logarithmic scale for clarity and illustrates that the binarization errors remain exceptionally small across the evaluated scenarios. Given that the maximum theoretical binarization error is 0.5, the observed results indicate that the proposed DNN model reliably produces discrete, binary outputs, thus fully complying with practical resource allocation constraints.

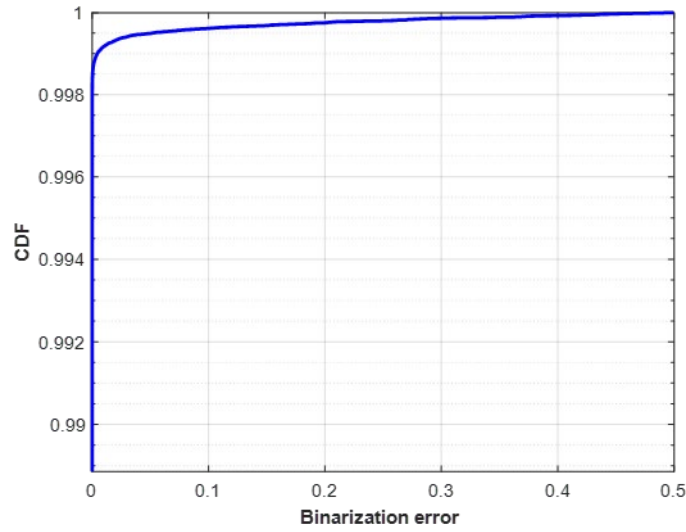


Figure 11: CDF of the binarization error for the RRM model, highlighting the effectiveness of the proposed framework in generating binary outputs for sub-band allocation with minimal deviation.

Furthermore, Figure 12 and Figure 13 illustrate the CDFs of average and individual SE values across all subnetworks, comparing the proposed DNN-based resource management method to the benchmark approach combining SISA and WMMSE under varying CSI delays ($\tau = 0, 1, 2, 3, 4$). Results indicate that increasing delays negatively impact the performance of both the proposed and benchmark methods due to reliance on outdated CSI for decision-making. Nevertheless, the proposed DNN-based model consistently demonstrates superior robustness and adaptability compared to the SISA-WMMSE approach, exhibiting notably less degradation as delay increases. Specifically, transitioning from zero delay ($\tau = 0$) to a 4-sample delay ($\tau = 4$) results in a median SE reduction of approximately 0.1 bps/Hz for the proposed method compared to approximately 0.2 bps for SISA-WMMSE.

Additionally, Figure 13, plotted on a logarithmic scale to highlight performance at lower percentiles, provides further insights into individual SE distributions under significant delays ($\tau = 4$). Notably, the minimum SE attained by the SISA-WMMSE benchmark decreases sharply to below 3 bps/Hz, whereas the proposed DNN-based approach maintains a significantly higher minimum SE of approximately 4.5 bps/Hz, representing a 50% improvement. This finding underscores the superior fairness and robustness of the DNN-based method, ensuring that even the most disadvantaged subnetworks maintain acceptable SE levels despite the challenges introduced by delays.

Overall, the analysis presented in Figure 12 and Figure 13 demonstrates clearly that the proposed DNN-based resource management model not only achieves higher average SE but also substantially improves fairness and robustness compared to the benchmark SISA-WMMSE method. These outcomes emphasize the model's suitability for realistic industrial scenarios, highlighting its ability to effectively mitigate the negative impact of CSI delays. The next subsection will explore how predictive capabilities further enhance performance under delay-induced impairments.

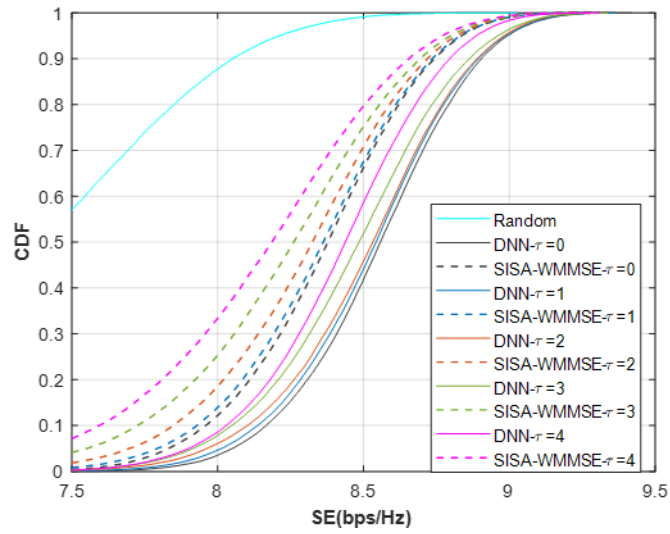


Figure 12: CDF of the average SE across all subnetworks, comparing the proposed DNN-based RRM and the benchmark under varying delay conditions

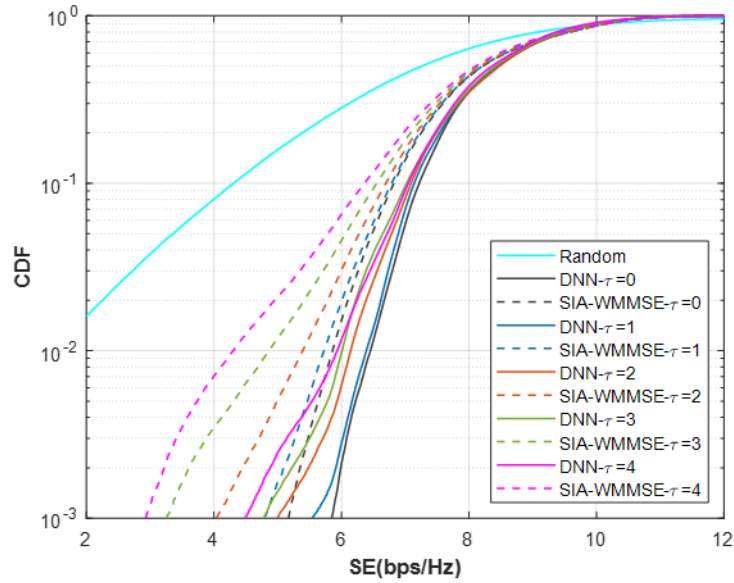


Figure 13: CDF of the individual SE across all subnetworks, comparing the proposed DNN-based RRM and the benchmark under varying delay conditions.

3.1.3.3 Comparison Results for Different CSI Prediction Methods

The primary aim of integrating channel predictors within this study is to enhance resource management by reducing the adverse effects caused by delays. To offer a comprehensive evaluation, Figure 14 compares the MSE loss trends during the training and validation phases for both the attention-enhanced LSTM and the standard LSTM predictors. These results provide valuable insights into the performance characteristics of the two predictive models. Initially, the dual-attention LSTM model exhibits faster

convergence in both training and validation losses compared to the standard LSTM, highlighting its superior learning capability during early training stages.

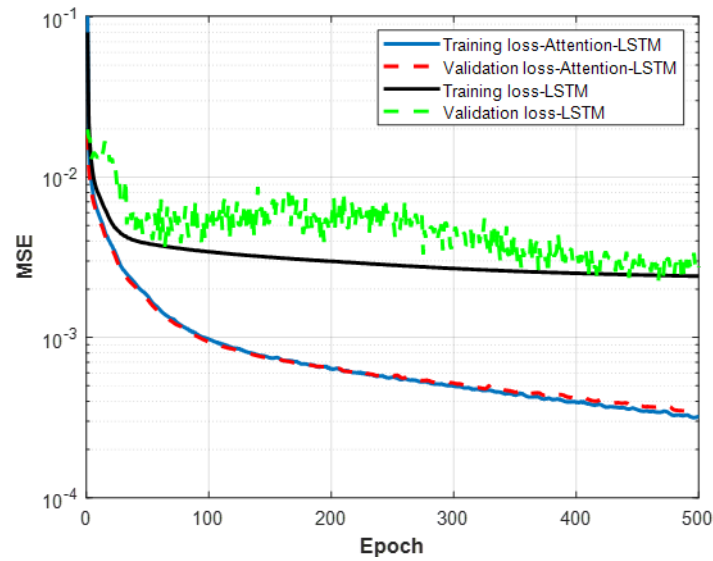


Figure 14: Training and validation loss trends for the dual attention LSTM and standard LSTM predictors.

Throughout the training period, the dual-attention LSTM consistently achieves lower loss values in both training and validation sets, demonstrating its enhanced effectiveness in capturing complex channel dynamics. Additionally, the narrow gap between training and validation losses indicates that the dual-attention LSTM model is less prone to overfitting, highlighting its robustness and generalizability. This superior performance is primarily attributed to the dual-attention mechanism's ability to effectively capture both spatial and temporal channel dependencies. Furthermore, the smoother loss curves observed for the dual-attention LSTM suggest a more stable and reliable training process, compared to the noticeable fluctuations associated with the standard LSTM.

To evaluate the effectiveness of the proposed prediction models in reducing delay-related performance degradation, we conducted a detailed analysis of SE across subnetworks. Figure 15 presents the CDF of the average SE for all subnetworks, while Figure 16 focuses on individual SE values under the maximum delay scenario ($\tau = 4$). The evaluations consider the DNN-based resource management framework and include baseline scenarios such as the ideal (no delay) condition and the sample-and-hold strategy, where the most recent CSI data is reused.

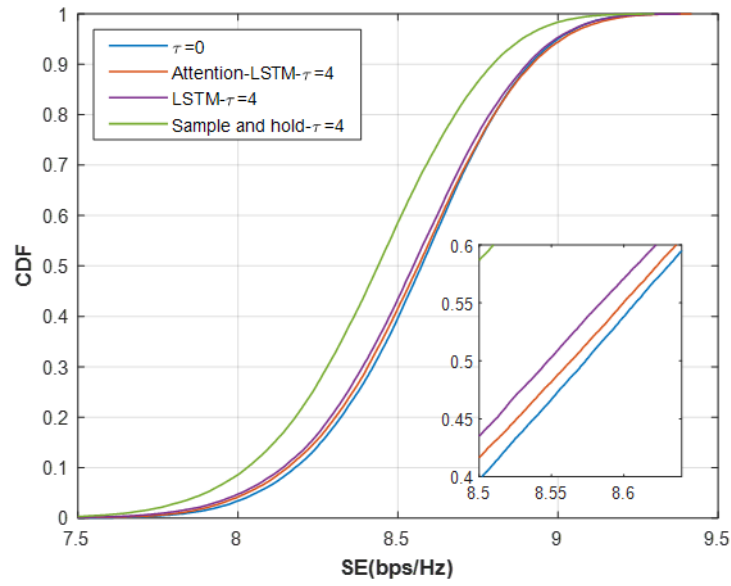


Figure 15: CDF of the average SE across all subnetworks for DNN-based RRM with a $\tau = 4$ -sample delay, comparing sample-and-hold, Attention-LSTM, and LSTM predictors.

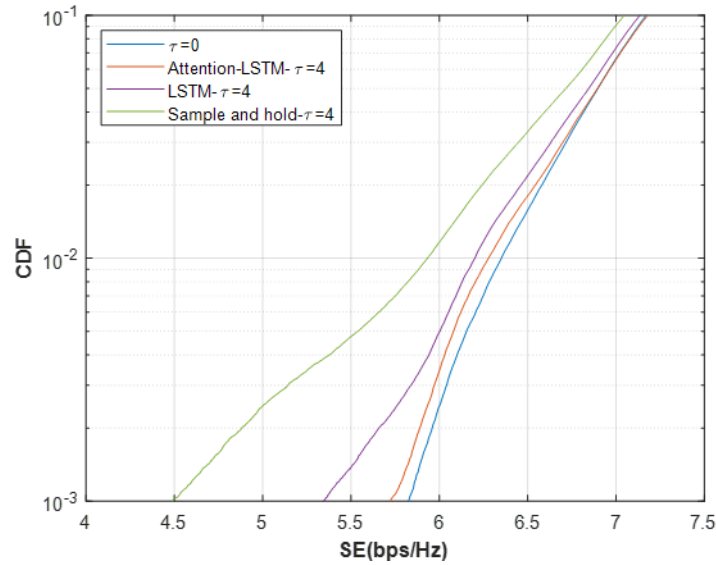


Figure 16: CDF of the SE for all subnetworks for DNN-based RRM with a $\tau = 4$ -sample delay, comparing sample-and-hold, Attention-LSTM, and LSTM predictors.

Figure 15 clearly demonstrates the substantial improvements achieved through the application of machine learning-based predictors, particularly the attention-enhanced LSTM. Under significant delay conditions ($\tau = 4$), the attention-enhanced LSTM consistently surpasses the standard LSTM in terms of average SE. Specifically, the CDF curve for the attention-based LSTM predictor is notably steeper and shifted towards higher SE values, approaching closely the ideal (no-delay) performance benchmark. This finding underlines its superior predictive capabilities, significantly reducing the negative impacts associated with outdated channel information.

Extending this analysis to individual SE distributions, Figure 16 further illustrates the substantial advantages of the attention-enhanced LSTM predictor. Not only does it deliver improved average SE, but it also ensures more balanced and equitable resource allocation among subnetworks. For instance, while the minimum SE with the sample-and-hold approach is approximately 4.5 bps, the attention-enhanced LSTM predictor significantly enhances this metric to approximately 5.7 bps/Hz. This improvement of over 25% underscores the predictor's effectiveness in mitigating delay-induced performance reductions, thus enhancing fairness by significantly reducing the occurrence of subnetworks experiencing low SE values due to delayed CSI.

These findings strongly validate the effectiveness of the attention-enhanced LSTM predictor in addressing delay-induced impairments. The observed improvements in both average and individual SE clearly demonstrate the advantage of incorporating advanced spatio-temporal attention mechanisms within the prediction framework, enabling efficient and fair resource allocation in delay-sensitive wireless communication systems. In typical deployments, the CSI reporting delay from nodes to the centralized RRM unit is significantly shorter than the 400 ms (4 samples $\times$ 100 ms) considered here. Therefore, the delay conditions examined in this work represent a conservative scenario, further underscoring the robustness of the proposed predictor.

3.2 Distributed RRM for in-X subnetworks

Solutions such as those presented in the previous section can only be applied in case subnetworks are able to communicate with the parent network and the central controller, which can perform centralized decision. Here, we are presenting a solution tailored instead to distributed deployments, where decisions are taken individually at each subnetwork.

In spectrum-constrained environments, wireless networks and subnetworks must share the same frequency bands, leading to interference between communication links. This is a key challenge addressed by the Subnetwork Co-existence in Factory Hall use case, where multiple subnetworks are deployed in closed proximity within a shared industrial space, requiring careful coordination to manage mutual interference.

Effective RRM is essential to mitigate this issue. Building upon the framework described in deliverable D4.1, which proposed a distributed AI-driven solution based on GNNs and over the air message passing, this section presents validation results obtained through simulations. This distributed power control solution employs over-the-air aggregation of pilot signals, enabling subnetworks to indirectly exchange interference-related information. This method dynamically optimizes transmission power allocation, minimizing inter-subnetwork interference while ensuring reliable communication. The solution has been integrated and tested within a 3GPP-compliant simulation environment.

3.2.1 System Model

3.2.1.1 General Architecture

Figure 17 presents a high-level depiction of the proposed RRM framework, where subnetworks (labelled as SN1 and SN2) operate in overlapping frequency bands, as they would if they had to coexist in a Factory Hall. Each subnetwork includes a HC node embedded with an AI/ML model within its radio protocol stack to manage RRM procedures. These HC nodes broadcast Channel State Information Reference Signals (CSI-RS) and information generated by the AI/ML model, namely Neural Network Information (NNI), to the whole network. The respective SNEs use the CSI-RS for channel measurements, while NNI is part of the over-the-air neural network computation. Using the channel feedback received from the SNEs, the AI/ML model computes the optimal transmission power for data communication towards the SNEs. The HC nodes dynamically adjust power control decisions by continuously processing updated CSI and the exchanged NNI to minimize interference across neighbouring subnetworks, thereby improving communication reliability and spectral efficiency.

As NNI is a scalar value presented in next section, we contained it in the CSI-RS by mapping it to its transmit power. Also, we configured all schedulers running on the HC nodes to broadcast the signal on the same resource elements on the resource grid. This results in over-the-air aggregation of all NNIs since they are transmitted simultaneously, and all SNEs receive the aggregated information. This implementation does not add any overhead or latency to the communication and allows the dynamic addition and removal of subnetworks. However, this technique requires a high level of synchronization in the time and frequency domain between the HC nodes because any misalignment can cause interference and corrupted NNI exchange. Additionally, it requires that each CSI-RS is orthogonal with each other for proper over-the-air message aggregation.

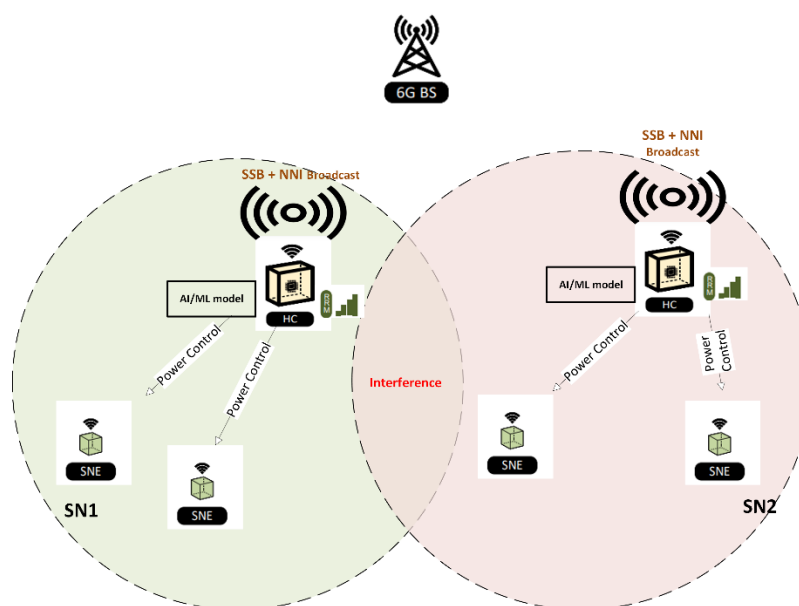


Figure 17: AI/ML RMM through power control

3.2.1.2 AI/ML Model for Distributed RRM

To implement a targeted RRM solution, a distributed power control mechanism for interference management was developed using the Air Message Passing Recurrent Neural Network (Air-MPRNN) paradigm, as introduced in [27]. This approach is built on a GNN-based Message Passing Neural Network (MPNN), which has been theoretically proven to enable efficient distributed power control. The MPNN framework scales with the number of vertices, where each vertex represents a subnetwork, and edges denote both direct and interference links within and between subnetworks. The distributed power control method exploits the temporal correlation of wireless channels, making the network recursive and reducing the time required for output generation. Each MPNN model follows a structured sequence of steps, some of which utilize Multi-Layer Perception (MLP) models. First step is the message generation phase, where each vertex transmits messages to its neighbouring vertices. Then follows the message aggregation phase, where vertices collect incoming messages. Finally, each vertex updates its state, generates an output, and the cycle repeats.

The mathematical representation of this framework is the following:

$$\text{Message generation: } \overline{p}_i(t) = \Phi(e_i(t-1), h_i(t-1)),$$

$$\text{Air message passing: } \overline{a}_i(t) = \sum_{j \neq i}^N (\overline{p}_j(t) |h_{j,i}(t)|),$$

$$\text{State Update: } e_i(t) = U(e_i(t-1), \overline{a}_i(t), h_i(t)),$$

$$\text{Output generation: } p_i(t) = \Omega(e_i(t))$$

Where:

- $\overline{p}_i$ is the generated message of vertex i
- Φ is the message generation MLP
- e_i is an embeddings vector to store the vertex's i state
- $h_i = h_{i,i}$ is the channel coefficient of the vertex's i link between SNE and serving HC
- $\overline{a}_i$ is the aggregated message received on the vertex's i SNE from neighboring HCs
- $h_{j,i}$ is the channel coefficient link of the vertex's j HC and vertex's i SNE
- U is the state update MLP
- p_i is the generated output of the vertex i
- Ω is the output generation MLP

As shown in Figure 18, the MPNN's vertices, embedded within the HC nodes, execute the following steps. The Φ MLP model manages message generation, using vertex embeddings e and channel coefficients h (estimated from channel measurements) as inputs to represent each vertex's internal state. During

the message passing phase, these messages - encoded as pilot transmit power $\bar{p}$ - regulate the CSI-RS transmission power of the HC nodes. Each SNE receives messages $\bar{a}$ from all HC nodes through reference signal measurements and, after extracting the relevant information, reports them back to its serving HC node. The U MLP model updates the vertex's internal state based on the received messages and channel measurements, feeding the updated state recursively into the Φ message-generation MLP for the next iteration. Finally, the Ω MLP model processes the vertex's internal state to generate the p network's output. This output determines the power control settings for data transmission from HC nodes to SNEs, optimizing transmit power to minimize interference between subnetworks.

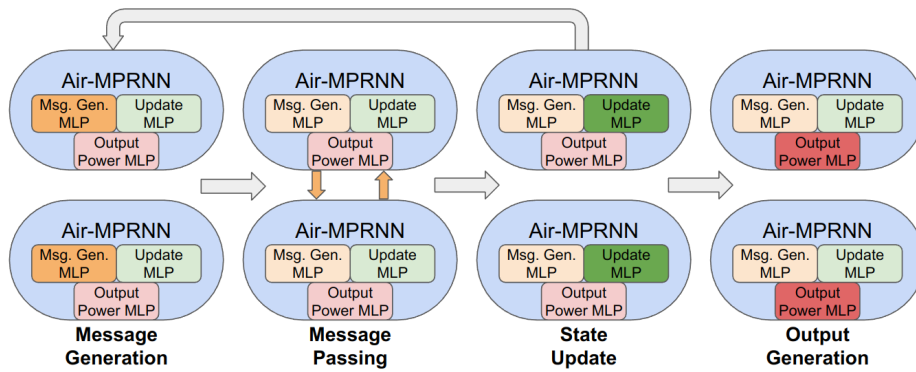


Figure 18: MPNN framework's phases.

Figure 19 illustrates two HC/SNE pairs utilizing MPNN to manage downlink transmit power in a distributed manner, collaboratively optimizing power control to maximize the network's overall sum rate. Each HC node integrates an AI/ML model that functions as an MPNN vertex, enabling decentralized decision-making. Power control operates in the downlink direction, facilitating message passing and interference management. Each SNE receives the serving HC node's signal combined with interfering signals and Gaussian noise. The channel estimation results are then fed back to the HC node via the uplink channel. A centralized manager, hosted at the 6G base station, coordinates the MPNN vertices. Even though the figure depicts only two HC/SNE pairs, the implementation is scalable seamlessly to any number of pairs. Additional HC nodes connect to the 6G parent network, with their serving SNEs receiving the downlink signal alongside interference and noise from other HC nodes.

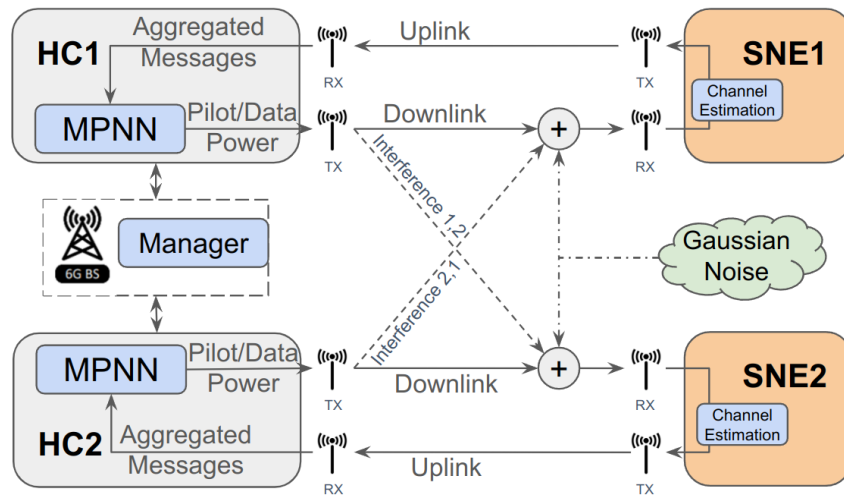


Figure 19: Distributed power control through MPNN for interference management in 6G subnetworks.

3.2.2 Simulation Setup

This section describes the network architecture, the training procedure, and the MPNN integration as a feature into an open-source RAN software application.

3.2.2.1 AI Model Definition and Training

Wireless channels are generated for direct and cross-links to train and evaluate the proposed implementation. Each channel attenuates signal power due to path loss and multipath fading. Channel realizations are based on standardized propagation model found in [27], suitable for short-range scenarios across a wide span of frequencies. Although industrial evaluations in D2.2 [3] rely on the 3GPP industrial channel model, this simulation tries to capture propagation conditions found also in other short-range deployments. This choice reflects the broader applicability of the proposed solution, which while relevant to industrial contexts, can be deployed for use across diverse scenarios. The generated channels span across various SINR values for complete model training. The dataset used to train the neural network comprises 1000 layouts, each with random SINR channels. For training, 80% of the dataset was utilized to update the network's weights, and the remaining 20% was reserved for validation in each epoch to prevent overfitting.

Regarding the network architecture and training parameters, various hyperparameters were tested, selecting those that yielded the best results. Table 2 provides details on the configurations of the message generation MLP, update MLP, and output MLP. The batch size is set to 10 layouts, representing the number of random layouts per epoch used to update the MLPs' weights to optimize the reward function, which is defined as the Sum-Rate of all pair's link Rate.

Table 2: Neural Network's hyper-parameters.

Parameter	Value
Message MLP Φ	{9, 32, 32, 1}
State Update MLP U	{10, 32, 8}
Output MLP Ω	{8, 16, 1}
Embedding size	8
Epochs	100
Batch Size	10
Initial Learning Rate (LR)	0.002
LR decay factor	0.9
LR decay step	10
Optimizer	Adam

The training was conducted in unsupervised manner on a desktop PC with an Intel i9 12th Gen processor, an NVIDIA GTX 1650 GPU, and 32GB of RAM, using the PyTorch library [28]. As shown in Figure 20, the model converged around the 30th epoch, achieving its peak performance with a 9.9% improvement of Sum-Rate for validation compared to the initial training phase, and approximately a 7.7% gain over an Equal Power Allocation (EPA) policy.

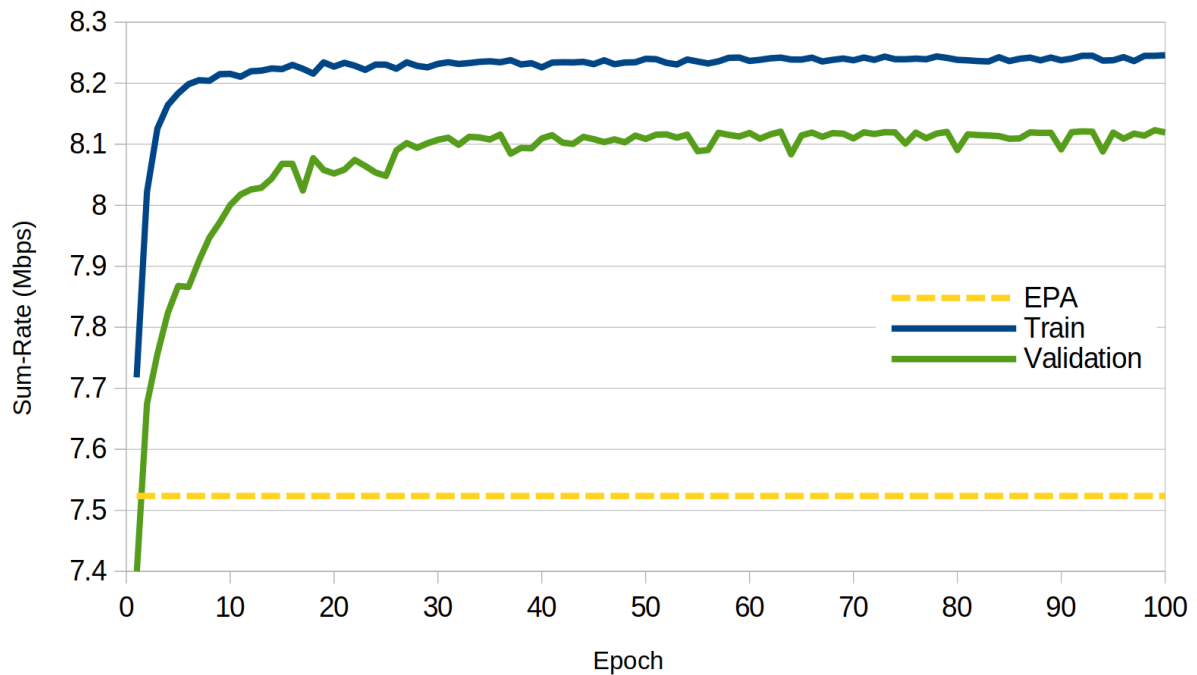


Figure 20: Overall rate during model training for train and test data

3.2.2.2 Software Setup

To deploy the proposed solution in a functional environment, we use the srsRAN software suite alongside GNURadio. srsRAN, an open-source platform, allows researchers, developers, and telecom enthusiasts to implement and experiment with LTE and 5G protocols using software-defined radios (SDRs) or in a fully software-based environment [29]. It provides a flexible framework for developing, testing, and deploying cellular network technologies, making it a valuable tool for advancing wireless communication research. With its modular design and extensive documentation, srsRAN is accessible to both academic and industry professionals. GNU Radio is an open-source software development toolkit that offers signal processing blocks for building communication systems [30]. It features a graphical user interface and enables the creation and deployment of complex radio frequency systems using general-purpose processors instead of specialized hardware. It is compatible with various SDR platforms but is also used in software-only simulations without requiring physical hardware.

Figure 21 illustrates the interaction between all applications within the setup, which consists of gNB-UE pairs implemented using srsRAN. In this context, we consider the gNB an HC node, while the UE represents an SNE. These pairs are interconnected and experience mutual interference through GNURadio, which simulates the wireless channels as previously described. Each pair is associated with its own MPNN vertex that executes MPNN functions. Additionally, we developed a "Centralized Scheduler" process to communicate with all MPNN vertices, coordinating them to enable synchronous message-passing broadcasts. Synchronization is crucial, as vertices may operate at different execution speeds, and an updated power control decision is only valid when all messages have been received and aggregated. Data exchange between processes is handled via inter-process communication using the

ZeroMQ library. ZeroMQ supports multiple programming languages and provides a lightweight, flexible messaging layer, simplifying communication while abstracting the complexities of socket programming [31]. Notable, even though the figure displays only two pairs, it can be easily expanded to more pairs.

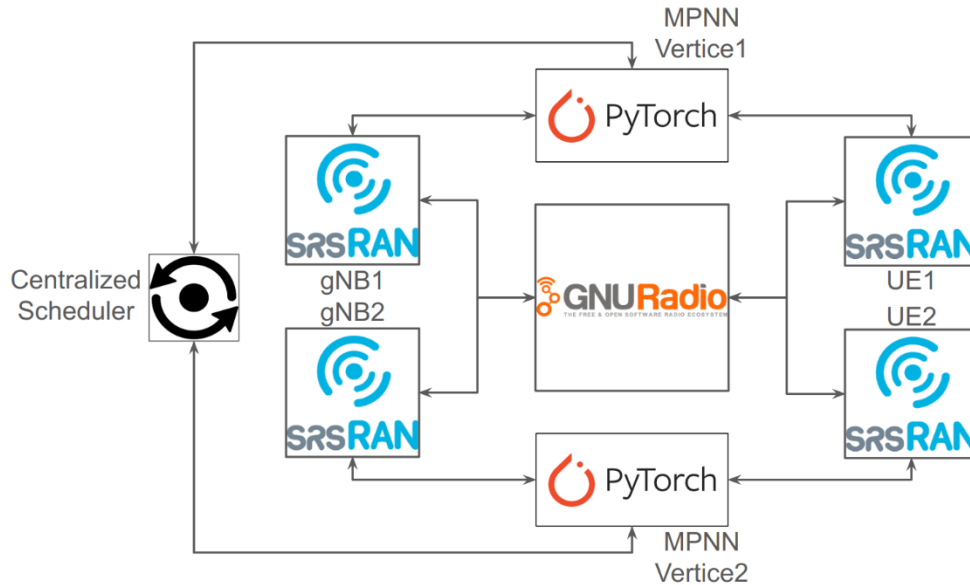


Figure 21: The block diagram of the implementation setup shows how the processes exchange data between them.

3.2.3 Results and Analysis

Measurements were carried out using two, three, and four gNB/UE pairs. All communication channels - including direct and interference links - were simulated using GNURadio. The SINR of each pair was controlled to explore different configurations. Pair-1's SINR was treated as a variable, while the SINR values for the remaining pairs were set to 5 dB, 10 dB, or 20 dB. Each scenario was simulated for at least 20 seconds, beginning with the EPA policy and transitioning to the GNN policy during runtime. During post-processing, the throughput reported by each pair was collected, averaged over time, and used to compute the system's overall Sum-Rate.

Figure 22 shows the total Sum-Rate across all measurement scenarios, varying by the number of gNB/UE pairs. Results using the MPNN approach are represented with solid lines, while the EPA baseline is shown with dashed lines. As expected, the Sum-Rate tends to rise by either increasing SINR values or using a higher number of pairs. When comparing MPNN to EPA, the proposed method generally matches or outperforms EPA - except in low-SINR scenarios with only two pairs. In situations where all pairs experience similar SINRs, there is typically no performance gain. However, the more the channel conditions differ among the pairs, the greater the advantage provided by the MPNN. As a notice, with two pairs, the Sum-Rate is capped at roughly 56 Mbps due to the numerology limiting the peak rate per pair to about 28 Mbps.

The scenarios where MPNN outperforms EPA were anticipated. When channel conditions are similar across all pairs, equal power allocation tends to be optimal, as no pair is disadvantaged. In contrast,

when some pairs have more favourable channels, they can reduce their transmit power - and consequently their SINR - to limit interference toward others, improving overall throughput. The best performance was observed in a three-pair scenario with SINRs of 28 dB, 5 dB, and 10 dB, where MPNN achieved a maximum gain of around 13.16% over EPA. Although the average improvement across all tests was around 7%, this seemingly modest gain can become substantial when aggregated across a larger network with multiple subnetworks, leading to significant overall performance benefits.

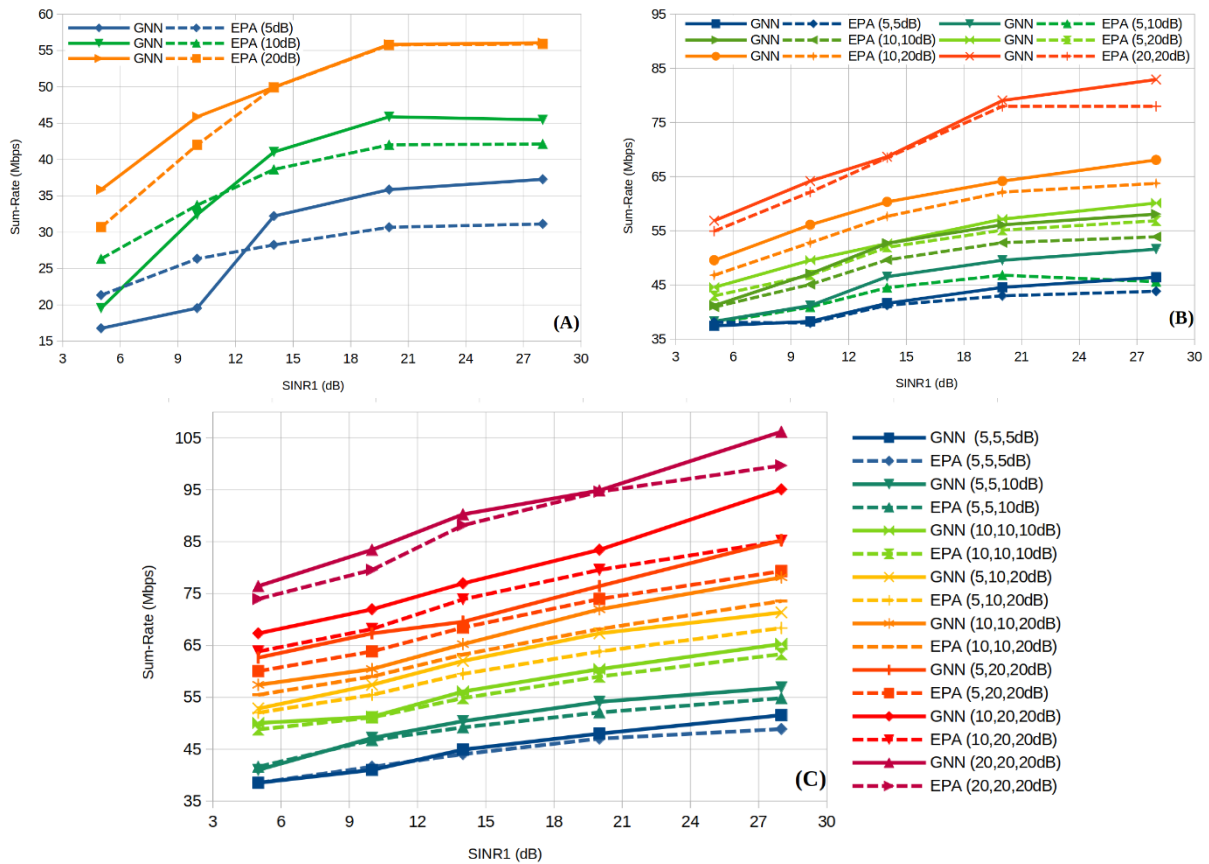


Figure 22: Sum-Rate of the considered 6G subnetwork over SINR1 (pair-1) for (A) one, (B) two, and (C) three interfering pairs with different SINR, such as 5, 10 and 20 dB considering the MPNN (GNN) and the EPA solutions.

3.3 Summary

In this chapter, we have addressed critical challenges in RRM for densely deployed and highly mobile in-X subnetworks, emphasizing the issues arising from outdated CSI and external interference. By developing a novel Spatio-Temporal Attention-Based LSTM model, we demonstrated substantial improvements in predicting future channel states, thereby enabling proactive and informed RRM decisions. Additionally, our proposed resilient DNN framework has shown significant benefits in jointly optimizing sub-band allocation and power control, effectively enhancing spectral efficiency, ensuring robust QoS, and mitigating performance degradation even under highly dynamic and interference-prone industrial environments. Overall, these contributions provide a strong foundation for reliable, efficient,

and adaptive resource allocation solutions, paving the way toward practical and robust 6G-enabled industrial wireless networks.

Furthermore, the simulation-based validation presented in the last section demonstrates that the distributed AI-driven power control solution, utilizing the Air-MPRNN framework, effectively mitigates interference between subnetworks operating in constrained spectral environments. The implemented framework specifically showcases the capability of HC nodes to dynamically coordinate their transmission power decisions by indirectly exchanging interference information through aggregated pilot signals, as part of the RRM role they are envisioned to support. This decentralized approach enhances overall network reliability and throughput performance, particularly under heterogeneous SINR conditions.

4 GOAL ORIENTED RRM

In this chapter, we address the challenge of inter-subnetwork coexistence in dense environments such as factory floors, where multiple robots - each equipped with its own subnetwork - can cause mutual interference. This problem, illustrated in Figure 23, builds on the formulation introduced in [2] and targets the factory floor use case proposed in [3]. Differently from the approaches presented in chapter 3, we investigate here a goal-oriented solution where radio resources are optimized with awareness of the underlying industrial actions and missions. The work presented here has been done in collaboration with the 5GSmartFact project [32].

4.1 Velocity Control for Inter-Subnetwork Interference Mitigation in Mobile Subnetworks

Conventional interference-mitigation methods in a subnetwork context, such as transmit power control or channel allocation, become less effective under tight spacing because transmitters remain in proximity, and interference grows rapidly with density. Instead, we propose a communication-aware dynamic speed control (CADSC), whereby each mobile robot adjusts its speed to maintain an acceptable distance from others, thus alleviating mutual interference. Crucially, signal-to-interference-plus-noise ratio (SINR) requirements must be upheld for ultra-reliable low-latency communication in each subnetwork. We employ reinforcement learning, specifically, proximal policy optimization (PPO), to learn an optimal control policy that balances travel-time minimization with stringent SINR constraints. Through simulations, we show that CADSC significantly improves SINR reliability (up to over 95% probability of meeting the SINR threshold) at the cost of only a modest increase in average travel time compared to a simple “maximum speed” control policy. It worths emphasizing that our results were published in [33].

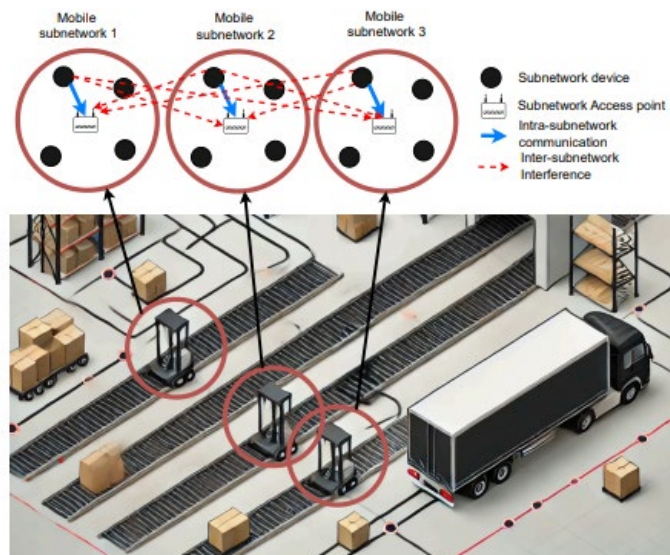


Figure 23: Each mobile robot carries a subnetwork to facilitate the wireless communication between devices in the robot. There is strong mutual interference between the subnetworks due to the proximity of the robots on the factory floor.

4.2 System Model

Each robot carries one subnetwork consisting of a device (e.g., a sensor or controller) and an onboard access point. All subnetworks share the same frequency spectrum. Let N denote the total number of mobile subnetworks (robots). At time t , subnetwork n transmits with power p_T^N and experiences interference from all other subnetworks $m \neq n$, whose transmit powers are p_T^M . The SINR for subnetwork n at time t is:

$$\text{SINR}_n^{(t)} = \frac{p_T^n h_n^{(t)}}{\sum_{m=1, m \neq n}^N p_T^m h_{m,n}^{(t)} + \sigma^2}$$

where:

- $h_n^{(t)}$ is the desired link gain from the subnetwork's device to its own AP,
- $h_{m,n}^{(t)}$ is the interference link gain from subnetwork m to AP n ,
- σ^2 is the noise power.

In realistic industrial deployments (e.g., 6G factory scenarios), channel factors include path loss, shadowing, and small-scale fading, often modelled via *3GPP indoor factory* channel models (as in [10], [34]). Because of *dense clutter* or line-of-sight conditions, interference intensifies whenever subnetworks move too close to one another, as demonstrated in [4].

4.2.1 Reliability Constraints (BLER Model)

Since each subnetwork may handle URLLC traffic [3], where reliability targets can be on the order of 10^{-5} or even 10^{-6} [3], we use the finite-block length *block error rate (BLER)* formulation from [35]. This approach captures the fundamental trade-offs inherent in URLLC and aligns well with the use cases outlined in [3]. Concretely, for a packet of b bits transmitted in a time slot of τ seconds over bandwidth B , the number of channel uses is $\psi = 2B\tau$. The BLER can be approximated as:

$$\text{BLER}_n^{(t)} = Q \left(\frac{\frac{\psi}{2} \cdot \log_2 \left(1 + \text{SINR}_n^{(t)} \right) - b + \frac{1}{2} \cdot \log_2(\psi)}{\sqrt{\psi \cdot V \left(\text{SINR}_n^{(t)} \right)}} \right)$$

where the function

$$V(\text{SINR}) = \frac{\text{SINR} \cdot (\text{SINR} + 2)}{2 \cdot (\text{SINR} + 1)^2} \cdot (\log_2 e)^2$$

captures the channel dispersion. Achieving a sufficiently low BLER implies maintaining an SINR above some minimum threshold ($\text{SINR}_{n,\min}$).

4.3 Robot Mobility and Speed Control Problem

4.3.1 Robot Motion Model

The subnetwork is installed in a mobile robot navigating a target course as shown in Figure 24. We assume each robot (carrying subnetwork n) navigates a preplanned trajectory of length (e.g., 25 meters). The robot state $x_n^{(t)}$ includes:

1. Lateral error e (distance from the center line of the path),
2. Yaw angular error θ ,
3. Derivatives of these errors $\dot{e}$ and $\dot{\theta}$,
4. The difference between the robot's current speed and its target speed v .

Hence $x_n^{(t)} = [e, \dot{e}, \theta, \dot{\theta}, v]^T$. The dynamics are discretized as:

$$x_n^{(t+1)} = A_n x_n^{(t)} + B_n u_n^{(t)}$$

Here, $u_n^{(t)}$ contains the robot's *angular* and *linear* acceleration inputs. A linear quadratic regulator (LQR) typically tries to drive e, θ , etc to zero errors by applying suitable accelerations. Crucially, that LQR relies on a *target speed* input $\bar{v}_n^{(t)}$, which we plan to adapt dynamically, to limit or increase each robot's speed based on interference conditions.

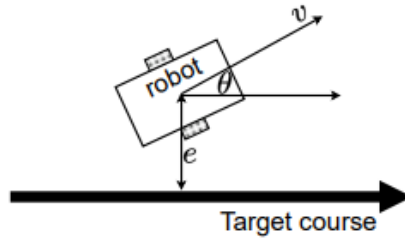


Figure 24: Mobility model of the robot. The state consists of the lateral error e , the yaw angular error θ , their derivatives and the velocity error v .

4.3.2 Optimization problem formulation

The overarching objective is to ensure that each robot completes its mission quickly while maintaining a minimum SINR for its intra-subnetwork communication. Minimizing travel time generally means choosing $\bar{v}_n^{(t)} \simeq v_{\max}$, but if multiple subnetworks cluster, the resulting interference can undercut the required SINR threshold. Hence, the problem is stated as:

$$\text{maximize}_{\bar{v}_n^{(t)}} \sum_{n=1}^N \sum_{t=0}^T \left\| \bar{v}_n^{(t)} - v_{\max} \right\|,$$

subject to:

$$SINR_n^{(t)} \geq SINR_{n,min} \quad , \quad v_{min} < v_n^{(t)} \leq v_{max}.$$

This reflects a *trade-off*: robots want to stay near v_{max} but must scale back whenever grouping too closely threatens the SINR constraint.

4.4 Reinforcement Learning Approach

4.4.1 PPO-Based Speed Control

Proximal Policy Optimization is used due to its robustness in continuous action spaces. We propose a centralized RL agent that, at regular intervals, observes the state of all subnetworks (channel gains, SINRs, current speeds) and outputs a continuous action vector $\{a_n^{(t)}\}$. Each component $a_n^{(t)}$ is then scaled to define the target speed $\bar{v}_n^{(t)} \in [0, v_{max}]$. The potential delay introduced by collecting state information from all subnetworks and transmitting it to the central agent for action computation is not accounted for in this study.

The observation vector is defined as following:

$$R^{(t)} = \sum_{n=1}^N \left(\|\bar{v}_n^t - v_{max}\| - R_{SINR,n}^{(t)} \right),$$

where:

$$R_{SINR,n}^{(t)} = K, \text{ if } SINR_n^{(t)} < SINR_{n,min}, \text{ and } 0, \text{ otherwise.}$$

A larger penalty constant K enforces stricter adherence to SINR constraints but may force the policy to reduce speeds more often. PPO updates an actor network (which selects actions) and a critic network (which estimates the value function) with a clipped objective, stabilizing learning by limiting large steps in policy space.

4.5 Evaluation and Results

4.5.1 Implementation and training

The neural networks (actor and critic) each have two hidden layers of 256 neurons and take as input a flattened state vector of dimension > 100 in the test scenario with 10 robots. During training:

1. Each episode simulates the robots traveling a set distance, gathering transitions $(s^{(t)}, a^{(t)}, R^{(t)}, s^{(t+1)})$.

2. The agent periodically updates its parameters using minibatch gradient descent on both the critic's value loss and the actor's clipped surrogate objective.
3. An exploration variance or standard deviation ρ for the action distribution is gradually decayed from a higher value to a smaller one.

After convergence (often hundreds of thousands of time steps), the trained policy can be deployed: at each sampling interval, the agent computes the speeds for all the different robots based on the latest channel/interference conditions, maximizing the reward. This behaviour is illustrated in Figure 25 for different CADSC training penalty constant K .

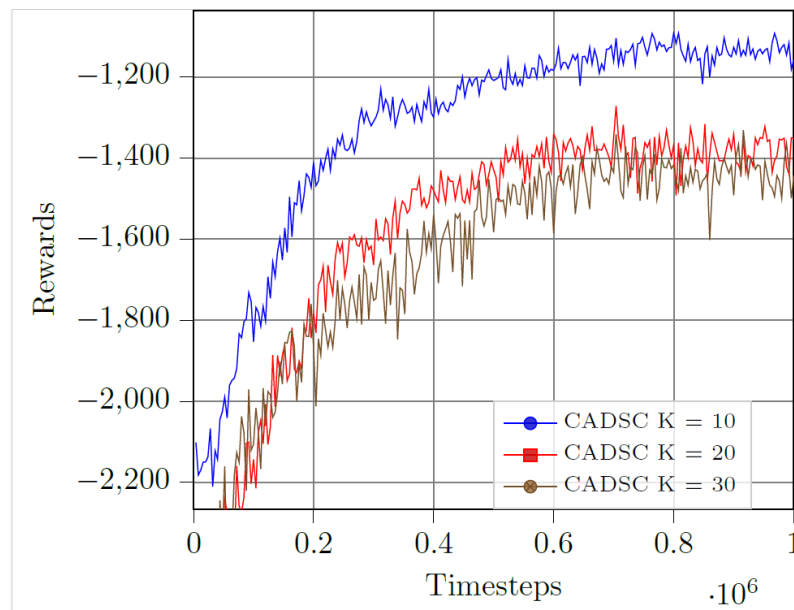


Figure 25: Cumulative reward per episode during the PPO training phase of the communication-aware dynamic speed control policy.

4.5.2 Simulation Setup

We simulate 10 robots traveling 25 meters on parallel tracks in a 30 m x 30 m factory environment. The main parameters are derived from the use-case analysis in [3] and can be summarized according to Table 3.

Table 3: Simulation Assumption

Parameter	Value	Parameter	Value
Factory area	30m x 30m	Number of mobile subnetworks	10

Subnetwork radius	1m	Number of devices per subnetwork	1
Max speed	2m/s	Travel distance	25m
Clutter density, Clutter size	0.2, 2m	Correlation distance	10m
Shadowing std (LOS, NLOS)	4dB, 5.7dB	Path loss exponent (LOS, NLOS)	2.15, 2.55
Maximum transmit power, Pmax	0 dBm	Bandwidth	1.3 MHz
Data size	32 bytes	Center frequency	10 GHz
Noise figure	10 dB	Traffic/Sampling interval, dt	20ms
Traffic type	Periodic	Latency	0.5ms

We compare, as baselines:

1. **Max-Speed (Fixed Power):** Robots move at full speed; each device transmits at 0 dBm0.
2. **Max-Speed + Transmit Power Control:** A non-convex optimization (SLSQP) tries to allocate powers to maximize throughput fairness.
3. **CADSC (PPO):** The proposed RL-based speed control, where transmit power is fixed but robot velocities are *adaptively* adjusted.

4.5.3 Performance Evaluation and Discussion

The evaluation results, as illustrated in Figure 26 and Figure 27, demonstrate the impact of communication-aware dynamic speed control (CADSC) on SINR reliability, travel time, and BLER distribution. In Figure 26, we observe that in scenarios where robots move at maximum speed without speed control, only about 75% of the robots achieve the required SINR threshold. By contrast, when CADSC is applied with a moderate penalty constant of $K=10$, the probability of meeting the SINR constraint increases to approximately 95%, with even higher reliability observed for larger values of K .

Introducing communication awareness through CADSC slightly increases the average travel time. For instance, while robots operating at maximum speed complete their journey in approximately 12.5 seconds, CADSC-adjusted robots take around 14 to 15 seconds. However, this minor delay is offset by a

significant improvement in SINR compliance, particularly in highly dense environments where interference is a critical issue.

The BLER distribution in Figure 27 further highlights the benefits of CADSC. Most robots operating under CADSC maintain BLER values in the range of 10^{-7} to 10^{-8} , demonstrating a high level of reliability. In contrast, purely adjusting transmit power proves ineffective in high-density scenarios, as interference from closely spaced subnetworks remains too strong when speed control is not applied. Additionally, the training curves referenced in Figure 25 show that the cumulative reward per episode stabilizes after a few hundred thousand training steps, indicating that the RL-based CADSC policy successfully learns an optimal strategy. The impact of penalty scaling is also evident: as the penalty value K increases, the RL agent prioritizes avoiding SINR violations more strictly, resulting in more conservative speed allocations. While this leads to better communication reliability, it also slightly prolongs travel time, reflecting the trade-off between network performance and mobility efficiency.

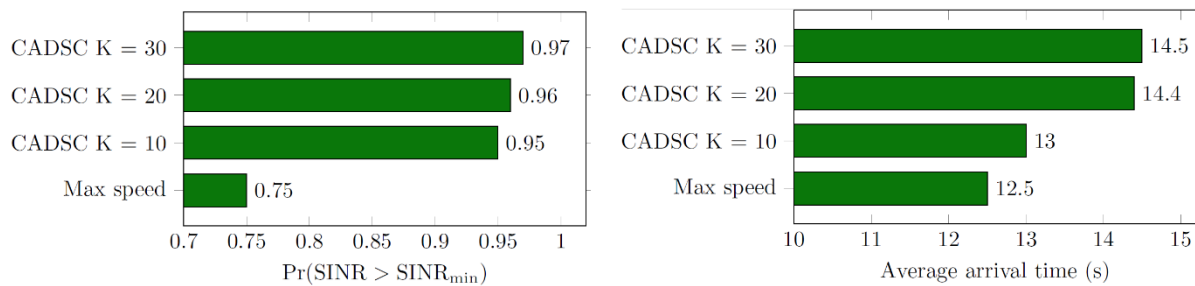


Figure 26: The average arrival time for a travel distance of 25m vs the probability.

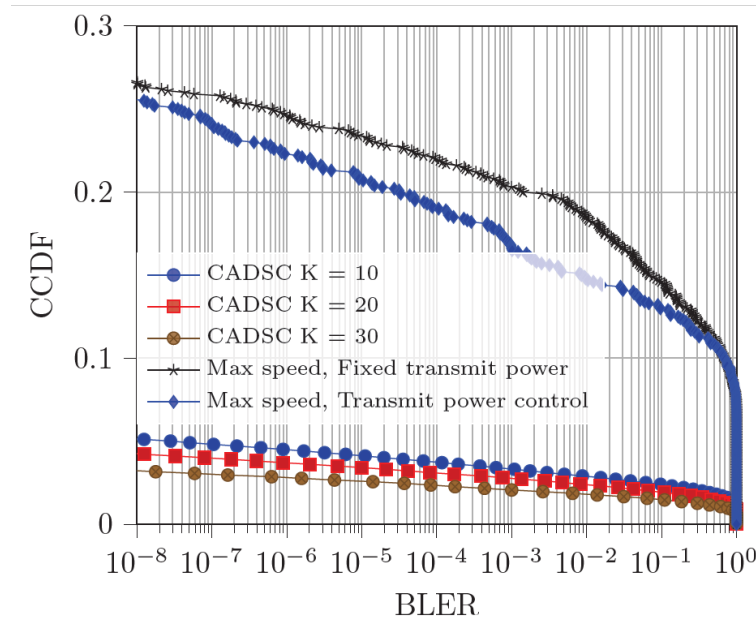


Figure 27: The Complementary Cumulative Distribution Function (CCDF) of the achieved block error rate.

4.6 Summary

We show that dynamic speed control adjusting the robots' velocities to maintain better spacing can significantly mitigate interference between adjacent subnetworks, where classical approaches (e.g., transmit power control alone) offer limited gains. By leveraging reinforcement learning, we build a flexible, data-driven policy that automatically balances navigation performance (travel time) and communication constraints ($\text{SINR} \geq \text{SINR}_{\min}$).

Key Takeaways

1. Controlling motion jointly with communication constraints yields superior interference management in future 6G subnetwork systems.
2. The penalty for slower speed is moderate, yet the SINR reliability improvement is large.
3. PPO manages the continuous action space of multi-robot speed settings and converges to stable solutions.

Looking forward, more advanced schemes could combine path planning (not just speed control) with interference-aware decision-making or jointly optimize transmit power and speed, potentially extending the benefits to even denser or more dynamic industrial scenarios. Moreover, future work should include feasibility studies that account for stale or delayed state information resulting from transmission latency and overhead. The proposed CADSC approach also underscores the importance of co-designing robotic mobility and wireless resource management in next-generation networks.

5 ENABLERS FOR RRM IN SUBNETWORKS

In this chapter, we discuss key aspects and potential solutions to enable radio resource management in both licensed and unlicensed spectrum bands applied for supporting the required communication needed in subnetwork solutions. The presented study and solutions are a continuation of the work on previous deliverable D4.1 [2]. Here as well, we assume that subnetworks can be built up on top of relevant features of Sidelink, specified in 3GPP standards. We will discuss potential enhancements needed for effective subnetwork operations, and revisit previous proposals and observations from deliverable D4.1. Additionally, we will provide further studies and considerations for improving subnetwork performance.

Subnetwork communication can be closely related to sidelink communication, particularly in scenarios where direct device-to-device (D2D) interactions are essential. For example, inter-subnetwork communication can occur between HC devices that act as access points for different subnetworks in an indoor interactive gaming scenario, as described in D2.2 [3], where multiple gaming devices such as VR headsets or gaming consoles of different players (each representing a subnetwork) may need to exchange information to synchronize eXtended Reality(XR) scenes. The HC devices on these consoles can communicate directly to share pose and orientation data, coordinate split-rendering operations, and ensure a seamless immersive experience. Similarly, intra-subnetwork communication can occur between SNE devices within a subnetwork or between a SNE and a HC device. In an XR setting, devices such as VR headsets and sensors (which may be categorized as LC or SNE devices) within a subnetwork may communicate directly to perform coordinated actions, such as synchronizing visual and sensory inputs. For example, in an indoor interactive gaming scenario, VR headsets worn by players can communicate directly with sensors and actuators attached to their bodies to track movements and provide real-time feedback. Additionally, these VR headsets can communicate directly with a central processing unit or gaming console (HC device) to provide video data and feedback, enabling responsive gaming experiences.

Sidelink communication, as specified in 3GPP standards [36], provides a rich framework which, depending on the configuration and deployment scenario, allows direct communication between devices to work independently of a central network infrastructure. This is particularly beneficial for subnetworks, which often require localized communication capabilities. The 3GPP sidelink specifications include several features that can be leveraged for subnetwork communication:

- In-coverage and out-of-coverage operation: Sidelink supports communication both within the coverage area of a base station and in scenarios where devices are outside the coverage area.
- Power saving features: Mechanisms such as Discontinuous Reception (DRX) help reduce power consumption, which is crucial for battery-operated devices.
- Inter-UE coordination (IUC): Sidelink includes features for coordinating transmissions between user equipment (UEs) to minimize interference and improve communication reliability.
- Support for licensed and unlicensed bands: Sidelink has been specified to operate in both licensed and unlicensed spectrum, providing flexibility for deployments.

- Centralized and distributed resource allocation modes: Sidelink supports two allocation modes. In Mode 1, the base station schedules the resources for the devices to communicate with other devices via sidelink. In Mode 2 the devices autonomously select the resources from a resource pool which can be pre-configured or configured by a base station.

Despite the advances in sidelink communication, several enhancements are needed to fully support the stringent requirements of subnetworks. These enhancements can include improvements to the air interface, such as enabling in-coverage or out-of-coverage subnetwork APs to sense the channel, acquire resources, and schedule those resources to subnetwork devices, extending beyond the capabilities of Sidelink Mode 2. Additionally, architectural enablers are needed to enhance authentication and policy enforcement in subnetwork scenarios, ensuring secure and compliant operations. Furthermore, enhancements in channel access mechanisms are recommended to better accommodate subnetworks operating in unlicensed bands, providing more efficient and reliable communication in these environments.

In deliverable D4.1, we proposed enhancement for subnetwork resource pool reservations as well as for in-band emission (IBE) mitigation, as summarized below:

- Subnetwork Resource Pool Reservations: We suggested that HC devices acting as access points should be capable of reserving shares of the sidelink resource pool for intra-subnetwork communication. This includes the use of enhanced Physical Sidelink Control Channel (ePSCCH) for HC-to-HC coordination and subnetwork-specific PSCCH (sPSCCH) for informing SNEs about available resources.
- IBE Mitigation for Subnetwork Resources: We highlighted the impact of IBE on subnetwork communication, particularly when using frequency domain multiplexing. Potential enablers for mitigating IBE include time-domain multiplexing, stricter IBE requirements in UE RF standards, adapting transmission starting points, and IBE-aware inter-UE coordination mechanisms.

In this deliverable, we will provide further studies on the operation of SNE-HC and HC-HC communication, both in separate bands (as also discussed in D4.1) and in shared bands (not treated in D4.1). We will explore further issues in relation to the subnetwork operation in unlicensed bands with respect to channel access leading to considerations for supporting semi-static channel access for sidelink. We also discuss further the impacts of IBE in these settings. Finally, we will discuss the opportunistic usage of licensed available resources for subnetworks to avoid sensitivity degradation issues for wide area communication.

5.1 HC-HC and SNE-HC sidelink communication in a shared band

In deliverable D4.1, we discussed the problem of IBE and its impact on subnetwork communication. IBE is the result of power leakage from the allocated transmission resource to the non-allocated transmission resource in the frequency domain (possibly being used by another device), primarily caused by transceiver impairments such as IQ imbalance, nonlinearity of RF components, quadrature imbalance,

and carrier leakage [37]. This leakage can significantly degrade the SINR for adjacent transmissions, leading to reduced communication reliability and efficiency.

IBE can cause substantial interference when two transmitters use intra-band adjacent resources, particularly in scenarios where one transmitter is close to a receiver while another transmitter is far from the receiver (near-far problem). This interference is exacerbated in unlicensed bands where interlaced resource allocation is used to meet regulatory requirements for occupied bandwidth (OCB) and power spectral density (PSD). In D4.1, we evaluated the impact of IBE in an indoor scenario, representing an immersive education use case. The evaluation considered both inter-subnetwork HC to HC communication (longer distance and higher power) and intra-subnetwork SNE to HC communication (shorter distance and lower power). The results showed significant performance degradation in terms of SINR, especially under high load conditions and with interlaced resource allocation. Some potential solutions discussed in D4.1 include enforcing stricter RF specifications for lower IBE, adjusting transmission starting points, and employing IBE-aware inter-UE coordination mechanisms.

Here, we provide further analysis using updated assumptions based on consumer subnetworks use cases, applicable for example to immersive education and interactive gaming scenarios, as described in D2.2 [3]. In this part (and in section 5.2), we consider the use of unlicensed spectrum for subnetworks, therefore the focus here is on the case where interlaced resource allocation is applied. In addition to the IBE aspect aforementioned, we discuss other limiting aspects in this kind of environment.

Setting the scene

In unlicensed spectrum bands, where multiple technologies like Wi-Fi, LTE-LAA, and 5G NR-U coexist without centralized coordination for using the spectrum resources, channel access procedures are critical to ensure fair and interference-free operation. Regulatory bodies (e.g., ETSI) mandate protocols such as Listen-Before-Talk (LBT) to prevent collisions and promote a fair spectrum sharing. These procedures require devices to check the channel availability through an energy sensing before transmitting, minimizing disruptions to incumbent systems like Wi-Fi. For cellular systems operating in sub-7 GHz unlicensed bands (e.g., NR-U), 3GPP specifies dynamic channel access, such as Type 1 and Type 2 procedures, as well as semi-static channel access procedures in order to comply with global regulations [38].

In this part we will assume that dynamic channel access is used for subnetworks based on Type 1 and Type 2 procedures. Type 1 involves a full contention-based LBT process with random backoff, requiring devices to sense the channel for a random duration to determine if it is idle before transmission. The process includes decrementing a counter based on idle sensing slots until the counter reaches zero, and after that the transmission is allowed. Type 2 procedures employ a predefined sensing period before transmission under restricted conditions, for example, Type 2A (single sensing in at least a 25 μ s channel idle gap) and Type 2B (single sensing in a 16 μ s channel idle gap) can only be used within a channel occupancy time initiated after a Type 1 procedure. Type 2C allows a short transmission without prior sensing, limited to 584 μ s transmission. These procedures are defined for what is known as load-based

equipment (LBE) channel access, and supported for DL, UL as well as for sidelink communication in unlicensed bands.

The evaluation scenarios which we assume here consider subnetwork deployments where SNE-HC and HC-HC communication can happen in separate bands (i.e., a channel of 20 MHz is dedicated to each type of communication) or in shared bands (i.e., a common channel of 20 MHz is shared by both type of communication). Figure 28 displays an example of the layout of such scenario where there may be possible simultaneous communication between HC to HC and between SNE to HC from different subnetworks coexisting in the same environment.

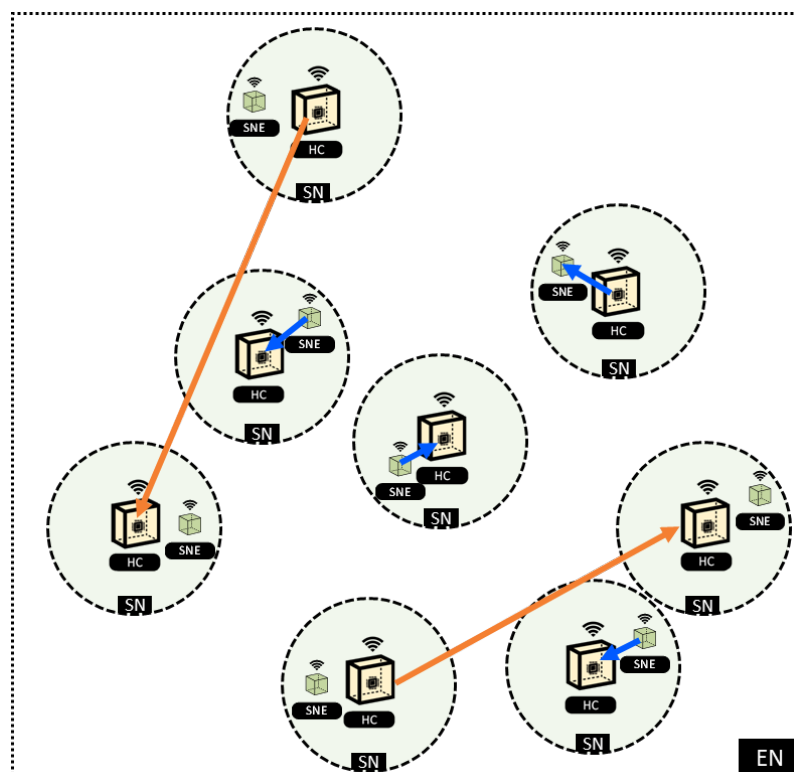


Figure 28: Example of subnetwork deployment where SNE-HC (blue) and HC-HC (orange) communication may coexist.

The HC devices are assumed to be of a higher power class, meaning they can transmit with a power 10 times higher than a SNE, when they are exchanging data with an HC device of another subnetwork. That may be for delivering low-latency position/orientation and command data between interacting users, as well as resource coordination signals for RRM based on sidelink control information and IUC control elements. While the SNEs are assumed to be of a lower power class for transmitting at the short distance to their local subnetwork HC device with a power of -10 dBm. The data may be, for example, local XR traffic such as video/audio and interaction data, as well as feedback signals and channel state reports used for RRM.

In order to mitigate the IBE impact when all these communications coexist, we adopt stricter requirements by considering an IBE general term 10 dB lower than the current limit in the 3GPP

specifications (defined in Table 6.4F.2.3-1 for interlaced RB allocation and Table 6.4.2.3-1 for contiguous allocation in [39]), and compare with the case where the devices just meet the exact minimum requirements. Additionally, in section 5.2, we will discuss other alternatives for mitigating IBE.

Evaluation methodology

The evaluation is performed using system level simulations with the assumptions mainly based on the evaluation methodology adopted for the 3GPP Rel-18 NR Sidelink evolution [40], with some adaptation for the subnetworks use case such as denser deployment in a smaller area and short distance low-power SNE-HC communication. Note that the existent channel model from 3GPP was assumed in this study as no updates has been recommended to the channel model for the indoor consumer scenario in the analysis from 6G-SHINE project's Deliverable 2.3: Radio Propagation Characteristics for IN-X Subnetworks [4].

Following the KPI aspects for the consumer use cases defined in D2.2, we consider the system capacity as described in TR 38.838 for XR traffic as the KPI for the study. The XR capacity can be defined as the maximum number of users per cell with at least 90% of UEs being satisfied. Here we associate a user as one subnetwork, where a SNE transmits to the HC or where an HC transmits to an HC of another subnetwork. The XR satisfaction ratio measures the percentage of XR users that receive at least 95% of their packets within the specified packet delay budget (PDB). For HC-HC communication, generating pose/control traffic, the PDB is assumed to be 10 ms. For SNE-HC communication, generating multi-stream traffic, the PDB is assumed to be 15 ms. Table 4 summarizes the evaluation assumptions.

Table 4: Summary of evaluation assumptions.

Parameter	Value
Scenario	A single room of 20 m x 20 m x 3 m (length x width x height) for consumer use case, e.g., indoor interactive gaming.
Subnetwork deployment	N subnetworks are deployed at a random location in the scenario. Each subnetwork consists of 1 HC acting as AP which communicates with 1 SNE at a time (multiplexing within subnetwork is not considered). SNEs are up to 2.5 m far apart from the HC device which they connect to. The subnetworks do not overlap in space.
Channel model	Indoor mixed office (InH) from 3GPP TR 38.901.
Traffic modelling	XR traffic based on 3GPP TR 38.838. HC-HC: Pose/control traffic (100 B periodic traffic with 4ms interval), PDB = 10ms.

	SNE-HC: Multi-stream traffic (Stream 1: Pose/control; Stream2: XR video frames following a truncated Gaussian with mean 20838 B/frame, minimum 10419 B/frame and maximum 31257 B/frame at 60 fps), PDB = 15 ms.
Antenna configuration	1 TX by 4 RX antenna configuration for SNE-HC. 2 TX by 4 RX antenna configuration for HC-HC.
Carrier frequency and bandwidth	5 GHz carrier frequency with 20 MHz channel bandwidth.
Slot structure	Orthogonal frequency division multiplexing (OFDM) with 15 kHz SCS. Assuming Sidelink slot configuration with 14 symbols per slot. • 4 out of 14 symbols are overhead (to account 2 DMRS, 1 AGC, 1 GP). • control channel (PSCCH) equivalent to 2 RBs of the sub-channel.
Sub-channel configuration	10 sub-channels of 10 RBs per sub-channel. For interlaced allocation, one sub-channel is equivalent to a 10-RB interlace. Up to 5 subchannels can be allocated at a time.
Scheduling	Sidelink mode 2 autonomous resource selection with 100 ms sensing window and 2 ms selection window.
Link adaptation	Link adaptation targeting 10% BLER, following MCS table 4 from TS 38.214 which includes 1024-QAM.
Power control	Fixed transmit power of 0dBm for HC in HC-HC communication and -10 dBm for SNEs in SNE-HC communication.
LBT	LBT procedure in unlicensed (Type 1 and Type 2 within a channel occupancy time according to TS 37.213) with energy detection threshold of -62 dBm.
IBE modelling	Based on minimal requirements from 3GPP TS 38.101-1. - Table 6.4.2.3-1 assumed with contiguous RB sub-channel allocation. - Table 6.4F.2.3-1 assumed with interlaced RB sub-channel allocation.

Analysis of unlicensed band performance

Figure 29 shows the XR satisfaction ratio results versus the number of deployed subnetworks generating the XR traffic in the room.

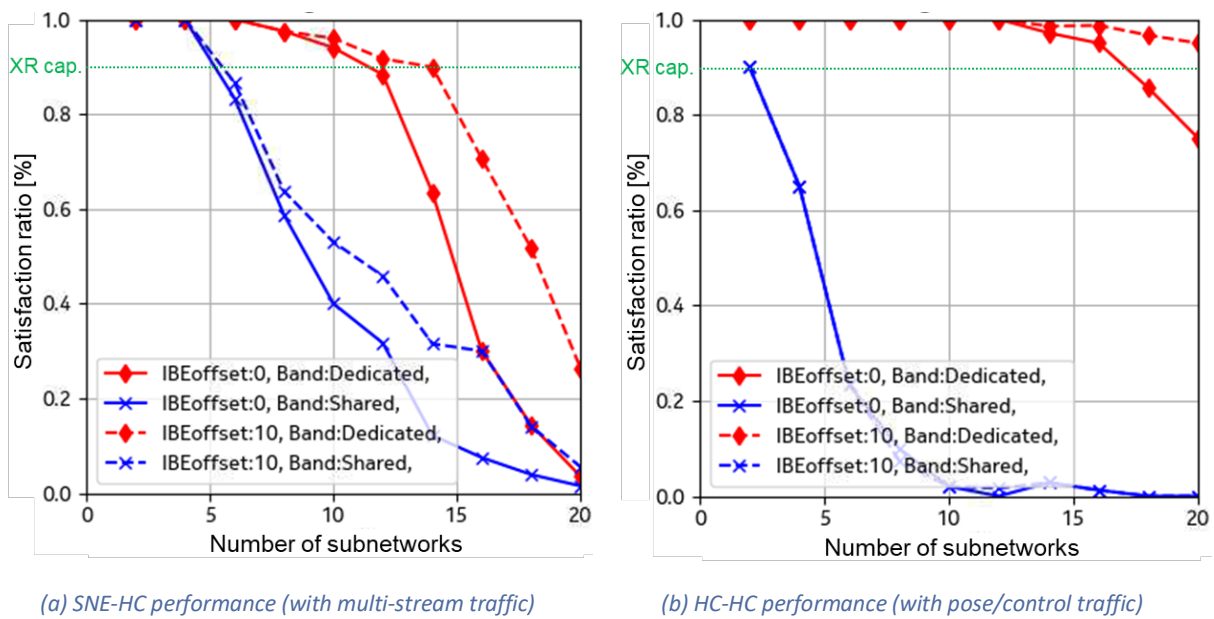


Figure 29: UE satisfaction ratio for different number of subnetworks in unlicensed band with Type1/Type2 channel access.

As expected, the satisfaction ratio decreases more rapidly with an increasing number of subnetworks when using a shared band compared to dedicated bands. A notable observation is the significant degradation in HC-HC communication within shared bands. The high-intensity multi-stream XR traffic of the SNE-HC communications predominantly utilizes the available resources in both frequency (limited to 5 subchannels) and time, resulting in prolonged channel occupancy. This blocks access for HC-HC traffic, leading to considerable performance degradation when they coexist in a shared band.

Regarding XR capacity, the results show that for SNE-HC communication, approximately 5 subnetworks can be supported in a shared band. In contrast, when operating in separate bands, the capacity increases to support 10 subnetworks without IBE mitigation and up to 14 subnetworks with IBE mitigation (i.e., with “IBEoffset: 10”, meaning that the IBE general term is reduced by 10 dB for subnetwork devices with enhanced front-end implementation). For HC-HC communication, the capacity is limited to only 2 subnetworks in a shared band. However, in separate bands, the capacity is significantly higher, supporting 16 subnetworks without IBE mitigation and at least 20 subnetworks with IBE mitigation.

These results highlight the importance of IBE mitigation in enhancing XR capacity and overall network performance. Moreover, they also make it evident that dynamic channel access remains a main bottleneck for performance, particularly in shared bands where resource contention and extended channel occupancy significantly impact communication efficiency. In the next part, we will discuss further enhancements for overcoming these issues.

5.2 Considerations for supporting semi-static channel access in unlicensed bands

In the previous part, it is assumed that the subnetworks are deployed using unlicensed spectrum bands, therefore the devices implement interlaced allocation to meet regulatory requirements. In addition, it

is assumed also that the devices need to perform a LBT procedure to gain access to the channel. This part discusses in more details about potential channel access enhancements for subnetworks.

In addition to the conventional dynamic channel access explained previously, 5G NR also specifies semi-static channel occupancy procedures defined for frame-based equipment (FBE) channel access, as shown in Figure 30. FBE is applicable for environments where the absence of other radio technologies sharing the spectrum is guaranteed, such as by regulation or private premises policies. This procedure involves initiating channel occupancy at fixed periods. The fixed frame period (FFP) can be configured from 1 ms to 10 ms. The base station or UE initiates a channel occupancy by transmitting a DL or UL transmission burst at the beginning of the FFP after a clear channel assessment (CCA) for at least a sensing slot duration (9 μ s). Once the channel occupancy is initiated, DL or UL transmission bursts can occur within the channel occupancy time without further sensing if the gap between bursts is at most 16 μ s. Additionally, no transmissions are allowed during an idle period before the start of the FFP. The idle period should be at least 5% of the FFP with a minimum of 100 μ s [38].

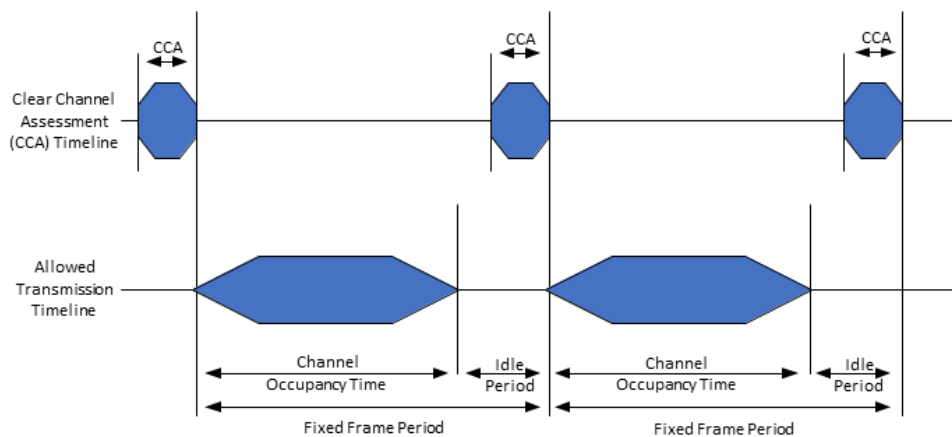


Figure 30: Example of timing of the channel access mechanism for FBE [41]

It should be noted that dynamic channel access is more suited for unpredictable environments, while semi-static channel access is more appropriate for controlled and predictable environments with guaranteed absence of other technologies. Previous works have shown that semi-static channel access is beneficial for low latency communications in controlled environments, due to its predefined channel access timing structure which allows better coordination whereas LBT procedures suffer with unpredictability of the random back-off algorithm and variable contention windows [42].

However, the described semi-static channel access is not specified for sidelink communications in 5G NR. Here we analyse the benefits of supporting semi-static channel access in the future for subnetworks.

For enabling semi-static channel access, we consider a new slot configuration including 2 guard-period symbols instead of the legacy sidelink slot with a single guard-period symbol, as illustrated in Figure 31. That obviously represent a capacity loss for the data channel, however, the extra gap is needed to meet the minimum idle period requirement for an FFP of 2 ms, as described above. Note that a 2 ms FFP is a

typical configuration for RRM test cases defined by 3GPP TS 38.133 [43] for UL and DL, though not yet defined for sidelink. We assume that the FFP configuration is common to all subnetworks. In practice, this configuration should be pre-configured or indicated, e.g., via radio resource control (RRC) signalling to the UEs in the room. Alternatively, the configuration could be broadcast, e.g., by a sync source UE using the sidelink synchronization signal block.

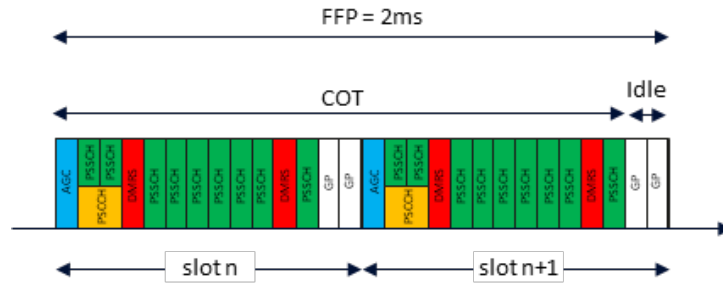


Figure 31: Modified sidelink slot configuration considered for semi-static channel access.

The remaining scenario settings and evaluation assumptions are the same as those of the previous part which are listed in Table 4.

Analysis of semi-static channel access

Figure 32 shows the XR satisfaction ratio results versus the number of deployed subnetworks applying semi-static channel access with the described configuration.

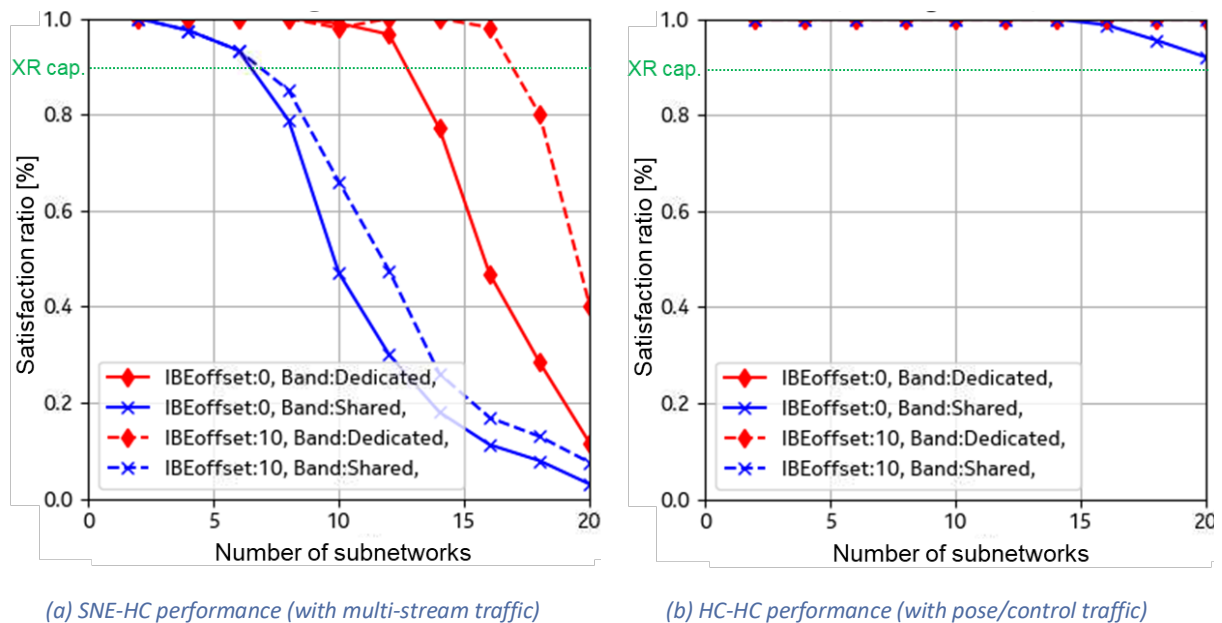


Figure 32: UE satisfaction ratio for different number of subnetworks in unlicensed band with semi-static channel access.

It can be noted that the performance has greatly improved in comparison to the results where dynamic channel access is used in the previous part within section 5.1, especially for the HC-HC communication. For SNE-HC communication, the XR capacity is of approximately 6 subnetworks in a shared band. With separate bands, the capacity ranges from 12 subnetworks without IBE mitigation to 16 subnetworks with IBE mitigation. For HC-HC communication, the XR capacity is increased to more than 20 subnetworks in both shared and separate bands. That is more than 10 times the capacity of dynamic access in shared bands. That is due to the almost negligible LBT blocking enabled by the semi-static channel access applied in the controlled environment. Since the CCA for both SNE-HC and HC-HC traffics are aligned in time during the idle period, they can access the channel simultaneously. The prevention of transmission collisions in this case relies mainly on the sidelink mode 2 distributed resource selection procedure. It should be noted that the FFP may still add a small delay for FFP alignment, since a channel occupancy can only start at the beginning of an FFP. Also, there is small capacity penalty due to the extra guard-period symbol needed in this case to meet the required idle period.

Analysis of IBE aware allocation

In the evaluation above, we show the potential performance improvement when IBE general component is reduced by an offset of 10 dB. However, pursuing such reduction via hardware improvements could translate in increased device cost. An alternative solution is to apply an IBE aware resource allocation targeting to diminish the leakage from the HC-HC communication to the SNE-HC communication and vice-versa when they operate in a shared band.

The solution can be enabled by an IBE aware inter-UE coordination mechanism. That includes the subnetwork devices exchanging information of power class or expected transmit power in the reserved resources of the sub-pools used for HC-HC and SNE-HC. The receiving devices sensing the reservations can then determine how severe the IBE will impact its reception and based on that it provides this information to the transmitting device, such that resources prone to suffer from IBE issues are indicated as non-preferable resources. As shown in Figure 33, the implementation of this solution can be based on enhancing the Sidelink IUC framework such that a UE can determine its preferred/non-preferred resources in IUC scheme 1 considering the impact of IBE from one interlaced RB sub-channel to another. The UEA and UEB in the figure could be two HC devices which coordinate the use of resources for inter- and intra-subnetworks communication.

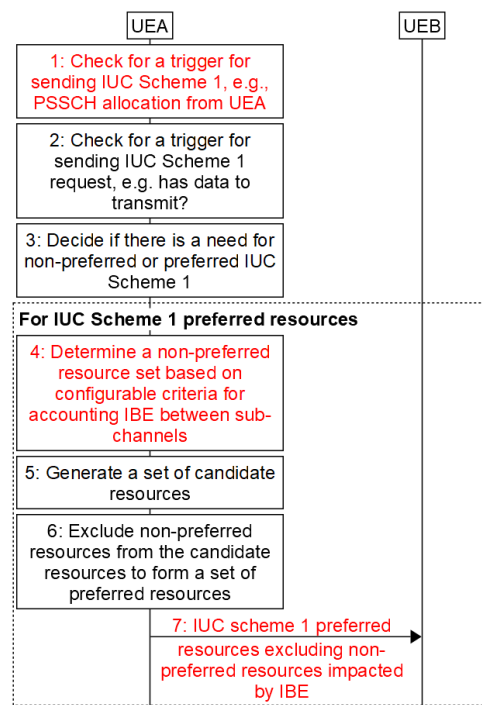


Figure 33: IBE-aware Inter-UE coordination scheme (Red steps highlight impact on existing IUC procedure).

For demonstrating the potential gain, we assume a scenario where both the SNE-HC and the HC-HC communications generate a pose/control traffic as described in Table 4. Note that this traffic assumption is different from the assumption in previous results where SNE-HC generates a high intensity multi-stream traffic, dominating the use of resources and becoming the main source of interference in the scenario. Here instead, with the common traffic model for SNE-HC and HC-HC, a higher number of subnetworks is assumed for the evaluation. This also means that the supported number of subnetworks is not directly comparable as to the earlier results.

For simplicity, the IUC signal in MAC layer is not explicitly modelled in the simulator, meaning that the transmitting devices have ideal knowledge of the interlaced sub-channels which should be excluded for avoiding IBE leakage from the neighbouring subnetworks towards its receiving device.

Figure 34 shows the XR satisfaction ratio results versus the number of deployed subnetworks using IBE aware resource coordination in a shared unlicensed band. Note that the semi-static channel access as described earlier is also applied here.

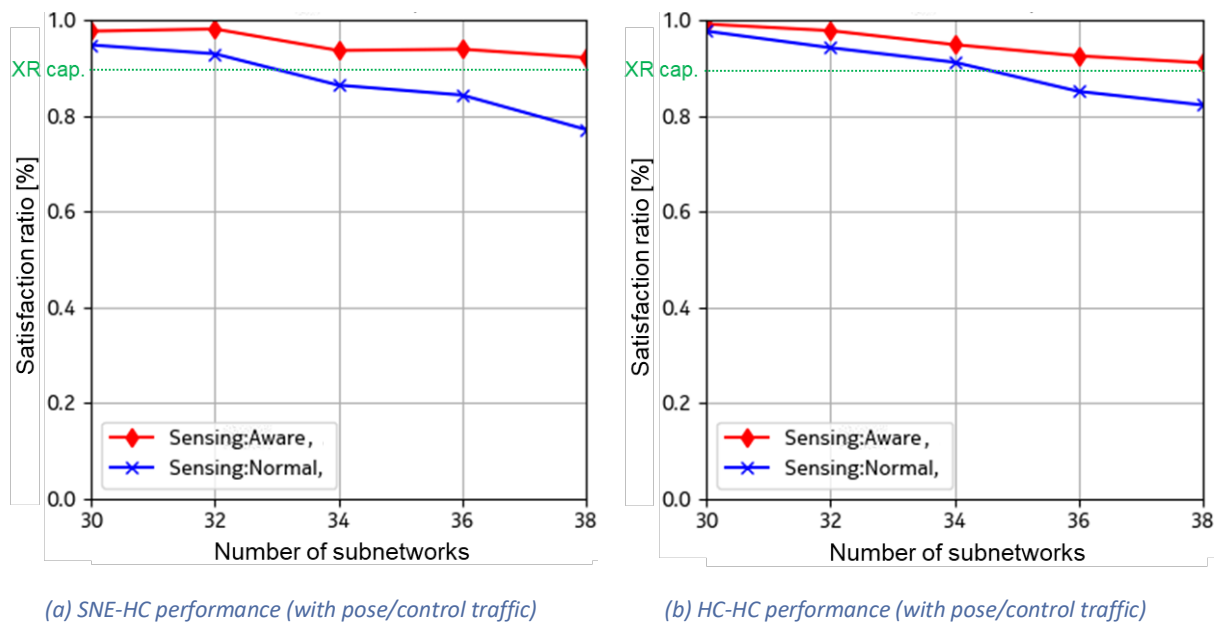


Figure 34: UE satisfaction ratio for different number of subnetworks using IBE aware resource coordination.

The results indicate that the XR capacity increases from 32 subnetworks performing SNE-HC communication when using the normal sidelink resource selection procedure to at least 38 subnetworks when the IBE aware allocation is introduced, i.e., about 19% improvement. Similarly, for HC-HC communication, the XR capacity increases from 34 subnetworks using the normal procedure to at least 38 subnetworks with the introduction of IBE aware allocation.

These results highlight the effectiveness of IBE aware allocation. However, it should be noted that this is mainly beneficial in the cases where the required allocation to convey the traffic is limited to a few interlaces, as is the case for the pose/control traffic, where 1 interlace is sufficient for transmitting the frequent but rather small 100B periodic payloads. That is because in such cases more transmissions are frequency multiplexed, which translates to more IBE leakage to interlaces of neighbouring transmissions.

5.3 Opportunistic usage of licensed available resources for subnetworks

A HC device should support multiple simultaneous active systems, allowing it to function both as a device within a wide area network (WAN) and as an AP that can create and manage subnetworks. To support future communication features with enhanced spectrum utilization, these UEs will include advanced front-end modules designed to handle multiple frequency bands simultaneously. This enables both inter-band and intra-band carrier aggregation (CA) and dual connectivity (DC) to boost coverage and capacity using new spectrum bands. These devices will also combine CA/DC and MIMO, using shared antennas for different frequency bands, facilitated by dual resonance and multiplexers for efficient signal separation.

However, the introduction of increased spectrum and band combinations in 6G increases the potential for self-interference and sensitivity degradation, posing significant challenges. Network operators must meticulously plan spectrum usage to mitigate these issues, which can increase roll-out costs and constrain spectrum utilization. Moreover, UE vendors face the pressure to develop robust designs capable of managing these complexities, resulting in more intricate hardware and stringent conformance testing requirements.

Self-interference and sensitivity degradation are critical challenges, especially when employing CA/DC. These issues arise from the simultaneous operation of multiple bands in different link directions, causing interference between different frequency bands within the same device, as illustrated in Figure 35.

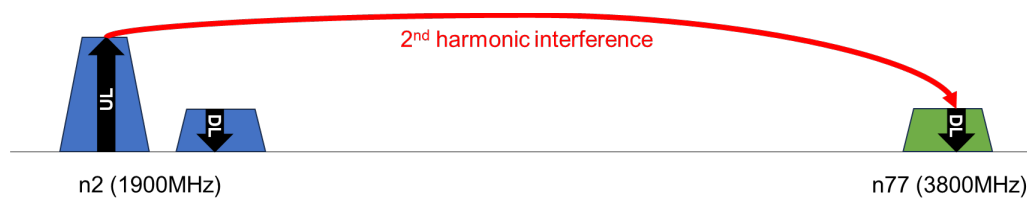


Figure 35: Example of self-interference which can lead to sensitivity degradation

This kind of interference can manifest as harmonics, harmonic mixing, intermodulation distortion (IMD), or cross-band isolation interference, all of which degrade the receiver's sensitivity. The impact of that is substantial, affecting WAN communication by causing dropped calls, reduced data rates, and increased latency. As an example from current 5G NR specifications TS 38.101-1 clause 7.3 [39], over 66.8% of the band combination cases for 3 CA can experience intermodulation issues, making it difficult to fully utilize the available spectrum. It is important to note that 3GPP permits a degree of degradation when the UE transmission emissions overlap with its own receiver band during CA/DC. This is achieved by relaxing the reference sensitivity requirements by up to a maximum sensitivity degradation (MSD) limit. However, that negatively affects the uplink link budget and throughput, making it a less desirable solution. Alternatively, using more advanced RF front-end could mitigate these issues, though this would significantly increase the overall cost and size of the UE implementation [44].

The specification of 6G subnetworks can offer a solution to overcome these issues and improve spectrum utilization. The solution consists of enabling subnetworks to opportunistically use spectrum resources that would otherwise be poorly utilized for WAN communication due to sensitivity degradation issues. That is further motivated by the short-range communication nature of subnetworks, meaning that they can operate at low power levels, therefore having reduced interference impact to the WAN. Figure 36 illustrates a procedure which can be executed between the HC device acting as AP for a subnetwork and the NW node.

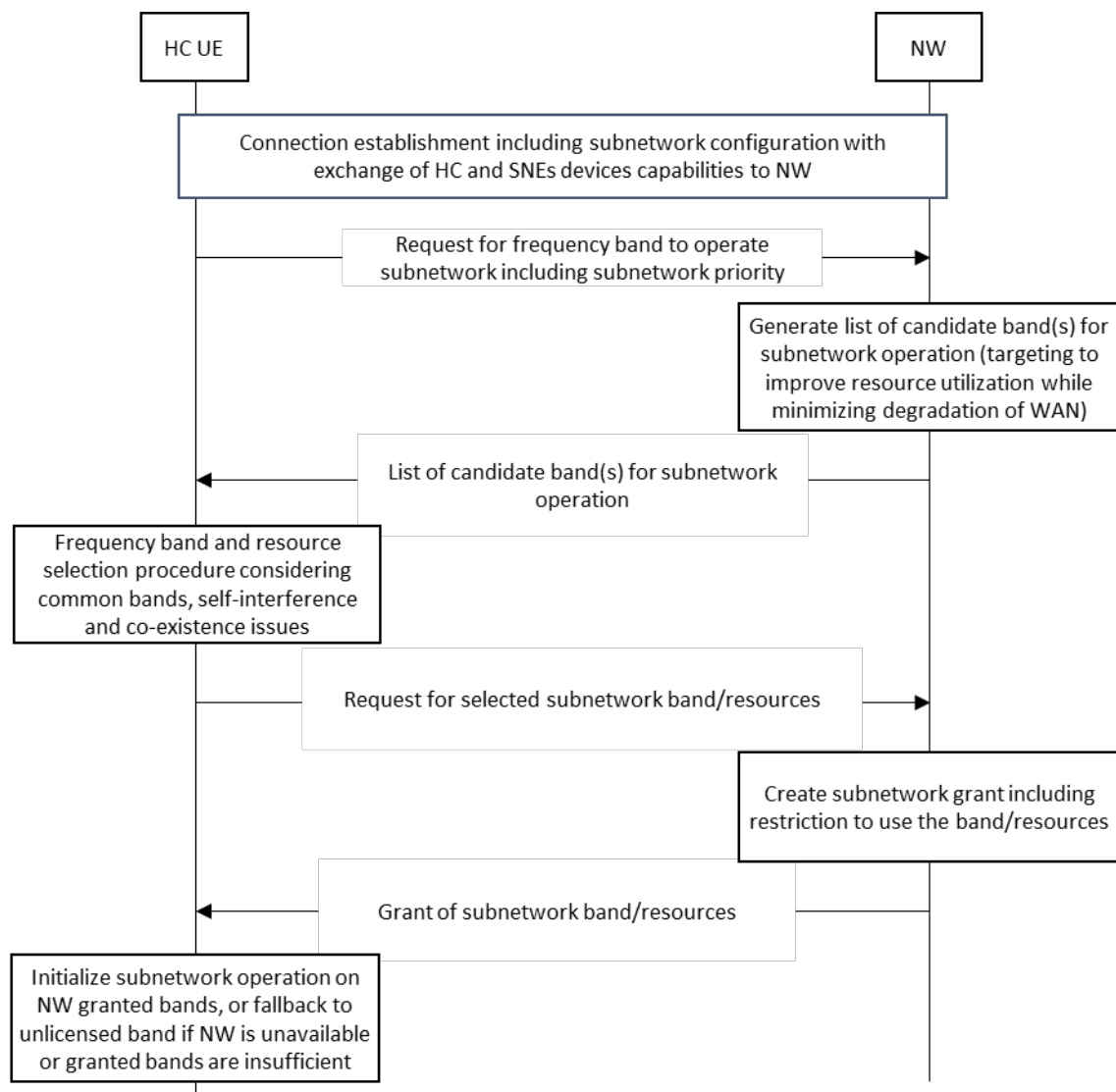


Figure 36: Example of signalling where the HC device acting as access point of a subnetwork obtains radio resource from NW.

The proposed procedure involves the UE requesting a frequency band to operate a subnetwork, and the NW responds with a list of candidate bands, often those underutilized in the WAN. The UE then analyses these options, prioritizing those which causes the least self-interference and least co-existence issues to determine the most suitable band for its subnetwork. The NW ultimately grants a band (or resources of a band) from the UE's preferred subset, ensuring quality of service (QoS) and preventing interference with neighbouring subnetworks. Constraints to the use of the resources may also be imposed, e.g., maximum transmit power allowed is -10 dBm. Notably, unlicensed bands can also be selected, offering additional bandwidth or a fallback option when coverage is limited, i.e., the UE should only use the licensed spectrum resources while it is allowed by a grant or semi-static configuration.

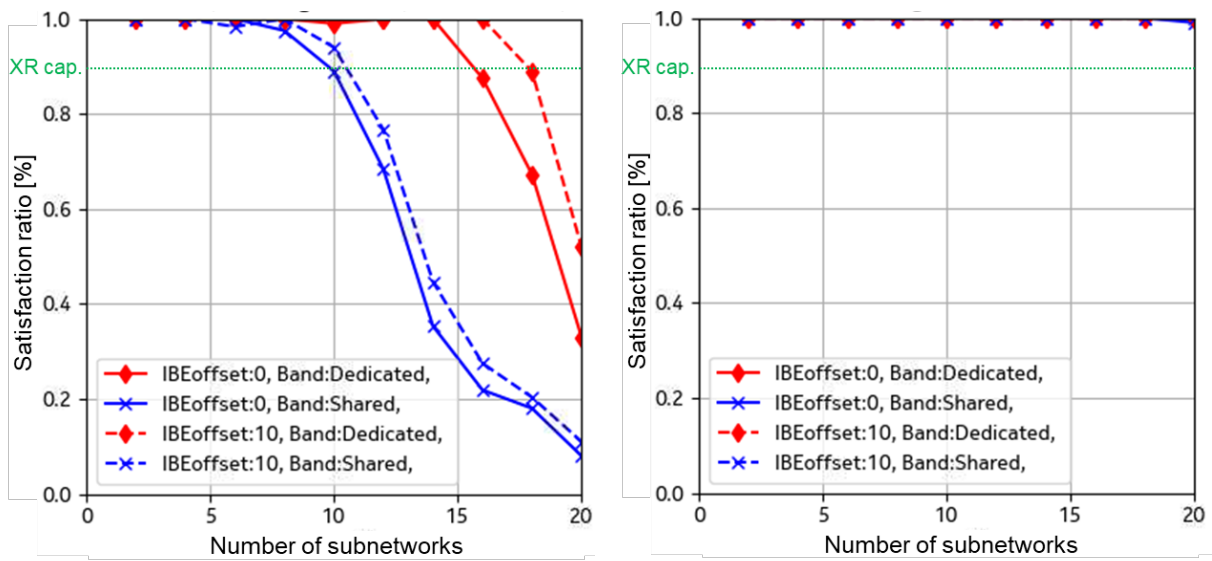
An advantage of the proposed procedure, in addition to the improved utilization of available spectrum and reducing self-interference, is in the use of licensed spectrum where devices are not required to

perform channel access procedures involving LBT. As discussed in previous part, LBT can be non-deterministic and introduce delays, hindering efficient communication. In licensed spectrum, there are less strict emission requirements since there is no incumbent technology operating outside the control of the parent network. This allows for more flexible resource allocation schemes, rather than the interlace-based approach required in unlicensed spectrum to meet OCB and PSD requirements.

Below, we analyse the performance of subnetworks following the same consumer use case assumptions as stated in previous parts. So here again, we follow the same slot structure and traffic assumptions as those listed in Table 4. However here we consider that the UEs use licensed spectrum, therefore no LBT procedure is needed. Also, in the license band, the sub-channel RBs can use the contiguous allocation.

Analysis of opportunistic use of licensed spectrum

Figure 37 shows the XR satisfaction ratio results versus the number of deployed subnetworks using a licensed band.



(a) SNE-HC performance (with multi-stream traffic)

(b) HC-HC performance (with pose/control traffic)

Figure 37: UE satisfaction ratio for different number of subnetworks in licensed band.

The results indicate a clear improvement in XR capacity when subnetworks use the available licensed band compared to the results from the semi-static channel access analysis. For SNE-HC communication, the XR capacity increases to between 15 and 17 subnetworks without and with IBE reduction, respectively, in separate bands, and to between 9 and 10 subnetworks without and with IBE reduction, respectively, in a shared band. In the latter case, it means almost 67% improvement relative to the unlicensed band performance with semi-static channel access shown before. For HC-HC communication, the XR capacity supports more than 20 subnetworks in both shared and separate bands.

This improvement is primarily due to the elimination of the need for LBT, which removes associated delays. Additionally, there is no requirement for the extra guard-period symbol, thus avoiding the capacity reduction it would cause. Furthermore, the use of contiguous RB allocation in the licensed band reduces the impact of IBE compared to the interlaced allocation in the unlicensed band. These results highlight the benefits of using licensed available spectrum for subnetworks whenever possible, for example, when they operate in coverage of a network which has underutilized spectrum.

5.4 Summary

In this chapter, we have studied various aspects to enhance RRM for subnetworks in licensed and unlicensed bands. The analysis covered intra- and inter-subnetwork communication based on sidelink in shared bands and in dedicated bands, considerations for semi-static channel access in unlicensed bands, and opportunistic usage of licensed available resources for subnetworks. The evaluation was conducted for a consumer use-case scenario, with performance assessed in terms of XR capacity.

The evaluation highlighted the significant impact of IBE on subnetwork communication. IBE, caused by transceiver impairments, leads to substantial interference, particularly in scenarios where adjacent resources are used by different transmitters. This interference is more pronounced in unlicensed bands with interlaced resource allocation, resulting in performance degradation under high load conditions. The analysis showed that IBE mitigation by an enhanced UE front-end with lower IBE general term or by applying an IBE-aware resource coordination can reduce these issues, improving the capacity.

In addition, the consideration of dynamic channel access in unlicensed bands revealed it as a main bottleneck for performance. Dynamic channel access suffers from resource contention and extended channel occupancy with high load traffic, significantly impacting performance of low delay budget communication. The evaluation indicated that dynamic channel access leads to rapid degradation in satisfaction ratios with increasing numbers of subnetworks, particularly in shared bands, calling for enhanced channel access schemes.

Further, the consideration of semi-static channel access for subnetworks in unlicensed bands demonstrated the benefits of predefined channel access timing structures. Semi-static channel access, unlike dynamic channel access, offers more deterministic communication and reduced latency in controlled environments. The evaluation indicated that semi-static channel access using a modified sidelink slot structure which satisfy regulatory requirements can significantly enhance the capacity, especially for HC-HC communication in shared bands, by minimizing LBT blocking and transmission collisions.

Lastly, the opportunistic usage of licensed available resources for subnetworks presented a promising solution to overcome sensitivity degradation issues and improve spectrum utilization. The analysis showed that the use of licensed spectrum resources for subnetworks, which may be underutilized resources for wide area communication, reduces delays and increases capacity, since it is not limited by LBT procedures. The contiguous RB allocation in licensed bands further reduces the impact of IBE

compared to interlaced allocation in unlicensed bands. Therefore, it is recommended that the opportunistic usage of licensed available resources should be prioritized for devices within coverage, while unlicensed bands should serve as a fallback solution for subnetworks unable to obtain licensed resources from a provider or operating out of coverage.

6 DETECTION AND MITIGATION MECHANISM OF EXTERNAL INTERFERENCE

Reliable communication in dense in-X subnetworks is increasingly threatened by external interference, which remains largely unaddressed in existing RRM research. While inter-subnetwork interference is typically managed through centralized or distributed coordination among subnetworks, external interference introduces a unique set of challenges due to its unpredictable, uncontrollable, and often malicious nature.

External interference in dense in-X subnetworks refers to unwanted signals or disruptions originating from sources outside the intended system, which degrade communication performance. These sources fall into three main categories:

- Natural interference, caused by environmental phenomena such as atmospheric noise, lightning, or solar flares. These can introduce random, often severe, disruptions.
- Unintentional interference, resulting from devices not part of the in-X system but operating in the same frequency band—e.g., industrial machinery such as motors, conveyor belts, and welding equipment can emit electromagnetic noise.
- Deliberate interference, commonly referred to as jamming, is introduced with the intention to disrupt network communications. This includes:
 - Constant jammers that emit continuous signals.
 - Reactive jammers that transmit only when activity is detected.
 - Pulsed jammers that intermittently emit disruptive bursts.
 - Intelligent jammers that adapt their strategies over time to avoid detection.

In dense and mission-critical deployments like in-factory environments or immersive consumer applications, in-X subnetworks are especially vulnerable to these forms of interference. Disruption in such environments can lead to system-wide communication breakdowns and operational failures. For this reason, effective detection and mitigation mechanisms for external interference are essential.

Several countermeasures can be employed to defend against deliberate jamming. Frequency hopping, for example, increases unpredictability, making it harder for jammers to disrupt transmissions. Additionally, Direction-of-Arrival (DoA) estimation using antenna arrays can identify the source of interference, enabling beamforming to nullify its effect.

Beyond deliberate jamming, cross-technology interference also poses a substantial challenge in dense, heterogeneous environments—especially in unlicensed or shared spectrum bands. Here, multiple wireless technologies (e.g., Wi-Fi, ZigBee, NR-U) coexist with incompatible protocols and potentially overlapping channels, leading to frequent collisions and packet loss.

To address these challenges, this chapter presents a comprehensive framework that spans from detection and modelling to resource management and receiver-level mitigation techniques. It includes:

- A stochastic external interference model capturing time-frequency dynamics and integrating realistic jammer behaviours.

- A centralized RRM algorithm (GDRA) using gradient descent for joint sub-band allocation and power control under interference-aware constraints.
- Modifications to two state-of-the-art benchmark algorithms, SISA and SIPA, to explicitly account for external interference power in ISR metrics.
- A performance evaluation of outage probability and SE across scenarios with and without interference awareness, demonstrating the superiority of the proposed methods.
- A dedicated section on robust receiver design under impulsive interference using LLR approximation techniques. The framework supports online, short-packet parameter estimation, enabling efficient demodulation under dynamic interference conditions.
- A practical protocol for managing LLR approximation in subnetwork nodes with constrained computing resources, enabling real-time selection and reporting of best-fitting functions for LLR estimation.

By combining network-level mitigation through intelligent RRM with device-level adaptability via efficient receiver approximation, the solutions presented offer a holistic and scalable response to the growing threat of external interference in in-X subnetworks.

6.1 Robust Radio Resource Management for In-Factory Subnetworks under External Interference

Unlike inter-subnetwork interference, external interference is inherently unpredictable and not directly manageable by network operators. In this section again we assume a centralized control framework, where a 6G base station coordinates and manages all subnetworks. This centralized approach effectively mitigates inter-subnetwork interference through strategic resource allocation and coordinated management. However, external interference remains a significant challenge due to its unpredictable nature and lack of operator control. To enhance the reliability and robustness of in-X subnetworks, we propose an approach that explicitly integrates external interference considerations into the resource allocation process.

Specifically, we address the joint sub-band allocation and power control problem within densely deployed InF-S in environments affected by external interference. By integrating external interference awareness into RRM decisions, our method ensures robust performance, significantly improving operational stability and QoS compliance in realistic industrial scenarios. The system model generally follows the description provided in Section 3.1.1.1, with the key difference that, in this scenario, small-scale fading varies across different sub-bands. This distinction arises due to the frequency-selective nature of the industrial wireless environment, where multipath propagation conditions differ significantly across frequency sub-bands. As a result, the wireless channels experience distinct fading patterns depending on their operating sub-band. Consequently, the channel gain matrix has dimensions $K \times N \times N$, explicitly capturing these sub-band-dependent variations. This refinement ensures a more accurate and realistic representation of channel characteristics, ultimately enhancing the effectiveness and reliability of the resource allocation strategies employed.

6.1.1 External Interference Model and Problem Formulation

External interference in the considered environment originates from multiple mobile interference sources, represented as the set $\mathcal{I} = \{I_1, I_2, \dots, I_J\}$. These sources operate within the same frequency bands as the InF-Ss and follow predefined trajectories across the factory floor.

Each sub-band is independently affected by external interference, which follows a Poisson traffic model. Packet arrivals in each sub-band adhere to a Poisson process characterized by an average arrival rate λ . Interference activation on a sub-band occurs whenever at least one packet arrives within a given time interval. Regardless of the number of arriving packets within that interval, interference power remains constant once activated. The external interference power level relative to the maximum transmission power is formally defined as:

$$I_{\text{ext(dB)}} = P_{\text{max(dB)}} + \alpha,$$

where α represents the interference power ratio in dB, indicating the strength of external interference sources relative to the maximum transmit power used by the subnetworks. Given the short-range nature of industrial wireless operations, HCs and SNEs typically utilize similar power levels. However, the actual received interference power varies depending on the distance between the HCs and interference sources, as well as the channel conditions at any given moment.

To accurately detect and measure external interference, we assume the implementation of a periodic detection mechanism using sounding reference signals and scheduled silent slots. Subnetworks periodically transmit reference signals, enabling neighbouring subnetworks to estimate CSI and measure internal interference. Complementing this, scheduled silent slots ensure HCs can precisely measure external interference without contamination from network transmissions, as subnetworks remain inactive during these intervals. Measurements from these silent periods are reported back to the CRM via dedicated backhaul links. Given limited temporal fluctuations in interference and channel conditions, these measurements remain valid and reliable for subsequent resource allocation cycles.

The main objective is to develop a resource allocation strategy that jointly optimizes sub-band selection and transmit power to minimize the outage probability while satisfying a predefined target SE, denoted as $\text{SE}_{\text{target}}$. This target ensures a baseline QoS, essential for mission-critical applications demanding reliable communication and high data throughput.

The problem formulation closely follows the approach described in Section 3.1.2.1, with the primary difference lying in the calculation of SE for subnetwork n on sub-band k . Specifically, the SE under consideration here includes the effects of external interference and is calculated as:

$$\text{SE}_n^k = \log_2 \left(1 + \frac{h_{n,n}^k a_n^k P_n}{\gamma_{m,n}^2 + \sum_{m \in \mathcal{N} \setminus \{n\}} h_{m,n}^k a_m^k P_m + I_n^k} \right),$$

where the additional term $\sum_{m \in \mathcal{N} \setminus \{n\}} h_{m,n}^k a_m^k P_m$ represents the inter-subnetwork interference caused by simultaneous transmissions from other subnetworks, and I_n^k denotes the external interference power, originating from external radio technologies operating in the same frequency band. These interference components differentiate this scenario from the one detailed in Section 3.1.2.1, necessitating tailored resource allocation strategies to effectively mitigate their impacts.

6.1.2 Gradient Descent-based Resource Allocation Algorithm

The proposed Gradient Descent-based Resource Allocation (GDRA) algorithm is described comprehensively in this section, outlining a detailed, two-stage approach for joint sub-band allocation and power control optimization.

In the first stage, the algorithm relaxes the original problem constraints by allowing continuous power distribution across multiple sub-bands. Specifically, it assigns power levels P_n^k to each sub-band k for subnetwork n , constrained by $\sum_{k=1}^K P_n^k \leq P_{\max}$, $\forall n \in \mathcal{N}$. This initial relaxation provides the flexibility necessary for gradient-based optimization methods, enabling the identification of optimal power allocations. During this process, a softmax function with a low-temperature parameter τ is employed to produce nearly binary (one-hot) power distributions across the available sub-bands, while still remaining differentiable for gradient updates. After this preliminary step, each subnetwork selects the sub-band achieving the highest SE, thus enforcing the single sub-band usage constraint.

Notably, this first stage alone can also function as an effective standalone sub-band allocation technique, as the Gradient Descent-based Sub-band Allocation with maximum transmit power (GDSA-maxPower). This alternative scenario, which involves assigning maximum allowed transmission power to the chosen sub-band, is independently evaluated in Section 6.1.4.

In the second stage, the algorithm optimizes the transmit power specifically for the sub-bands selected in stage one. By fixing sub-band allocation, the algorithm concentrates solely on power control adjustments. Power levels are continuously tuned within the range $0 \leq P_n \leq P_{\max}$ using gradient descent, guided by a sigmoid-based differentiable representation. This approach enables end-to-end optimization, facilitating efficient gradient descent.

The detailed GDRA algorithm for joint sub-band allocation and power control is summarized below:

GDRA Algorithm for Joint Sub-band Allocation and Power Control:

Inputs: Channel gain matrix H

Initialization: Initialize power $p^{(0)}$ and selection variable $\theta^{(0)}$ as zero matrices

Stage 1: Sub-band Selection

1. Compute the power distribution across sub-bands using:

$$P = \text{softmax}\left(\frac{\theta^{(l-1)}}{\tau}\right) \cdot P_{\max}$$

2. Calculate SE using the equation:

$$SE_n^k = \log_2 \left(1 + \frac{h_{n,n}^k a_n^k P_n}{\gamma_{m,n}^2 + \sum_{m \in \mathcal{N} \setminus \{n\}} h_{m,n}^k a_m^k P_m + I_n^k} \right)$$

3. Optimize sub-band allocation using gradient descent by maximizing the minimum SE across subnetworks.
4. Identify optimal sub-band $k^*(n)$ for each subnetwork, based on maximum SE.

Stage 2: Power Control Optimization

1. With the selected sub-bands fixed, apply power control optimization:

for fixed sub-band selection $P_n = P_{\max} \cdot \sigma(\rho_n)$ for fixed sub-band selection

2. Recalculate SE using the previously defined SE formula.
3. Adjust power values iteratively through gradient descent, ensuring constraints are respected.

Output: Final optimized sub-band allocations and transmit power levels for all subnetworks.

6.1.3 Modified SISA-SIPA Algorithm

In this section, we describe the necessary modifications to two SoA resource management algorithms: SISA, originally proposed in [25], and SIPA, as introduced in [45]. These algorithms aim to minimize the sum Interference-to-Signal Ratio (ISR) across the entire network. Originally, they assume scenarios free from external interference. However, in practical deployments, subnetworks are subjected to additional external interference sources, which must be explicitly considered.

6.1.3.1 External-Interference-Aware SISA

We first discuss modifications to the SISA algorithm, which iteratively allocates sub-bands to subnetworks. The original algorithm begins with an arbitrary allocation, progressively refining sub-band assignments by sequentially analysing each subnetwork. At iteration d , subnetwork n selects a sub-band k based on the minimization of mutual ISR, defined as:

$$k^* = \arg \min_{k \in \mathcal{K}} \sum_{m \in \mathcal{A}_k^d} (W_{nm}^k + W_{mn}^k),$$

where $\mathcal{A}_k^d$ denotes the set of subnetworks allocated to sub-band k after $d - 1$ iterations, and $(W_{nm}^k = \frac{\omega_{nm}^k}{\omega_n^k})$ represents the ISR from subnetwork m to subnetwork n . Here, ω_{nm}^k and ω_n^k denote the channel gain powers of interfering and desired channels, respectively.

To incorporate external interference into the SISA algorithm, we introduce an additional term representing external interference ISR. Consequently, the modified selection criterion becomes:

$$k^* = \arg \min_{k \in \mathcal{K}} \left(\sum_{m \in \mathcal{A}_k^d} (W_{nm}^k + W_{mn}^k) + W_{n,\text{ext}}^k \right),$$

where $W_{n,\text{ext}}^k = \frac{I_n^k}{\omega_n^k}$ captures the ISR caused by external interference sources. By adding this term, the modified algorithm proactively selects sub-bands with lower external interference, thus enhancing robustness and reliability.

6.1.3.2 External-Interference-Aware SIPA

Similarly, the SIPA algorithm is adapted to explicitly account for external interference. SIPA iteratively selects transmission (Tx) power levels from a discrete set of available powers, denoted as $\mathcal{P}$. Initially, random power levels are assigned, and the algorithm sequentially optimizes each subnetwork's Tx power to minimize mutual ISR while keeping other subnetworks' power levels fixed. With external interference incorporated, the modified SIPA selection criterion at iteration d becomes:

$$(\phi_n^{(d)}(P) = \sum_{m \in \mathcal{A}_k} \left(\frac{[\mathbf{P}^{(d-1)}]_m}{P} W_{nm}^k + \frac{P}{[\mathbf{P}^{(d-1)}]_m} W_{mn}^k \right) + \frac{1}{P} W_{n,\text{ext}}^k),$$

where $\mathbf{P}^{(d-1)}$ denotes the vector of Tx powers selected by subnetworks after $d - 1$ iterations, and $W_{n,\text{ext}}^k = \frac{I_n^k}{\omega_n^k}$ quantifies the ISR from external interference. This modification explicitly encourages higher Tx power allocations under scenarios experiencing stronger external interference, ensuring improved communication reliability.

The modified SIPA algorithm operates iteratively, updating each subnetwork's Tx power L times, identical to SISA. The comprehensive algorithm is outlined below:

Algorithm: External-Interference-Aware SIPA for Sub-band k

Input:

- Set $\mathcal{A}_k$ of subnetworks on sub-band k .
- Mutual ISR values ($W_{nm}^k, \forall n, m \in \mathcal{A}_k$)
- Discrete set of transmission powers, $\mathcal{P}$

Initialization:

Initialize Tx power levels $P^{(0)}$ randomly from $\mathcal{P}$.

Procedure:

For each iteration $l = 1$ to L :

For each subnetwork $n = 1$ to N :

Set iteration number: $d = N(l - 1) + n$

For each power level $P \in \mathcal{P}$:

Compute $\phi_n^{(d)}(P)$ using the modified criterion.

Update Tx power level of subnetwork n to minimize $\phi_n^{(d)}(P)$.

Update the Tx power level for subnetwork n : $[\mathbf{P}^{(d)}]_n = \arg \min_{P \in \mathcal{P}} \phi_n^{(d)}(P)$

Output: Optimized power allocation $\mathbf{P}^{(d)}$ after completing all iterations.

6.1.4 Simulation results and analysis for RRM in the presence of external interference

In this section, we evaluate the performance of the proposed GDRA algorithm for joint sub-band allocation and power control. The GDRA algorithm is benchmarked against SoA algorithms, specifically the modified SISA and SIPA algorithms. Additionally, to provide a broader assessment, we include comparisons with two additional strategies: SISA combined with maximum transmit power and the GDSA approach paired with maximum transmit power.

The external interference model utilized is detailed in Section 6.1.1. In this evaluation, a single mobile external interference source is considered, generating independent interference across different sub-bands. Unless explicitly stated otherwise, the simulations employ parameters summarized in Table 5. Specific parameter adjustments are highlighted where necessary to assess their impact on algorithm performance. All simulations were conducted using a custom-built simulator implemented in Python, specifically developed to model in-X subnetwork behaviour, interference dynamics, and RRM algorithm execution in a controlled and flexible environment.

The key performance metric evaluated is the outage probability, defined as the probability that the achieved SE falls below the target SE, SE_{target} .

Table 5: Simulation Parameters for RRM under external interference

Parameter	Value
Factory area	20 m×20 m
Number of subnetworks	20
Number of sub-bands	4
Subnetwork radius	0.5 m
Number of devices per subnetwork	1
Minimum distance between HCs	1 m
SNE-to-HC minimum distance	0.3 m
Shadowing standard deviation	4 dB
DL clutter density, clutter size	0.6, 2
De-correlation distance	5 m
Maximum transmit power	0 dBm
Interference power ratio	-20 dB
Number of sub-bands with active interference	2
Average arrival rate for external interference	0.3
Sounding reference signal period	100 ms

Maximum velocity of InF-S and external interferer	10 m/s
Sub-band bandwidth	30 MHz
Center frequency	10 GHz
Noise figure	5 dB
Temperature parameter for softmax	0.001
Batch size	20000
Number of epochs	1000

Figure 38 illustrates the outage probability of individual links across all subnetworks as a function of the target SE, SE_{target} . We compare the performance across three distinct scenarios: *No External Interference*, *External Interference-Unaware*, and *External Interference-Aware*.

In the *No External Interference* scenario, the only interference considered is from other subnetworks within the system. The *External Interference-Unaware* scenario introduces external interference sources that the algorithms are not configured to address or mitigate explicitly. Conversely, in the *External Interference-Aware* scenario, algorithms explicitly adapt their resource allocation decisions to account for external interference.

In the *No External Interference* case, all algorithms demonstrate low outage probabilities, effectively handling inter-subnetwork interference. However, under the *External Interference-Unaware* scenario, outage probabilities rise significantly, particularly at higher SE_{target} , highlighting the consequences of neglecting external interference. The *External Interference-Aware* scenario shows notable improvements across all algorithms. Nevertheless, GDRA consistently provides superior performance. For instance, at $SE_{\text{target}} = 5$, GDRA reduces outage probability to 0.008, representing around a 90% reduction compared to 0.077 for SISA-SIPA.

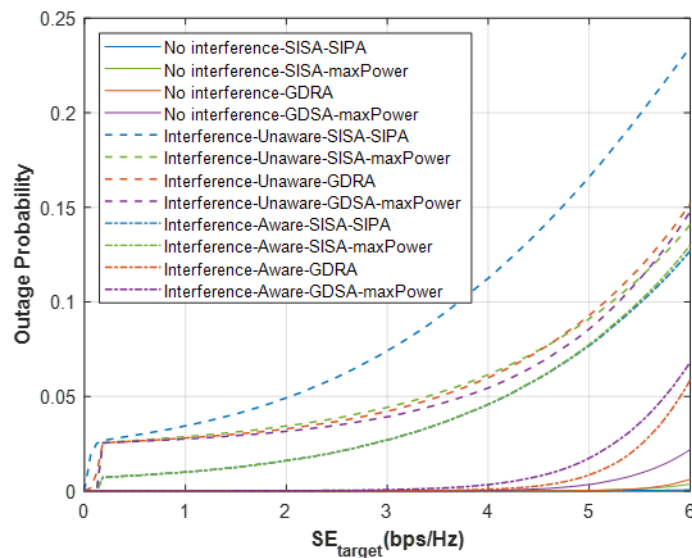


Figure 38: Outage probability of individual links across all subnetworks, under three scenarios: No Interference, Interference-Unaware, and Interference-Aware.

Figure 39 presents the CDF of the average SE across subnetworks. This figure indicates that while optimized power control significantly reduces outage probabilities, methods employing maximum transmit power achieve higher SE at upper percentiles. Overall, gradient descent-based algorithms such as GDRA consistently outperform benchmarks, particularly in interference-aware scenarios, emphasizing their suitability for mission-critical communication systems and other scenarios demanding high spectral efficiency.

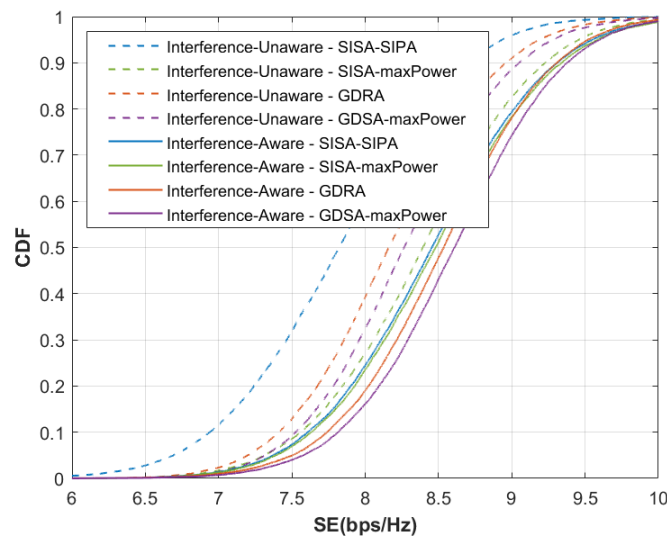


Figure 39: CDF of the average SE across all subnetworks, under two scenarios: Interference-Unaware and Interference-Aware.

Figure 40 further evaluates the algorithms under varying external interference power levels. GDRA consistently yields the lowest outage probability, notably ensuring 99% reliability at $SE_{\text{target}} = 3$ bps/Hz, substantially outperforming benchmark methods that exhibit around 25% outage probabilities. At lower external interference power levels, GDRA outperforms GDSA-maxPower due to its optimized power

allocation. However, as external interference approaches the maximum transmit power $P_{\max}$, GDRA and GDSA-maxPower performances converge, indicating reduced effectiveness of power optimization under extreme interference conditions.

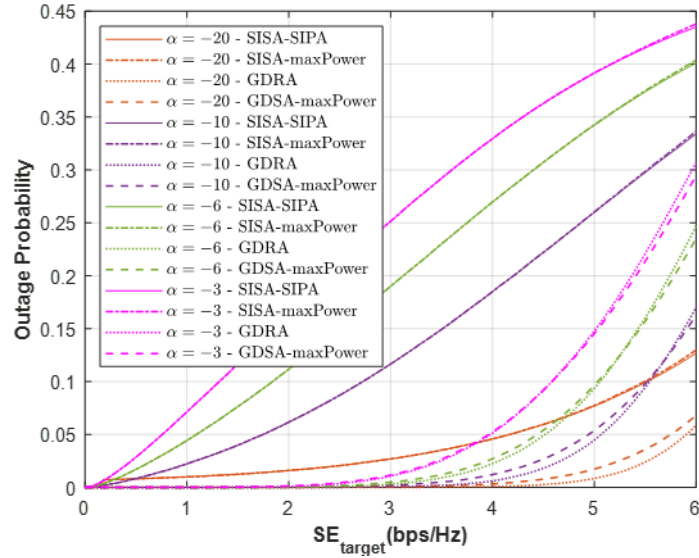


Figure 40: Outage probability of individual links across all subnetworks, under varying levels of external interference power.

Figure 41 analyses the performance impact when interference affects multiple sub-bands. With an increasing number of interfered sub-bands K_{intf} , the outage probability rises across all methods. Yet, GDRA maintains a consistently lower outage probability than SISA-SIPA, demonstrating superior adaptive interference management. The performance advantage of GDRA becomes even more pronounced as interference conditions worsen, underscoring its robustness.

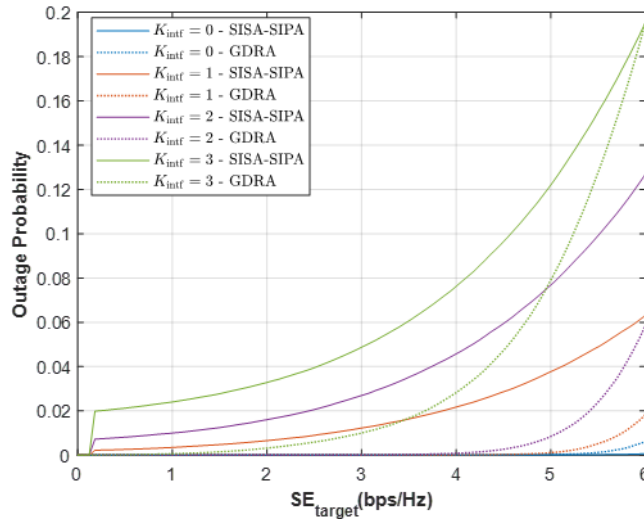


Figure 41: Outage probability of individual links across all subnetworks, under scenarios with no interference and with interference activated on 1, 2, and 3 sub-bands.

Figure 42 shows the CDF of transmit powers utilized across subnetworks, highlighting that algorithms tend to employ higher power levels to maintain target SE under increased external interference. As interference intensifies, all methods shift towards higher transmit powers.

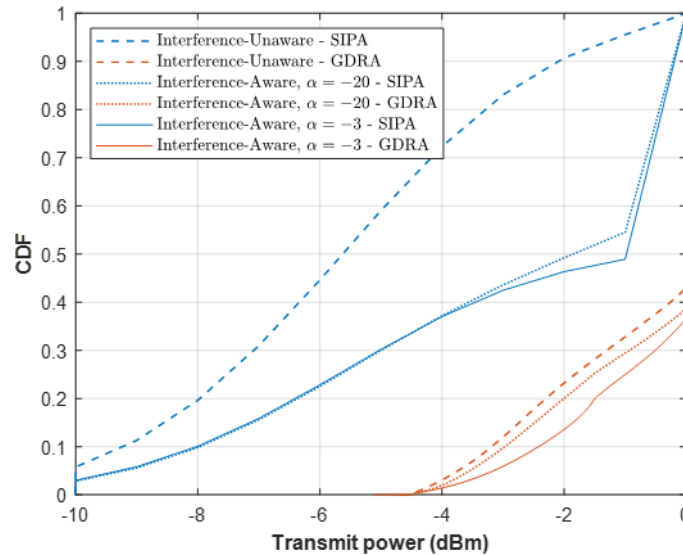


Figure 42 CDF of transmit powers across all subnetworks, under varying levels of external interference power.

While the asymptotic complexity for SISA, SIPA, and GDRA algorithms simplifies to $O(KN^2)$ with fixed parameters [25],[45],[46], GDRA distinguishes itself through efficient parallelization and GPU-based computation. Table 6 Execution Time per Sample for Different Batch Sizes and Algorithms summarizes the execution time per sample, demonstrating GDRA's superior computational efficiency, particularly at increased batch sizes, due to effective batch processing capabilities.

Table 6: Execution Time per Sample for Different Batch Sizes and Algorithms

Algorithm	Batch Size	Execution Time (s)
GDRA	200	0.1925
GDRA	2000	0.0509
GDRA	20000	0.0213
SISA-SIPA	-	0.0243

6.2 Performance evaluation framework for efficient receiver adaptation over subnetworks.

In this section, we complement the radio resource management approaches tailored to external interference, with receiver-level strategies, targeting non-Gaussian impulsive noise.

6.2.1 Receiver Design approaches

To achieve reliable and efficient communications, it is imperative to consider the impulsive nature of interference when designing receivers. Interference modelling frequently exhibits impulsive characteristics, which can be represented using various statistical methods and probability distributions. However, designing a dedicated receiver for each specific scenario is impractical due to the significant temporal and spatial variability of interference characteristics. Consequently, there is a strong need for a receiver architecture capable of adapting to a broad spectrum of interference models, encompassing both impulsive and non-impulsive behaviours, and accommodating varying degrees of impulsiveness.

In the realm of receiver design, we discussed in D4.1 several approaches, each with its own set of advantages and disadvantages. Among these, we focused mainly on direct LLR approximation which stands out due to its simplicity and ability to facilitate online learning. For further information we refer the reader to D4.1.

Recalling that achieving exact LLR values is computationally prohibitive, to address this, multiple approximations can be utilized. These approximations can belong to different families, such as piecewise functions, rational functions, and more. Selecting the optimal LLR approximation typically involves extensive Bit Error Rate (BER) or Frame Error Rate (FER) simulations to identify the best-performing function. While effective, this process is time-consuming and computationally intensive.

To overcome these challenges, we introduce in the following? a novel framework that derives a new metric for selecting the best LLR approximation. This metric is designed to adapt well to varying channel conditions while maintaining low complexity. By leveraging this framework, we can achieve efficient online learning without the need for exhaustive simulations, thus streamlining the selection process and enhancing overall performance.

6.2.2 Approximation functions

Different functions may apply, continuous functions (e.g., Identity functions, Constant functions, polynomial functions, quadratic functions, cubic functions, etc.), or non-continuous functions (e.g., rational functions, modulus functions, Dirichlet functions, step functions, piecewise-defined functions). Combination of one or more of these sets of functions will form a pool of functions to select from the best that can describe the channel. By enriching this pool, the probability to reach the ideal function representation will increase but with the trade-off with additional searching complexities. It is worth noting that, by enabling piecewise functions this will give additional level of control in terms of complexity and accuracy where, the more parameters are defined or included to the piecewise function, the closer the function will be to the truth or the best representation of the samples. Furthermore, one can have a combination of simple functions (e.g., linear functions) that each can be mapped to a segment.

Functions that best approximate the channel interference are crucial to the LLR estimation process, as they can capture accurate representations of the channel interference, saving in complexity. As discussed, this is a two-step process where first, a function that best adapts to the channel interference level and type is selected and applied. Once done, the parameter estimation for that function takes place.

So, we consider parametric approximation L_θ of the LLR. The family of functions is L_θ chosen for its simplicity and flexibility to represent the LLR in different channel types. To narrow down the search, we consider the estimated LLR L_θ is an odd piece-wise function. We consider both demappers L_{ab} and L_{abc} [47],[48] as shown in Figure 43, that outperforms other LLR approximations as shown in [48], in terms of performance.

$$L_{ab}(y) = \text{sgn}(y) \min(a|y|, b/|y|),$$

$$L_{abc}(y) = \text{sgn}(y) \min(a|y|, b/|y|, c).$$

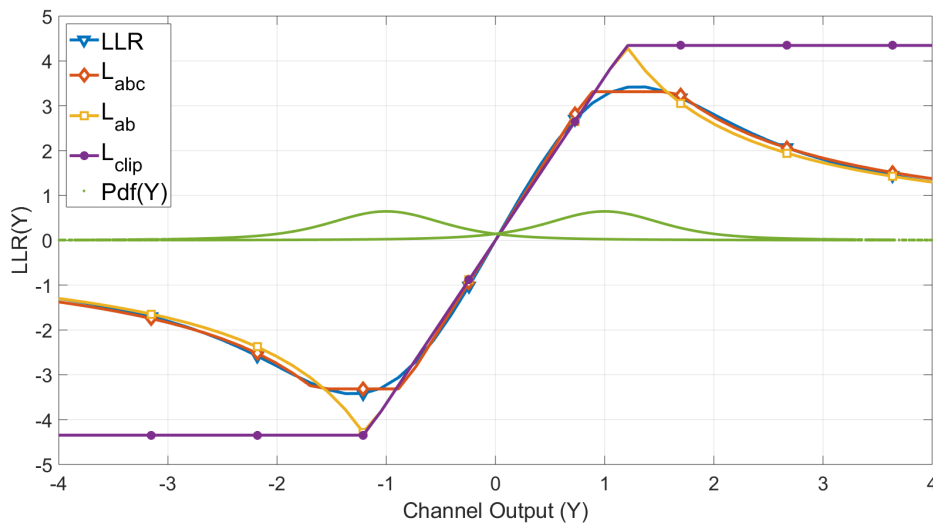


Figure 43: Comparison of the optimal LLR shape with different approximations

The LLR approximations depend on several parameters, which must be optimized to make the approximation as close as possible to the LLR. Several parameter estimation methods are considered in the literature. In [47],[49],[50], the authors proposed a framework to enable online real-time parameter estimation, but they consider long block length regime. For short packets as in in-X Subnetworks the proposed framework suffers from significant performance degradation due to the lack of availability of large number of samples. To solve this problem, authors in [48] proposed a solution that enables unsupervised learning in the short block length regime which is suitable for in-X Subnetworks.

6.2.3 Parameter estimation

The LLR approximations depends on two to three parameters, grouped here under the variable θ which must be optimized to make the approximation as close as possible to the LLR. In [39],[40], authors proposed a method for supervised learning of θ . The receiver is looking for θ that maximizes

$$C_{L_\theta} = 1 - E[\log_2(1 + e^{-X L_\theta(Y)})].$$

This search is to approach the capacity of the channel as closely as possible with the approximate likelihoods.

However, it should be noted that an actual implementation cannot be based directly on expectation due to the lack of the pdf. From a learning sequence $(x_1, \dots, x_n)$ and the corresponding output $(y_1, \dots, y_n)$, the receiver optimizes a version of the aforementioned Equation where the expectation is replaced by an empirical mean,

$$C_{L_\theta} = 1 - \sum_{i=1}^n [\log_2(1 + e^{-x_i L_\theta(y_i)})].$$

The LLR data $L(y)$ is equivalent to that of the probability *a posteriori* $p(x|y) = \Pr[X = x|Y = y]$ because $p(x|y) = \frac{1}{(1+e^{xL(y)})}$. Similarly, the LLR approximation $L_\theta(y_i)$ provides an approximation of the *a posteriori* $q(x|y) = \frac{1}{(1+e^{xL_\theta(y_i)})}$. It is then possible to show that the C_{L_θ} criterion is bounded and the bound is reached when $q(x|y) = p(x|y)$. More precisely, the difference between C_{L_θ} and capacity is the Kullback-Leibler distance between a posteriori and its approximation.

$$D(q(x|y)||p(x|y)) = \int \log_2 \frac{q(x|y)}{p(x|y)} p(x|y) dx p(y) dy,$$

Where $p(y)$ is the density of the channel output.

Receivers using LLR approximation should be compared to identify the best balance between simplicity and performance. However, performance should be evaluated based on the error rate, which can be computationally intensive. For example, to analyse the robustness of the previous approximation L_{ab} for a noisy α -stable channel with parameters ($\alpha = 1.4$) and ($\gamma = 0.4$), authors in [39] show, in the binary error rate for a regular LDPC code (3,6) of size 20000 using this approximation for different parameter values (a) and (b). These contours are superimposed on the zone of parameters that optimize the criterion (θ) to verify the adequacy between the receiver and this type of channel.

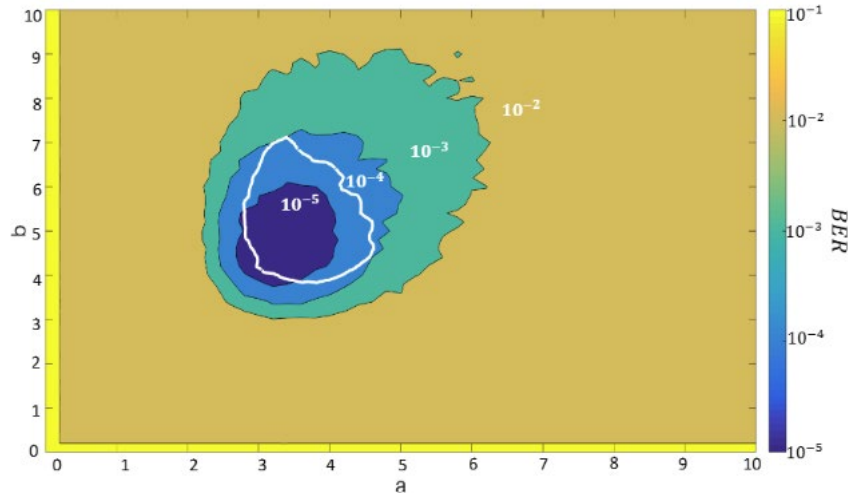


Figure 44: Optimal region and BER as a function of a and b

This approach of comparing LLR approximation stemmed by using the KL divergence is crucial for selecting the best LLR approximation, however, it can quickly deplete computational resources, especially for receivers that are not well-suited to the channel model. This inefficiency arises because extensive simulations and evaluations are required to identify the optimal LLR approximation, which is computationally burden.

To address this challenge, we propose a new criterion based on measuring the MSE between the true LLR and the approximated LLR to rank the approximated LLR, as shown in the following:

$$MSE = \int_{-\infty}^{\infty} [\Lambda(y) - L_{\theta}(y)]^2 p(y) dy.$$

Here, $\Lambda(y)$ represents the true LLR, and $L_{\theta}(y)$ represents the approximated LLR. The term $[\Lambda(y) - L_{\theta}(y)]^2 p(y)$ ensures that the approximation error is weighted by the likelihood of y , emphasizing accuracy in regions where y is more probable. This proposed metric is justified by the need for an $L_{\theta}(y)$ receiver to best approximate the likelihood of $\Lambda(y)$. This approximation must be more accurate for the most likely values of y , which is why different regions need to be weighted differently. This weighting is represented by the factor $p(y)$, which denotes the probability distribution of y .

To formalize this, consider the integral over x in the previous equation. We aim to find two constants, K and K' , that can frame this integral, ensuring that the approximation remains within acceptable bounds,

$$K[\Lambda(y) - L_{\theta}(y)]^2 \leq \int \log \frac{q(x|y)}{p(x|y)} p(x|y) dx \leq K'[\Lambda(y) - L_{\theta}(y)]^2,$$

which shows the equivalence between the Kullback-Leibler distance and the MSE criterion:

$$K \text{ MSE} \leq D(q(x|y) \parallel p(x|y)) \leq K' \text{ MSE},$$

Thus, comparing receivers according to the KL criterion and the MSE criterion is a first coherent approach and allows for quickly selecting the best LLR approximation. This criterion offers a more efficient method for selecting the best LLR approximation. By focusing on the MSE, we can directly assess the accuracy of the approximated LLR without the need for exhaustive BER simulations. This selection criterion is straightforward to implement and can be executed online by edge devices with limited computational capabilities. In the following section, we are going to introduce how such a criterion can be exploited in the subnetwork context.

6.2.4 In subnetwork receiver approximation

In this section, a solution for LLR approximation for nodes within a subnetwork is proposed. It consists of an example protocol and details on how a subnetwork element could performed the described approximations.

A first consideration is on the capabilities of the network element. Due to the compute resource constraint nature of an SNE and LC, when compared to an HC, there could be scenarios where SNEs and LCs would not be able to compute certain functions for approximations. This needs to be communicated to the parent 6G network, so that proper function management can be applied. This can be done with the first two steps in Figure 45, where LC is a presentation of either an LC or an SNE. Worth noting that, if a subnetwork node cannot perform this approximation, then it will not be able to perform BER estimation either.

A node in the subnetwork may have a set of preconfigured functions to apply to the channel and further be configured to only apply a subset only of the functions based on, e.g., current levels of computational delays. This is important to note, especially in cases where the latency requirements are strict, and because the computational delay of each function may be known, but it is also a function of the current CPU usage at the subnetwork receiver. The pool of functions will be designed based on modulation type used by the subnetwork node, channel conditions, interference type, etc. Thresholds for computational delay may be configured at the subnetwork node, under the form of a max time, a max CPU load, max power spent in computing the function, etc. This constitutes the configuration of pool of functions and estimation rules by the parent 6G network, in this case, a 6G-BS. This configuration is then delivered to the subnetwork management node.

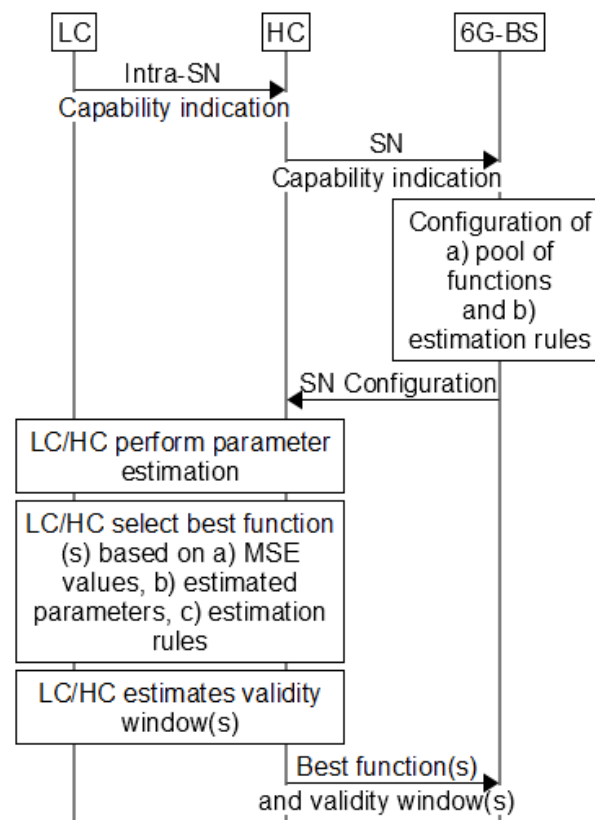


Figure 45: Example procedure for reporting the best function approximating LLR and its validity window

Once this is complete, subnetwork nodes can start performing parameter estimation for the configured function(s). Once parameter estimation is complete, the subnetwork node can now determine what the best function(s) are. This is done based on the MSE values, estimated parameters, and the configuration estimation rules received from the 6G-BS.

The subnetwork node can also estimate for how long the current functions are valid for, i.e., for how long the subnetwork node will continue to use them for LLR estimation, which may be estimated as a function of channel variations and interference level, in addition to other metrics.

With this procedure, the subnetwork node will perform MSE computation instead of BER, which could be more efficient in terms of computation, latency, power consumption, overhead, etc., by incurring in much lower costs for these metrics. The subnetwork node can select the best N functions with lowest MSE value (where N is a positive integer). This can be useful for the parent 6G network to be able to test its functions in the subnetwork domain, as they would be assessed in parallel for the same channel conditions.

6.3 Summary

This chapter addressed the often-overlooked problem of external interference in radio resource management for dense in-X subnetworks. A key contribution is the modelling and incorporation of stochastic external interference - originating from both benign and malicious sources - into a centralized RRM framework. A novel gradient descent-based RRM algorithm (GDRA) was proposed for joint sub-band and power allocation. The GDRA algorithm integrates interference-aware utility functions and is validated through simulations that show up to 90% reduction in outage probability compared to interference-unaware baselines.

The study further introduced enhanced versions of the benchmark SISA and SIPA algorithms by integrating an external interference ISR term into their respective optimization objectives. Simulation results confirmed the importance of interference awareness, especially in high-SE and multi-band interference scenarios.

To complement RRM, this chapter also explored receiver-level strategies for mitigating performance loss due to non-Gaussian impulsive noise. Various LLR approximation techniques were reviewed, with emphasis on piecewise parametric functions offering high accuracy and low complexity. A procedure was proposed for subnetwork nodes to dynamically select and validate the best LLR approximation functions based on the MSE metric, enabling adaptive receiver operation under limited computational budgets.

Together, these solutions form a robust and flexible framework for managing external interference in in-X subnetworks. They enable not only reliable RRM under adverse conditions but also efficient real-time receiver adaptation at the subnetwork edge. The results and proposed methodologies support the robustness objective of the 6G-SHINE project, especially in environments with significant external interference.

7 CONCLUSIONS

The final results of the 6G-SHINE project on radio resource management (RRM) for dense and dynamic in-X subnetworks are presented in this document. Through an integrated set of centralized, distributed, and goal-oriented solutions, substantial progress has been made toward meeting the project's ambitious targets on reliability, scalability, latency, spectral efficiency, and resilience to external interference.

A centralized RRM framework combining spatio-temporal attention-based LSTM prediction with resilient deep neural network (DNN)-based resource allocation was developed to address the impact of outdated CSI. With a 4-sample CSI delay, the proposed centralized method achieves a minimum spectral efficiency (SE) that is 53% higher than the SoA without any predictor and 94% higher than the SoA with a standard LSTM predictor. These gains were validated for dense deployments of 25,000 subnetworks per km², aligning with the project's objective of achieving approximately ten times the density of current 5G ultra-dense networks.

Complementing the centralized solution, a distributed RRM approach based on Graph Neural Networks (GNNs) was proposed to enable autonomous power control in scenarios where global coordination is limited. The GNN-based strategy improves spectral efficiency by approximately 7% under uniform conditions and up to 13.16% under heterogeneous channel conditions compared to equal power allocation, while relying on realistic over-the-air message passing mechanisms compatible with 3GPP protocols.

We also introduced a goal-oriented RRM solution for mission-critical industrial automation, where communication quality is jointly optimized with application-specific metrics, such as minimizing robot mission completion time. By employing a Proximal Policy Optimization (PPO) reinforcement learning method, the proposed mobility control algorithm achieves a 20% higher probability of maintaining the same block error rate (BLER) as the SoA under a 0.5 ms latency constraint, significantly enhancing URLLC performance in motion-intensive scenarios.

Recognizing the growing importance of shared-spectrum operations, the project developed enablers for supporting dense subnetwork deployments in unlicensed or hybrid licensed-unlicensed bands. Semi-static channel access techniques were shown to enable up to ten times higher XR capacity compared to dynamic access methods, while licensed-assisted operation led to a 67% increase in the number of supported subnetworks compared to semi-static access alone. Furthermore, mitigation of in-band emissions (IBE) through device front-end improvements and coordination strategies contributed up to 40% additional capacity gain under high-density unlicensed operation.

The management of external interference, a critical challenge for in-X subnetworks operating in real-world environments, was addressed through a two-pronged strategy combining robust resource allocation and advanced receiver design. The Gradient Descent-based Resource Allocation (GDRA) algorithm successfully limited spectral efficiency degradation to 9.7% in the presence of external interference, outperforming state-of-the-art benchmarks, which experienced a 13.3% loss. On the

receiver side, low-complexity yet robust likelihood ratio approximation methods were proposed, allowing resilient decoding even in the presence of impulsive noise and jamming, with adaptations based on real-time channel observations and minimal computational overhead.

The results demonstrate that the developed RRM strategies are capable of supporting dense, autonomous, and resilient subnetwork deployments across industrial, consumer, and vehicular domains. They deliver substantial improvements in reliability, scalability, spectral efficiency, and interference robustness over current benchmarks, meeting or surpassing the technical targets established for the 6G-SHINE project. These contributions form a strong foundation for future advancements in enabling 6G subnetworks to operate effectively under the highly dynamic and interference-prone conditions anticipated in next-generation wireless networks.

REFERENCES

- [1] M. Series, "IMT Vision–Framework and overall objectives of the future development of IMT for 2020 and beyond," *Recomm. ITU*, vol. 2083, no. 0, 2015xxx.
- [2] Saeed Hakimi and et al, "D4.1 - PRELIMINARY RESULTS ON THE MANAGEMENT OF RADIO RESOURCES IN SUBNETWORKS IN THE PRESENCE OF LEGITIMATE AND MALICIOUS INTERFERERS," 6G-SHINE, Jun. 2024.
- [3] Basuki Priyanto and et al, "D2.2. – REFINED DEFINITION OF SCENARIOS, USE CASES AND SERVICE REQUIREMENTS FOR IN-X SUBNETWORKS," 6G-SHINE, Feb. 2024.
- [4] R. Adeogun, G. Berardinelli, P. E. Mogensen, I. Rodriguez, and M. Razzaghpour, "Towards 6G in-X Subnetworks With Sub-Millisecond Communication Cycles and Extreme Reliability," *IEEE Access*, vol. 8, pp. 110172–110188, 2020.
- [5] M. Ding, D. Lopez-Perez, H. Claussen, and M. A. Kaafar, "On the Fundamental Characteristics of Ultra-Dense Small Cell Networks," *IEEE Netw.*, vol. 32, no. 3, pp. 92–100, May 2018.
- [6] S. Hakimi, P. Koteswar Srinath, S. Bagherinejad, R. Adeogun, and G. Berardinelli, "Robust Resource Management for Mission-Critical In-Factory Subnetworks under External Interference," presented at the 2025 IEEE 101st Vehicular Technology Conference (VTC2025-Spring), Jun. 13, 2025. doi: 10.5281/zenodo.15657826.
- [7] S. Hakimi, R. Adeogun, and G. Berardinelli, "Resilient DNN for Joint Sub-Band Allocation and Power Control in Mobile Factory Subnetworks," Accepted for Publication in *EURASIP Journal on Wireless Communication and Networking*, May 2025.
- [8] S. Hakimi, G. Berardinelli, and R. Adeogun, "Attention-Aided Channel Prediction for Efficient Resource Management in Industrial IoT Subnetworks," Manuscript submitted for publication in *IEEE Internet of Things Journal*, 2024.
- [9] S. Bagherinejad, T. Jacobsen, N. Kiilerich Pratas, and R. O. Adeogun, "DRL-based Distributed Joint Sub-band Allocation and Power Control for eXtended Reality over In-Body Subnetworks: IEEE Wireless Communications and Networking Conference 2025," 2025 IEEE Wirel. Commun. Netw. Conf. WCNC2025, 2025.
- [10] 3GPP TR 38.901, 5G; Study on channel model for frequencies from 0.5 to 100 GHz, v17.1.0.
- [11] S. Lu, J. May, and R. J. Haines, "Effects of correlated shadowing modeling on performance evaluation of wireless sensor networks," in 2015 IEEE 82nd Vehicular Technology Conference (VTC2015-Fall), IEEE, 2015, pp. 1–5.
- [12] Enrico M. Vitucci and et al, "D2.3: Radio Propagation Characteristics for IN-X Subnetworks," 6G-SHINE, Nov. 2024.
- [13] L. Su, C. Yang, and S. Han, "The Value of Channel Prediction in CoMP Systems with Large Backhaul Latency," *IEEE Trans. Commun.*, vol. 61, no. 11, pp. 4577–4590, Nov. 2013.
- [14] W. Jiang and H. D. Schotten, "Neural Network-Based Fading Channel Prediction: A Comprehensive Overview," *IEEE Access*, vol. 7, pp. 118112–118124, 2019.
- [15] K. E. Baddour and N. C. Beaulieu, "Autoregressive modeling for fading channel simulation," *IEEE Trans. Wirel. Commun.*, vol. 4, no. 4, pp. 1650–1662, Jul. 2005.

- [16] C. Wu, X. Yi, Y. Zhu, W. Wang, L. You, and X. Gao, "Channel Prediction in High-Mobility Massive MIMO: From Spatio-Temporal Autoregression to Deep Learning," *IEEE J. Sel. Areas Commun.*, vol. 39, no. 7, pp. 1915–1930, Jul. 2021.
- [17] H. Kim, S. Kim, H. Lee, C. Jang, Y. Choi, and J. Choi, "Massive MIMO Channel Prediction: Kalman Filtering Vs. Machine Learning," *IEEE Trans. Commun.*, vol. 69, no. 1, pp. 518–528, Jan. 2021.
- [18] X. Wen and W. Li, "Time series prediction based on LSTM-attention-LSTM model," *IEEE Access*, vol. 11, pp. 48322–48331, 2023.
- [19] X. Yuan, L. Li, Y. A. W. Shardt, Y. Wang, and C. Yang, "Deep Learning With Spatiotemporal Attention-Based LSTM for Industrial Soft Sensor Model Development," *IEEE Trans. Ind. Electron.*, vol. 68, no. 5, pp. 4404–4414, May 2021.
- [20] J. Zhao et al., "Do RNN and LSTM have Long Memory?," Jun. 10, 2020, arXiv: arXiv:2006.03860.
- [21] S. Hochreiter and J. Schmidhuber, "Long Short-Term Memory," *Neural Comput.*, vol. 9, no. 8, pp. 1735–1780, Nov. 1997.
- [22] M. Sundermeyer, R. Schlüter, and H. Ney, "LSTM neural networks for language modeling," presented at the Proc. Interspeech 2012, 2012, pp. 194–197.
- [23] S. Hakimi, R. Adeogun, and G. Berardinelli, "Rate-conforming Sub-band Allocation for In-factory Subnetworks: A Deep Neural Network Approach," in 2024 Joint European Conference on Networks and Communications & 6G Summit (EuCNC/6G Summit), IEEE, 2024, pp. 729–734.
- [24] H. Sun, X. Chen, Q. Shi, M. Hong, X. Fu, and N. D. Sidiropoulos, "Learning to optimize: Training deep neural networks for interference management," *IEEE Trans. Signal Process.*, vol. 66, no. 20, pp. 5438–5453, 2018.
- [25] D. Li, S. R. Khosravirad, T. Tao, and P. Baracca, "Advanced frequency resource allocation for industrial wireless control in 6G subnetworks," in 2023 IEEE Wireless Communications and Networking Conference (WCNC), IEEE, 2023, pp. 1–6.
- [26] Q. Shi, M. Razaviyayn, Z.-Q. Luo, and C. He, "An Iteratively Weighted MMSE Approach to Distributed Sum-Utility Maximization for a MIMO Interfering Broadcast Channel," *IEEE Trans. Signal Process.*, vol. 59, no. 9, pp. 4331–4340, Sep. 2011.
- [27] Y. Gu, C. She, Z. Quan, C. Qiu, and X. Xu, "Graph neural networks for distributed power allocation in wireless networks: Aggregation over-the-air," *IEEE Trans. Wirel. Commun.*, 2023.
- [28] PyTorch, "Pytorch." [Online]. Available: <https://pytorch.org/>
- [29] "srsRAN Project with GNU Radio, Copyright 2019-2024, Software Radio Systems." [Online]. Available: https://docs.srsran.com/projects/4g/en/latest/app_notes/source/zeromq/source/index.html
- [30] "GNU Radio Project, GNU Radio: The Free and Open-Source Toolkit for Software Radio, 2024." [Online]. Available: <https://www.gnuradio.org/>
- [31] ZeroMQ, "Zeromq." [Online]. Available: <https://zeromq.org/>
- [32] "Future Wireless Connected and Automated Industry Enabled by 5G (5GSmartFact)," MSCA-ITN, Grant Agreement ID 956670.
- [33] D. Abode, P. M. de Sant Ana, R. Adeogun, and G. Berardinelli, "Communication-Aware Dynamic Speed Control for Interference Mitigation in 6G Mobile Subnetworks," in 2024 IEEE Conference on Standards for Communications and Networking (CSCN), IEEE, 2024, pp. 1–7.

- [34] R. Adeogun and G. Berardinelli, "Distributed Channel Allocation for Mobile 6G Subnetworks via Multi-Agent Deep Q-Learning," in 2023 IEEE Wireless Communications and Networking Conference (WCNC), Mar. 2023, pp. 1–6.
- [35] P. Popovski et al., "Wireless Access in Ultra-Reliable Low-Latency Communication (URLLC)," IEEE Trans. Commun., vol. 67, no. 8, pp. 5783–5801, Aug. 2019.
- [36] 3GPP TS 38.300, NR; NR and NG-RAN Overall description; Stage-2, TS, v18.0.0.
- [37] D. Li and Y. Liu, "In-Band Emission in LTE-A D2D: Impact and Addressing Schemes," in 2015 IEEE 81st Vehicular Technology Conference (VTC Spring), May 2015, pp. 1–5.
- [38] 3GPP TS 38.214, 5G; NR; Physical layer procedures for data, clause 8.1.
- [39] 3GPP TS 38.101-1, NR; User Equipment (UE) radio transmission and reception; Part 1: Range 1 Standalone, v18.4.0.
- [40] 3GPP TSG RAN WG1, Final Report of 3GPP TSG RAN WG1 #110 v1.0.0, August 2022.
- [41] European Telecommunications Standards Institute (ETSI), ETSI EN 301 893 V2.2.1 (2024-11), 5 GHz RLAN; Harmonised Standard for access to radio spectrum.
- [42] R. Maldonado, C. Rosa, and K. I. Pedersen, "A Fully Coordinated New Radio-Unlicensed System for Ultra-Reliable Low-Latency Applications," in 2020 IEEE Wireless Communications and Networking Conference (WCNC), May 2020, pp. 1–6.
- [43] 3rd Generation Partnership Project (3GPP), "NR; Requirements for support of radio resource management," 3GPP TS 38.133 v18.4.0.
- [44] S. Ahmadi, 5G NR: Architecture, technology, implementation, and operation of 3GPP new radio standards. Academic Press, 2019.
- [45] D. Li, S. R. Khosravirad, T. Tao, P. Baracca, and P. Wen, "Power Allocation for 6G Sub-Networks in Industrial Wireless Control," in 2024 IEEE Wireless Communications and Networking Conference (WCNC), IEEE, 2024, pp. 1–6.
- [46] D. Abode, R. Adeogun, and G. Berardinelli, "Power Control for 6G Industrial Wireless Subnetworks: A Graph Neural Network Approach," in 2023 IEEE Wireless Communications and Networking Conference (WCNC), IEEE, 2023, pp. 1–6.
- [47] Y. Mestrah, A. Savard, A. Goupil, L. Clavier, and G. Gellé, "Blind Estimation of an Approximated Likelihood Ratio in Impulsive Environment," in 2018 IEEE 29th Annual International Symposium on Personal, Indoor and Mobile Radio Communications (PIMRC), Sep. 2018, pp. 1–5.
- [48] Y. Mestrah et al., "Unsupervised Log-Likelihood Ratio Estimation for Short Packets in Impulsive Noise," in 2022 IEEE Wireless Communications and Networking Conference (WCNC), Apr. 2022, pp. 944–949.
- [49] Y. Mestrah, A. Savard, A. Goupil, G. Gellé, and L. Clavier, "An Unsupervised LLR Estimation with unknown Noise Distribution," EURASIP J. Wirel. Commun. Netw., vol. 2020, no. 1, p. 26, Dec. 2020.
- [50] Y. Mestrah, A. Savard, A. Goupil, G. Gelle, and L. Clavier, "Robust and simple log-likelihood approximation for receiver design," in 2019 IEEE Wireless Communications and Networking Conference (WCNC), IEEE, 2019, pp. 1–6.