



Performance evaluation of synergic operation of algorithms enabling opportunistic networks – D4.3

Project Number:	ICT-2009-257385
Project Title:	Opportunistic networks and Cognitive Management Systems for Efficient Application Provision in the Future Internet - OneFIT
Document Type:	Deliverable

Contractual Date of Delivery:	31.12.2012
Actual Date of Delivery:	14.01.2013
Editors:	Seiamak Vahid (UNIS)
Participants:	Please see the list of contributors
Workpackage:	WP4
Estimated Person Months:	51
Nature:	PU ¹
Version:	1.0
Total Number of Pages:	208
File:	OneFIT_D4.3_20121231.doc

Abstract

This deliverable presents final set of the algorithmic solutions and evaluation results for enabling opportunistic networks. The results are based on matured algorithms and performance assessments include aspects related to the practicality of the solutions proposed. Finally, an integrated approach is adopted to showcase synergic and dynamic operation of different algorithms in support of main ON scenarios.

Keywords List

Opportunistic networks, algorithms, infrastructure-less networks, traffic aggregation, suitability determination, route selection, node selection, spectrum selection.

¹Dissemination level codes:

PU = Public

PP = Restricted to other programme participants (including the Commission Services)

RE = Restricted to a group specified by the consortium (including the Commission Services)

CO = Confidential, only for members of the consortium (including the Commission Services)

Executive Summary

This document presents descriptions of the final set of algorithmic solutions and results for enabling opportunistic networks developed in the OneFIT project [1]. The first part of the document describes the matured final design of algorithms that have been proposed to cope with the technical challenges arising in the ON management stages. Proposed algorithms include the ON suitability determination in different scenarios, the spectrum opportunity identification and selection, the nodes and routes selection, and the capacity extension through femtocells. This part of the document also includes further performance evaluation results of each of the specified algorithms, according to Key Performance Indicators in representative scenarios. In the following, a list is presented with the main conclusions that are obtained for each of the considered main topics:

- **ON Suitability determination:** The possibility of direct communication between devices to establish the ON has been evaluated in the “infrastructure supported device-to-device communication” and “coverage extension” scenarios. It is revealed that this probability depends on the probability of the users being in the same “Area of Interest” and on the range of the wireless interface.
- **Spectrum opportunity identification and selection:** Different approaches are proposed depending on the specific scenario considerations. The fittingness factor concept has been proposed as a novel metric that captures the time-varying suitability of available spectrum resources to different applications supported in each ON link. Advanced statistics of this metric capturing the dynamic radio environment are stored in a Knowledge Database and used to support the spectrum selection and spectrum mobility algorithms. Performance evaluation shows that the strategy introduces significant gains with respect to a random selection and performs very closely to an upper-bound optimal scheme. The algorithm has been extended to cope with non-stationary environments, through inclusion of a novel Reliability Tester (RT) component, and, to support different operator preferences in the spectrum pools that can be assigned to each radio link, reflecting e.g. different operator policies or different regulatory constraints related to the band that each specific pool belongs to. The new set of results that rely on the spatial and temporal dimensions are captured, by analysing the robustness to non-stationary environments, indicate that proposed strategy does lead to substantial increases of the global efficiency under varying traffic loads. The possibility of jointly considering the frequency, bandwidth and radio access technique selection has also been studied. The algorithm is used to solve operator network congestion or coverage problems by selecting band, RAT, and bandwidth for creating a short range communication links, whilst guaranteeing a proper QoS to the users. The algorithm makes the decision taking into account users’ capabilities, velocity and application requirements. The inclusion of machine learning in the knowledge acquisition on spectrum usage is also addressed in this document. Results reveal the capability of learning by means of successful interactions with the environment (rather than relying on spectrum usage history). It is shown that the proposed strategies can identify unused resource blocks with high degree of accuracy, resulting also in significantly higher average cell throughput and satisfaction probabilities in the considered scenarios, while co/cross-layer interference can be reduced significantly, thus resulting in average cell throughput increases of 30% and satisfaction probabilities increase of 47%, in the considered scenarios.. Finally, another field of study in this area is the inclusion of the spectrum aggregation in the spectrum selection. The final set of results based on system level simulations, confirm the original analytical performance results and show that aggregation in deployments with small size cells can result in significant levels of channel switching by cognitive radios.
- **Node and route selection:** Different strategies have been proposed and analysed regarding the capability to select the adequate nodes to form the ON and the most suitable routes between them.

An algorithm on knowledge-based suitability determination and selection of nodes and routes is proposed for the capacity extension scenario in which a congestion situation occurs in the infrastructure. The route selection is based on the Ford-Fulkerson maximum flow algorithm while the node selection is based on a fitness function that weights different parameters of the candidate nodes. Extended evaluation results based on new use-cases reveal that the energy consumption of congested base stations is reduced within a range from 15 to 25% as more terminals switch to ON. Also, the quality of communication is benefited, as delay of successfully delivered messages drops approximately 15-35%, compared to the congested situation. Furthermore, an average decrease of 15-40% has been achieved in the load of the congested base stations. In another approach, a network coding optimization is used for multi-flow route co-determination, based on the introduction of delegated nodes. Results show good improvements on the throughputs in terms of number of data packets to be exchanged between pair nodes.

Another work in this area addresses the creation and maintenance of network topology through enabling coordination in decision making among nodes taking into account different parameters to establish links/topology with a set of desired global properties and constraints. Updated results based on combined consideration of k-connectivity, power and interference, in the final formulation indicate that the proposed models can provide practical, yet optimal power levels to minimize the power utilization while maintaining connectivity. It is shown that the applied approximation technique (column generation), together with PSO heuristic are able to produce near optimal power levels (min. total network power) within guaranteed connected topologies.

The establishment of a WLAN network between different devices using the connections allowed by another technology is considered in the so-called UE-to-UE trusted direct path activity. It deals with the selection of the candidate access point and the operating channel to fulfil the QoS and power minimization requirements. The extended formulations propose a QoS-constrained, novel joint AP and channel allocation technique for WLAN configuration. Challenges specific for the OneFIT scenario 5 "Opportunistic resource aggregation in the backhaul network" are addressed with three algorithmic solutions. First, the application cognitive multipath routing algorithm provides the complete solution for route identification and selection in the backhaul of WMNs for purpose of opportunistic BW sharing and aggregation. Related performance evaluation results show that the algorithm provides better load balancing and BW resource utilization than the QoS aware single path routing solutions while maintaining or increasing QoS level offered to the end users. In addition the autonomic trigger recognition feature of the algorithm was presented together with proof of concept results. The "Content conditioning and distributed storage virtualization/aggregation for context driven media delivery" algorithm provides context aware, knowledge based and opportunistic aggregation of backhaul storage resources in WMNs. Corresponding mathematical model for node selection is thoroughly evaluated in the custom built simulator. In addition, the GW selection variant of the algorithm, showing algorithm's ability to make proper decision making regarding selection of WMN nodes, is implemented in the open platform WMN test-bed and practically validated. These two algorithms provide complete solution for node and route identification and selection from perspective of the scenario 5. The spectrum selection challenge is successfully covered with the fitness factor based spectrum selection algorithm. These three algorithms provide complete algorithmic solution for the scenario 5 related challenges throughout all ON management phases. Achieved context aware, opportunistic and knowledge based resource management in the backhaul can be implemented in operator's network together with solutions which are addressing capacity extensions in the access side in order to provide pervasive resource management resulting in the highest level of resource utilization while maintaining requested level of QoS.

- Capacity extension through femtocells: A novel algorithm is presented addressing the optimization problem of allocation of network resources to reroute macro-terminals to deployed femtocells in congestion situations, while keeping the minimum transmit power levels and switching off unnecessary femtocells and in this way extending the capacity of the congested networks in an energy-efficient manner. The problem is mathematically formulated and solved by means of a novel greedy algorithm (Energy-Efficient Resource Allocation - ERA). The evaluations prove that the utilization of femtocells in a congested area can reduce the average delay in packet delivery by 15% in average as well increase the delivery probability. Furthermore, with respect to user distribution within the congested area, when users are distributed in a centralized manner energy benefits exist due to the fact that only 25% of the femtocells in average need to be switched on, while in sparse distributions femtocells are switched on but they operate at minimum transmission power levels.
- In the support activity to validate ON algorithms on an offloading-oriented real-deployment testbed, the objective has been to evaluate the performance of an ON suitability determination algorithm on a specific femtocell-populated urban environment, in order to identify some recommendations for the implementation of an ON-based macro-to-femto offloading mechanism in LTE. The Results indicate noticeable performance gains due to ONs.

The second part takes a closer look at the specific approaches adopted to address the two main themes of spectrum selection and nodes/route selection. More specifically, the proposed spectrum selection techniques, for instance, have taken advantage of a wide range of methods and techniques including, utility based approach for aggregation, fitness-based method for spectrum opportunity identification, knowledge based selection, machine-learning for spectrum opportunity identification and fuzzy techniques. The second part then presents an examination of the value-add provided by these approaches, in terms of benefits and increased adaptability.

Finally, the last part of the deliverable deals with aspects related to dynamic operation of ONs and outlines the overall relationship between algorithmic solutions, different ON phases and individual scenarios, based on integrated and synergic operation of algorithms. With a particular focus on the solutions for the technical challenges of spectrum opportunity identification and selection and node and route selection, identified as two of the major challenges in the ON formation and management stages, on per scenario basis, a mapping of algorithms to challenges and algorithm execution sequence are provided, as examples that illustrate how dynamic operation of ONs can be realized.

Contributors

First Name	Last Name	Affiliation	Email
Jordi	Pérez-Romero	UPC	jorperez@tsc.upc.edu
Oriol	Sallent	UPC	sallent@tsc.upc.edu
Faouzi	Bouali	UPC	faouzi@tsc.upc.edu
Ramon	Ferrús	UPC	ferrus@tsc.upc.edu
Ramon	Agustí	UPC	ramon@tsc.upc.edu
Jens	Gebert	ALUD	Jens.Gebert@alcatel-lucent.com
Rolf	Fuchs	ALUD	Rolf.Fuchs@alcatel-lucent.com
Panagiotis	Demestichas	UPRC	pdemest@unipi.gr
Andreas	Georgakopoulos	UPRC	andgeorg@unipi.gr
Vera	Stavroulaki	UPRC	veras@unipi.gr
Kostas	Tsagkaris	UPRC	ktsagk@unipi.gr
Yioulis	Kritikou	UPRC	kritikou@unipi.gr
Lia	Tzifa	UPRC	etzifa@unipi.gr
Nikos	Koutsouris	UPRC	nkouts@unipi.gr
Dimitris	Karvounas	UPRC	dkarvoyn@unipi.gr
Marios	Logothetis	UPRC	mlogothe@unipi.gr
Asimina	Sarli	UPRC	Mina.sarli@gmail.com
Aimilia	Bantouna	UPRC	abantoun@unipi.gr
Louisa-Magdalene	Papadopoulou	UPRC	lpapadop@unipi.gr
Aristi	Galani	UPRC	agalani@unipi.gr
Panagiotis	Vlacheas	UPRC	panvlah@unipi.gr
Petros	Morakos	UPRC	pmorakos@unipi.gr
Alexandros	Antzoulatos	UPRC	alexant@unipi.gr
Óscar	Moreno	TID	omj@tid.es
Milenko	Tosic	LCI	milenko.tosic@lacidelleing.com
Dragan	Boskovic	LCI	dragan.boskovic@lacidelleing.com
Mirko	Cirilovic	LCI	mirko.cirilovic@lacidelleing.com
Heli	Sarvanko	VTT	heli.sarvanko@vtt.fi
Xianfu	Chen	VTT	Xianfu.chen@vtt.fi
Stéphane	Pega	TCS	Stephane.pegas@thalesgroup.com

Michel	Bourdellès	TCS	Michel.bourdelles@thalesgroup.com
Seiamak	Vahid	UNIS	s.vahid@surrey.ac.uk
Abdoulaye	Bagayoko	NTUK	abdoulaye.bagayoko@nectech.fr
Christian	Mouton	NTUK	christian.mouton@nectech.fr
Thomas	Delsol	NTUK	thomas.delsol@nectech.fr
Dorin	Panaitopol	NTUK	dorin.panaitopol@nectech.fr
Maryam	Riaz	UNIS	m.riaz@surrey.ac.uk
Ghassan	Alnwaimi	UNIS	ghassan.almwaimi@surrey.ac.uk
Haeyoung	Lee	UNIS	haeyoung.lee@surrey.ac.uk

Table of Acronyms

Acronym	Meaning
3GPP	Third Generation Partnership Project
ADTL	Average Delivery Tree Length
AOI	Area of Interest
AP	Access Point
BS	Base Station
CR	Cognitive Radio
CSG	Closed Subscriber Group
D2D	Device-to-Device
DH	Digital Home
ERA	Energy-efficient Resource Allocation
GW	GateWay
ISM	Industrial, Scientific and Medical
KD	Knowledge Database
KM	Knowledge Manager
KPI	Key Performance Indicator
MNO	Mobile Network Operator
OF	Objective Function
ON	Opportunistic Network
OSG	Open Subscriber Group
RT	Reliability Tester
SM	Spectrum Mobility
SpHO	Spectrum HandOver
SS	Spectrum Selection
TVWS	TeleVision White Space
UE	User Equipment
WMN	Wireless Mesh Network
SU	Secondary User
PU	Primary User

Table of Contents

1.	Introduction.....	18
2.	Algorithmic Solutions (final version)	19
2.1	Suitability determination for direct device-to-device (D2D) communication.....	19
2.1.1	Problem formulation and algorithm concept.....	19
2.1.2	Algorithm specification.....	19
2.1.3	Integration in OneFIT architecture.....	21
2.1.4	Further performance evaluation results	21
2.2	Suitability determination for the coverage extension scenario	21
2.2.1	Problem formulation and algorithm concept.....	21
2.2.2	Algorithm specification.....	22
2.2.3	Integration in OneFIT architecture.....	24
2.2.4	Further performance evaluation results	24
2.3	Modular decision flow approach for selecting frequency, bandwidth and radio access technique for ONS24	
2.3.1	Problem formulation and algorithm concept	24
2.3.2	Algorithm specification.....	24
2.3.3	Integration in OneFIT architecture.....	26
2.3.4	Further performance evaluation results	26
2.4	Fittingness factor-based spectrum selection.....	28
2.4.1	Problem formulation and algorithm concept.....	28
2.4.2	Algorithm specification.....	28
2.4.3	Integration in OneFIT architecture.....	32
2.4.4	Further performance evaluation results	32
2.5	Machine Learning based Knowledge Acquisition on Spectrum Usage	33
2.5.1	Problem formulation and algorithm concept.....	33
2.5.2	Algorithm specification.....	33
2.5.3	Integration in OneFIT architecture.....	39
2.5.4	Further performance evaluation results	39
2.6	Techniques for Aggregation of Available Spectrum Bands/Fragments	47
2.6.1	Problem formulation and algorithm concept.....	47
2.6.2	Algorithm specification.....	47
2.6.3	Integration in OneFIT architecture.....	48
2.6.4	Further performance evaluation results	48
2.7	Knowledge-based suitability determination and selection of nodes and routes	50
2.7.1	Problem formulation and algorithm concept	50
2.7.2	Algorithm specification.....	51
2.7.3	Integration in OneFIT architecture.....	54
2.7.4	Further performance evaluation results	56
2.8	Route pattern selection in ad hoc network	59
2.8.1	Problem formulation and algorithm concept.....	59
2.8.2	Algorithm specification.....	60
2.8.3	Integration in OneFIT architecture.....	62
2.8.4	Further performance evaluation results	64
2.9	Multi-flow routes co-determination	65
2.9.1	Problem formulation and algorithm concept.....	65

2.9.2	Algorithm specification.....	66
2.9.3	Integration in OneFIT architecture.....	66
2.9.4	Further performance evaluation results.....	67
2.10	Techniques for Network Reconfiguration – topology Design.....	69
2.10.1	Problem formulation and algorithm concept.....	69
2.10.2	Algorithm specification.....	70
2.10.3	Integration in OneFIT architecture.....	70
2.10.4	Further performance evaluation results.....	72
2.11	Application cognitive multi-path routing in wireless mesh networks.....	73
2.11.1	Problem formulation and algorithm concept.....	73
2.11.2	Algorithm specification.....	74
2.11.3	Integration in OneFIT architecture.....	77
2.11.4	Further performance evaluation results.....	78
2.12	UE-to-UE Trusted Direct Path.....	85
2.12.1	Problem formulation and algorithm concept.....	85
2.12.2	Algorithm specification.....	87
2.12.3	Integration in OneFIT architecture.....	94
2.12.4	Further performance evaluation results.....	94
2.13	Content conditioning and distributed storage virtualization/aggregation for context driven media delivery.....	94
2.13.1	Problem formulation and algorithm concept.....	94
2.13.2	Algorithm specification.....	94
2.13.3	Integration in OneFIT architecture.....	97
2.13.4	Further performance evaluation results.....	97
2.14	Capacity Extension through Femto-cells.....	101
2.14.1	Problem formulation and algorithm concept.....	101
2.14.2	Algorithm specification.....	102
2.14.3	Integration in OneFIT architecture.....	104
2.14.4	Further performance evaluation results.....	105
2.15	Support activity to validate ON algorithms on an offloading-oriented real-deployment testbed.....	108
2.15.1	Problem formulation and algorithm concept.....	108
2.15.2	Algorithm specification.....	110
2.15.3	Integration in OneFIT architecture.....	111
2.15.4	Further performance evaluation results.....	112
3.	Functional Components of the comprehensive OneFIT Solution.....	146
3.1	OneFIT solution for Spectrum selection.....	146
3.1.1	Description.....	146
3.1.2	Evaluation.....	147
3.2	OneFIT solution for Selection of Nodes and Routes.....	161
3.2.1	Description.....	161
3.2.2	Evaluation.....	165
4.	Dynamic Operation of the ON.....	168
4.1	Scenario 1"Opportunistic Coverage Extension".....	168
4.1.1	ON Suitability.....	170
4.1.2	ON Creation.....	171
4.1.3	ON Maintenance.....	171
4.1.4	ON Termination.....	172

4.1.5	Performance evaluation	172
4.2	Scenario2 "Opportunistic Capacity Extension"	177
4.2.1	ON Suitability	178
4.2.2	ON Creation	179
4.2.3	ON Maintenance.....	180
4.2.4	ON Termination	182
4.2.5	Performance evaluation	183
4.3	Scenario 3 "Infrastructure supported D2D communication"	183
4.3.1	ON Suitability	185
4.3.2	ON Creation	185
4.3.3	ON Maintenance.....	186
4.3.4	ON Termination	186
4.3.5	Performance evaluation	186
4.4	Scenario 5 "Opp. resource aggregation in the backhaul network"	189
4.4.1	ON Suitability	191
4.4.2	ON Creation	192
4.4.3	ON Maintenance.....	193
4.4.4	ON Termination	194
4.4.5	Performance evaluation	195
5.	Conclusions	201
6.	References	207

List of Figures

Figure 1: Infrastructure supported opportunistic networking via direct device-to-device communication	19
Figure 2: Suitability Determination for an Infrastructure supported direct device-to-device communication	20
Figure 3: Terminating the direct device-to-device communication upon certain events	21
Figure 4: Opportunistic coverage extension when a UE is going out of infrastructure coverage	22
Figure 5: Suitability Determination for the coverage extension scenario	23
Figure 6: Termination of the ON when the device earlier being out of coverage gets into coverage ..	24
Figure 7: Modular decision flow for band, RAT, channel, and bandwidth selection.....	25
Figure 8: Band selection for different traffic types.....	27
Figure 9: Capacity of different traffic type users using enhanced algorithm and former algorithm. ..	28
Figure 10: Functional architecture of the proposed Fittingness Factor-based Spectrum Management Framework.....	29
Figure 11: Functional proposed multi-objective CODIPAS-HRL model for FCs self-configuration in HNs deployment.....	35
Figure 12: Cell layout and SINR map of the LTE heterogeneous network HNs	41
Figure 13: The selected strategy of FC number 20 using the GF learning strategy.....	42
Figure 14: The average received utility of the proposed learning model for FCs number 5, 10, 15, and 20, with initial strategy profile $\{=1\}$	42
Figure 15: The dissatisfaction probability and the average cell throughout of the proposed leaning model as compared to the GDC.....	43
Figure 16: SINR cumulative distributed function (CDF) for the eNBs and FCs with and without the proposed learning model.....	43
Figure 17: The normalised interference level and the utilisation index for the SMU and SSU versus the learning time.....	44
Figure 18: FCs RBs selection and avoidance behaviour and normalised interference level of the GF learning strategy.	45
Figure 19: FCs RBs selection and avoidance behaviour and normalised interference level of the MRE and MBM learning strategies.....	46
Figure 20: Framework of the proposed spectrum aggregation approach.....	47
Figure 21: (a) Normalized average Throughput (b) Normalized number of channel switching (c) Normalized number of sub-channels.....	49
Figure 22: Channel switching performance for different ISD	50
Figure 23: Outline of the knowledge-based suitability determination approach	51
Figure 24: Flowchart of the proposed solution for the suitability determination problem	52
Figure 25: Flowchart of the proposed algorithm.....	53
Figure 26: Flowchart of the proposed algorithm.....	54
Figure 27: Integration of the “knowledge-based suitability determination” in the OneFIT Functional Architecture	54
Figure 28: Integration of the “selection of nodes and routes through a fitness value evaluation” in the OneFIT Functional Architecture.....	55
Figure 29: Integration of the “Capacity extension of the infrastructure through neighbouring terminals” in the OneFIT Functional Architecture	56

Figure 30: Mean delay (in seconds) for each case	57
Figure 31: Delivery probability for each case.....	57
Figure 32: Mean power of intermediate terminals	58
Figure 33: Mean power of edge terminals (terminals that switch to ON).....	58
Figure 34: Mean power of congested BS	59
Figure 35: Mean power of non-congested BSs	59
Figure 36 : Flow chart for signalling packet route selection	60
Figure 37 : Flow chart for conversational packet route selection	61
Figure 38 : Flow chart for streaming packet route selection.....	61
Figure 39 : Flow chart for background packet route selection.....	62
Figure 40 : Mapping of the route selection algorithm in the OneFIT functional architecture	62
Figure 41: Sequence Diagram for Route pattern selection algorithm during ON Suitability determination	63
Figure 42: Sequence Diagram for Route pattern selection algorithm during ON maintenance	64
Figure 43:Example of route pattern selection	65
Figure 44: Integration of the algorithm in the functional architecture	66
Figure 45: Integration of the algorithm in the functional architecture	67
Figure 46	67
Figure 47	68
Figure 48	68
Figure 49: Algorithm flow charts	70
Figure 50: Placement of the TC algorithm & components in OneFIT Functional Model.....	71
Figure 51: Placement of the TC algorithm in OneFIT Functional Architecture.....	71
Figure 52: EER vs #Nodes	72
Figure 53: total network power vs # Nodes.....	73
Figure 54: Application cognitive multipath routing algorithm mapped onto ON management phases	74
Figure 55: Suitability determination steps of the application cognitive multipath routing algorithm.	75
Figure 56: Autonomic trigger detection feature of the application cognitive multipath routing algorithm.....	75
Figure 57: MSC of the multipath routing algorithm	77
Figure 58 – Recognition of end user’s request for service	78
Figure 59 – Recognized trigger for ON creation	78
Figure 60 – Recognized trigger for ON reconfiguration/termination	79
Figure 61 – Total number of end users at the UNS WMN during 8 month period.....	79
Figure 62 – The number of end users at one AP during one week.....	80
Figure 63 – The number of end users on one AP during one month	80
Figure 64 – Standard approach for ETX value gathering with detected threshold breaches.....	83
Figure 65 – Impact of local data pre-processing on accuracy of ETX awareness	84
Figure 66: AP and Channel Selection during the creation of a WLAN	86
Figure 67: Flowchart of SNR_{req} derivation.....	88
Figure 68: Links mean SINR for given AP and channel assumptions.	89
Figure 69: Flowchart of SNR_{req} derivation, and of the per-link PER achievability checking.....	90

Figure 70: Flowchart of the joint AP and Channel selection for WLAN establishment: 1 ^{er} algorithm. This algorithm allows finding one compliant couple (AP, Wi-Fi channel) that meets the per-link PER requirements.	91
Figure 71: Flowchart of minimum required transmit power derivation	92
Figure 72: Flowchart of the joint AP and Channel selection for WLAN establishment: 2 nd algorithm. This algorithm allows finding the most power efficient couple (AP, Wi-Fi channel) that meets the per-link PER requirements associated to a particular 3GPP QCI.....	93
Figure 73: Mapping of the algorithm onto ON management phases.....	96
Figure 74 – Set up of the open platform WMN test-bed for validation of the GW selection algorithm	98
Figure 75 – User mobility results in change in request distribution and GW selection	98
Figure 76 – Jitter levels measured at U4 side	100
Figure 77 – Packet loss percentage measured at U4 side	100
Figure 78: The concept of the Energy-efficient Resource Allocation	101
Figure 79: Flow-chart of the ERA algorithm.....	104
Figure 80: Mapping of the ERA concept functional entities to the OneFIT functional architecture ..	105
Figure 81: Considered topology.....	105
Figure 82: Delivery probability for each considered case.....	106
Figure 83: Average delay for each considered case.....	107
Figure 84: Power allocation to femtocells for each grouped case	107
Figure 85: Number of terminals acquired by femtocells for each grouped case	107
Figure 86: ERA runtime for each grouped case	108
Figure 87: Base LTE test scenario and focus area	109
Figure 88 : Mapping of the algorithmic procedure to the OneFIT functional architecture.....	112
Figure 89: Baseline test scenario. Green: macros. Cyan: femtos. Red: UEs.	113
Figure 90: Per user DL throughput.....	114
Figure 91: Per user DL throughput.....	115
Figure 92: Best server map before and after a macro cell is turned off.....	116
Figure 93: Best server map before and after a macro site is turned off.	117
Figure 94: Best server map before and after 3 macro cells are turned off.	118
Figure 95: Best server maps as all 5 macro cells are subsequently turned off.....	119
Figure 96: RSRP maps as all 5 macro cells are subsequently turned off.	120
Figure 97: Number of disconnected macro users.....	121
Figure 98: Distribution of connected users.....	121
Figure 99: Average DL throughput.....	122
Figure 100: Total DL traffic.....	122
Figure 101: Throughput enhancement for ON users.....	123
Figure 102: Number of disconnected macro users.....	124
Figure 103: Distribution of connected users.....	124
Figure 104: Number of disconnected macro users.....	125
Figure 105: Average DL throughput.....	125
Figure 106: Macro DL throughput.	126
Figure 107: Total DL traffic.....	126

Figure 108: Total DL traffic.....	127
Figure 109: Throughput enhancement for ON users.....	127
Figure 110: Number of disconnected macro users.....	128
Figure 111: Distribution of connected users.....	128
Figure 112: Average DL throughput.....	129
Figure 113: Total DL traffic.....	129
Figure 114: Throughput enhancement for ON users.....	130
Figure 115: Number of disconnected macro users.....	131
Figure 116: Number of users transferred to the ON.	132
Figure 117: Average DL throughput in the macro layer.	132
Figure 118: Average DL throughput in the femto layer.	133
Figure 119: Total DL macro traffic.	133
Figure 120: Total DL femto traffic.	134
Figure 121: Achievement rate for macro UEs.....	134
Figure 122: Achievement rate for femto UEs.	135
Figure 123: Achievement rate for UEs in the ON.....	135
Figure 124: Number of disconnected macro users.....	136
Figure 125: Number of users transferred to the ON.	137
Figure 126: Average DL throughput in the macro layer.	137
Figure 127: Average DL throughput in the femto layer.....	138
Figure 128: Total DL macro traffic.	138
Figure 129: Total DL femto traffic.....	139
Figure 130: Achievement rate for macro UEs.....	139
Figure 131: Achievement rate for macro UEs in traffic mixes A, B and C.....	140
Figure 132: Achievement rate for femto UEs.	141
Figure 133: Achievement rate for femto UEs in traffic mixes A, B and C.	142
Figure 134: Achievement rate for UEs in the ON.....	143
Figure 135: Achievement rate for UEs in the ON for traffic mixes A, B and C.....	144
Figure 136: General framework for spectrum selection.....	146
Figure 137: Achievable bit rates (Mbit/s) with the different configurations (a) Pool #1, (b) Pool #1 in the presence of an external interference, (c) Pool #2, (d) Pool #3.....	148
Figure 138: Location of the receivers for the two considered links	148
Figure 139: Global efficiency level for the considered approaches.....	151
Figure 140: Regret of link #2 in using pool #3 for the considered approaches	152
Figure 141: Spectrum HO rate	153
Figure 142: Sensitivity to R in terms of (a) Global satisfaction level, (b) Regret of link #2 in using pool #3, (c) Spectrum HO rate per session	154
Figure 143: Strategy dynamics with different initial value of and 	155
Figure 144: Considered DH environment for evaluating capability to adapt to non-stationary environments	155
Figure 145: Evolution of dissatisfaction probability in the presence of Change 1	157
Figure 146: Evolution of the SpHO rate in the presence of Change 1	158

Figure 147: Evolution of the relative regret $Usage(2,3) \times Regret(2,3)$ for pool #3 in the presence of Change 1	158
Figure 148: Evolution of dissatisfaction probability in the presence of Change 2	159
Figure 149: Evolution of the SpHO rate in the presence of Change 2	159
Figure 150: Evolution of the relative regret $Usage(2,3) \times Regret(2,3)$ for pool #3 in the presence of Change 2	159
Figure 151: RAT selection for different data types.	160
Figure 152: Comparison of the accumulated utilities corresponding to different channel selection schemes.	161
Figure 153: Mapping of “selection of nodes and routes” relevant algorithms to ON phases.....	162
Figure 154: Mapping of “selection of nodes and routes” relevant algorithms to ON phases.....	163
Figure 155: Synergic flowchart of “selection of nodes and routes” relevant algorithms.....	163
Figure 156: Synergic inputs/outputs of “selection of nodes and routes” relevant algorithms.....	164
Figure 157: Mean power of congested BS	165
Figure 158: Mean power of non-congested BSs.....	166
Figure 159: Mean power of edge terminals (terminals that switch to ON).....	166
Figure 160: Network average delay before and after solution enforcement.....	167
Figure 161: Opportunistic Coverage Extension	169
Figure 162: Mapping of the Algorithms to the different phases for Scenario 1.....	169
Figure 163 – Example algorithm execution sequence for the OneFIT scenario 1	170
Figure 164: Scenario 1 - Mapping of algorithms for ON Suitability phase.....	170
Figure 165: Scenario 1 - Mapping of algorithms for ON Creation phase.....	171
Figure 166: Scenario 1 - Mapping of algorithms for ON Maintenance phase	172
Figure 167: (a) SNR (dB) with the direct link, (b) area (in green) where the direct communication can be established	174
Figure 168: Probability of having direct coverage as a function of the distance D to the infrastructure transmitter with and without shadowing.....	174
Figure 169: Coverage probability as a function of distance D for the different spectrum band alternatives, with and without shadowing.	175
Figure 170: Average value of the surface S_i with direct coverage that can be reached through an ON link for the different options.....	175
Figure 171: Coverage probability as a function of the node density η at distance $D=4\text{km}$ when shadowing is considered.....	176
Figure 172: Indicative duration of connections by taking into account the number of accepted ON nodes.....	176
Figure 173: Percentage of aborted messages by taking into account the number of accepted ON nodes.....	177
Figure 174: Opportunistic Capacity Extension	177
Figure 175: Mapping of algorithms to ON phases in Scenario 2	178
Figure 176: Scenario 2: Flow diagram for ON Suitability phase.....	179
Figure 177: Scenario 2: Flow diagram for ON Creation phase.....	180
Figure 178: Scenario 2: Flow diagram for ON Maintenance phase	181
Figure 179: Scenario 2: Flow diagram for ON Termination phase.....	182

Figure 180: Power allocation to femtocells with respect to the fittingness factor of the available spectrum	183
Figure 181: Infrastructure supported opportunistic device-to-device networking.....	184
Figure 182: Mapping of the Algorithms to the different phases for Scenario 3.....	184
Figure 183: Scenario 3: Mapping of the Algorithms to the Suitability phase and to the Creation Phase	185
Figure 184: Location of the receivers for the two considered links	186
Figure 185: Global efficiency as a function of the offered traffic.....	188
Figure 186: Dissatisfaction probability of link #2 as a function of the offered traffic	188
Figure 187: Relative regret of link#2 when using pool #3 as a function of the offered traffic.....	189
Figure 188: Spectrum HO rate	189
Figure 189 – Mapping of the OneFIT algorithms onto scenario 5 challenges	190
Figure 190 – Example algorithm execution sequence for the OneFIT scenario 5	191
Figure 191 – Trigger detection and initial decision making	192
Figure 192 – Selection of nodes and routes in suitability determination phase for OneFIT scenario 5	192
Figure 193 – Selection of nodes, routes and spectrum in ON creation phase of the OneFIT scenario 5	193
Figure 194 – Maintenance trigger recognition and decision making	193
Figure 195 - Selection of nodes, routes and spectrum in ON maintenance phase of the OneFIT scenario 5.....	194
Figure 196 – Recognition of ON termination triggers and termination execution.....	194
Figure 197: Backhaul links considered in the evaluation.....	196
Figure 198: Global efficiency.....	198
Figure 199: Dissatisfaction probability for link #1	199
Figure 200: Relative regret in using pool #3 by link#1 (defined as $Usage(1,3) \times Regret(1,3)$).....	199
Figure 201: Rate of spectrum handovers per session.....	200

List of Tables

Table 1: Simulation Parameters.....	40
Table 2: Initial strategy profiles $\{\Pi_0\}_1$ of selecting action $a_{j,t}=1$, for the gradient follower (GF) learning strategy	45
Table 3: Comparisons learning strategies: the gradient follower (GF), the modified Roth-Erev (MRE) and the modified Bush and Mosteller (MBM)	45
Table 4: Simulation parameters.....	48
Table5: Studied cases.....	56
Table 6: Considered cases.....	106
Table 7 : Grouped cases	108
Table 8: Baseline test results	114
Table 9: Test case 1 for Scenario 1 results.....	116
Table 10: Test case 2 for Scenario 1 results.....	117
Table 11: Test case 3 for Scenario 1 results.....	118
Table 12: Average achievable bit rates (Mbit/s) in the different pools.....	149
Table 13: Average achievable bit rates (Mbit/s) in the different pools.....	156
Table 14: Studied cases.....	165
Table 15: Mapping of Algorithms to Scenarios and Phases.....	168
Table 16: Scenario parameters.	173
Table 17: Input traffic (Mbit/s) for the different APs	196
Table 18: Required configurations for additional links #1 and #2	196
Table 19: Average achievable bit rates (Mbit/s) in the different pools.....	197

1. Introduction

This deliverable describes final set of algorithmic solutions for enabling opportunistic networks developed in the OneFIT project [1]. Each algorithm is accompanied by performance evaluations, according to proper KPIs in representative scenarios. Performance assessment also includes aspects related to the practicality of the proposed solutions.

The deliverable is organized in three main chapters. Chapter 2 presents further results from the different enabling algorithms that have been developed to cope with the technical challenges arising in the ON management stages. Two algorithms in particular address the ON suitability determination challenge in different scenarios. Then the chapter focuses on algorithms that based on different approaches, specifically address the spectrum opportunity identification and selection technical challenge. Next the node and route selection challenge is addressed, including also different strategies that consider specific problems within this technical challenge. Finally, specific solutions are presented to the capacity extension problem by means of femto-cells and to the macro-to-femto offloading mechanism.

Chapter 3 presents fresh results that highlight benefits of adopted approaches i.e. utility functions and knowledge management, machine-learning and multi-objective formulations, to spectrum selection challenge and cognitive multipath routing, network-coding for multi-flow route selection, dynamic resource allocation and direct D2D communication for node and route selection.

The dynamic operation of ONs based on integration and synergic operation of algorithms is addressed in chapter 4. On a per scenario basis, mapping of algorithms to different ON phases is shown as well as examples of execution sequences. For selected set of scenarios and cases, based on the identified execution sequence, further results are provided validating the approach adopted.

2. Algorithmic Solutions (final version)

2.1 Suitability determination for direct device-to-device (D2D) communication

2.1.1 Problem formulation and algorithm concept

This algorithm is for the case that a user sends a request to the infrastructure to establish a direct device-to-device (D2D) session as described in the OneFIT scenario 3 on Infrastructure supported device-to-device (D2D) communication [2][3]. This algorithm can also be used by the infrastructure to modify and optimize an existing session where the traffic is routed via the infrastructure towards a direct D2D connection. Figure 1 shows such a situation where the user data is directly exchanged between the two devices while there is still a control path from each device to the infrastructure to enable the infrastructure to provide further guidance and control to the D2D connection..

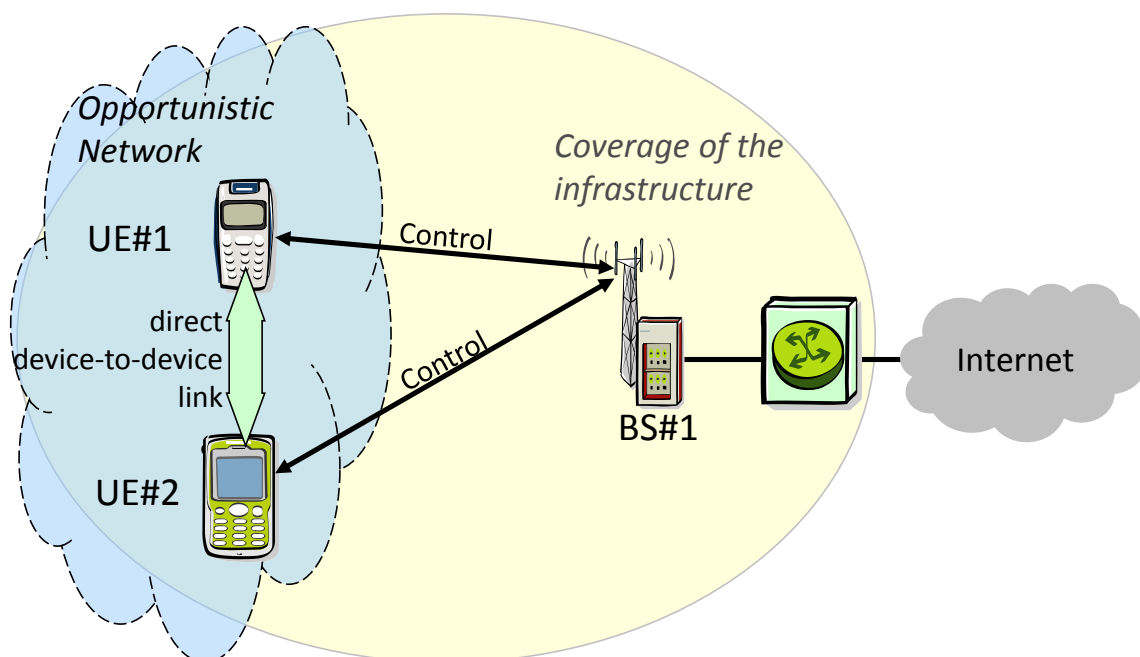


Figure 1: Infrastructure supported opportunistic networking via direct device-to-device communication

2.1.2 Algorithm specification

Figure 2 shows the Flow-chart of the principle procedure to be executed on infrastructure side to determine if a direct device-to-device link is feasible. If a connection setup request is received, either indicating that a direct D2D connection is requested, or, if the infrastructure initiates suitability determination based on other criteria as described e.g. in D2.2.2 [4], different decisions must be made. First, policies must be checked if opportunistic networking (direct D2D) is allowed and the service requirements must be checked if the service can benefit from an opportunistic network. For example, for a video conference, the users can benefit from a possibly good performance of the direct D2D link and the infrastructure network benefits from the traffic offloading. For other services, e.g. the exchange of an SMS, an ON shall not be created because the overhead for establishing an ON would be much higher than any savings for the traffic offloading.

Second, it must be checked if both endpoints support direct D2D links. In the last step it must be reviewed if both endpoints are attached to the same network. At last it must be controlled if the two endpoints are in the geographical neighbourhood.

If these conditions are met, then the direct link is seen as feasible and a discovery process shall be started and the connection is established. If the direct D2D connection setup of the two devices is successful, then an ON with direct D2D shall be established, else the ON with direct D2D is not feasible.

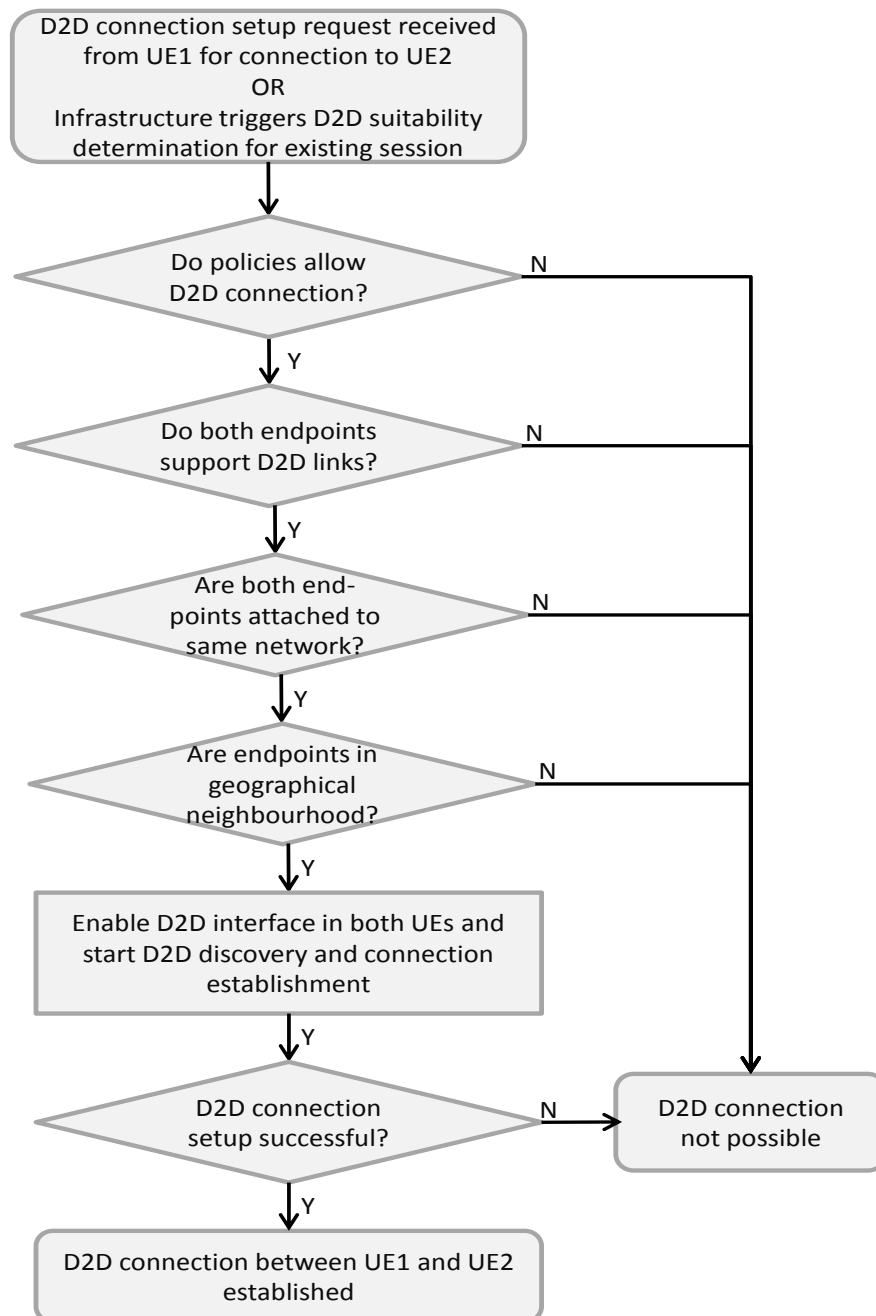


Figure 2: Suitability Determination for an Infrastructure supported direct device-to-device communication

When the direct D2D connection is longer needed (e.g. because the user terminates the D2D communication sessions), then the ON can be released.

Further on, as also illustrated in Figure 3, it may also be possible that the D2D can no longer be maintained because e.g. the devices are moving out of proximity of each other. In that case, a handover towards the infrastructure shall be executed so that that the traffic is then routed via the infrastructure.

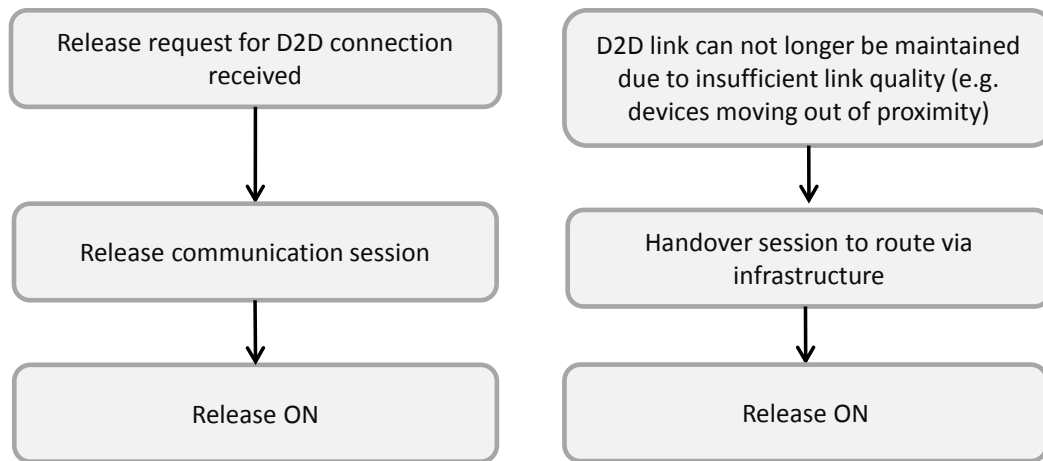


Figure 3: Terminating the direct device-to-device communication upon certain events

2.1.3 Integration in OneFIT architecture

This algorithm runs in the CMON on the infrastructure side.

2.1.4 Further performance evaluation results

The probability for being able to have a direct D2D communication between two devices is evaluated in section 3.2 of D4.2 [7] in dependency of the distribution and the density of the users as well as independency of the range of the direct device-to-device interface. This algorithm has been verified with the opportunistic networking demonstrator as described in OneFIT Deliverable D5.2 [12] section 2.2. An example on the throughput and coverage gains for coverage extension scenario achieved with the demonstrator is described in the OneFIT Deliverable D5.3 [13].

2.2 Suitability determination for the coverage extension scenario

2.2.1 Problem formulation and algorithm concept

This section describes an algorithm for a scenario where a UE is first inside the coverage of an infrastructure network but is then moving outside the coverage. This scenario is shown in Figure 4 and described in more detail as scenario 1 “opportunistic coverage extension” in [2][3].

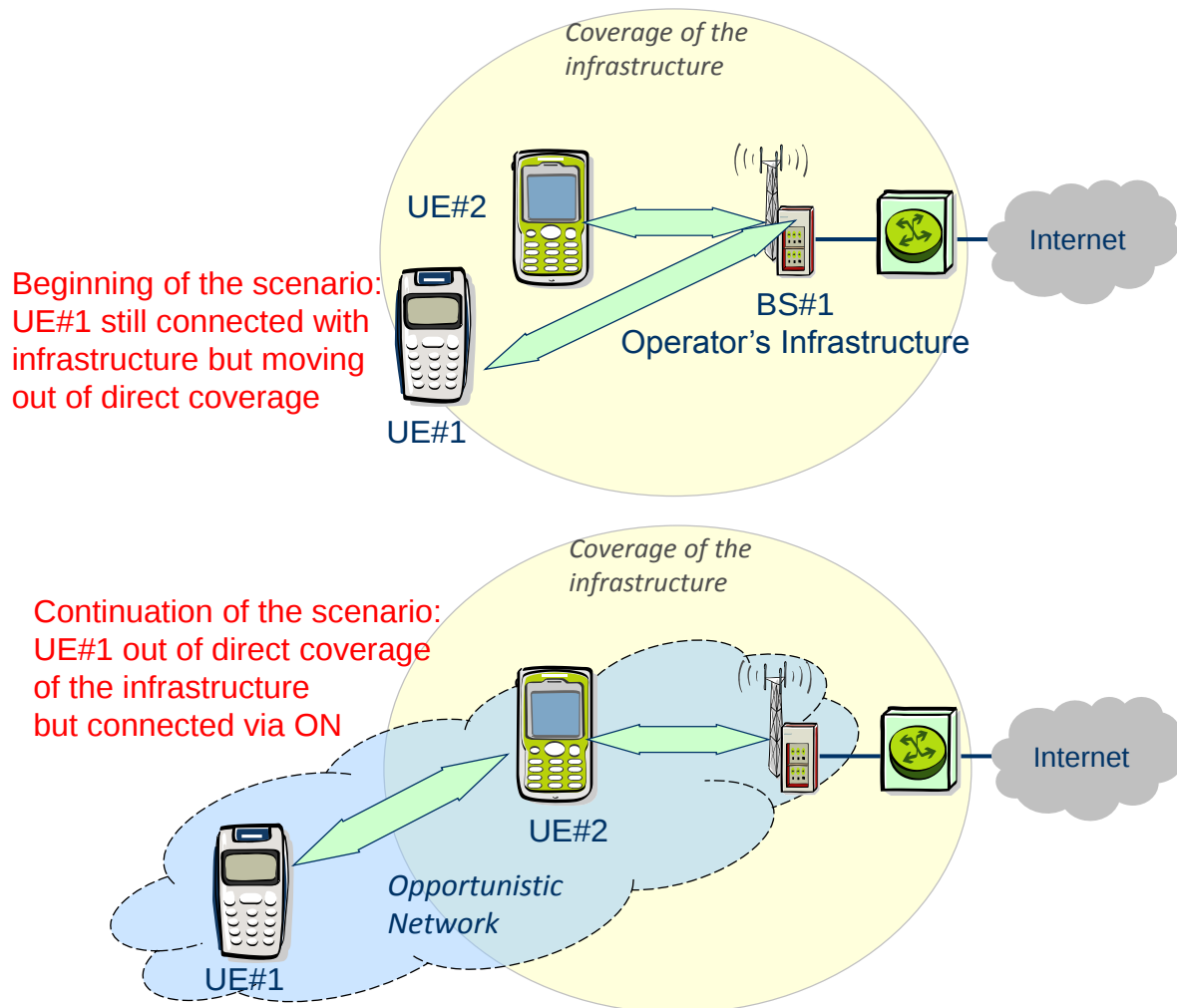


Figure 4: Opportunistic coverage extension when a UE is going out of infrastructure coverage

2.2.2 Algorithm specification

Figure 5 shows the flowchart for the procedure to be executed on infrastructure side when a device is going out of infrastructure coverage. When a terminal is moving out of the infrastructure coverage, the radio conditions are getting worse. If one or more thresholds are crossed (e.g. link quality or signal strength going below a threshold), then the Radio Resource Management (RRM) is informed about this event. As with traditional RRM procedures, it is checked if a handover to another cell of the infrastructure is possible. If so, then the handover will be executed.

If such a handover to another cell is not possible, then the suitability determination to find a supporting device for an ON creation will be started. The algorithm searches through the list of devices attached to the current cell or neighbour cells for at least one device in the geographical neighbourhood which supports opportunistic networking and which has the capabilities to provide a direct D2D link to the terminal going out of coverage.

If several candidate supporting devices are found in the neighbourhood, then a ranking of the devices must be made to decide on the most suitable device. This can be made using e.g. the fitness as described in section 2.7 "Knowledge-based suitability determination and selection of nodes and route" of this document.

In the next step, the D2D interface for the selected supporting device is activated (if not already active) and direct D2D discovery procedures are initiated. If the devices discover each other and the

direct D2D link is estimated to provide sufficient quality, then an ON can be created to provide coverage extension. If the link establishment with sufficient quality cannot be established, then the ranking of devices is updated. In the following a search for another supporting device in the neighbourhood is initiated and the procedure is repeated until a coverage extension can be provided or a decision is made that a coverage extension cannot be provided.

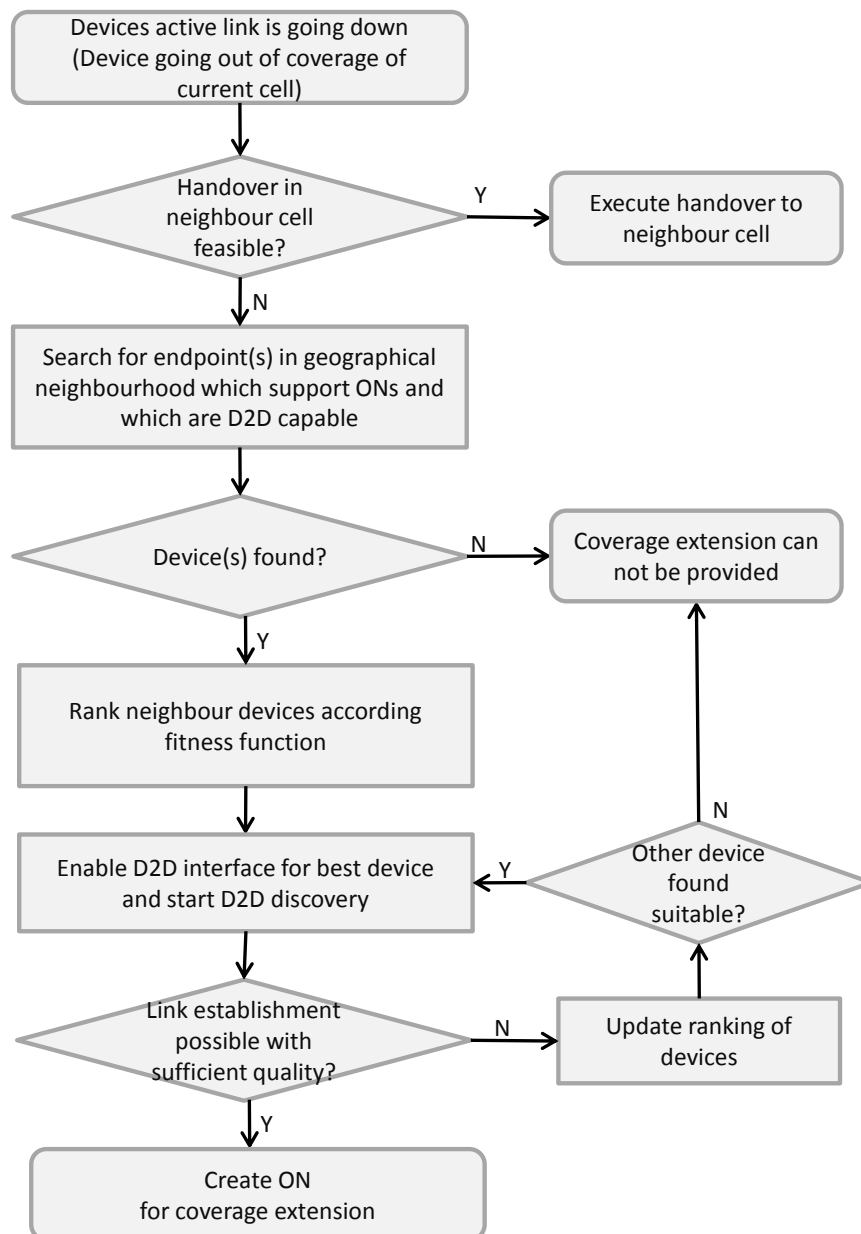


Figure 5: Suitability Determination for the coverage extension scenario

After ON creation, the device being out of direct infrastructure coverage still executes infrastructure discovery procedures. In the case that this procedure detects that the device is now in coverage of the infrastructure, a handover to the infrastructure can be performed and the ON is released as illustrated in Figure 6.

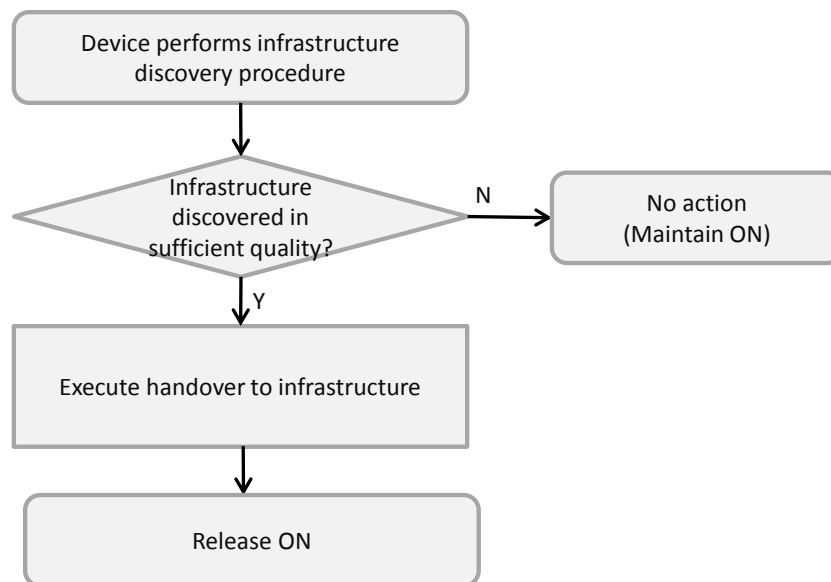


Figure 6: Termination of the ON when the device earlier being out of coverage gets into coverage

2.2.3 Integration in OneFIT architecture

This algorithm runs on the infrastructure side. The first steps where triggers like “active link is going down” are evaluated and where a check is made if a handover to another cell is feasible are executed in the legacy Radio Resource Management (RRM) or in the Joint Radio Resource Management (JRRM). In the case a handover is not feasible, the CSCI is triggered to search of candidate endpoints in the neighbourhood and to rank the devices according their suitability. The creation of the ON will then be executed by the CMON.

2.2.4 Further performance evaluation results

The probability of being able to solve an out of cellular coverage scenario with an opportunistic coverage extension in dependency of the number of users supporting opportunistic networking in that area and in dependency of the coverage range of the direct device-to-device radio interface is evaluated in section 3.3 of D4.2 [7].

Further on, D5.3 [13] describes example capacity extension results achieved with the opportunistic networking demonstrator.

2.3 Modular decision flow approach for selecting frequency, bandwidth and radio access technique for ONs

2.3.1 Problem formulation and algorithm concept

This algorithm is used to solve operator network congestion or coverage problems by selecting band, RAT, and bandwidth for creating a short range communication link between different users. The algorithm makes the decision taking into account users’ capabilities, velocity and application requirements. More detailed description of problem formulation can be found in [7] in section 2.3.1.

2.3.2 Algorithm specification

In [7] section 2.3.2 detailed description of algorithm is provided. A new feature is added to algorithm: selected band is divided into channels and reciprocity based reinforcement learning is used in channel selection. We also use channel selection to decrease delay so that during maintenance phase the channel is changed rather inside same band than between different bands. Figure 7 shows improved flow chart of modular decision flow.

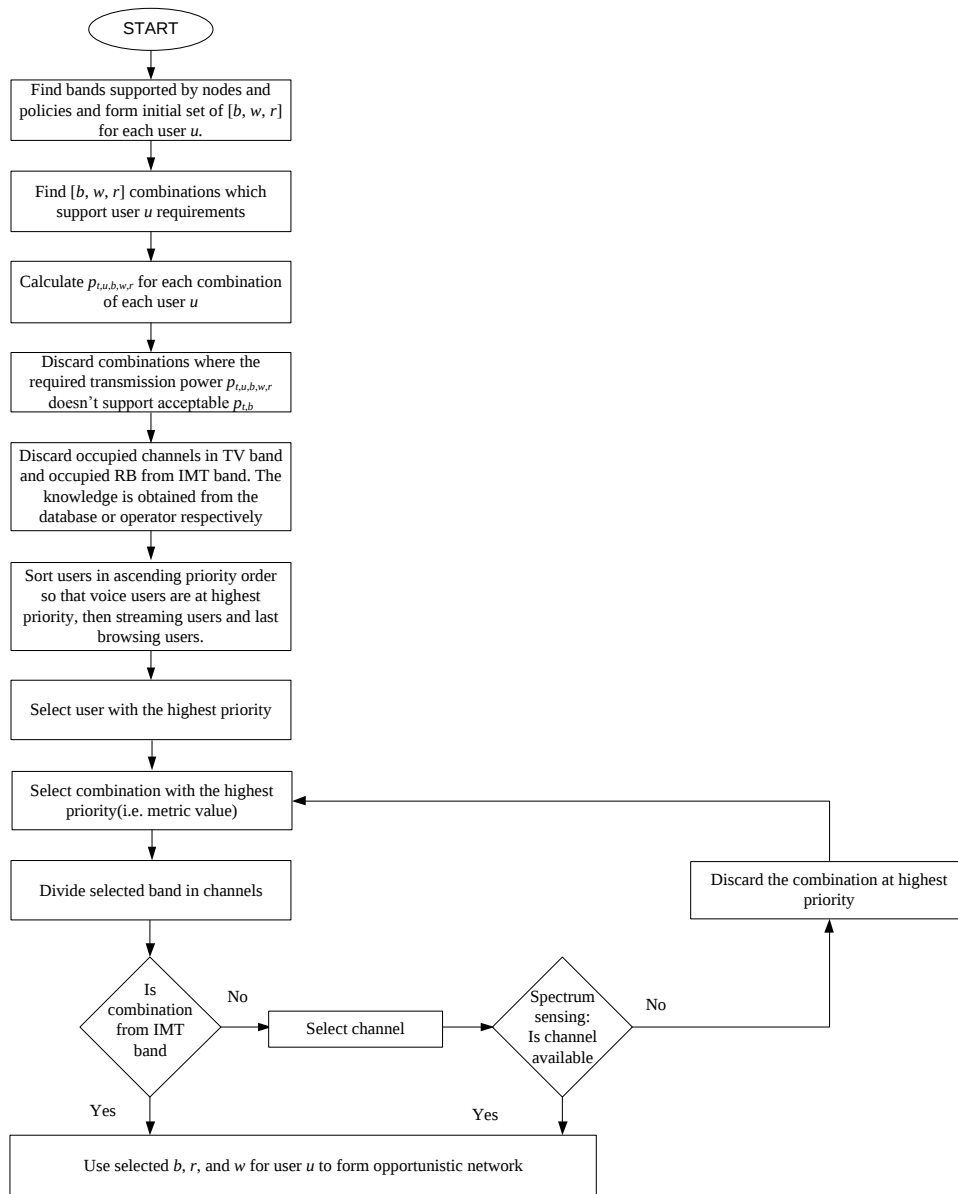


Figure 7: Modular decision flow for band, RAT, channel, and bandwidth selection.

The new learning based channel selection algorithm goes as follows:

Define a strategy $\mathbf{p}_i^t$ to be a probability distribution over the set of channels N , $\mathbf{p}_i^t = (p_i^t(1), \dots, p_i^t(N))$, where $p_i^t(n)$ means the probability with which SU i selects channel n at time slot t . The objective of each SU i is to learn the channel selection strategy $\mathbf{p}_i^t$ at every time slot t that maximizes its expected utility, which is given by

$$U_i(\mathbf{p}_i^t, \mathbf{p}_{-i}^t) = \sum_{n \in N} \theta_n p_i^t(n) \prod_{j \in I \setminus \{i\}} (1 - p_j^t(n)), \quad (1)$$

where θ_n is the idle probability of channel n , and $\mathbf{p}_{-i}^t = (p_1^t, \dots, p_{i-1}^t, p_{i+1}^t, \dots, p_I^t)$.

Nash equilibrium (NE) is one solution for the non-cooperative game. To encourage the possible cooperation among the SUs, we propose to use the conjecture concept which is based on 'reciprocity' in economic. Specifically, each SU can estimate $b_i^t(n) = \prod_{j \in I \setminus \{i\}} (1 - p_j^t(n))$ using

$$\tilde{b}_i^t(n) = b_i^{t-1}(n) - \delta_{i,n} (p_i^t(n) - p_i^{t-1}(n)), \quad (2)$$

with belief factor $\delta_{i,n} > 0$. Substituting (2) into (1), we have

$$U_i(p_i^t, \tilde{b}_i^t) = \sum_{n \in N} \theta_n p_i^t(n) \tilde{b}_i^t(n). \quad (3)$$

Based on (3), we propose a best response learning algorithm, where the best response strategy is

$$p_i^t(n) = \begin{cases} \left[\frac{p_i^{t-1}(n)}{2} + \frac{1}{\delta_{i,n}} \left(b_i^{t-1}(n) + \frac{\gamma_i^t}{\theta_n} \right) \right]_0^1, & \text{if } \theta_n > 0; \\ 0, & \text{otherwise,} \end{cases} \quad (4)$$

where γ_i^t is a parameter that satisfies $\sum_{n \in N} p_i^t(n) = 1$. Herein, $[x]_a^b$ denotes the Euclidean projection of x on the interval $[a, b]$. Detailed description of the proposed algorithm is as follows.

Algorithm: A best response learning algorithm for channel selection

Initialization:

$t=1$, initialize the channel selection strategies $p_i^t(n)$ and the belief factor $\delta_{i,n}$ for all users and all channels.

Learning:

- (1) Set $t \leftarrow t + 1$.
- (2) For all users and all channels, do (4).
- (3) Each SU i selects channel n at time slot t with probability $p_i^t(n)$.

End Learning

2.3.3 Integration in OneFIT architecture

In [7] section 2.3.3 the algorithm mapping onto OneFIT functional architecture is described in detail. The decision on used spectrum, bandwidth and RAT is done in CMON side in decision making block where the output parameters are sent to corresponding nodes to form link between ON nodes.

2.3.4 Further performance evaluation results

In simulations five TV channels within the TV band, each having bandwidth of 8 MHz, one 2.4 GHz band with 20 MHz bandwidth and one 60 GHz band with 100 MHz bandwidth are considered. And as an operator own band we consider one IMT band with 20 MHz bandwidth. Used RATs are IEEE 802.11a with subcarrier spacing 312.5 kHz, IEEE 802.15.3c with subcarrier spacing 1.5625 MHz, and LTE with subcarrier spacing 15 kHz.

Simulations have been made to three traffic types which each have different data rate and delay requirements. We have used single carrier frequency division multiplexing (SC-FDMA) and we consider only contiguous resource blocks per user. We reserve different amount of resource blocks for different data type users when using LTE. In case of IEEE 802.11a and IEEE 802.15.3c one subcarrier per user is reserved for the sake of simplicity. Simulations have been carried out for voice users which demand at least data rate of 60 kbps to maintain adequate user experience and reserve 3 resource blocks when using LTE. Streaming users require at least data rate of 384 kbps and 5 resource blocks when using LTE. Browsing users require 13 kbps data rate and 2 resource block requirement when using LTE.

Simulations have been run in a scenario where 27 browsing users, 27 streaming users, and 27 voice users exists and their velocity is randomly selected to be either 0 m/s, 1.11 m/s, and 5.55 m/s. Link range varies randomly from 1 m up to 30 m. In addition on each band there exist also users not

involved with ON operation. Such users in TV band are primary users like TV broadcasting and wireless microphones. In ISM band they are users with same kind of priority as ON users. In IMT bands they are other users within the operator's network. We model traffic caused by users not involved with ON using birth death process where death rate is α and birth rate is β . α and β are uniformly distributed between 0.1 and 0.5. D_{shift} is on our simulations 1.5 ms when changing channels between different bands and $2k|c_{new} - c_{old}|$ when channel is changed inside one band. k is constant [8], c_{new} is selected new channel and c_{old} is previously used channel.

In Figure 8 the bandwidth usage for each data type is shown and it can be seen that streaming users favour TV band and IMT band while usage in 60 GHz band is quite rare due to limitations set by the transmission range. On the other hand 60 GHz band provides high data rates whenever its usage is sufficient and we can see that all data types utilize the band. The usage drops on 60 GHz band refer to situations where transmission range is high causing severe path loss and band usage is not sufficient. Since browsing users have least demanding data rate and delay requirements their usage is directed to ISM more often to offload traffic from operator own band (IMT).

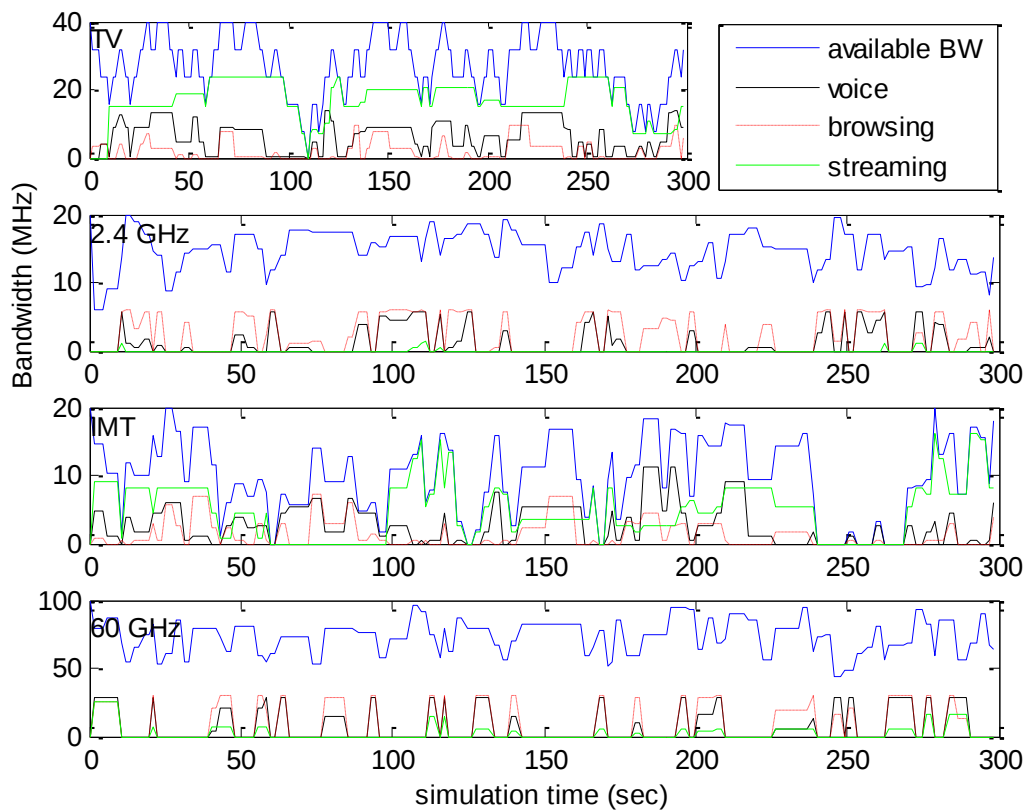


Figure 8: Band selection for different traffic types.

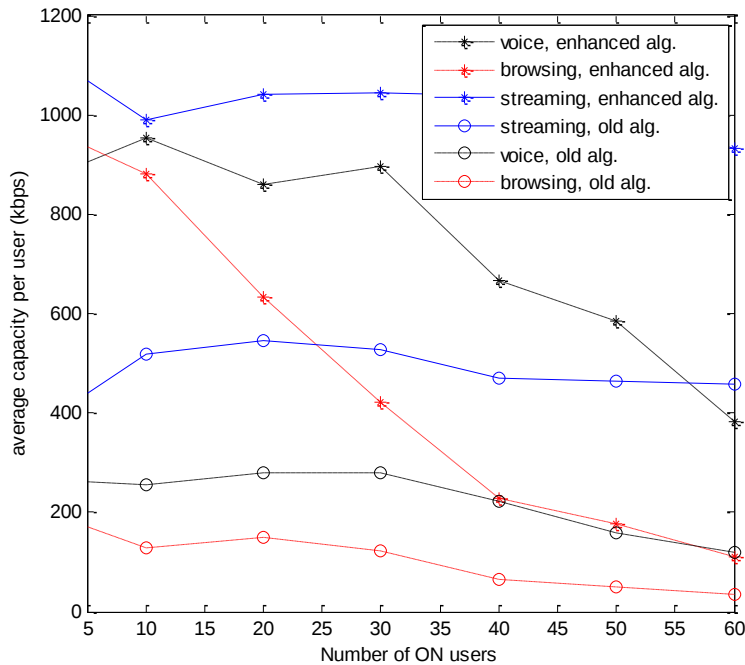


Figure 9: Capacity of different traffic type users using enhanced algorithm and former algorithm.

For Figure 9 the simulation setting has changed so that the number of ON users varies from 5 to 60 users, and link length is fixed to 10 m. We have calculated the average capacity per user with a normalized capacity model introduced in [9]. We assume perfect sensing and thus the channel idle time can be expressed as $1/\beta$. Figure 9 compares the average user capacity when using algorithm introduced in [10] and the new enhanced algorithm where channel selection aspects are also taken into account. Clear improvements in the performance levels of each data type user can be seen when the new algorithm is used.

2.4 Fittingness factor-based spectrum selection

2.4.1 Problem formulation and algorithm concept

The problem considered here is the selection of the spectrum to be assigned to a set of radio links between pairs of terminals and/or infrastructure nodes. Each radio link belongs to one ON, and one ON can be composed of one or several radio links. The purpose of each radio link is to support a certain CR application with certain bit rate requirements. The available spectrum is modelled as a set of spectrum blocks (denoted here as “pools”), each one with a given bandwidth, that can belong to different spectrum bands subject to different regulatory constraints. Based on radio link requirements and spectrum pool characteristics, the general aim is to efficiently assign a suitable spectrum pool for each of the L radio links. For further details on this problem formulation the reader is referred to section 2.4.1 of deliverable D4.2 [7].

2.4.2 Algorithm specification

For details on the algorithm specification the reader is referred to section 2.4.2 of deliverable D4.2 [7]. Specifically, the algorithm is based on the fittingness factor $F_{l,p}$ concept as a metric that captures how suitable a given spectrum pool p is for a certain radio link l depending on the bit rate that can be achieved in this pool in relation to the bit rate required by the link.

In the following, details are given only for the developed algorithm extensions to cope, on the one hand, with non-stationary environments, and, on the other hand, to introduce different operator

preferences in the spectrum pools that can be assigned to each radio link, reflecting e.g. different operator policies or different regulatory constraints related to the band that each specific pool belongs to.

The modified functional architecture of the algorithm to deal with non-stationary environments is presented in Figure 10. The main difference with respect to the previous architecture presented in D4.2 [7] is the inclusion of the so-called Reliability Tester (RT) component. Its objective is to detect relevant changes in the statistics stored in the Knowledge Database (KD) that may occur whenever the environment is non-stationary (e.g. when significant changes in positions of the radio link receivers occur, or when the behaviour of the interference sources is modified, etc.). When these changes are detected, the statistics of the KD are regenerated.

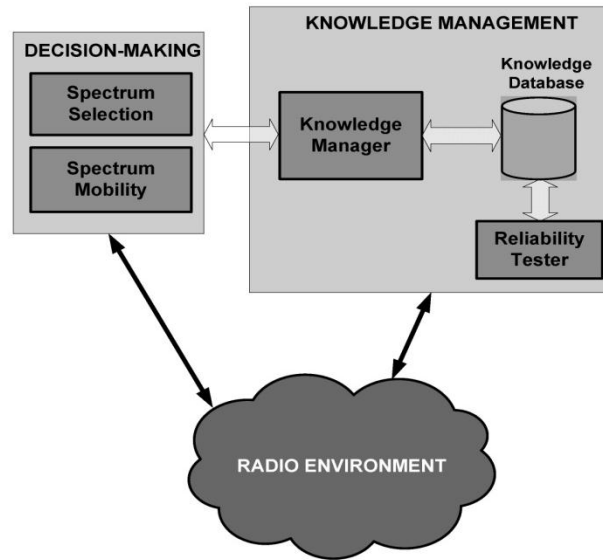


Figure 10: Functional architecture of the proposed Fittingness Factor-based Spectrum Management Framework

2.4.2.1 Reliability tester specification

In order to detect changes in the radio and interference conditions of the different spectrum pools, the RT monitors the reliability of the stored KD data. Whenever a relevant change is detected, the KD statistics are regenerated under the new conditions.

The RT detects potential changes in pools currently assigned to active link sessions based on monitoring a set of Key Performance Indicator (KPIs) of interest from the radio environment. To carry out this procedure, the RT considers first an initial estimate for each KPI x computed based on a sequential analysis of its observed mean over a number of samples obtained from the established link sessions (let denote this initial mean estimate as $\bar{x}$). The sequential procedure gradually increases the number of samples S as new link sessions are established until the following stopping rule is met:

$$\bar{x}_{\max} - \bar{x}_{\min} < \beta \bar{x} \quad (5)$$

where $0 < \beta < 1$ and $[\bar{x}_{\min}, \bar{x}_{\max}]$ denotes the γ -confidence interval of $\bar{x}$. The limits of this confidence interval are updated each time a new sample is available based on the estimates of the mean and the variance of the KPI. After meeting the stopping rule the initial estimate $\bar{x}$ and the corresponding value of the window size S are kept unchanged. Then, the RT keeps comparing, for each KPI, its initial estimate ($\bar{x}$), to a new estimate ($\tilde{x}$) and its corresponding γ -confidence interval $[\tilde{x}_{\min}, \tilde{x}_{\max}]$ calculated online over a sliding window of the same size S based on hypothesis testing [24]. Specifically, the RT infers whether the null hypothesis (H_0), meaning that no relevant changes in the KD statistics have occurred, or its alternative (H_1) holds, meaning in the latter case that relevant changes in the radio

environment have been detected and thus the KD statistics are no longer valid.

The proposed hypothesis testing procedure is based on the following tests:

- Test for $H_0: SUP_{min} \leq INF_{max}$

- Test for $H_1: SUP_{min} > INF_{max}$

Where $[SUP_{min}, SUP_{max}]$ is the interval, either $[\bar{x}_{min}, \bar{x}_{max}]$ or $[\tilde{x}_{min}, \tilde{x}_{max}]$ with the highest maximum value, that is:

$$[SUP_{min}, SUP_{max}] = \begin{cases} [\bar{x}_{min}, \bar{x}_{max}] & \text{if } \bar{x}_{max} \geq \tilde{x}_{max} \\ [\tilde{x}_{min}, \tilde{x}_{max}] & \text{otherwise} \end{cases} \quad (6)$$

Correspondingly, $[INF_{min}, INF_{max}]$ is the interval with the lowest maximum value, that is:

$$[INF_{min}, INF_{max}] = \{[\bar{x}_{min}, \bar{x}_{max}], [\tilde{x}_{min}, \tilde{x}_{max}]\} \setminus [SUP_{min}, SUP_{max}] \quad (7)$$

The intuition behind the test is that if KPI x has already achieved a good level of convergence by the time the initial estimate $\bar{x}$ was obtained and provided that no subsequent relevant changes occur, the confidence interval of any new estimate $\tilde{x}$ is likely to overlap that of the initial estimate $\bar{x}$. If such condition does not hold, H_1 is selected. Otherwise, H_0 is selected. It is worth mentioning that the above conditions for the hypothesis testing have been formulated with the general objective to minimize the risk of detecting false changes of the KD statistics (the so-called Type I error in hypothesis testing terminology [25]), while keeping the error of not detecting a real change (the so-called type II error) at a reduced level. This reduces the risk of useless regeneration of the KD.

Finally, the obtained hypothesis testing results for each of the considered KPIs are combined to decide about the reliability of the whole KD data. Specifically, if H_1 is selected for any of the considered KPIs, the whole KD data is judged as unreliable, so it is regenerated and all the initial estimates for the different KPIs are recalculated under the new conditions. Otherwise, the KD data is kept unchanged.

Note that the above procedure is adequate to detect changes that occur in the pools that are regularly used by the system. However, it is not valid to detect changes occurring in pools that are never assigned (e.g. this may happen if the actual KD statistics reveal that a pool always has low fittingness factor $F_{l,p}$ values). To overcome this issue, a specific procedure is carried out for each link/pool that remains inactive for more than T_{Inact} . When this occurs, a total of N forced measurements are performed to obtain the value of the actual $F_{l,p}$ that is compared against the estimate $\hat{F}_{l,p}$ stored in the KM to detect possible changes.

2.4.2.2 Modification of the selection criterion to incorporate operator preferences

An additional input parameter has been introduced in the spectrum selection criterion to capture the fact that the different spectrum pools can belong to different bands and thus be subject to different regulatory constraints, leading to different preferences in using them. More specifically, the following new algorithm input parameter has been introduced:

- $\lambda_{l,p}$: It is a preference factor reflecting the operator preference regarding the use of spectrum pool p for link l . It takes values in the range $\lambda_{l,p} \in (0,1]$ and the higher the value the higher the preference.

The definition of the fittingness factor, the associated statistics stored in the knowledge database and the Knowledge Manager (KM) remain exactly the same as in deliverable D4.2 (see sections 2.4.2.1, 2.4.2.2 and Algorithm 1 for the KM).

In turn, with the inclusion of the operator preferences, the spectrum selection criterion for link l is defined as follows:

$$p^*(l) = \arg \max_{p \in Av_Pools} \left(g(\hat{F}_{l,p}) \right) \quad (8)$$

$\hat{F}_{l,p}$ is the estimated fittingness factor for pool p and link l , provided by the KM and Av_Pools is the set of available pools. In turn, function $g(\hat{F}_{l,p})$ jointly considers the current preference factor ($\lambda_{l,p}$) and the average preference factor observed while in the HIGH state of the fittingness factor up to the end of link session $T_{req,l}$.

$$g(\hat{F}_{l,p}) = \begin{cases} \lambda_{l,p} + \frac{1}{T_{req,l}} \sum_{k=1}^{T_{req,l}} P_{H,H}^{l,p} (\Delta k_{l,p} + k, \delta_{l,p}) \times \lambda_{l,p} & \text{if } \hat{F}_{l,p} \text{ is HIGH} \\ 0 & \text{if } \hat{F}_{l,p} \text{ is LOW} \end{cases} \quad (9)$$

Based on the above criterion, the associated spectrum selection and spectrum mobility algorithms of the decision-making process are described in the following.

In the Spectrum Selection (SS) algorithm (see Algorithm 1), based on the fittingness factor values determined by the KM, a suitable spectrum pool is allocated for each new radio link that is activated (i.e. at the ON creation stage). Upon receiving a request for establishing link l , the request is rejected if the set of available pools (Av_Pools) is empty (line 3). Otherwise, an estimation of $\{\hat{F}_{l,p}\}_{1 \leq p \leq P}$ values

is obtained from the KM (line 5). Based on provided $\hat{F}_{l,p}$ values, the algorithm selects the available spectrum pool $p^*(l)$ with the largest $g(\hat{F}_{l,p})$ is performed (lines 6 and 7).

Algorithm 1: Fittingness Factor-based Spectrum Selection

```

1: if (service  $l$  request) do
2:   if (  $Av\_Pools = \emptyset$  ) do
3:     Reject request;
4:   else
5:     Get  $\{\hat{F}_{l,p}\}$  from the KM;
6:     Compute  $g(\hat{F}_{l,p})$  for all pools
7:      $p^*(l) = \arg \max_{p \in Av\_Pools} \left( g(\hat{F}_{l,p}) \right)$ 
8:   end if
9: end if

```

The Spectrum Mobility (SM) functionality pursues a highly efficient allocation of spectrum pools to active applications. Therefore, it is executed during the ON Maintenance stage whenever an event that might have influence into the spectrum selection decision-making process occurs. Such events can be (1) a spectrum pool has been released due to finalization of an application, or (2) changes in suitability of the available spectrum pools have been detected.

As detailed by Algorithm 2, the proposed algorithm first gets the current $\{\hat{F}_{l,p}\}_{1 \leq l \leq L, 1 \leq p \leq P}$ estimates from the KM. Then, it explores the list of currently active links ($Active_Links$) in decreasing order of the required throughputs ($R_{req,l}$) in order to prioritize the neediest links. The decision to reconfigure or not each active link is based on a comparison between the actually used pool ($p^*(l)$) and the currently best pool in terms of $g(\hat{F}_{l,p})$ ($new_p^*(l)$) (line 7). Specifically, if $\hat{F}_{l,p^*}$ is LOW and $\hat{F}_{l,new_p^*}$ is

HIGH, a Spectrum HandOver (SpHO) from $p^*(l)$ to $new_p^*(l)$ is performed since $new_p^*(l)$ fits better link l . The same SpHO should be performed if both $\hat{F}_{l,p^*}$ and $\hat{F}_{l,new_p^*}$ are either LOW or HIGH with $\lambda_{l,new_p^*(l)} > \lambda_{l,p^*(l)}$ since $new_p^*(l)$ is preferred for link l (line 8). Finally, the SpHO should also be performed in case $p^*(l)$ is no longer available to link l after being pre-empted during previous reassignments (line 9). Once all active links have been explored, the list of assigned pools is updated to consider all SpHOs that need to be performed as a result of the algorithm (line 16).

Algorithm 2: Fittingness Factor-based Spectrum Mobility

```

1: if (service  $\hat{l}$  ends OR change in any active  $F_{l,p}$ ) do
2: Get  $\{\hat{F}_{l,p}\}$  from the KM;
3:  $New\_Assigned \leftarrow \emptyset$ ;
4: Sort  $Active\_Links$  in the decreasing order of  $R_{req,l}$ ;
5: for  $l=1$  to  $|Active\_Links|$  do
6:  $new\_p^*(l) = \arg \max_{p \in New\_Assigned} (g(\hat{F}_{l,p}))$ 
7: if  $(\hat{F}_{l,new\_p^*} \geq \delta_{l,p}) \&\& (\hat{F}_{l,p^*} < \delta_{l,p})$  OR
8:  $(\hat{F}_{l,new\_p^*} - \delta_{l,p}) \times (\hat{F}_{l,p^*} - \delta_{l,p}) > 0 \&\& (\xi_{l,new\_p^*} > \xi_{l,p^*})$  OR
9:  $((p^*(l) \in New\_Assigned) \text{ do}$ 
10:  $p^*(l) = new\_p^*(l)$ ;
11:  $New\_Assigned \leftarrow New\_Assigned \cup \{new\_p^*(l)\}$ ;
12: else
13:  $New\_Assigned \leftarrow New\_Assigned \cup \{p^*(l)\}$ ;
14: end if
15: end for
16:  $Assigned \leftarrow New\_Assigned$ ;
17: end if
18: end if

```

2.4.3 Integration in OneFIT architecture

Proposed spectrum selection framework is mapped onto the decision making and knowledge management entities at the CMON of the OneFIT architecture. For further details, including the C4MS signalling used by the proposed framework the reader is referred to section 2.4.3 of deliverable D4.2 [7].

2.4.4 Further performance evaluation results

In deliverable D4.2 [7] different evaluations of the proposed fittingness factor-based framework were presented. In particular, the robustness of the spectrum selection algorithm to the interference dynamic variations was analysed. For that purpose, the main focus of the scenario was on the temporal component without considering explicitly the spatial dimension regarding the specific positions of the existing links. This prior analysis is extended in this deliverable by considering a more realistic environment in which both the spatial and temporal dimensions are captured, by analysing the robustness to non-stationary environments and by including the effect of operator preferences for the band selection. In that respect, detailed results to analyse the impact of the different components of the framework will be presented in section 3.1 in the context of the comprehensive OneFIT solution for spectrum selection. Moreover, general performance results will also be presented in sections 4.3.5 and 4.4.5 specifically addressing OneFIT scenarios 3 and 5, respectively.

2.5 Machine Learning based Knowledge Acquisition on Spectrum Usage

2.5.1 Problem formulation and algorithm concept

Interference is one of the most limiting factors when trying to achieve high spectral efficiency in the deployment of heterogeneous networks (HNs). In this work, the HN is modelled as a layer of closed access LTE femtocells (FCs) overlaid upon a LTE radio access network. In the context of dynamic learning games, this work proposes a novel heterogeneous, multi-objective, fully distributed strategy based on reinforcement learning (RL) model (CODIPAS-HRL) for the FCs self-configuration. The self-configuration capability enables the FCs to autonomously and opportunistically sense the radio environment using different learning strategies and tune their parameters accordingly, in order to operate under restrictions of avoiding interference to both network tiers and satisfy a certain QoS requirements. The proposed model takes the advantage of calculating the learning cost associated with each learning strategy. We also derive a new accuracy metric in order to provide comparisons between the learning behaviour of different strategies. The simulation results show the convergence of the learning model that reflects the fairness and optimality of the played game between participating FCs, under the uncertainty of the HN environment. We show that Co/cross-layer interference can be reduced significantly, thus resulting in higher cell throughputs.

2.5.2 Algorithm specification

Dynamic Games Model

Let's assume that $j = \{1, \dots, J_f\}$ is the set of all active FCs, and $i = \{1, \dots, N\}$ is the number of available resource blocks. Here $\mathbf{a}_{j,t}$ is noted as the set of actions taken by the j_{th} node on all resource blocks at each time step t , defined as a binary $1 \times N$ vector $\mathbf{a}_{j,t} = \{a_{j,t}^1, \dots, a_{j,t}^N\}$, where $a_{j,t}^i \in \{0, 1\}$, $\forall i$ and $\forall j$. Action $a_{j,t}^i = \{1\}$ means that the i_{th} subchannel is available and can be used by the j_{th} FC, while $a_{j,t}^i = \{0\}$ means that the subchannel is utilised by other FCs/eNBs, and the j_{th} FC should avoid using this subchannel. Here, the former type is considered as being "ON" for opportunistic access, while the latter is designated as being "OFF" for opportunistic access. These parameters compose a class of dynamic game, given by:

$$\mathcal{G}_t(j \in J_f, \mathbf{a}_{j,t}, r_{j,t}(\mathbf{A}_{t,f}, \mathbf{A}_{t,m}, S_t)) \quad (10)$$

where G_t is the stage of the game in a discrete time space, $\mathbf{A}_{t,f}$ is the action profile: the actions of all femtocells, defined as the $N \times J_f$ matrix. $\mathbf{A}_{t,m}$ is the action profile for the macrocells, defined as the $N \times J_m$ matrix. The variable $r_{j,t}$ is the received realisation (payoff) at a time t when the j_{th} FC interact with the environment selecting action $a_{j,t}$ which depends on the state of the network S_t (i.e. network topology, FCs' locations, user distribution etc.), and the action profiles $\mathbf{A}_{t,f}$ and $\mathbf{A}_{t,m}$ for the femtocells and macrocells' network, respectively.

From a practical point of view, femtocells with opportunistic access are more often required to guarantee a certain QoS level and fulfil operation constraints rather than achieving the highest performance (i.e. global nash equilibrium (GNE), dominated strategy, etc.), which is not necessarily true for all femtocells simultaneously, as well as pursuit to the global optimum solution can be costly in terms of terminals' reconfiguration (channel switching) and energy consumption. Therefore, using the satisfaction equilibrium as a solution concept is a more relevant assumption in such game. A satisfaction equilibrium can be observed when all femtocells simultaneously satisfy their constraints and have no interest to change their strategy. In this work, we assume the following as the optimality criterion for each femtocell:

$$\mathbf{a}_j^* \in \arg \max \mathbb{E} \bar{u}_j(\mathbf{a}_j(\mathbf{a}_{-j}), \mathbf{a}_{-j}, \mathbf{A}_m, S_t) \quad (11)$$

$$\forall j \text{ and } \forall q, \mathbf{a}_j^* \in f_{q,j}(\mathbf{a}_{-j}^*) = \{\mathbf{a}_j \in \mathbf{A}_f : \mathbb{E} \bar{u}_j(\cdot) \geq \zeta_j\} \quad (12)$$

where $\mathbb{E} \bar{u}_j$ is the expected average utility model of the j th FC, q is the network's conditions, $f_{q,j}$ is a satisfaction function, and ζ_j is the minimum utility value required by femtocell j . Note that equation (4) describes the strategy's behaviour of each femtocell, while equation (5) states the satisfaction equilibrium of the played game (conditions/constraints). The argument of the optimisation of any selected FC is a function of all other FCs/eNBs' actions. Thus, in order to maximize their own received utility during the dynamic game, the main task of the active FC is to predict current/future actions of nearby FCs/eNBs over the shared spectrum.

The main assumptions for the game considered in this scenario are summarised as follows:

- The space of actions $\mathbf{a}_{j,t}$ is not necessarily known in advance by the j th FC, but will be explored and exploited by the probability $p_{j,t}^i$ at each time step t .
- At each time step, the j th FC knows only its own action $\mathbf{a}_{j,t}$, but not the action selected by other nodes $\mathbf{a}_{-j,t}$. It is assumed that the action profile $\mathbf{A}_{t,t}$ is not known by each individual FC.
- At each time step, each FC knows its received reward signal $r_{j,t}$ but not the other FCs received payoff $r_{-j,t}$.
- No cooperation is required between the two network tiers' FCs and eNBs, nor between FCs.
- It is only assumed that each FC has the knowledge of the instantaneous value of its utility function $u_{j,t}$ at each stage of the game.

For the given problem formulation, we considered following notations:

- $\omega_{j,t}^i$: defined as the propensity (probability weights) of the selected action $\mathbf{a}_{j,t}^i \in \{0, 1\}$ at the i th subchannel.
- $p_{j,t}^i$: defined as the probability distribution of the j th FC over its actions on the i th subchannel, known as the strategy, where $p_{j,t}^i \in \Delta(\mathbf{a}_{j,t}^i)$, $\forall i$, and $\Delta(\mathbf{a}_{j,t}^i) = \Delta\{0, 1\} = \{(p_{j,t}^i(0), p_{j,t}^i(1)) \in \mathbb{R}^+, \sum p_{j,t}^i = 1\}$. Here, it is assumed mixed-strategy where femtocells independently make $\forall a$ their decision on each subchannel.
- $\{\pi_{j,t}\}_a$: defined as the strategies set taken by the j th FC to select action a over all sub channels, represented by $1 \times N$ vector $\pi_{j,t} = \{p_{j,t}^1(1), \dots, p_{j,t}^N(1)\}$. For example, $\{\pi_{j,t}\}_1$ is the strategy (probability distribution) of taking action 1 for all subchannels.
- Π_t : is the strategy profile of $\forall j$, defined as $N \times J_f$ matrix.
- $u_{j,t}$: is the instantaneous value of the utility model of the j th FC, at each time step t .
- $\bar{u}_{j,t}$: is the instantaneous average utility model at time step t .

Learning Model

In this work we propose a multi-objective CODIPAS heterogeneous reinforcement learning (CODIPAS-HRL) model for femtocells' self-configuration. Under the proposed multi-objective approach, each femtocell j selects a learning strategy LS_l , $l = \{1, \dots, L\}$ to self-configure its terminals, based on its capability and target objective obj. In addition, the proposed model takes advantage of calculating the learning cost associated with each node's learning strategy. The proposed model is in the following form:

$$\begin{cases} \bar{u}_{j,t+1} = y_{j,t} \left(\lambda_{j,t}, r_{j,t}, (u)_{j,t}^{obj}, \bar{u}_{j,t} \right) \\ \pi_{j,t+1} = x_{j,t} \left(\mathcal{L} S_l(\cdot), r_{j,t}, (u)_{j,t}^{obj}, \bar{u}_{j,t}, \pi_{j,t} \right) \end{cases} \quad (13)$$

Where the functions $\mathcal{U}_{j,t}$ and $x_{j,t}$ depend on each node j selected learning strategy LS_t , perceived noisy reward signal $r_{j,t}$, and estimated utility value $(u)_{j,t}^{obj}$. The variable $\lambda_{j,t}$ is an averaging constant step-size parameter, to weights recent rewards more heavily than long-past ones.

The main objectives of the learning process are to:

- Acquire spectrum knowledge awareness and identify spectrum occupancy states for opportunistic access.
- Select subchannels from the spectrum pools and configure femtocell terminals to operate under restrictions of avoiding interference and meeting QoS conditions.

From a general perspective, the execution time and the convergence behaviour of the learning process depends on the structure of the utility function and the stated objectives. So, having a complex utility function or a tight objective could lead to a longer learning process, and therefore introduce more interference to the primary network. Remarkably, splitting the learning of the FCs into two sequential levels speeds up the discovery of the forbidden subchannels and the convergence of the learning model, as well as helping the FCs to identify channel occupancy states which could be used later on to monitor the activity of the primary users over the spectrum. The proposed learning paradigm is illustrated in Figure 11. Detailed descriptions of main building blocks are given in the following sections.

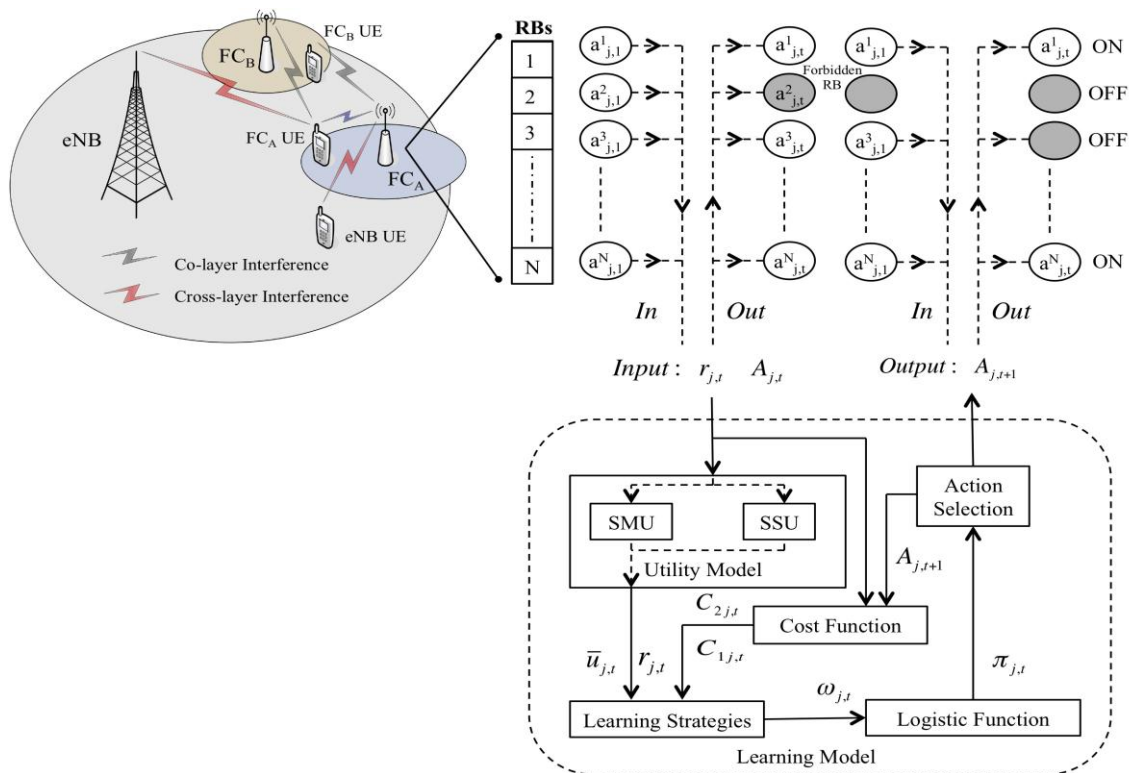


Figure 11: Functional proposed multi-objective CODIPAS-HRL model for FCs self-configuration in HNs deployment.

a. Utility Model

The utility model block is used to shape the reward signal into a desired form. Two definitions of the utility model are given here to cope with the above-mentioned objectives. First, the spectrum modelling utility (SMU), where the task of this model is to identify spectrum pools for opportunistic

access by generating a probabilistic data base of the channel availability, ranking the resources by a metric and eliminating all resources that do not achieve an adequate value.

The spectrum modelling utility function is given by:

$$(u)_{j,t}^{sm} = \begin{cases} \frac{r_{j,t} - \tau_j}{r_{j,t}} & \text{if } r_{j,t} \geq \tau_j \\ 0 & \text{if } r_{j,t} \leq \tau_j \end{cases} \quad (14)$$

Where the reward signal $r_{j,t}$ is given by means of the average measured SINR and the τ_j is the minimum SINR threshold for the j_{th} FC to be able to start a communication link. The value of the utility model is bounded within the interval $[0, 1]$.

The second model is the spectrum selection utility (SSU). The goal of the SSU is to configure the FC to transmit on certain subchannels such as to meet the QoS requirements, based on the spectrum pools generated by the SMU and the required number of resources by the FC. The SSU model is given by:

$$(u)_{j,t}^{ss} = \psi_{j,t} e^{\left(\frac{- \left(\sum_{\forall} a_{j,t} - N_{reqj} \right)^2}{Q_j} \right)} \quad (15)$$

Where N_{reqj} is the number of requested resources by the j_{th} FC, and $\psi_{j,t}$ is the ratio of the average achieved FC throughput $Th_{j,t}$ to requested/threshold throughput Th_{reqj} , given by:

$$\psi_{j,t} = \begin{cases} \frac{Th_{j,t} - Th_{reqj}}{Th_{j,t}} & \text{if } Th_{j,t} \geq Th_{reqj} \\ 0 & \text{for } Th_{j,t} \leq Th_{reqj} \end{cases} \quad (16)$$

Where the average achieved FC throughput $Th_{j,t}$ is calculated based on the received reward signal $r_{j,t}$. The exponential function in equation (9) optimises the FCs' resources selection hence it has a value equal to $\{1\}$ when the number of assigned resources is equal to the requested one, otherwise it is exponentially decreased based on the value of Q_j . The value of the SSU utility is also bounded between the interval $[0, 1]$.

b. Learning Strategies

Based on the problem formulation, we present a dynamic solution using methods from reinforcement learning; the gradient follower (GF), the modified Roth-Erev (MRE), and the modified Bush and Mosteller (MBM) learning algorithms [34]. The proposed methods are within the actor-critic methods.

At each incremental learning step, the j_{th} FC receives a reward signal $r_{j,t}$ as a consequence of selecting action $a_{j,t}$. Then, accordingly, it estimates/predicts the average utility value $\bar{u}_{j,t+1}$ and adjusts the weights $\omega_{j,t+1}^i$ of selecting action $a_{j,t+1}$ for the next learning step using one of the learning strategies LS_l , $l = \{1, \dots, L\}$, in order to find the optimal spectrum configuration toward maximising the utility function. The main considered learning strategies are described in the following

i. The gradient follower **LS₁**

$$\bar{u}_{j,t+1} = \bar{u}_{j,t} + \lambda_{j,t} (u_{j,t} - \bar{u}_{j,t}) \quad (17)$$

$$(\omega_{j,t+1})_{a=1} = \omega_{j,t} + \alpha_{j,t} (u_{j,t} - \bar{u}_{j,t}) F(\omega_{j,t}, \tau_1, a_{j,t}^i) \quad (18)$$

$$F(\omega_{j,t}, \tau_1, a_{j,t}^i) = \nabla_{\omega} \text{Sigm}(\omega_{j,t}, \tau_1, a_{j,t}^i) = \frac{\xi(a_{j,t}^i)}{\frac{\omega_{j,t}}{e^{\tau_1}} \left(1 + e^{\tau_1}\right)^2} \quad (19)$$

$$\xi(a_{j,t}^i) = \begin{cases} 1 & \text{for } a_{j,t}^i = 1 \\ -1 & \text{for } a_{j,t}^i = 0 \end{cases} \quad (20)$$

where $\alpha_{j,t}$ is the j_{th} FC's learning rate bounded between the value, $(0 < \alpha < 1)$, τ_1 is a positive Boltzmann cooling parameter, and $\nabla_{\omega} \text{Sigm}(\omega_{j,t}, \tau_1, a_{j,t}^i)$ is the gradient of the sigmoid function with respect to the probability weights $\omega_{j,t}^i$.

ii. The modified Roth-Erev **LS₂**

$$\bar{u}_{j,t+1} = \bar{u}_{j,t} + \lambda_{j,t} (u_{j,t} - \bar{u}_{j,t}) \quad (21)$$

$$(\omega_{j,t+1})_{\forall a, a \in \{0,1\}} = (1 - \phi) (\omega_{j,t})_{\forall a, a \in \{0,1\}} + H(\epsilon, u_{j,t}, a_{j,t}^i, \omega_{j,t}^i) \quad (22)$$

$$H(\epsilon, u_{j,t}, a_{j,t}^i) = \begin{cases} \frac{u_{j,t}(1 - \epsilon)}{v_j} & \text{for } a_{j,t}^i = a \\ \frac{(\omega_{j,t}^i)_a \epsilon}{v_j} & \text{for } a_{j,t}^i \neq a \end{cases} \quad (23)$$

Where v_j is weight parameter, ϕ is called the recency parameter $(0 < \phi < 1)$, as ϕ approaches 0, approximately equal weights placed on all received rewards $r_{j,t}$ over time. Hence, a high value ϕ is desirable in the time-changing environment. While ϵ is the experimentation parameter $(0 < \epsilon < 1)$, as ϵ approaches 0, the reward resulting from the chosen action $\{a_{j,t}^i = a\}$ is attributed only to a , implying that only the weight of the selected action $\{a_{j,t}^i = a\}$ is updated [1]. The use of ϵ is to encourage the exploitation of new actions and avoid premature convergence on a suboptimal action.

iii. The modified Bush and Mosteller **LS₃**

$$\bar{u}_{j,t+1} = \bar{u}_{j,t} + \lambda_{j,t} (u_{j,t} - \bar{u}_{j,t}) \quad (24)$$

$$\bar{\mu}_{j,t} = \frac{u_{j,t} - M_j}{\sup_a |U_j(a) - M_j|} \quad (25)$$

$$(\omega_{j,t+1}^i)_{a=1} = \begin{cases} \omega_{j,t}^i + \beta_{j,t} \bar{\mu}_{j,t} (\kappa_j - \omega_{j,t}^i) & \text{for } \{a_{j,t}^i = 1 \text{ and } \bar{\mu}_{j,t} \geq 0\} \\ \omega_{j,t}^i & \text{for } \{a_{j,t}^i = 0 \text{ and } \bar{\mu}_{j,t} \geq 0\} \\ \omega_{j,t}^i + \beta_{j,t} \bar{\mu}_{j,t} & \text{for } \{a_{j,t}^i = 1 \text{ and } \bar{\mu}_{j,t} < 0\} \\ \omega_{j,t}^i & \text{for } \{a_{j,t}^i = 0 \text{ and } \bar{\mu}_{j,t} < 0\} \end{cases} \quad (26)$$

where $\bar{\mu}_{j,t}$ is the average stimulus corresponding to the selected action $\{a_{j,t}^i = a\}$, $\sup |\cdot|$ is the superior limiting of the Cesro-mean payoff of the j th FC, M_j is an aspiration level, $\beta_{j,t}$ is the learning rate, and κ is a clipping logit function to set an upper bound for $\omega_{j,t}^i$ as the probability $p_{j,t}^i$ approaches $\{1\}$, given by $\kappa_j = \log(pmax) - \log(1 - pmax)$. Note that $\bar{\mu}_{j,t}$ is always bounded between $[-1, 1]$.

c. Logistic Function

Logistic functions are often used in machine learning to "smooth" the variant and transform the full-range variables into a limited range of a probability, restricted to the specific range $[0, 1]$, by applying a bounded logistic function to the weights. In essence, the logistic function works as a graded function of the selected actions where they are ranked and weighted according to the value estimates. The most common model uses the sigmoidal distribution, which expects a relatively linear function in the middle of the probability range, while it is stretching exponentially towards the extremes as it approaches the range's bounds. The learning strategies LS1 and LS3 update the probability of selecting action $\{a_{j,t}^i = a\}$ according to the following rule:

$$(p_{j,t+1}^i)_{a=1} = \frac{1}{e^{\frac{(\omega_{j,t+1}^i)_{a=1}}{\tau_1}}} \quad (27)$$

The strategy profile for the j th FC on the i th subchannel is given by:

$$p_{j,t}^i = \{(p_{j,t+1}^i)_{a=0}, (p_{j,t+1}^i)_{a=1}\} \quad (28)$$

The set of strategies taken up by the j th FC for the selected action over all subchannels is given by:

$$(\pi_{j,t+1})_{a \in \{0,1\}} = \{(p_{j,t+1}^1)_{a \in \{0,1\}}, \dots, (p_{j,t+1}^N)_{a \in \{0,1\}}\} \quad (29)$$

A generalisation and extension of the logistic function to multiple inputs is the softmax activation function. The actions' probability of the learning strategy LS2 is updated according to the following rule:

$$(p_{j,t+1}^i)_{a \in \{0,1\}} = \frac{e^{\frac{(\omega_{j,t+1}^i)_{a \in \{0,1\}}}{\tau_2}}}{\sum_{\forall a \in \{0,1\}} e^{\frac{(\omega_{j,t+1}^i)_a}{\tau_2}}} \quad (30)$$

d. Action Selection

The action vector $a_{j,t+1}^i$ of the j th FC for the next time step $(t + 1)$ will be generated based on a Bernoulli binary generator given by means of each node's updated strategy $(\pi_{j,t+1})_{a \in \{0,1\}}$, as shown below:

$$(a_{j,t+1})_{a \in \{0,1\}} = \text{Ber} \left((\pi_{j,t+1})_{a \in \{0,1\}} \right) \quad (31)$$

e. Cost Function

In the previous section, it is shown that in order for the FC to derive its own transmission policy, it needs to know how its decision process is coupled to that of other FCs/eNBs, where the goal of the proposed game framework, is to find the policy such that FCs' utility is maximised during the learning process. Hence, such greedy interaction often involves a "cost" to both network tiers.

For comparison and performance optimisation, two cost functions are considered in this work, namely, the collision cost and the average reconfiguration cost during the learning process. For instance, the learning algorithms should be aware of the previous state of the network, so that the next state will be selected, taking into account the minimum number of reconfigurations. This has the effect of ensuring that fewer configurations will be required and signalling overheads for all the network elements are minimized. The collision cost is defined as the normalised total interference experienced by femtocell j at each time step t , given by:

$$C_{1,j,t} = \frac{\sum_{n=1}^{n_{f,j}} \sum_{i \in \mathcal{V}_i} \left((X_c)_{n,j}^{f,i} + (X_o)_{n,j}^{f,i} \right)}{n_{f,j} \bar{\delta}_j} \quad (32)$$

Where the parameter $\bar{\delta}_j$ is a normalised factor, such that $(0 \leq C_{1,j,t} \leq 1)$.

The reconfiguration cost is defined as the average switching between ON/OFF states of all resources during the learning process, given by:

$$C_{2,t} = \frac{\sum_{i=1}^N \mathbb{1}_{a_{i,t+1} \neq a_{i,t}}}{N} \quad (33)$$

The average reconfiguration cost can be calculated in an incremental step as follows:

$$\bar{C}_{2,t} = \bar{C}_{2,t-1} + \gamma_{j,t} (C_{2,t} - \bar{C}_{2,t-1}) \quad (34)$$

Where $\gamma_{j,t}$ is an averaging parameter.

Note that fast discovery of forbidden resources could lower the mutual interference experienced by the FCs during the learning period, which is the target of the SMU learning phase. On the other hand, a fast convergence strategy $\pi_{j,t}$ towards selecting actions $\{0, 1\}$ lowers the reconfiguration cost (SSU learning strategy behaviour).

2.5.3 Integration in OneFIT architecture

Proposed spectrum selection framework is mapped onto the decision making entities at the CMON of the OneFIT architecture. For further details on the proposed framework the reader is referred to section 2.4.3 of deliverable D4.2 [7].

2.5.4 Further performance evaluation results

a. Evaluation Platform and Simulation Parameters

The considered scenario in this work consists of an LTE radio access cellular network, comprising tri-sectorised eNBs located at the centre of the macrocell, and a layer of closed access, randomly-located LTE FCs overlaid upon the macrocell system. Each macrocell's sector contains 10 pedestrian users, while 5 indoor users are located uniformly inside each femtocell. The simulation parameters and assumptions are summarised in Table (I).

The simulations carried out are based on the Downlink Long Term Evolution (LTE) System Level simulator. A spectrum configuration toolbox, with different learning strategies and optimisation techniques had been added to the baseline simulator. The accuracy of the presented learning strategies is evaluated using the cross-validation technique for the holdout set of actions, to ensure a fair comparison of different learning strategies' behaviour on the detection of the ON/OFF subchannels (selection/avoidance behaviour). Here, we use the history of each learning strategy $h_{j,t}^{\ell} = (a_{j,0}, a_{j,1}, \dots, a_{j,t-1})$ to measure the accuracy. The accuracy can be expressed as the model's ability to correctly classify ON/OFF in the holdout dataset (same as the definition of the detection and false alarm probabilities in spectrum sensing).

We calculate the accuracy of the algorithms according to the following definition:

$$Pr_{ON} = \frac{T_{on}}{T_{on} + F_{on}} \quad (35)$$

where Pr_{ON} is the probability of correctly classify ON subchannels during the learning process. The parameter T_{on} is defined as how many times the learning algorithm correctly identified ON. On the other hand, F_{on} is the fault identification of ON (collision cost occurs when the decision of the FC is ON but the resource is in use by other FCs/eNBs).

The probability of correctly identified OFF subchannels is given by:

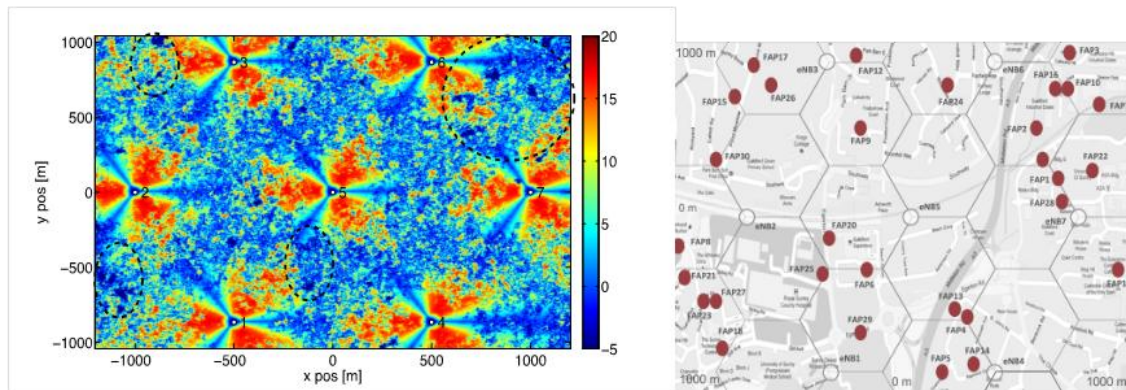
$$Pr_{OFF} = \frac{T_{off}}{T_{off} + F_{off}} \quad (36)$$

where the same definition applies to T_{off} and F_{off} .

Table 1: Simulation Parameters

Parameter		Value
Macrocells Network	No. of eNodeB	7
	Inter site distance (ISD)	1 Km
	Number of sectors	3 sectors per site
	eNBs transmission mode	SISO
	Number of attached UEs	10/sector
	Minimum coupling loss	70 dB
	Pathloss model (TS25814)	$PL_{mc} = 128.1 + 37.6 \log_{10}(d)$
	Macro environment setting	urban
	Max. eNBs transmission power	43 dBm
	Shadow fading mean	mean 0 dB, standard deviation 10 dB
	Shadowing correlation distance	50 m
	Small Scale Fading Model	Typical Urban (TU)
	Antenna gain pattern	Directional
	UE receiver noise figure	9 dB
	UE thermal noise density	-174 dBm/Hz
	UE speed	3 Km/h
Femtocells Network	No. of Femtocells	30
	Cell Radius	20 m
	Max. FCs transmission power	10 dBm
	Pathloss model	$PL_{fc} = 37 + 32 \log_{10}(d)$
	Wall Penetration loss	10 dB
	Number of attached UEs	5
	Power allocation	always ON, homogeneous
LTE Spectrum	Scheduler	Round Robin
	Frequency	2.14 GHz
	Channel BW	5 MHz
	No. of Resource Blocks	25 RBs

b. Performance Evaluation Results



(a) The SINR map of the LTE cellular network, the circular dashed line shows low coverage area. (b) Residential area covered by 7 tri-sectored macrocells and 30 Femtocells.

Figure 12: Cell layout and SINR map of the LTE heterogeneous network HNs

Figure 12 depicts the cell layout and the SINR coverage of a layer of one tier cellular network overlaid with a layer of 30 FCs randomly located over the area. Here, each eNB optimised its spectrum utilisation to create spectrum opportunities for the secondary network tier, adopting algorithms proposed in [35]. Meanwhile, FCs randomly select carriers based on the predicted demand and the method of gradual deployment of carriers (GDC). The main purpose of Figure 12 is to illustrate the impact of the secondary FCs' coexistence within a legacy macro-cellular system, where dead zones (poor SINR) often appear within the coverage of cellular network due to the excessive cross-layer interference. Indeed, this emphasises the need for smart terminals (FCs) that have the capability to self-configure their operational parameters under restrictions of avoiding/causing interference and meeting QoS conditions.

To sum up, each FC runs the multi-objective CODIPAS-HRL model, in order to first discover channel occupancy states, then accordingly, maximise the average cell throughput, with the necessary support of reconfigurability, self-organisation and learning characteristics. The outcome of the SMU learning model for the 20th FC is illustrated in Figure 13. It shows the selected strategy $\pi_{i,t}$ of taking action $a_{20,t}^i = 1$ for all subchannels after 2.5×10^3 , 5×10^3 , and 10×10^3 learning steps. The initial strategy profile $\{\pi_0\} \forall j$ is equal 0.5 for all femtocells. As shown, after 5×10^3 and 10×10^3 learning steps, higher probability is assigned to subchannels $\{9 - 13\}$ and $\{19 - 24\}$ as compared to the remaining sub channels. Indeed, these subchannels with a higher probability are identified as the resources that are most likely to be free for opportunistic access by the secondary node, labelled as ON, while other subchannels will be labelled as OFF. Also, it is clear that as the number of learning steps increase, the nodes (FCs) acquire better knowledge about the spectrum usage and primary network activity.

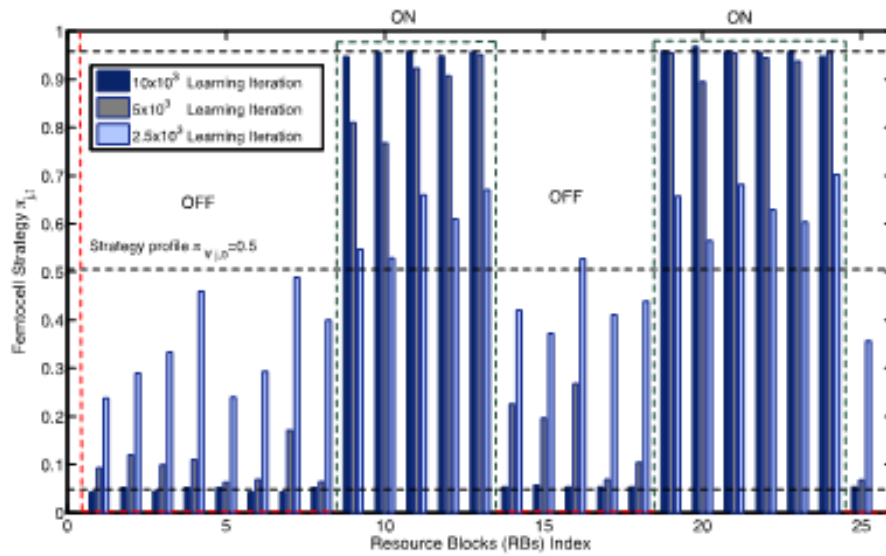


Figure 13: The selected strategy of FC number 20 using the GF learning strategy.

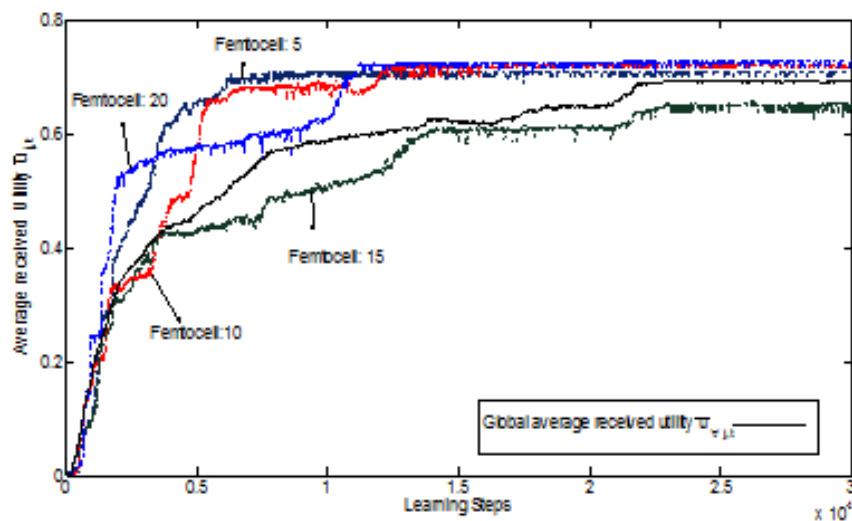
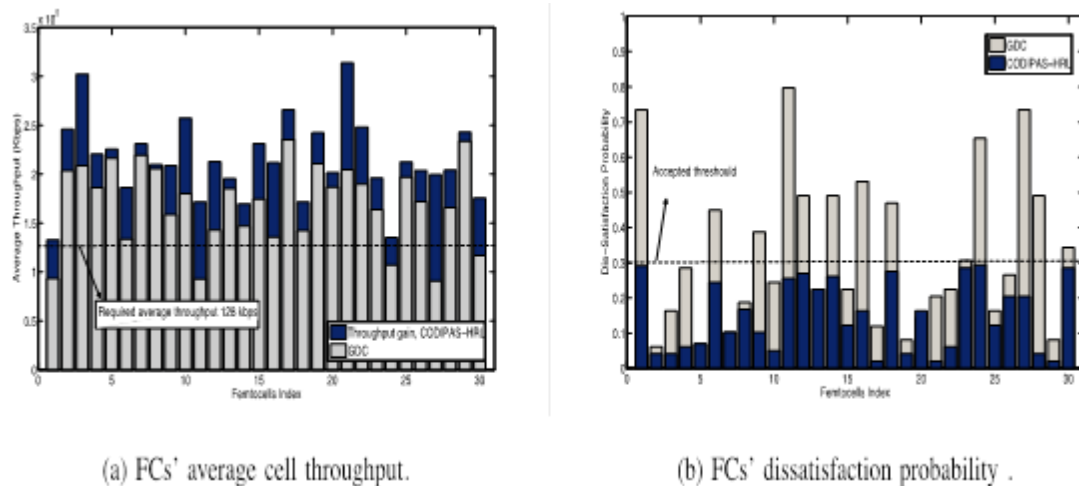


Figure 14: The average received utility of the proposed learning model for FCs number 5, 10, 15, and 20, with initial strategy profile $\{=1\}$.

Subsequently, the obtained strategy π_j from the SMU model is used to configure the initial strategy $\pi_{0,j}$ of the SSU learning model. Figure 14 shows the normalised average FCs' reward signal $r_{j,t}$ versus the execution time of the learning process. The convergence behaviour of the learning model's trajectories is shown clearly which reflects the fairness and optimality of the played game among participated FCs. As expected, the convergence behaviour differs according to the adopted learning strategies and network states (i.e. femtocells' location, macrocells' traffic, etc.).



(a) FCs' average cell throughput.

(b) FCs' dissatisfaction probability .

Figure 15: The dissatisfaction probability and the average cell throughput of the proposed leaning model as compared to the GDC.

In this work, we used the average cell throughput and the dissatisfaction probability as performance metrics. As shown in Figure 15.a, the proposed model improves the average cell throughput by approximately 30% and achieves a satisfactory solution that meets the network conditions, as compared to the random GDC spectrum configuration. In addition, the dissatisfaction probability is used here to ensure the fairness of access among all FCs' users.

From Figure 15.b we can note that the optimised configuration did enhance the dissatisfaction probability by an average of 47% and meets the satisfactory threshold that could be set up in advance by the FCs' provider. Indeed, Figure 14 and Figure 15 show the algorithm converges to a satisfactory equilibrium or stable solution that explicitly fulfils femtocells constraints.

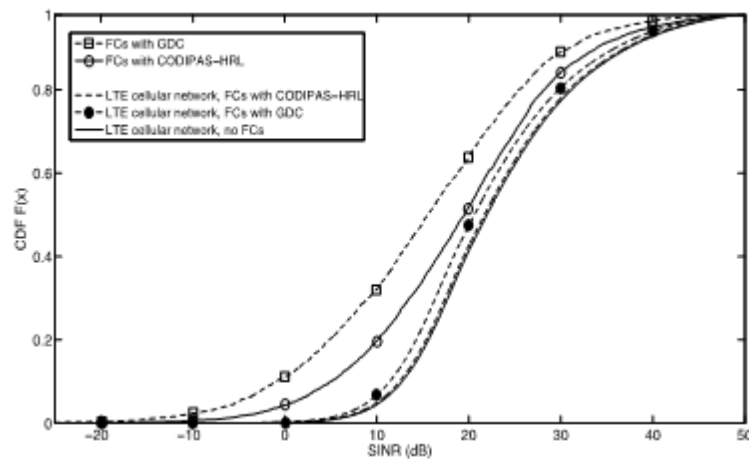


Figure 16: SINR cumulative distributed function (CDF) for the eNBs and FCs with and without the proposed learning model.

Moreover, Figure 16 shows the CDF of the SINR of all FCs/eNBs obtained with the optimised spectrum configuration, for the coexistence scenario and for the case when there are no FCs. It is seen that there are improvements in the SINR for all network tiers. Although the mean improvement is around 6dB for the FC tier, it is around 2dB for the macrocell tier. This is due to the fact that the eNBs are considered to be the main aggressors with high power transmission. Furthermore, we found that the SINR of the macrocells' tier is nearly similar to the case when there is no FCs overlay upon them, thanks to the learning model that aims to minimise the cross-layer interference between network tiers.

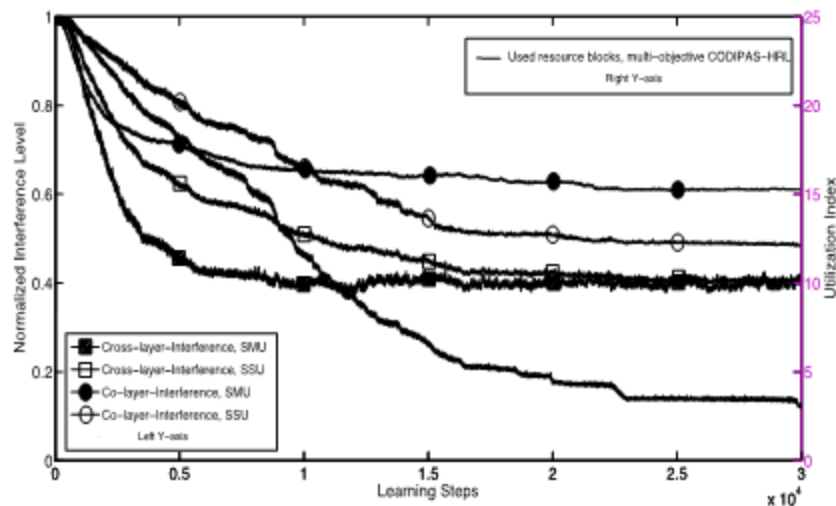
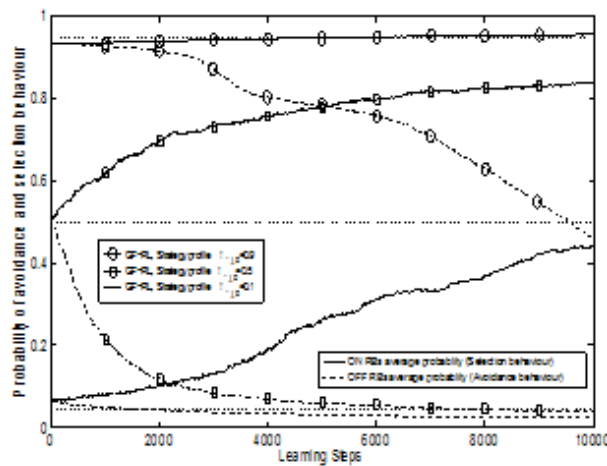


Figure 17: The normalised interference level and the utilisation index for the SMU and SSU versus the learning time.

The advantage of dividing the learning process into two sub-sequential levels is illustrated in Figure 17. It shows the normalised co/cross-layer interference level and the utilisation index versus the number of learning steps. The utilisation index is defined as how much of the available spectrum has been used by each FC. It is assumed that each FC requires 3 subchannels to deliver its own users' requested traffic. Initially, each FC begins with a full spectrum configuration. Thereafter, it starts to release harmful spectral resources according to the learning update rules of SMU/SSU model, until it reaches the best 3 subchannels that fulfil the FC users' requirements. During the initial stage of the learning interactions, it is shown clearly in Figure 17 that FCs with an SMU model experience lower inference levels from both networks' tiers, as compared to the FCs that adopted an SSU model. Although the SSU model has slower convergence, it leads to a better spectrum configuration as we approach the end of the learning steps. Hence, a hybrid learning process i.e. SMU model to discover forbidden subchannels during initial learning phase, preceded by an SSU for terminal configuration is outperforming individual learning utility.

The amount of perceived interference during the SMU learning process also depends on the initial configured strategy of each FC $\pi_{i,t}$. Figure 18 illustrates the selection/avoidance (ON/OFF) convergence behaviour and the normalised experienced interference during the resource configuration for 10×10^3 iteration steps using LS_1 . The tendency to select action $\{a_{j,t} = 1\}$ starts with the initial strategy $\pi_{0,j} = \{0.9\}$, $\{0.5\}$ and $\{0.1\}$ as shown in Figure 18.a.

Clearly, the initial strategy $\pi_{0,j} = \{0.1\}$ is the most conservative approach, with accuracy of around 88% to identify OFF subchannels. It assigns low probability to all subchannels, thereafter it gradually increases the probability of subchannels that maximise the FC's utility function. However, it has a slow convergence trajectory toward detecting the ON subchannels and higher reconfiguration cost, as compared to the strategies $\pi_{0,j} = \{0.5, 0.9\}$ (low accuracy, so long learning steps needed), per Table (II). On the other hand, the strategy $\pi_{0,j} = \{0.9\}$ has low accuracy, around 55%, to discover OFF subchannels, which implies higher interference perceived during the learning steps, as depicted in Figure 18.b. FCs with strategy profile $\pi_{0,j} = \{0.5\}$ have the fastest convergence of the ON/OFF channel discovery with slightly lower accuracy, experience lower interference levels, and have the minimum reconfiguration cost as compared to other, per Table (II). Indeed, a mixed-value initial strategy profile could outperform the single-value strategy profile, since it is assigning high probability for channels that are most likely to be free and low probability for channels that are most likely to be used by other FCs/eNBs. Thus, prediction to calculate the prior probability for the initial strategy profile should be considered for further research.



(a) RBs selection and avoidance behaviour (ON/OFF probability).

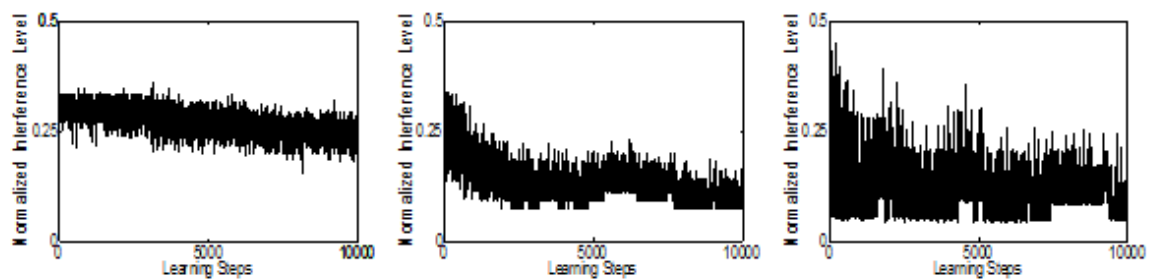
(b) Interference level, with $\{\Pi_0\}_1 = 0.9$. (c) Interference level, with $\{\Pi_0\}_1 = 0.5$. (d) Interference level, with $\{\Pi_0\}_1 = 0.1$.

Figure 18: FCs RBs selection and avoidance behaviour and normalised interference level of the GF learning strategy.

Table 2: Initial strategy profiles $\{\Pi_0\}_1$ of selecting action $a_{i,t}=1$, for the gradient follower (GF) learning strategy

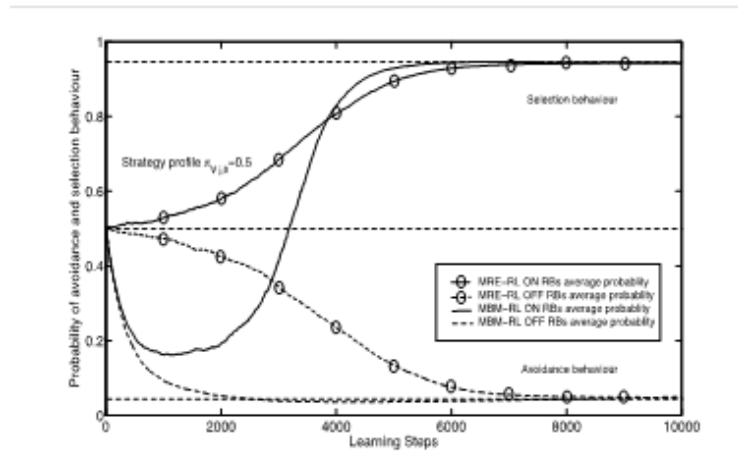
Π_0	Pr_{off}	Pr_{on}	C_2
$\Pi_0 = 0.1$	0.88 %	0.56 %	0.85 %
$\Pi_0 = 0.5$	0.85 %	0.78 %	0.76 %
$\Pi_0 = 0.9$	0.55 %	0.81 %	0.83 %

Table 3: Comparisons learning strategies: the gradient follower (GF), the modified Roth-Erev (MRE) and the modified Bush and Mosteller (MBM)

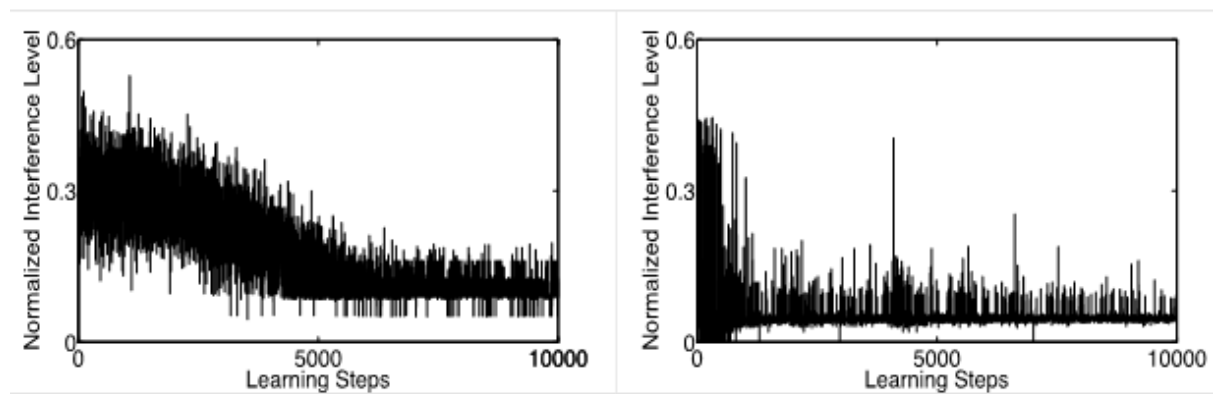
Case	Pr_{off}	Pr_{on}	C_2
LS ₁ : GF	0.85 %	0.78 %	0.76 %
LS ₂ : MRE	0.69 %	0.69 %	0.77 %
LS ₃ : MBM	0.95 %	0.70 %	0.80 %

Finally, Figure 19 extends the results of the selected strategy $\pi_{0,j} = \{0.5\}$ to compare the performance of the proposed learning algorithms. Algorithm LS3 exhibits the advantages of fast convergence, high accuracy of correctly identified OFF, and experiences low interference, compared to LS₁ and LS₂. However, the low accuracy of detecting ON subchannels and the high reconfiguration cost could be the payoff, per Table (III). Indeed, the algorithm LS₃ is suitable for the SMU learning as the target is to readily and correctly identify the forbidden harmful carriers while LS₁ and LS₂ are more suited to

the SSU learning, with the high accuracy of selecting subchannels that maximise the objectives from the spectrum pools.



(a) RBs selection and avoidance behaviour (ON/OFF probability).



(b) FC's interference level, MRE learning strategy.

(c) FC's interference level, MBM learning strategy.

Figure 19: FCs RBs selection and avoidance behaviour and normalised interference level of the MRE and MBM learning strategies.

In summary, a multi-objective CODIPAS-HRL model to solve the interference coordination problem in the downlink of LTE FCs in heterogeneous networks has been provided. The proposed model enables the FCs to autonomously sense their environment and to tune their parameters accordingly, in order to operate under restrictions of minimizing interference to both network tiers. The communication environment has been modeled as a dynamic, non-cooperative game with incomplete/imperfect information, and with heterogeneous players. The simulation results show the convergence behaviour of the proposed model which reflects the fairness and optimality of the played game among participating FCs. The total experienced interference and the reconfiguration cost has been also studied in order to provide a reasonable comparison between different learning strategies and to highlight the cost associated with the learning process. Specifically we show that the proposed strategies can identify unused resource blocks with high degree of accuracy, co/cross-layer interference can be reduced significantly, thus resulting in average cell throughput increases of 30% and satisfaction probabilities increased by 47%, in the considered scenarios. In this research, we proposed the learning model as a practical solution to configure/optimize spectrum selection for FCs. Given this context, future work could consider the potential use of the proposed learning scheme as an approach to facilitate spectrum sensing, particularly, for wide band channels where the use of the acquired spectrum knowledge is highly valuable.

2.6 Techniques for Aggregation of Available Spectrum Bands/Fragments

2.6.1 Problem formulation and algorithm concept

The problem considered here is the approach for spectrum selection with spectrum aggregation (SA) capabilities. When it is difficult to support enough bandwidth with contiguous available spectrum, multiple small spectrum fragments (sub-channels) can be exploited to yield a (virtual) single channel by spectrum aggregation. When spectrum is allocated to users, there are some factors which should be considered. First, since maximizing the total throughput by efficient spectrum utilization is still important, the sub-channels with the high signal quality can be preferred. Second, in the opportunistic spectrum access, since channel switching caused by a returning primary users (PUs) leads to system overhead, selection of the channel with least-likelihood for the appearance of a PU is advantageous. Finally, the aggregate channel which is composed of less number of sub-channels with the type of spectrum aggregation with less complexity can be selected. The proposed algorithm makes use of the utility-based approach for these three objectives. For further details on this problem formulation, the reader is referred to section 2.6.1 of deliverable D4.2 [7].

2.6.2 Algorithm specification

The proposed spectrum aggregation approach considering multi-objectives depicted in Figure 20 consists of two parts.

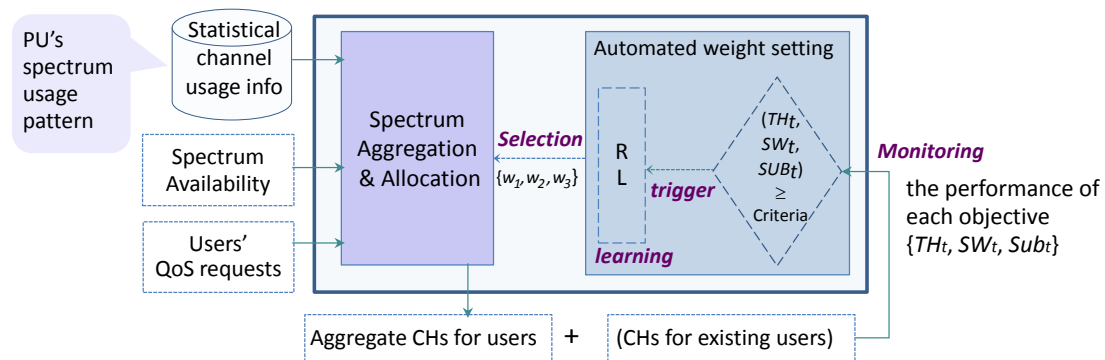


Figure 20: Framework of the proposed spectrum aggregation approach

- **Multi-objective utility-based spectrum aggregation and allocation:** The proposed algorithm considers three objectives: 1) to maximize the total throughput, 2) to minimize the channel switching, and 3) to reduce the spectrum aggregation complexity. The utility value for each objective is calculated from three different utility function and these three objectives are integrated into a weighted sum utility function. Multiple sub-channels can be aggregated to satisfy the secondary user's throughput requirement.
- **Automated weight setting by the learning technique:** While the multi-objective spectrum aggregation approach exploits the weighted sum approach, the optimal weights setting can be different depending on the interested performance metric. The system is assumed to have set/pre-defined thresholds for each performance metric based on system level key performance indicators (KPIs) to be met. Then, the interested performance metric can be changed depending on the environment changes. In the framework of the proposed algorithm, the MTLs process (i.e. *Monitoring* → *Triggering* → *Learning* → *Selection*) repeats to adapt to the different environments. That is, when the spectrum is allocated to users, the performance of each objective is monitored (*Monitoring*). In case that the monitored performance is not met, the learning module is triggered (*Triggering*). As a result of the

learning process (Learning), the set of optimal weight vector is found (Selection) and used by the spectrum aggregation algorithm.

For further details on this algorithm including utility functions and weighted setting module, the reader is referred to section 2.6.2 of deliverable D4.2 [7]

2.6.3 Integration in OneFIT architecture

The proposed spectrum aggregation approach is mapped onto the function entity of the decision making at the CMON of the OneFIT architecture. For further details, the reader is referred to section 2.6.3 of deliverable D4.2 [7].

2.6.4 Further performance evaluation results

In order to evaluate the performance of the proposed algorithm; we used the LTE system level simulator [34]. The main simulation parameters are shown in TABLE below.

Table 4: Simulation parameters

ISD	50m	100m	300m	500m	1000m	1732m
Layout	Hexagonal grid, 19 cell sites, 1 sectors per site			Hexagonal grid, 19 cell sites, 3 sectors per site		
# of UEs	10 UEs (constant UEs distribution)					
Traffic	Full-buffer traffic model					
Scheduler	RR / PF					
Channel Model	$L[dB]=\begin{cases} 39+20\log_{10}(d[m]) & 10m < d \leq 45m \\ -39+67\log_{10}(d[m]) & d > 45m \end{cases}$			L=l + 37.6log ₁₀ (.R), R in kilometres, l=128.1 @ 2GHz		
Tx Pwr	9 dBm		41dBm	49dBm (@ 20MHz)		
	Based on 41dBm(@ 300m ISD), to calculate the minimum receiver sensitivity (-65.dBm) at the cell edge . Then by considering the min. rx. sensitivity and the channel model, other ISD Tx Pwrs are calculated			(43dBm @5MHz, 46dBm @10MHz in LTE → thus to set two times larger Tx Pwr for two times wider bandwidth)		
Shadowing mean, s.d.	0, 10dB			0, 8dB		
Correlation distance	25m			50m		
Antenna	1x1 SISO					
Carrier Frequency/Bandwidth	2GHz /20MHz					
Secondary network	• Single network (with only one sector) • Varying number of users (2,4,6,8,10) • The proposed utility-based resource allocation • One sub-channel: 180kHz RB unit • The rest settings are the same with LTE network (e.g. Full-buffer traffic model, channel model, 1x1 SISO)					

It is assumed that the opportunistic network operates as the secondary user while LTE system operates as the primary user. In the LTE system level simulator, we integrate the opportunistic network as the secondary user and the proposed algorithm for its spectrum aggregation and allocation. For the opportunistic network, the traffic is modelled as the full buffer model. Since the proposed algorithm aims three objectives, the performance evaluation is based on the following metrics, which try to reflect each objective.

- Average throughput performance, average number of channel switching, average number of sub-channels

For evaluation of the proposed algorithm's performance, we use two reference approaches.

- Random aggregation algorithm: It selects sub-channels randomly among the available sub-channels
- Single objective aggregation algorithm: It can be implemented by adjusting the weight of each objective
 - For the weight of three objectives, $\{W_1, W_2, W_3\} = \{1, 0, 0\}$, $\{0, 1, 0\}$, $\{0, 0, 1\}$

The same priority to three different objectives i.e. the weight-vector of multi-objective utility $\{w_1, w_2, w_3\}$ is set as $\{1/3, 1/3, 1/3\}$ for the first stage simulation. Performance evaluations in the following sections compare performance of the random aggregation method (labelled as "Random") with different settings of the weight vector as shown in section 3.7.2 of deliverable D4.2 [7].

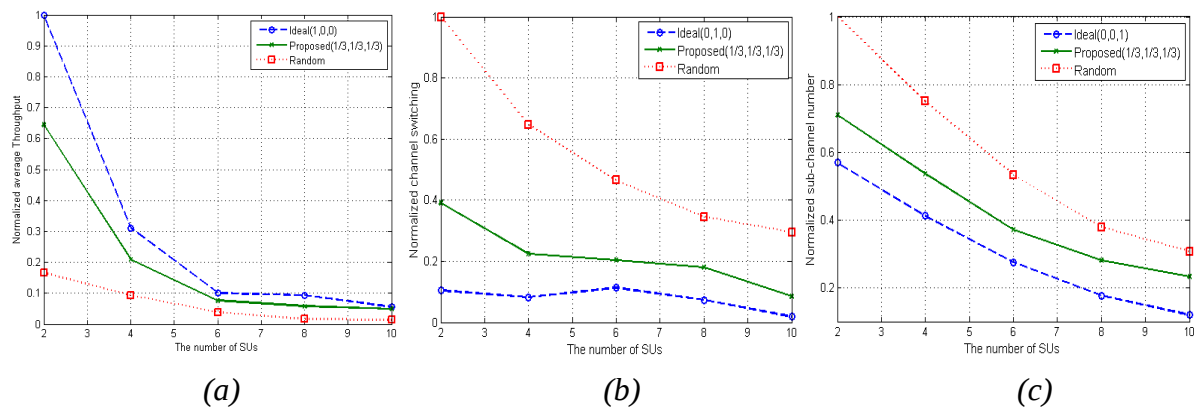


Figure 21: (a) Normalized average Throughput (b) Normalized number of channel switching (c) Normalized number of sub-channels

Figure 21 (a) presents the normalized average throughput experienced by each secondary user. While the throughput performance of the proposed equal-weight setting algorithm outperforms Random algorithm's and falls behind the optimal setting algorithm's. Since the full buffer traffic model is exploited, average throughput per users decreases for the increase in the number of users. Figure 21 (b) evaluates the channel switching performance. Similarly, the equal-weight setting results in more channel switching but presents better performance than the Random aggregation. Lastly, the number of sub-channels is evaluated. It is shown that as the number of users increases, the less number of sub-channels becomes to allocate to each user. When two figures in (b) and (c) are compared, it is analysed that the times of channel switching is reduced for the large number of users since the less number of sub-channels is allocated to users. All three graphs show the same pattern in the Figure 77 of deliverable D4.2 [7].

We evaluate the effect of inter-site distance (ISD) on the channel switching. For the simulation setting, the values of ISD $\{ 50\text{m}, 100\text{m}, 300\text{m}, 500\text{m}, 1000\text{m}, 1732\text{m} \}$ are considered. Based on the

value of ISD, the base stations of PU are assumed to be located and a secondary network is also located. In the case of a large ISD setting, the possibility of channel switching of SUs becomes to decrease. Since a fixed number of users (10 users) are distributed in wider area, the distance between a PU user and a SU user is longer compared to the case of a small ISD. Thus, varying spectrum usage by the primary user influences less the spectrum usage by the secondary user and it results better channel switching performance for secondary network (i.e. reductions). The channel switching performance for different ISDs is shown in Figure 22.

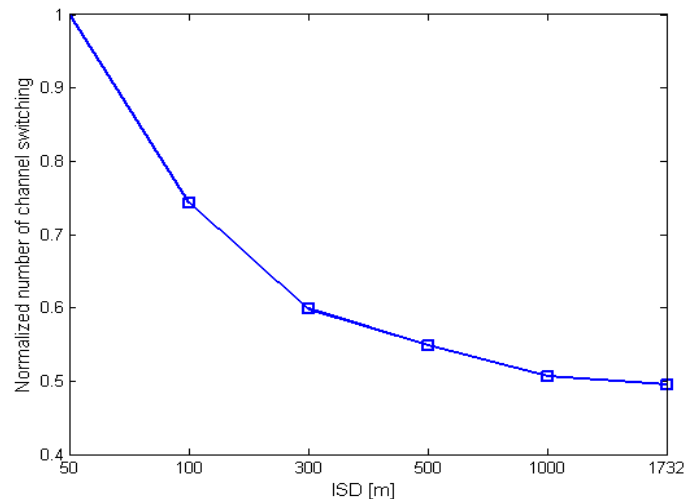


Figure 22: Channel switching performance for different ISD

2.7 Knowledge-based suitability determination and selection of nodes and routes

2.7.1 Problem formulation and algorithm concept

During the OneFIT project the following algorithmic solutions have been developed and evaluated:

- Knowledge-based suitability determination
- Selection of nodes through a fitness value evaluation
- Capacity extension of the infrastructure through neighbouring terminals

This document summarizes and highlights main features of these algorithms which have been extensively analysed in [7], [11].

2.7.1.1 Knowledge-based suitability determination

The algorithm will be responsible for making decisions upon the feasibility of the creation of ONs when judged as appropriate. The delineation of such an algorithm-strategy is the objective of this section. In particular, at the input level, a properly defined algorithm will need to read context information, which according to the preceding analysis, comprises (i) the number and/or spatial distribution of terminals, (ii) the type (requirements) of applications requested, (iii) mobility levels and (iv) access point and terminal capabilities and characteristics, such as supported applications, routing protocols etc. In the output, the algorithm must select (i) the transmission power (range) of the Access Points (APs) and/or terminals and (ii) the ad hoc routing protocol that will be used for routing traffic between the infrastructure and the ONs (Figure 23).

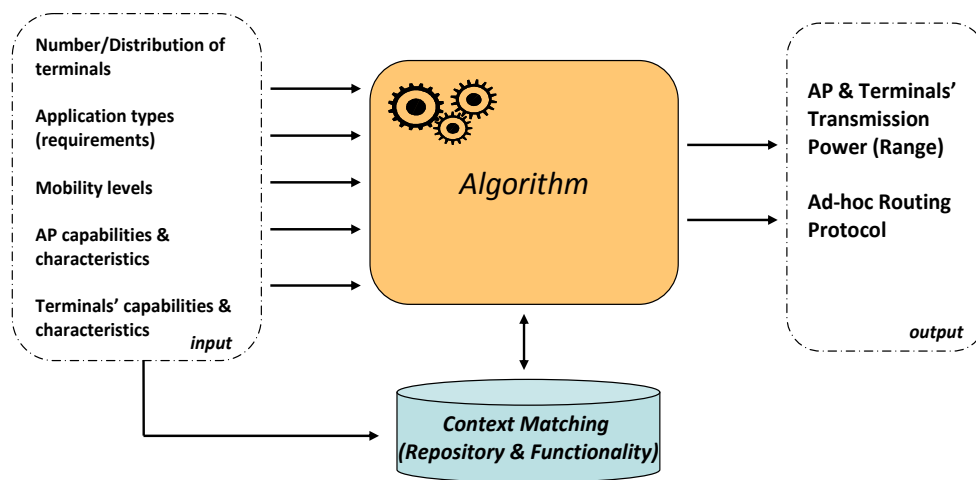


Figure 23: Outline of the knowledge-based suitability determination approach

2.7.1.2 Selection of nodes through a fitness value evaluation

According to the OneFIT concept, ONs are created in an infrastructure-less manner under the supervision of the operator and include numerous network-enabled elements. To ensure application provisioning in an acceptable QoS level, the selection of the proper nodes, among all discovered nodes in the vicinity is rather essential. As a result, this approach proposes a mechanism for selecting nodes that will participate in the ON. The selection mechanism is based on a fitness function which is a weighted, linear formula which takes into consideration specific parameters as defined in [7], [11], [15], [16].

2.7.1.3 Capacity extension of the infrastructure through neighbouring terminals

The algorithm is executed in the infrastructure side and more specifically in a congested base station which needs to solve the problem of congestion. It is triggered whenever a congestion situation occurs in the infrastructure and makes decisions on the establishment of routes which will redirect traffic from the congested service area into non-congested ones. The proposed solution of the algorithm is based on the Ford-Fulkerson maximum flow algorithm and it is assumed that each terminal in the congested service area creates one application flow which can be redirected through neighbouring terminals (e.g., which are in range of a typical Wi-Fi connection) to alternate, non-congested BSs.

2.7.2 Algorithm specification

2.7.2.1 Knowledge-based suitability determination

Figure 24 illustrates the proposed solution for the “suitability determination” problem by using a flowchart [14]. Initially, a monitoring process (point 1 in Flowchart) is taking place in the network so that the decision maker (the operator) will be aware of the nodes’ related information, capabilities and set of characteristics. Under certain circumstances/criteria, the suitability determination process will be triggered. Such criteria could be (i) the bad coverage of nodes inside the cell or (ii) the existence of nodes out of the coverage. Interference problems among APs or/and nodes could also trigger suitability determination process but this is not covered by the results in this section. The suitability determination process (point 2 in Flowchart) should take mainly into account multiple input parameters consisting, among others, the requested traffic/application types and respective requirements and node mobility levels. At the same time it should envisage anticipated benefits for the whole network such as the provision of adequate QoS levels to applications/ users or the decrease of the transmission powers, which can also mean lower energy consumption. The suitability determination process is going to decide whether or not it is appropriate to set up an ON,

at specific time and place and this is going to trigger the creation of these networks. Additionally, it will give as an output a request for the creation process of the ON, associated with a pre-selected set of candidate nodes that offer at least on radio path to an infrastructure AP and at least one radio path between each pair of nodes.

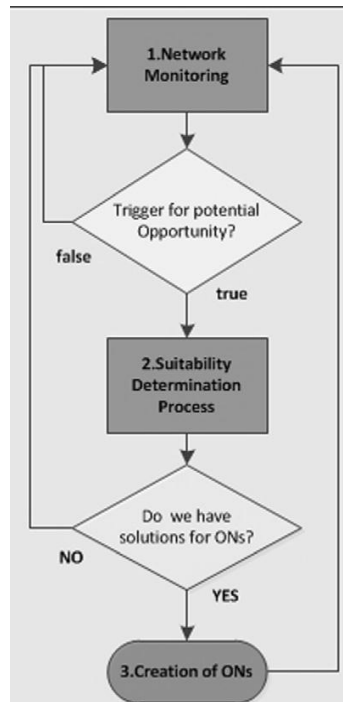


Figure 24: Flowchart of the proposed solution for the suitability determination problem

2.7.2.2 Selection of nodes through a fitness value evaluation

The input of the algorithm consists of the set of candidate nodes which are located in a specific service area that have ON capabilities (i.e., they have the ability to participate in an ON) and they are legitimate to participate in an ON according to the operator policies, in order to stress on the operator-governance during the ON formation. The output of the algorithm consists of the selected nodes that will form the ON. Then, the ON creation procedure is triggered in order to allow the nodes/terminals to negotiate between each other and establish links. The aforementioned procedure is illustrated in Figure 25.

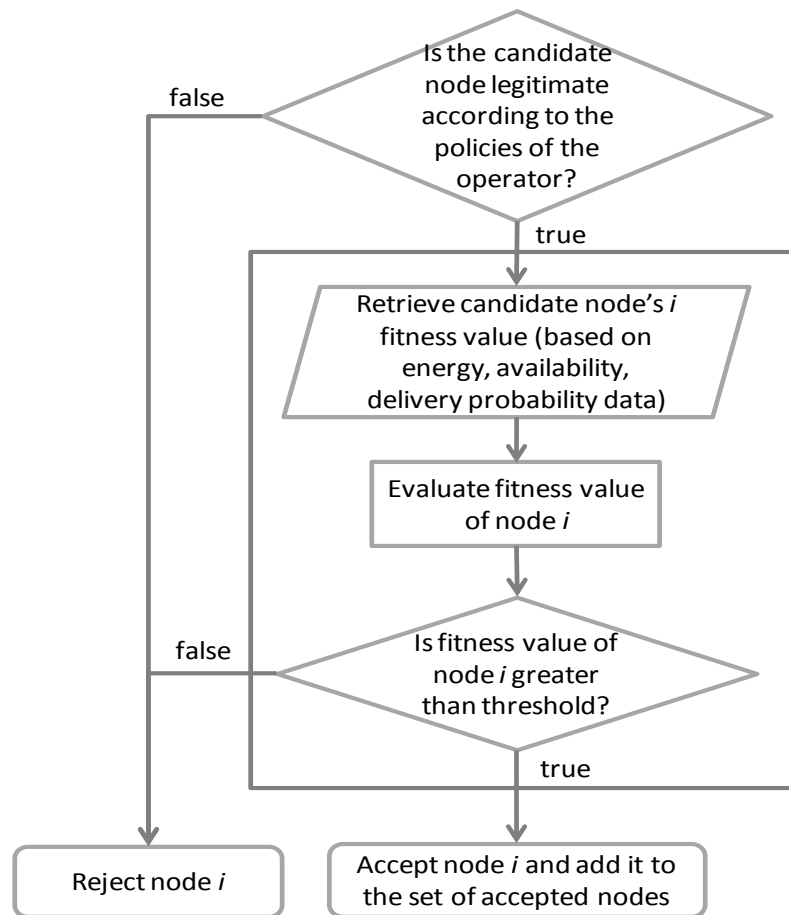


Figure 25: Flowchart of the proposed algorithm

2.7.2.3 Capacity extension of the infrastructure through neighbouring terminals

The input of the algorithm consists of the following:

- Sets of congested and non-congested BSs
- Set of terminals in the congested area that have ON capabilities and can be redirected to alternate BSs
- Paths from source to destination. Each path originates from a 'virtual' source where terminals with ON capabilities in the congested area are connected and ends to a 'virtual' destination where BSs are connected.

At the initialization phase the link capacities are estimated. Their values derive from parameters such as the RAT of the link or the quality of the link. In addition, the flows of all links are initialized to zero. Then the set P is created by the paths that are being discovered by the breadth-first search algorithm which yields the shortest path. The algorithm runs for each path $p \in P$ for which there is a link where more flow can be directed. As soon as all possible paths are saturated (i.e. there is no path with available capacity anymore and flow could not be pushed to it) the algorithm finalizes. The following figure provides the flowchart of the proposed algorithm. Detailed formulation and analysis is available in [7], [11], [17], [18].

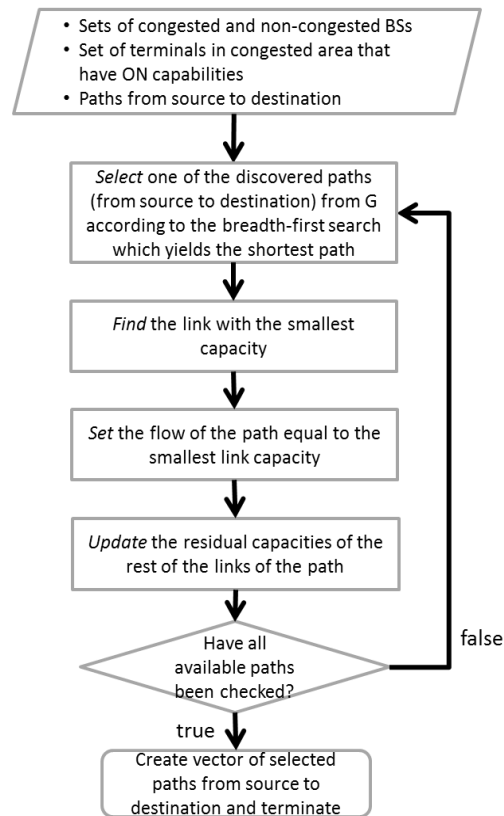


Figure 26: Flowchart of the proposed algorithm

2.7.3 Integration in OneFIT architecture

2.7.3.1 Knowledge-based suitability determination

The integration of the “knowledge-based suitability determination” in the OneFIT architecture has been documented since [11]. Figure 39 depicts the mapping to CSCI functional entity, in terms of parameters used by the knowledge based suitability determination algorithm, since the CSCI is responsible for the execution of the suitability determination phase as described in [3], [4]. The CSCI involves the following entities: context awareness, policy derivation and management, profile management, decision making mechanism and knowledge management.

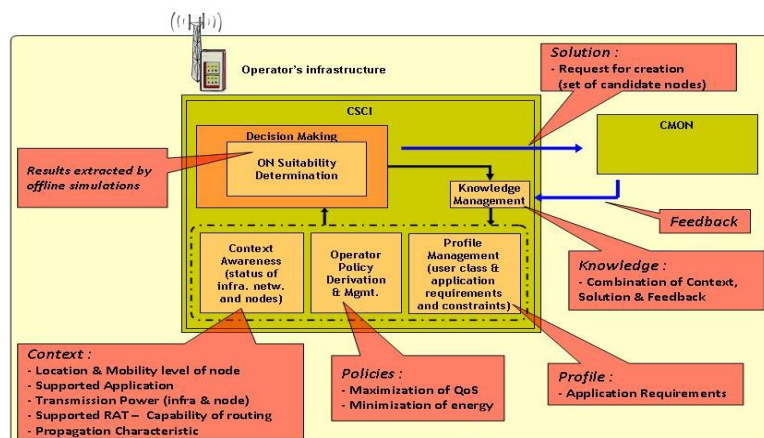


Figure 27: Integration of the “knowledge-based suitability determination” in the OneFIT Functional Architecture

2.7.3.2 Selection of nodes through a fitness value evaluation

Also, Figure 40 provides the mapping of the functionalities of the selection of nodes through a fitness value evaluation to the OneFIT functional entities. Details have been discussed in [7].

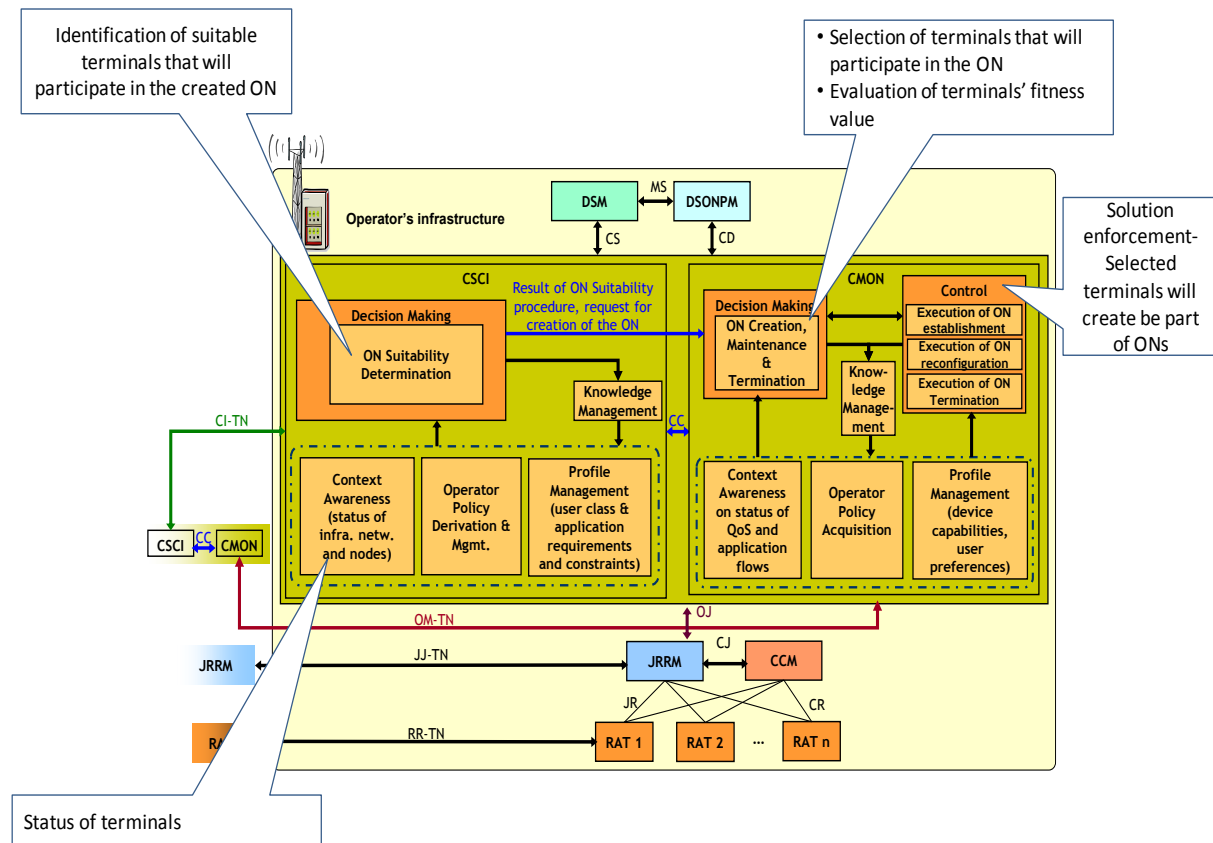


Figure 28: Integration of the “selection of nodes and routes through a fitness value evaluation” in the OneFIT Functional Architecture

2.7.3.3 Capacity extension of the infrastructure through neighbouring terminals

Figure 41 provides a mapping of the functionalities of the capacity extension through neighbouring terminals algorithm to the OneFIT functional architecture and the functionalities entities as defined in D2.2 [3].

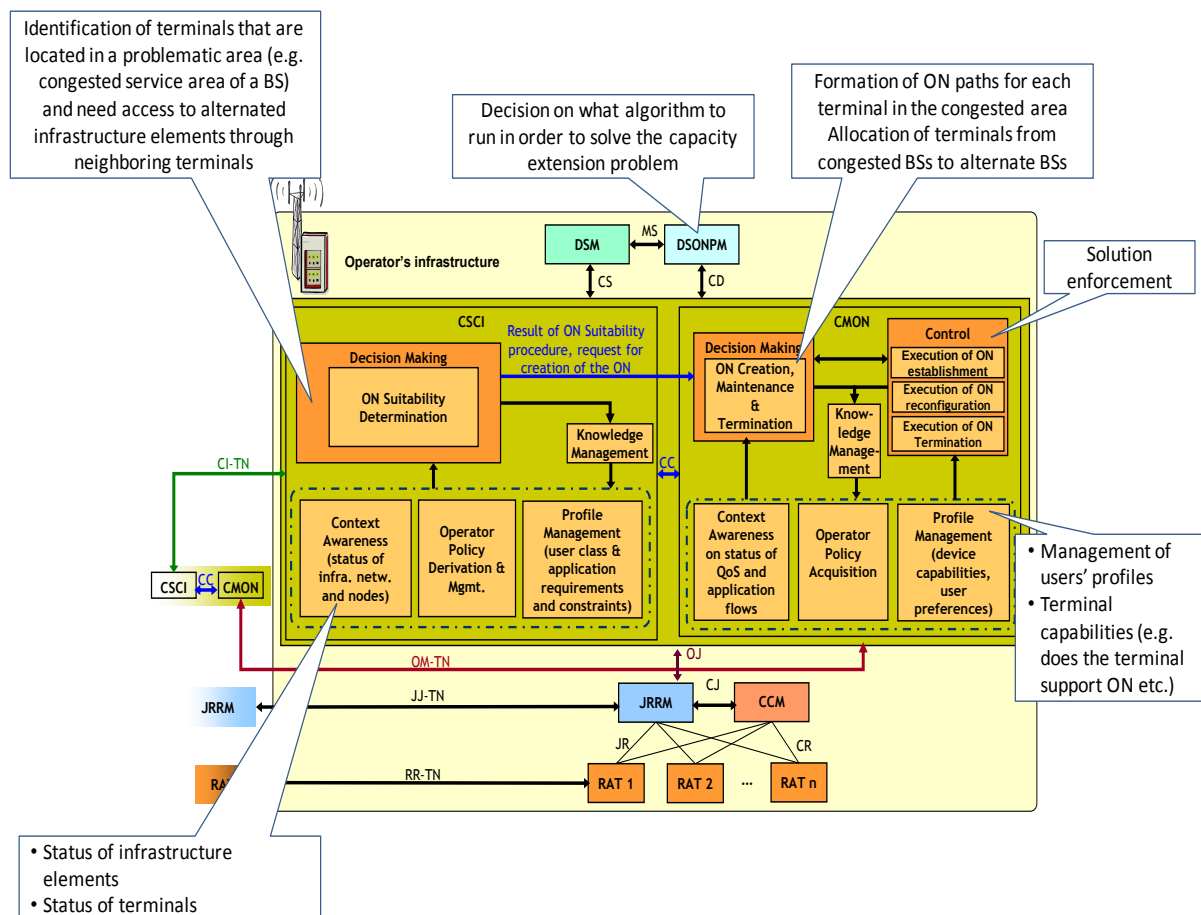


Figure 29: Integration of the “Capacity extension of the infrastructure through neighbouring terminals” in the OneFIT Functional Architecture

2.7.4 Further performance evaluation results

Since D4.1 and D4.2, specific results to each algorithm have been provided and analysed. However, an extra result set is available in D4.3 which gives a more detailed insight to the investigated algorithms. Specifically, seven new use cases are considered in the simulations (simulations have been conducted in a modified version of the ONE simulator [22]). In these use cases, 6 BSs are considered as non-congested and 1 as congested. Moreover, each non-congested BS serves 15 terminals and the congested BS serves 40 terminals. In the first use case, all terminals are stationary. From use cases 2 to 4, from each BS, 6 terminals are moving with a speed of 1, 2 or 4 m/s accordingly. In the final set of use cases (5 to 7), 3 terminals from each BS are moving with a speed of 1, 2 or 4 m/s accordingly.

Table5: Studied cases

Case	Moving terminals of each BS	Speed (m/s)
1	0	0
2	6	1
3	6	2

4	6	4
5	3	1
6	3	2
7	3	4

As Figure 42 suggests, the mean delay (i.e., time to transmit a whole message from source to destination), tends to drop after the creation of operator-governed ONs takes place.

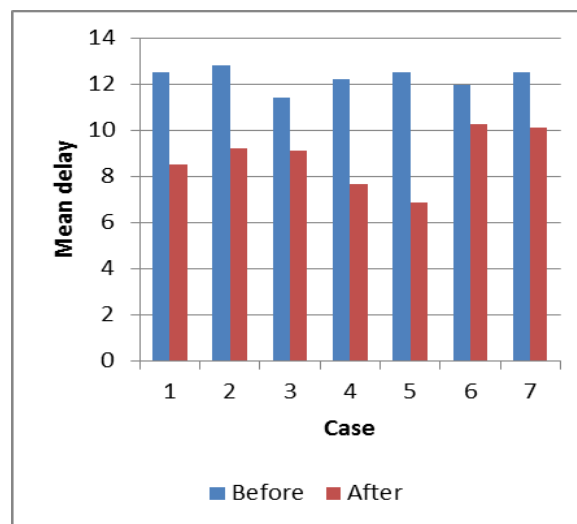


Figure 30: Mean delay (in seconds) for each case

Regarding the delivery probability, as the speed increases, the links among the nodes of the ONs are not stable and more messages are dropped. Thus, the delivery probability after the solution decreases as the speed increases as Figure 31 suggests.

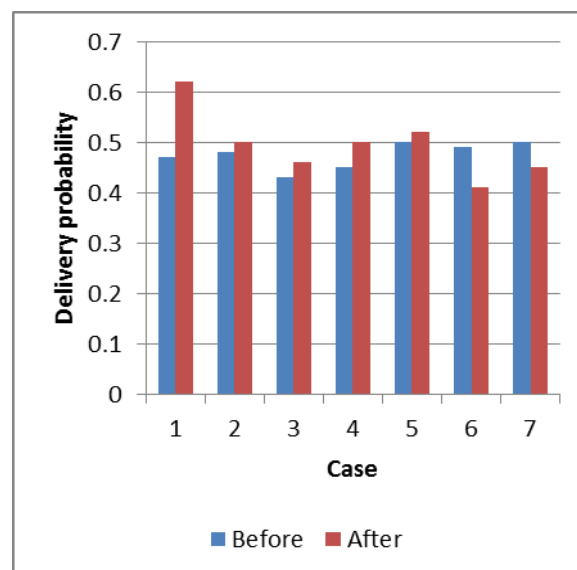


Figure 31: Delivery probability for each case

Moreover, according to conducted simulations it is worth mentioning that for low mean speed values (i.e. 0 and 1 m/s) the links of the ON nodes are stable and therefore the mean power of the

intermediate nodes increases due to the fact that they serve additional nodes. When the speed value is high, the links among the ON nodes begin to break and therefore the intermediate nodes don't serve additional nodes for long. In addition, after the ON creation they begin to move and the distance between the intermediate nodes and their serving BS decreases. As a result the mean power decreases (Figure 32).

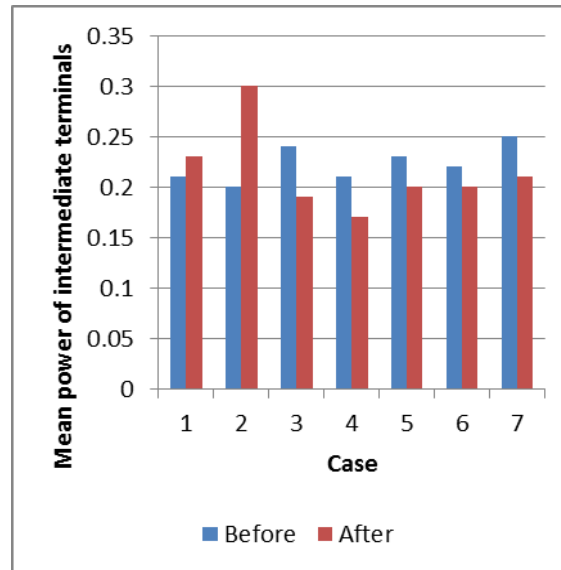


Figure 32: Mean power of intermediate terminals

Regarding the edge terminals (these are the terminals that switch to ON), after the ON creation their traffic is transferred through the neighboring terminals and the distance is very small compared with the distance from the BS. Therefore, the power that is needed for the communication among the terminals is lower than the one that is needed for the communication with the BS (Figure 33).

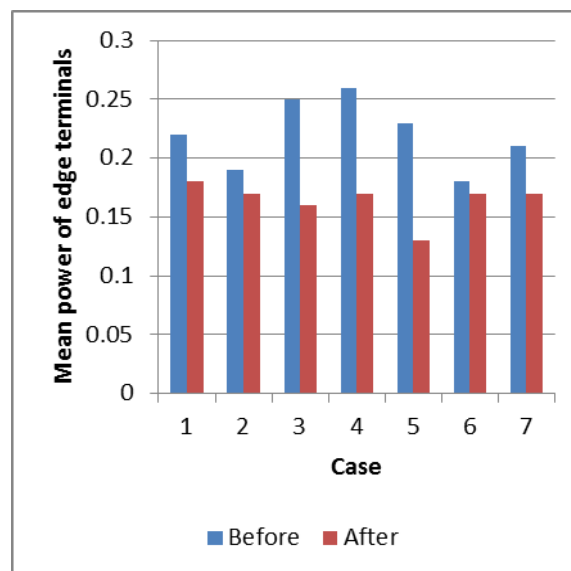


Figure 33: Mean power of edge terminals (terminals that switch to ON)

As far as the BSs are regarded, the mean power is the one that is needed for the communication with a terminal. When an ON is created, the edge terminals are offloaded to the neighboring BSs, therefore the remaining terminals have small distance from the BS and lower power is needed for the communication as Figure 34 illustrates.

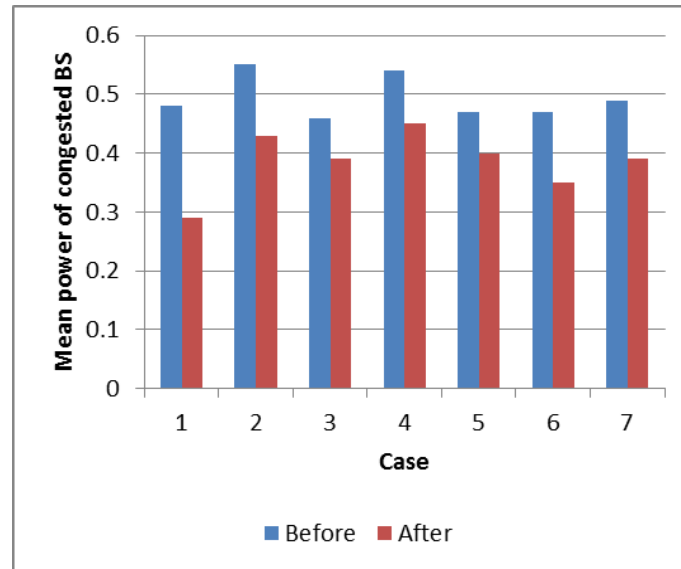


Figure 34: Mean power of congested BS

Finally, regarding the non-congested BSs, for low mean speed values that lead to stable ONs, the mean power increases after the solution due to the fact that the non-congested BSs acquire a proportion of the traffic of the congested BS. For higher mean speed values, the ONs after some point it is more likely to be destroyed (due to mobility), hence the traffic acquired by non-congested BSs is lower. Therefore, the mean power is more or less the same before and after the solution. Figure 35 illustrates the aforementioned findings.

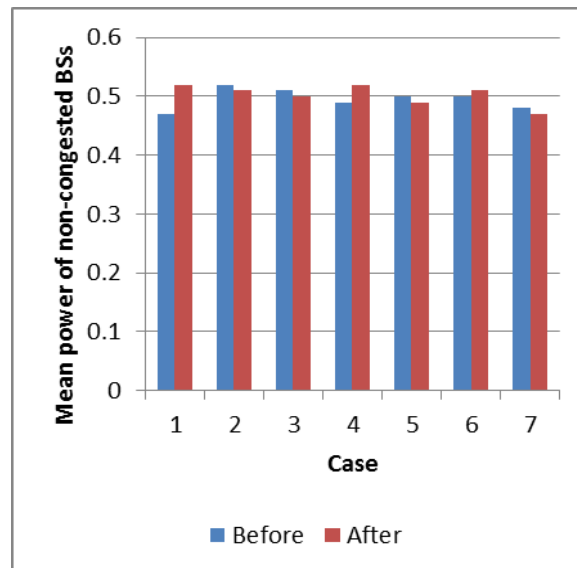


Figure 35: Mean power of non-congested BSs

2.8 Route pattern selection in ad hoc network

2.8.1 Problem formulation and algorithm concept

The algorithm focuses mainly on the Suitability Determination and Maintenance as the technical challenges as they have been analysed during WP2. The OneFIT network supporting ad hoc configuration is operator governed.

The diversity of user applications implies different constraints for data transfer in order to guarantee the QoS. In an opportunistic network composed of heterogeneous equipment having different

characteristics, it is necessary to take into account the specificities of these equipment and the characteristics of the supported radio access technologies to determine the way to transmit and to receive data. Depending on the service requested, the routing behaviour shall be adapted, and the routing protocol shall consider particular metrics.

2.8.2 Algorithm specification

The proposed algorithm is an enhancement of the routing protocols to take into account the constraints associated to user applications, by selecting appropriate metrics for each service class and to compute these metrics in order to determine the most adapted route to exchange data. The route pattern selection algorithm acts as a cognitive router to determine the most adapted route to reach the destination in an ad-hoc opportunistic network. The algorithm is able to change the selected pattern depending to the end user application requesting data to transmit. The algorithm determines the set of patterns to be used.

It manages 4 different service classes:

- Conversational class
- Streaming class
- Interactive class
- Background class

It gets also a particular behaviour for the transmission of the control and signalling information.

The algorithm selects the distance (number of hops to reach the destination), to determine the best route for the transmission of control and signalling packets (Figure 36). This pattern applies also for C4MS messages.

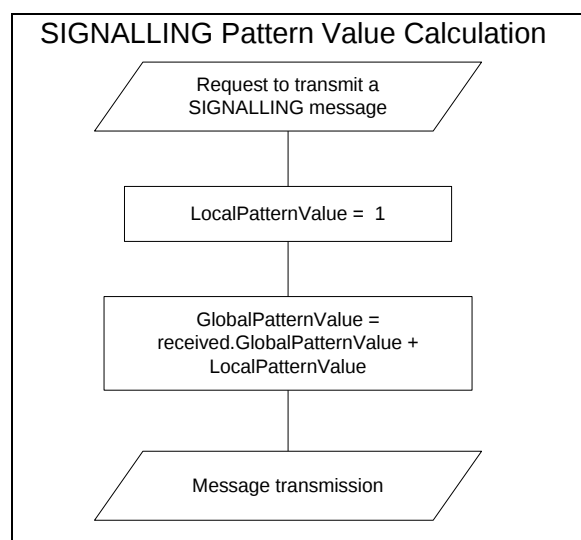


Figure 36 : Flow chart for signalling packet route selection

To transmit user packets related to conversational application, the algorithm selects the route depending on the distance (number of hops), because the distance introduces a certain latency. The second pattern used to determine the most adapted route is the age of the considered route in order to use the most stable routes (Figure 37).

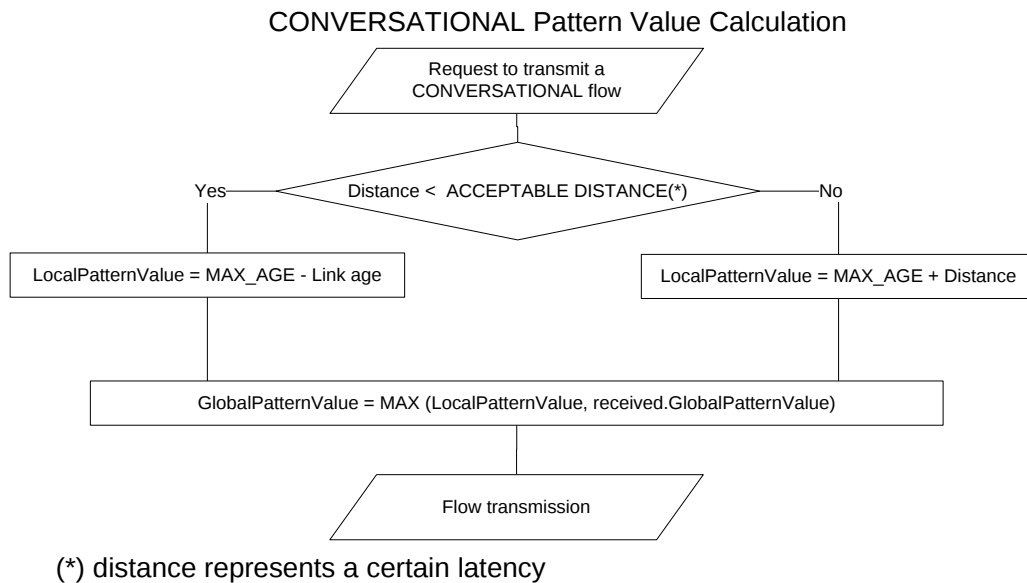


Figure 37 : Flow chart for conversational packet route selection

To transmit user packets related to streaming application, the algorithm selects the route depending on the available throughput, and it uses also as the link quality as secondary pattern (Figure 38). The available throughput is calculated by for all the possible routes connected to reach the destination.

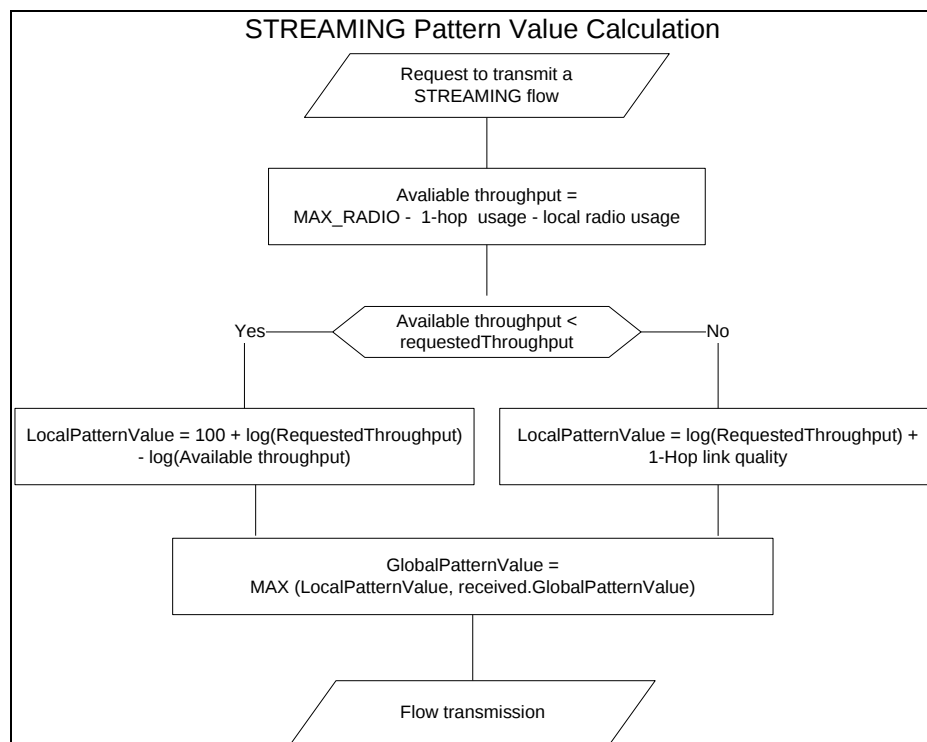


Figure 38 : Flow chart for streaming packet route selection

The background class applications are usually identified to use a “best effort” mechanism. The algorithm uses only the pattern distance (as for the signalling and control packet) to determine the most adapted route (Figure 39).

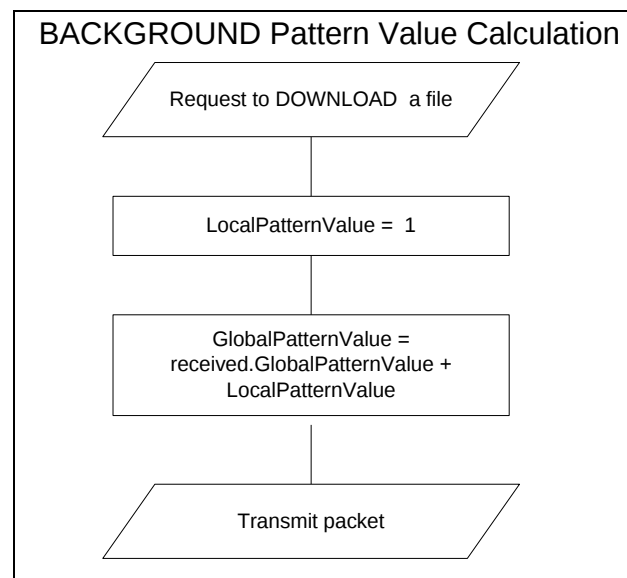


Figure 39 : Flow chart for background packet route selection

2.8.3 Integration in OneFIT architecture

The route pattern selection algorithm has to be seen as a functional module integrated in the both CSCI and CMON. The purpose of the algorithm is to be triggered any time a packet has to be transmitted over the air, to be able to establish and to select the most adapted route to transfer the data to the destination (i.e. to reach the infrastructure).

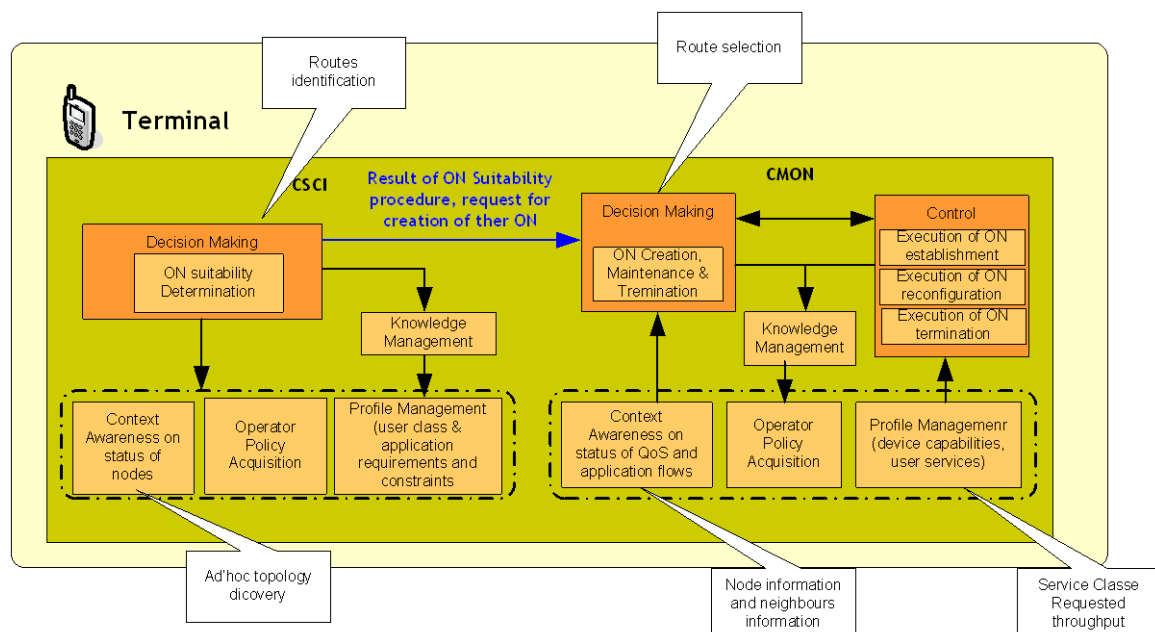


Figure 40 : Mapping of the route selection algorithm in the OneFIT functional architecture

During the ON suitability determination stage, it participates to the discovery of the topological environment, of each node involved in the opportunistic network. The figure below depicts precisely the moment of the algorithm processing during the ON suitability determination stage.

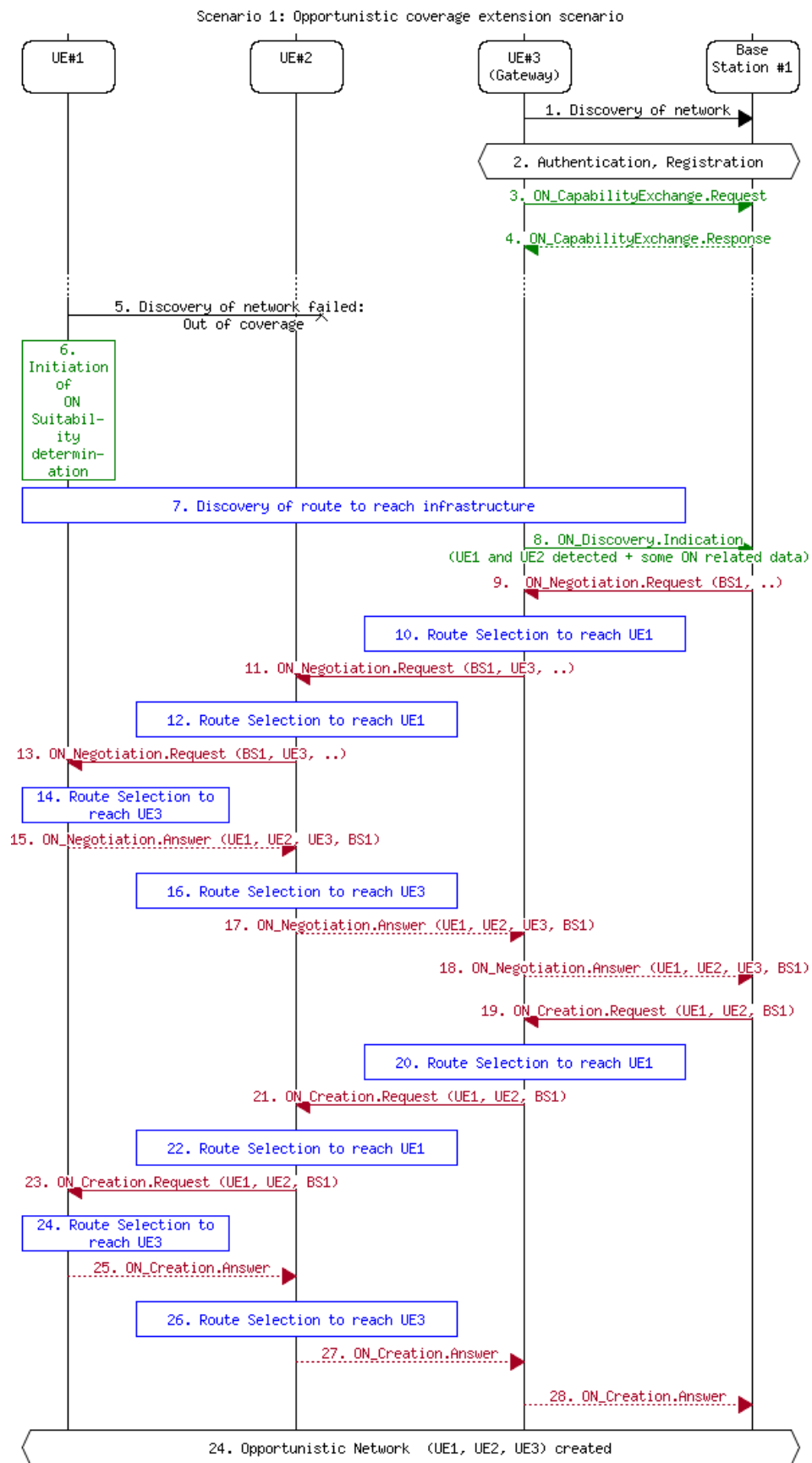


Figure 41: Sequence Diagram for Route pattern selection algorithm during ON Suitability determination

During the ON maintenance stage, the algorithm is applied for each packet of user data to be transmitted over the air, and according to the type of application (i.e.: the DSCP field located into the IP frame header), it determines, which route is allowed to be activated in order to satisfy the QoS related to the application. The algorithm requires also information about the context of the nodes and the context of its neighbours provided by the C4MS messages to apply its decision rules.

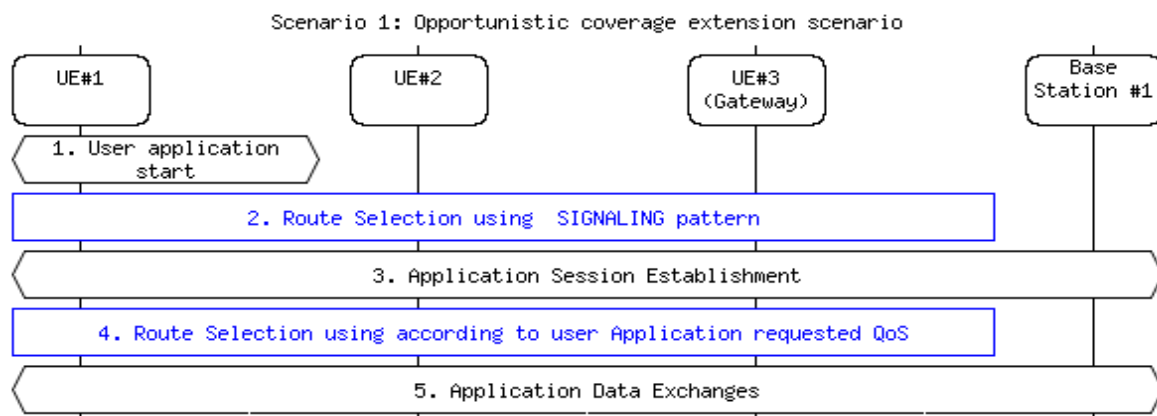


Figure 42: Sequence Diagram for Route pattern selection algorithm during ON maintenance

2.8.4 Further performance evaluation results

The main purpose of the algorithm is to select the most adapted pattern to select a route according to end user requested service and with the constraint to satisfy the Quality of Experience (QoE). The routing pattern selection algorithm is a qualitative algorithm.

There is no meaning to perform any quantitative analysis.

The reason, why quantitative results are not applicable, is that the algorithm opportunistically adapts according to the end user service requested, and the metrics that are necessary to evaluate a quantitative result vary from an application family to another. It would be necessary to tune finely every threshold of the used metrics, this is out of the scope of the project

The validity of the algorithm has been proved through by using OMNET simulation. The figure below (Figure 44) shows one example of the algorithm result for applied to streaming video application where the different links and radio paths are depicted in different colors:

- In GRAY, the available radio connections
- In BLUE , the signalling route (RTSP protocol) using SIGNALING pattern
- In GREEN, the data flow using STREAMING pattern

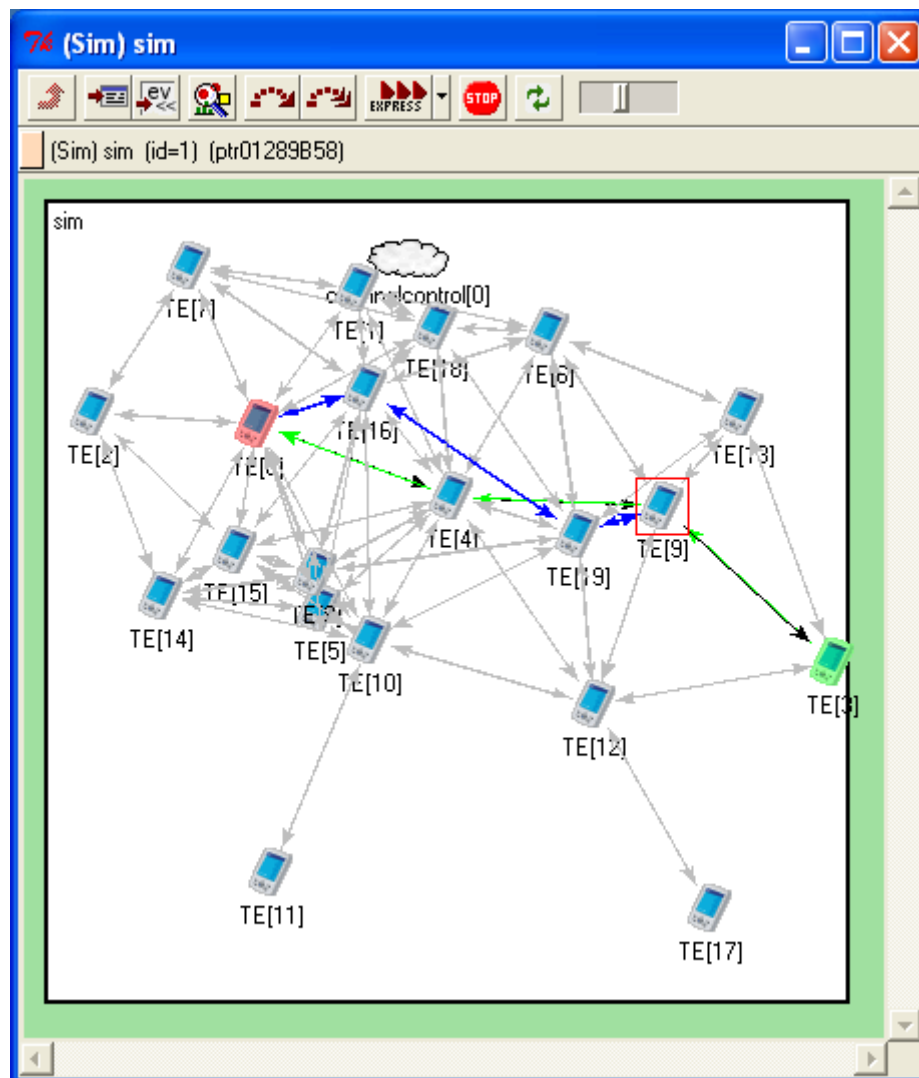


Figure 43: Example of route pattern selection

2.9 Multi-flow routes co-determination

2.9.1 Problem formulation and algorithm concept

As referred in the D4.2 deliverable, the class of Mobile Ad-Hoc Network (MANET) is particularly prominent in Private Mobile Radios networks management, as these networks have to be deployable in harsh environments, with infrastructure-less dependency.

The main characteristics of these networks are the self-neighbourhood detection and routing paths creation and the dynamic management of links/nodes.

The issues to be raised in terms of routing enhancement are manifold. One of these issues is the Quality of Service management, and in particular the routing based throughput optimization including resource allocation optimization. Algorithms have been proposed to optimize routing protocols for MANETs. Extensions integrate to these protocols a quality of service management.

In D4.2 a multi-flow routes co-determination algorithm is proposed, in addition to these optimizations. This algorithm combines the (re)routing of traffic flows on ad-hoc network with a throughput optimization technique called network coding. We present in the following simulation results on the multi-flow routes co-determination, and on extensions to initial and terminal delegated nodes to extend the topology situations the algorithm may be applied to.

2.9.2 Algorithm specification

The multi flow route co-determination algorithm aims at completing routing algorithm in the flooding and information recovery phases from the egress nodes of flows to be established. The detailed description of the algorithm may be found in [10]. In Figure 9, as laptops demonstration results, a sequence diagram of the algorithm message exchanges on 3 nodes bidirectional flows establishment is shown.

In a first step, during the flooding phase of the flow route establishment, the following information is stored on the nodes crossed. For each flow, the distance measured in number of hops to the initial node and the precedent node in the path are considered. The links are assumed bidirectional and stable during the traffic establishment. A simple bounded Dijkstra algorithm may be used to implement this phase. This phase may be applied at traffic setup or by polling after traffic establishment to optimize the global traffic workload with new independent flows establishment.

In a second step the egress nodes send information to the ingress nodes to create the flow routes, catching in the relay nodes the information needed to define at the ingress node level if a network coding optimization with one other flow is applicable. This step consists of sending information on these flows, periodically from candidate egress nodes by sending specific messages.

The last phase consists on identifying the best route to send the flows at the initial node level, from all the messages received. The information of these messages contains in particular the information needed to know the traffics from other flows, to identify pivot nodes (which initiates and relay the network coding of flows), and if these pivot nodes are bidirectional.

In the flows initialization, the establishment messages set-up the paths to transmit the flows, and inform relay nodes of these path network coding is applied, identifying bidirectional ones, on which memorization, coding and decoding phases are applied.

2.9.3 Integration in OneFIT architecture

As referred in the D4.2 deliverable, the algorithm is located in the CMON module (see Figure 44). It is a distributed algorithm located on each terminal (see Figure 45). It runs as a distributed algorithm. It is activated during the ON maintenance procedure.

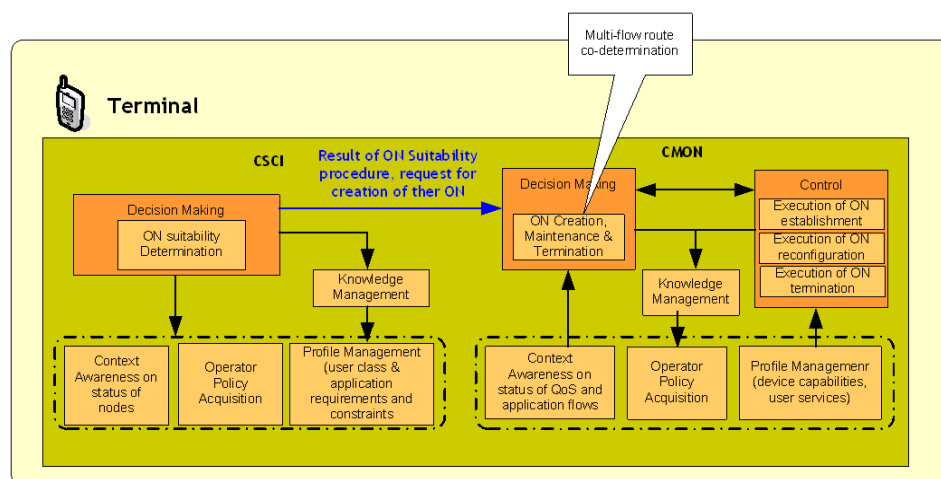


Figure 44: Integration of the algorithm in the functional architecture

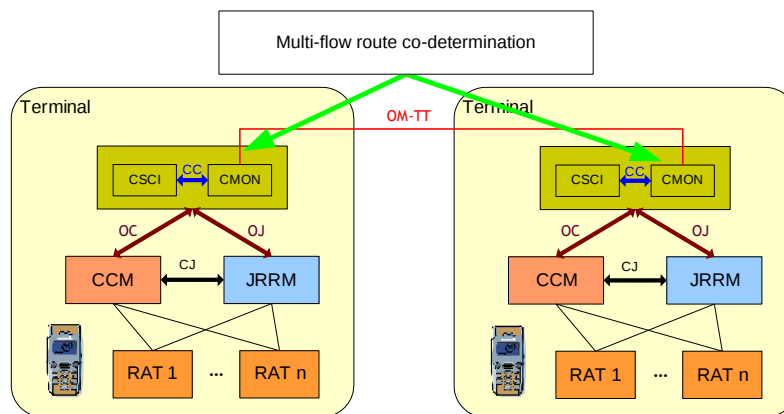


Figure 45: Integration of the algorithm in the functional architecture

2.9.4 Further performance evaluation results

The following figures recap further experiments applied on a C based simulation of the algorithms on randomized topologies. These results are the object of the publication [36]. Further details may be also found on the previous D4.2 deliverable.

Figure 46 gives for the three kind of flows (two traffics 1to2, two traffics 1to1 and one traffic 1to1), on randomized samples the traffics can be set (without optimizations), the percentage an optimization is applicable to with respect to standard routing protocols. Samples of randomized topologies of 6 to 10 nodes, on a 4×4 grid, with the capability for the (x, y) located node to communicate to the $x \pm 1$ and/or $y \pm 1$ located nodes. For each of these sets of samples have been randomized 2 samples of traffic flows between nodes, with a distance of at least 2 hops. These 2 traffic flows may be both with one or two destination nodes. In abscissa is indicated the number of nodes, and in ordinate the percentage of optimization applicability. We may see that this percentage decreases with the density of the network. The explanation is that there are less situations of flow paths crossing in a randomized draw of the samples.

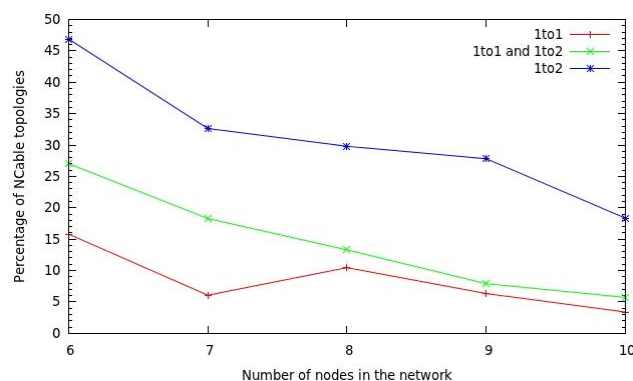


Figure 46

Figure 47 presents, for all the situations where the traffic is possible (with or without optimization), the global gain of the network coding optimizations. In abscissa is indicated the number of nodes, and in ordinate is indicated the global percentage of gains with the use of network coding optimization with delegated nodes on the global samples the traffics were effective.

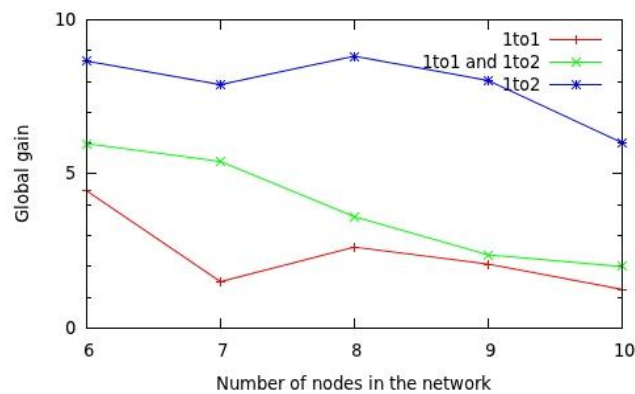


Figure 47

Figure 48 presents the gain of the delegated node optimization with respect to the initial network coding optimization algorithm proposed.

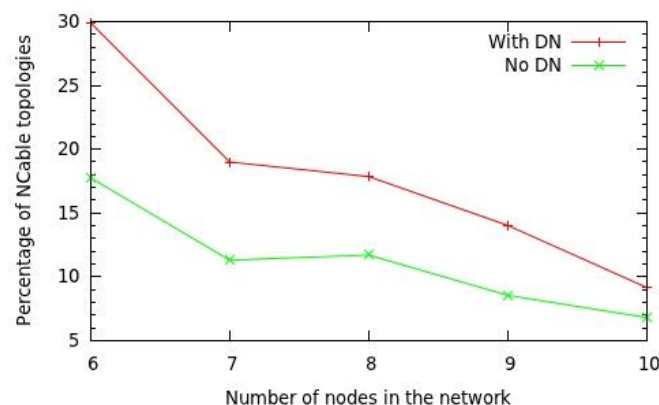


Figure 48

In the 6 last months of the project, evaluation results were provided on the evaluation of the estimate of the control signalling overhead of the use of the proposed solutions. These evaluations are studied on a 3 nodes 2 bidirectional flows example..

In this example, 20 IP packets of 1324 bytes are transmitted per second, which gives a global traffic of received and sent packets per second for the three nodes of 211840 bytes without optimization, and 185360 bytes with optimization.

The overhead in terms of IP packets header is of 40 bytes per packets (identifier of the 2 packet flows, size of the data, information in the case the packets have different size) which means an overhead of **3%**.

Per polling period, set here as 4 seconds, the three nodes send 4 flooding, 4 Topo and 4 Establish messages of respective 54, 116, 64 bytes size, for a global throughput for these messages of 936 bytes every 4 seconds.

The added throughput per second in this example is of $936 \div 4$ (extra signaling) + 40×20 (overhead in packet headers sent by node 2) each second, which equals 1034 bytes, for a gain of 26480 bytes (**ratio 3.9%**).

The overhead in terms of messages exchanged is of 12 control messages sent any 4 seconds, when 640 data messages (in the case of no network optimization applied in node 2) are exchanged (**ratio of 1.88%**),.

As the control messages already exist in routing protocol (as RREQ, RREP messages for AODV), and the size of the information added in the IP packet header may be optimized, the overheads are

overestimated, and may be optimized in further studies. However, gains obtained from this experimentation justify the use of such protocol optimizations.

2.10 Techniques for Network Reconfiguration – topology Design

2.10.1 Problem formulation and algorithm concept

This section presents further continuation of work started in D4.2 [7] section 2.11. The challenge addressed is the creation and maintenance of network topology through enabling coordination in decision making among the nodes taking into account impact of various parameters (e.g. Transmission range) to establish links/topology with a set of desired global properties and constraints. The desired properties are:

- K connectivity (for reliability, reduction in rerouting)
- Interference minimization (Capacity improvement)
- Energy efficiency (Reduced power consumption and increased network lifetime)

In order to attain the optimal solution for topology control different approaches of Network Optimization Theory are adopted; for instance convex optimization, linear programming, geometric programming and stochastic programming. The optimal scheduling of nodes and links as addressed in [9], however, achieving the optimal minimum power allocation remains an infeasible problem. The optimal solutions so far provided either are applicable to the networks where number of nodes is under six [10] or the suboptimal or approximation approach is adopted [11]. The reason of infeasibility is the non-linear nature of interference constraint (SINR) and the presence of conflicts in power assignments as the number of nodes increases. In order to address this problem, different approximation approaches are used. The column generation approach is one of the significant method opted for optimal power allocation while considering interference [12]. Besides, column generation, recently, the interference is mapped as knapsack optimization problem [13] however, it is still applied to achieve optimal scheduling instead of optimal power allocation. In [9] the optimal power allocation is provided under k connectivity and flow constraints but the interference is not taken into account.

To date, among the optimal solutions, approximation approaches and other practical methods such as graph theory, prediction and learning methods, evolutionary algorithms [14] etc. have been used to address individual problems, however combination of the constraints of k connectivity, interference handling and energy efficiency has not been addressed. The proposed algorithm captures under a single formulation the requirements on k-connectivity, interference and energy efficiency. As outlined in D4.2, section 2.11.1, the implementation of the proposed TC algorithmic approach is conducted in two phases. In phase 1, the Mixed Integer Linear Programming was used to obtain the optimal solution/values of network parameters (max. TX power per node) under k-connectivity, energy and maximum power constraints. The MILP formulation was done in three sets: continuous power levels, discrete power levels and incremental power levels.

In the second phase, the “interference” constraint, as shown below, is introduced into previous MILP formulation, making the linear formulation infeasible for the instances larger than 6 nodes.

Objective Function:

$$\min \sum_i P_t(i)$$

Subjected to:

$$\begin{aligned}
 p_t(i) &\leq p_{max} \\
 p_t(i) &\geq d_{i,j}^\alpha RX_{min}, \forall edge_{i,j} \in G \\
 \beta d_{i,j}^\alpha p_t(k) - d_{k,j}^\alpha p_t(i) &\leq -\beta N d_{i,j}^\alpha d_{k,j}^\alpha \\
 \text{if } edge_{i,j} &\in G \text{ and } edge'_{k,j} \in G'
 \end{aligned}$$

- Implementation of approximation technique: column generation, to produce near optimal guaranteed topology control solutions.
- Enhancing the search and performance of column generation through initialization from particle swarm optimization algorithm.

As, the performance of heuristic techniques cannot be guaranteed while the approximation techniques can provide a performance guaranteed, the column generation approximation is implemented here. The column generation decompose the constraint linear programming formulation into two parts that are: master problem and pricing problem. However, the main challenges of the problem formulation considered here are: first the possibility of feasible solution, second the time to find that feasible solution and the size of the problem instance.

2.10.2 Algorithm specification

The particle swarm optimization algorithm first assign the frequencies to the links , and the solution is then passed on to the constraint particle swarm formulation. The given solution is then provided to the column generation to improve the solution (by bring it closer to the possible optimal solution). The improved solution is fed again to the PSO algorithm which acts as the feedback to enhance the algorithm capability to get near optimal solution.

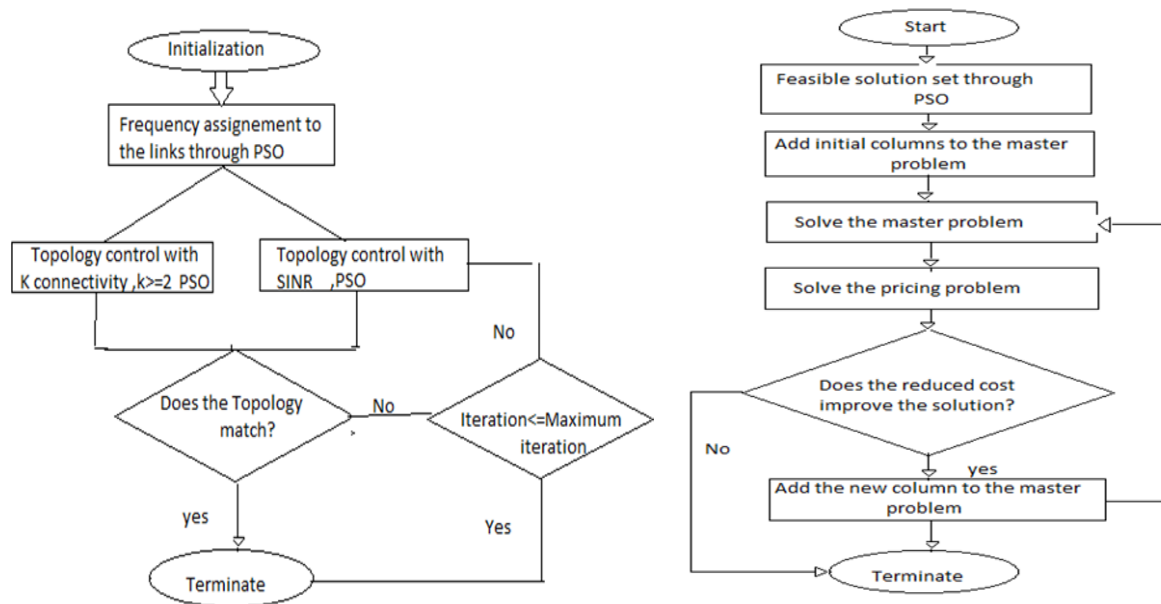


Figure 49: Algorithm flow charts

2.10.3 Integration in OneFIT architecture

The topology configuration algorithm addresses the creation and maintenance phase of the ON network. The scenario 1 , for the user case 1-4, the topology control algorithm will output the power levels which are nearest to the possible optimal solution. The algorithm will reside in the APs or any other central unit in the scenario. As, the CMON is the main entity involved in the decision making and control for ON creation, maintenance and termination phases, it will be the main entity where

topology control will reside. CMON-to-CMON communication will be via OM –TT to enable relaying, while OM-TN interface will be used in case of direct BS service. The CSCI will provide CMON with candidate nodes once ON suitability decision has been made by CSCI. The JRRM will provide neighbourhood information (through C4MS) and conditions regarding network for the second stage of distributed topology control.

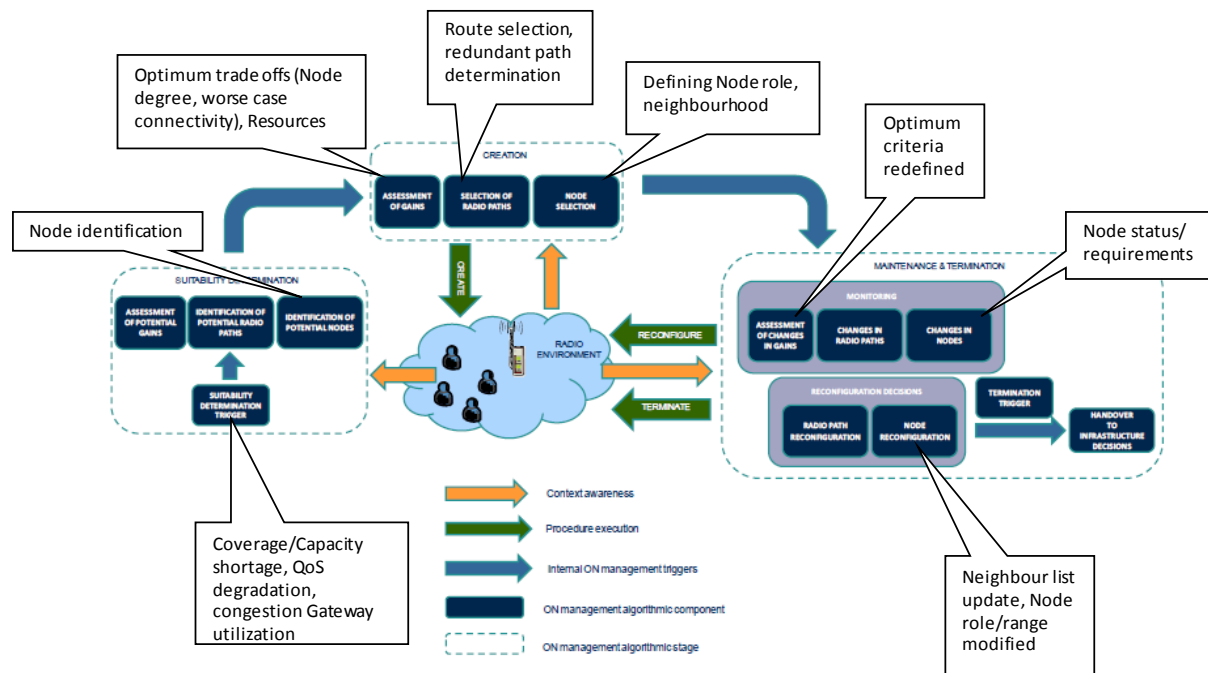


Figure 50: Placement of the TC algorithm & components in OneFIT Functional Model

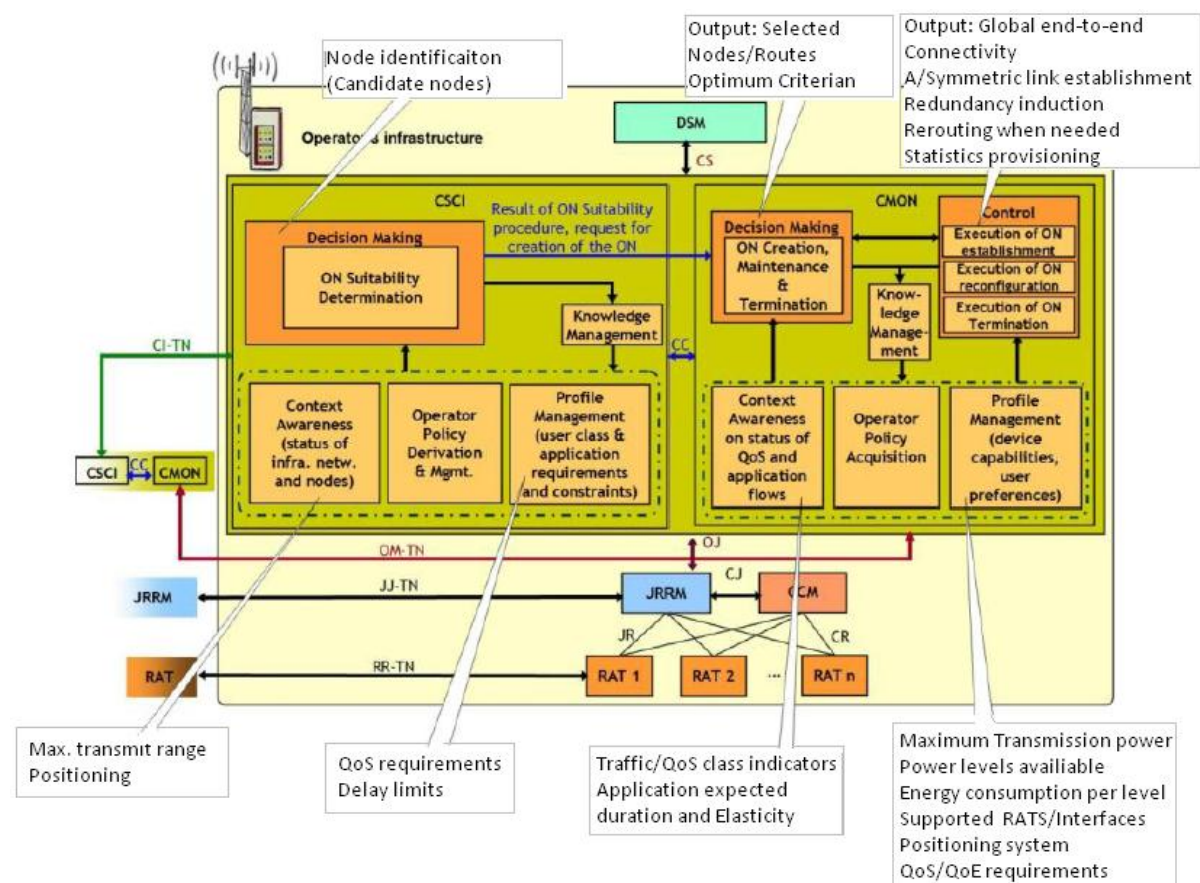


Figure 51: Placement of the TC algorithm in OneFIT Functional Architecture

2.10.4 Further performance evaluation results

The mathematical analysis is done for the phase one using IBM ILOG CPLEX. Here, optimal minimum power per node is determined under path loss and flow constraints with symmetric topology (to enable RTS/CTS procedure of 802.11 MAC). For the particle swarm optimization Matlab programming is used. The assumptions and data set include: Number of nodes 5,10,20,30, Random distribution. Free space path loss model with noise 10^{-7} and path loss parameter as 3, number of frequencies per AP is 6, full mesh initial topology graph, K demand per flow where $k \geq 2$, Maximum power per node ≤ 1 and cannot be negative, different power levels per node, stationary network and SINR Threshold per node is set as 10.

The optimal fault tolerant power levels for various values of K, as provided in phase one, will be referred here as the optimal bench mark. The heuristic algorithm that is constraint particle swarm optimization provide the power levels to column generation (the approximation technique).The Continuous power model provided least power values however it considers infinite power levels to be available at all network nodes. The Incremental power model provided minimum power levels as compared to discrete model while taking into account a finite set of available power levels in heterogeneous network. The constraint particle swarm optimization here is based on the basic particle swarm optimization and greedy approach. While the search space is explored by the particles , the unfeasible solutions are reduced by the greedy approach. The quality of our algorithm is that unlike the approximation and exact techniques where SINR constraint had to be relaxed, the PSO applied here take the un-relaxed constraints and tend to explore the space for suboptimal or pareto optimal solutions. The algorithm under SINR interference and K connectivity constraints has reduced the energy consumption by almost 80% as compared to the maximum power transmission and 60% to the random power allocation. The expended energy ratio (ERR): that is average power over all nodes to the Maximum transmission power is displayed as in the figure below.

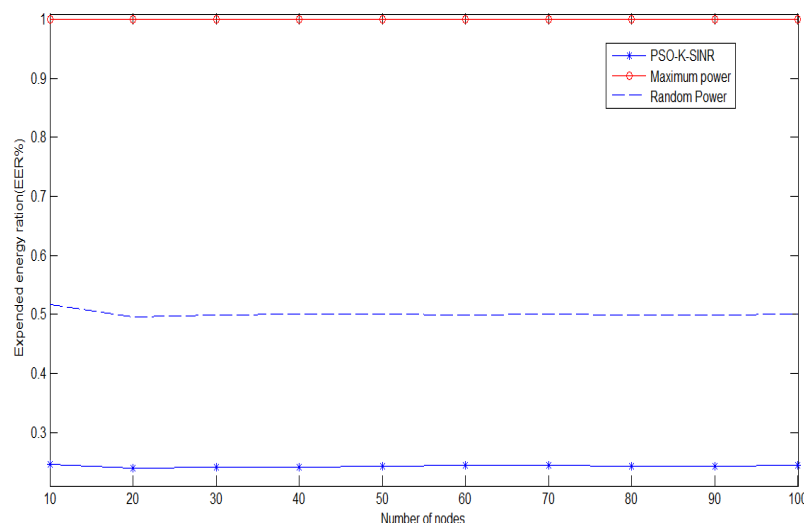


Figure 52: EER vs #Nodes

The power needed for maintaining the K connectivity tends to reduce as the number of the nodes in the network increase, however in case of SINR and K connectivity constraints together, the power required per node is more than what is required by only k connected network. On the contrary, for the smaller networks, the power required to fulfil SINR an k connectivity constraints tends to be less than k connected network, as shown in Figure 53.

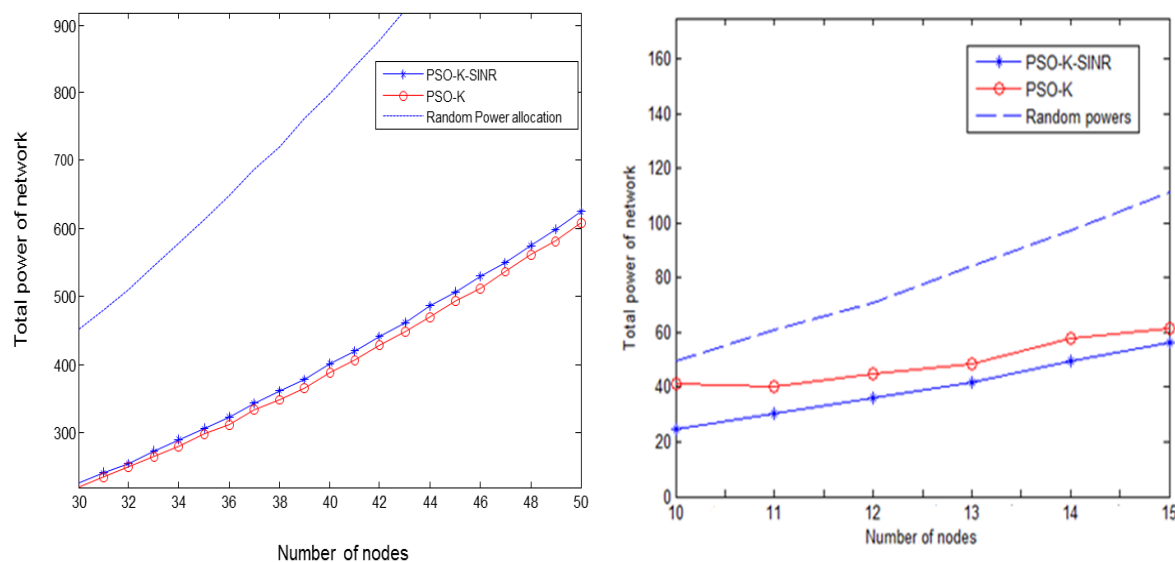


Figure 53: total network power vs # Nodes

In summary,

- In comparison to the optimal power levels (incremental model) provided by the mathematical formulation (CPLEX), the gap is at most three times. For instance the optimal power levels in case of a network of 15 nodes with k connectivity, where $k=6$, is 6.80 while PSO-K provides the power level of 27.40.
- The algorithm is robust against the initialization parameters. The initial parameters are randomly chosen and on the other hand set of extreme possible values is implemented. In both cases the curves remained in cordiality with each other. This is due to the impact of optimized adaptive inertia, social and cognition factors, maximum velocity, global and local best values.
- The implementation of column generation provides simplifications due to reduction in variables and fast computing due to separate consideration of constraints, but it also provides near optimal solutions (better approximation) for otherwise infeasible problems through exact methods.

2.11 Application cognitive multi-path routing in wireless mesh networks

2.11.1 Problem formulation and algorithm concept

This algorithm falls into the group of route selection and management algorithms. Its main function is selection and establishment of appropriate set of multiple paths in the wireless backhaul side of the wireless mesh networks (WMNs) in order to opportunistically provide aggregation of the backhaul bandwidth on the access side of struggling APs. Algorithm takes into account topology of the underlying WMN, backhaul traffic patterns, status of the WMN backhaul links and bandwidth requests at access side of the WMN APs. By providing bandwidth aggregation over multiple backhaul paths, higher levels of backhaul bandwidth utilization and load balancing can be achieved. This algorithm addresses OneFIT Scenario 5. More precisely the use case "Opportunistic backhaul bandwidth aggregation in unlicensed spectrum" is addressed.

Aspects of the algorithm are included into LCI's patent application with EU patent office: EP12181759.7 - *A method and apparatus for enabling context aware and opportunistic sharing and aggregation of resources across plurality of wireless networks.*

2.11.2 Algorithm specification

Developed algorithm makes decision on whether or not to create opportunistic network for providing multi-path routing as a solution for the detected problem. Besides having responsibility to decide when to kick in with the multipath routing, the algorithm also selects the multiple paths set which is able to provide required level of bandwidth aggregation. The net effect of opportunistic backhaul bandwidth aggregation is to match the access bandwidth of modern wireless technology with the adequate transport bandwidth in the backhaul/core network. The proposed solution makes use of OLSR as underlying routing protocol in the WMN. Necessary contextual data is gathered from WMN nodes with simple network monitoring protocol (SNMP). Decision making algorithm resides on the centralized management server, which monitors network status/state with SNMP protocol and stores gathered contextual data into a database. This database contains current and historical contextual data as well as history of previous decision instances. For more detailed description of the decision making process regarding identification of candidate paths and selection of the appropriate multipath solution please refer to section 2.12 in D4.2 [7] document.

Developed algorithm covers all ON management phases, as described within D4.2 [7] and D5.2 [12]. Figure 54 shows mapping of the algorithm onto ON management phases while Figure 55 depicts the ON suitability phase in more details.

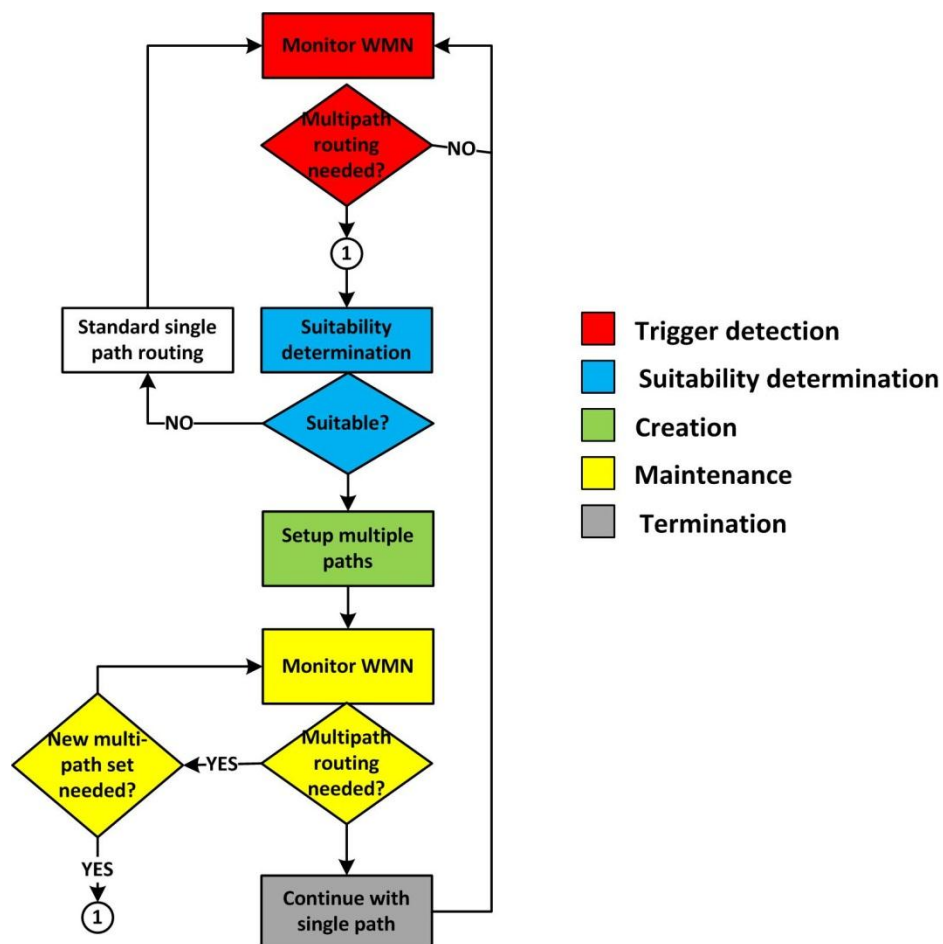


Figure 54: Application cognitive multipath routing algorithm mapped onto ON management phases

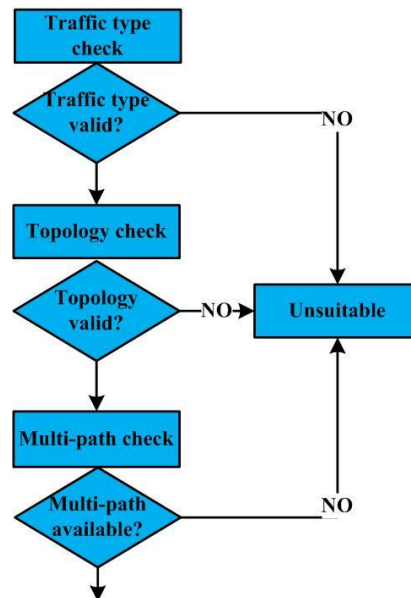


Figure 55: Suitability determination steps of the application cognitive multipath routing algorithm

In order to enable proper decision making, with respect to opportunistic management of backhaul resources, appropriate trigger recognition procedures need to be developed. Practical implementation of the multipath routing algorithm includes several techniques for autonomic trigger detection and classification. This autonomic trigger recognition process is presented in Figure 56.

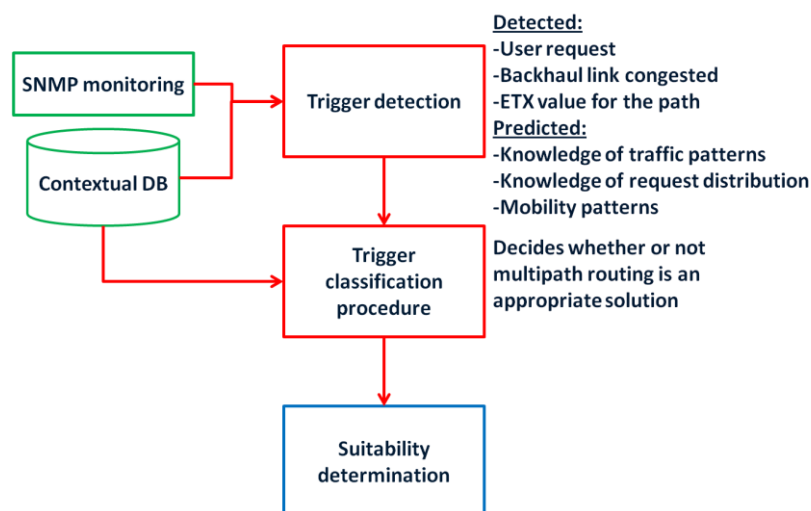


Figure 56: Autonomic trigger detection feature of the application cognitive multipath routing algorithm

Trigger detection procedure uses inputs from the network monitoring system (Simple network management protocol – SNMP, which is one variant of the C4MS implementation) for direct detection of triggers for ON suitability determination. Three distinctive triggers are taken into account as triggers which are directly detected with the monitoring system:

- User's request for new service/application;
- Congestion of the path/link in the backhaul of the WMN;
- High ETX value of a backhaul path.

In the D5.2 document [12], the web interface of the implemented algorithm shows a table containing the last obtained values of monitored parameters such as ETX, throughput and available BW on a path. These parameters are obtained every one minute over the mentioned SNMP protocol. Previous values of the monitored parameters are stored into the contextual database for further pattern recognition. Values of the parameters are constantly compared against predefined trigger levels (i.e. available BW over a backhaul path reaches less than 1Mbit/s, then every new request should be served over additional backhaul path). If QoS capabilities of certain backhaul paths drops below predefined level, then WMN backhaul paths are reconfigured (this is often supported by QoS aware routing protocols). However, for decision making from perspective of the multipath routing algorithm, measured levels of QoS related parameters are combined with recognized user request to form a trigger for ON suitability determination phase.

User's request is detected and through packet inspection approach a protocol used by the requested service is recognized. For now, the implemented backhaul resource management system, which utilizes the multipath routing algorithm, recognizes the following protocols based on packet header: HTTP, BitTorrent, VoIP, FTP and RTP/RTSP. These protocols are recognized by custom built monitoring solution based on the layer 7 monitoring protocol and Net Filter. Existence of a video delivery service is assumed. The resource management system obtains information from the video delivery service about requested level of quality for particular video request (requested video resolution).

The network monitoring system uses detected level of QoS related parameters of currently established backhaul paths and QoS requirements of the recognized end user request for service/application. Ability of the currently established WMN backhaul paths to meet QoS requirements of requested services represents trigger for ON suitability determination phase. The ON providing the multipath routing solution will be created when triggering situation is met. Triggering events can also be predicted if certain patterns in network and user behaviour are detected/recognized through analysis of the history of contextual data. This context history is available from the contextual database which is constantly filled by the monitoring system. Patterns which can be used for successful prediction of challenging situations (triggers for ON creation) are:

- Backhaul traffic patterns;
- Users' request distribution;
- Users' mobility patterns.

Contextual database is constantly updated with new values of monitored contextual parameters (QoS related parameters, status of nodes and user requests). These parameters form a contextual history which is used for recognition of patterns in changes of the monitored contextual parameters and detection of their interdependencies. Pattern recognition goes from utilization of simple statistical analysis of gathered data to applying more complex machine learning tools. Backhaul traffic pattern is recognized through correlation of gathered contextual data related to topology changes (temporal and spatial) and levels of available BW (or achieved throughput) on WMN backhaul links (these values are gathered every one minute by the monitoring system). User spatial and temporal request distribution is derived by correlating detected user's request for specific service/application, time at which the request is made, duration of service provision and AP to which the user was connected while making the request. In this way, request distributions can be made for every service/application/content. User mobility pattern can be extracted from gathered data about service provision to the user and APs to which the user connects while using the service. The users' mobility patterns affect the request distribution and backhaul traffic patterns.

In certain WMN deployments and with provision of certain services there are easily detectable and recognizable patterns in changes of related contextual parameters. However, some WMN deployments require significant calculations and data mining in order to detect and recognize any

pattern in resource utilization. Recognized triggers will kick off the ON suitability determination phase. If certain triggers find ON constantly being unsuitable for solving the challenging situation, these triggers will not be used for starting the ON decision making procedure in the future (detected situation cannot be solved through utilization of the multipath routing in WMN backhaul).

Further on, the trigger recognition process is constantly active during the ON maintenance phase. Reconfiguration triggers are defined as QoS capabilities of established WMN backhaul paths (opportunistically established multipath resource sharing and aggregation). If established multipath solution doesn't meet required QoS levels it can be reconfigured. A new suitability determination phase will be executed before reconfiguration is made. If the multipath routing solution cannot be reconfigured (ON doesn't provide proper solution for the problem), the ON will be terminated. Also, the ON will be terminated if there is no longer the need for backhaul BW aggregation/sharing through multipath routing.

2.11.3 Integration in OneFIT architecture

The algorithm described here is implemented into the open platform WMN test-bed. Test-bed and algorithm's implementation are described in detail in [12]. All performance evaluation results presented in D4.2 are obtained through experiments conducted with this algorithm implementation. The MSC of the algorithm is presented in Figure 57.

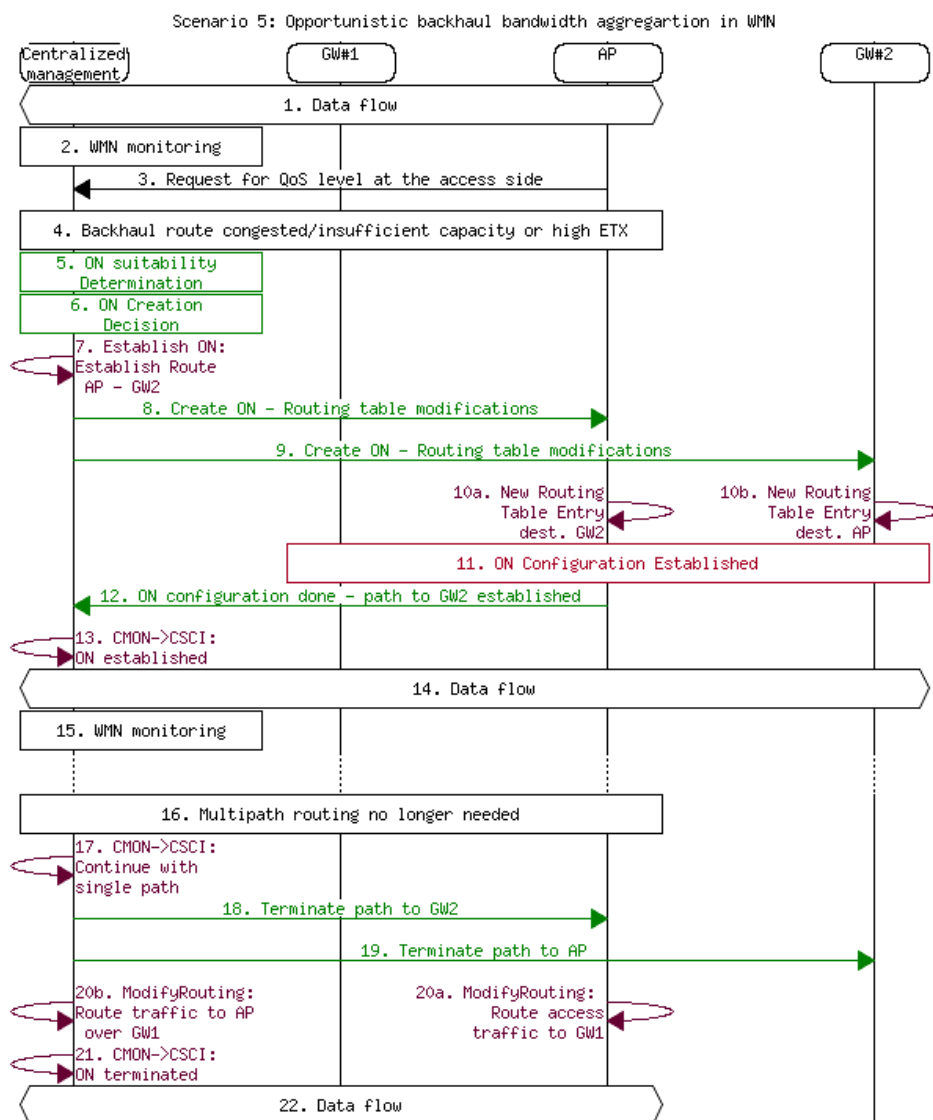


Figure 57: MSC of the multipath routing algorithm

2.11.4 Further performance evaluation results

The performance evaluation results presented in the D4.2 document give insight into performance of the suitability determination procedures of the algorithm (ability to choose appropriate multipath set for the problem at hand). Also, impact of multipath routing on QoS perceived by the end users is analysed.

Autonomic trigger recognition capability of the WMN resource management system, which utilizes the presented multipath routing algorithm, will be presented here in short. The result of user's request for service (in this example Skype call utilizing the VoIP protocol) is recognized (see *Figure 58*). This recognition is used, in combination with the detected levels of QoS indicators on the current backhaul path and QoS requirements of the recognized service/application (see *Figure 59*), for indicating that certain trigger is met. In this example, user's request for a VoIP based service is recognized. The currently established WMN backhaul path between AP, to which the user is connected, and the WMN GW has the following QoS capabilities: available BW = 2.118 Mbit/s and ETX= 2.0. The VoIP is defined as a service profile with the following QoS requirement: available BW = 0.1 Mbit/s and ETX = 1. Since the current backhaul path cannot offer required QoS level (ETX value bigger than requested), additional backhaul path needs to be established (see *Figure 59*). This path needs to be able to support requested service (ETX value less or equal to 1). The ETX value of the new path is estimated to the level equal to the average ETX over this path while it was established (if this backhaul path wasn't established before, the ETX value is configured to be 1). After the ON is created (the additional backhaul path is established), the maintenance phase starts (see *Figure 59*). After the user's request for service is finished, the packet inspection system will detect that there is no longer the VoIP service provision over the established path (see *Figure 60*). This will result in ON termination. For detailed description of the WMN test-bed, the web interface of the implemented algorithm and naming of the WMN nodes please refer to section 6.4 of the D5.2 [12].

Protocol detected
End user request for service recognized

```

root@CAMRRI:~# iptables -t mangle -L POSTROUTING -n -v
Chain POSTROUTING (policy ACCEPT 19 packets, 2065 bytes)
  pkts bytes target      prot opt in     out     source    destination
    0    0          all -- *      *       0.0.0.0/0 0.0.0.0/0
root@CAMRRI:~# iptables -t mangle -L POSTROUTING -n -v
Chain POSTROUTING (policy ACCEPT 1087 packets, 463K bytes)
  pkts bytes target      prot opt in     out     source    destination
   61 7020          all -- *      *       0.0.0.0/0 0.0.0.0/0
root@CAMRRI:~#
  
```

LAYER7 17proto skypetoskype

LAYER7 17proto skypetoskype

Figure 58 – Recognition of end user's request for service

Recognized trigger **Selected solution**

Node	Time	Required_bandwidth	Required_ETX	Available_bandwidth	Current_ETX	ON_phase	Selected_path	SP_available_bandwidth	Total_available_bandwidth	SP_ETX	FLAG
OUTDOOR	2012-05-18 15:57:46 UTC	0.1	1.0	2.118	2.0	Suitability determination			3.696		false
OUTDOOR	2012-05-18 15:57:57 UTC	0.1	1.0	3.696	2.118	Creating	BEGEJ-TISA	3.691	7.386	1.0	false
OUTDOOR	2012-05-18 15:59:42 UTC	0.1	1.0	1.66	3.253	Maintaining	BEGEJ-TISA	3.614	5.274	1.976	false

Set requirements for candidate nodes:

Node: Application: Submit

PIVA: Submit

Figure 59 – Recognized trigger for ON creation

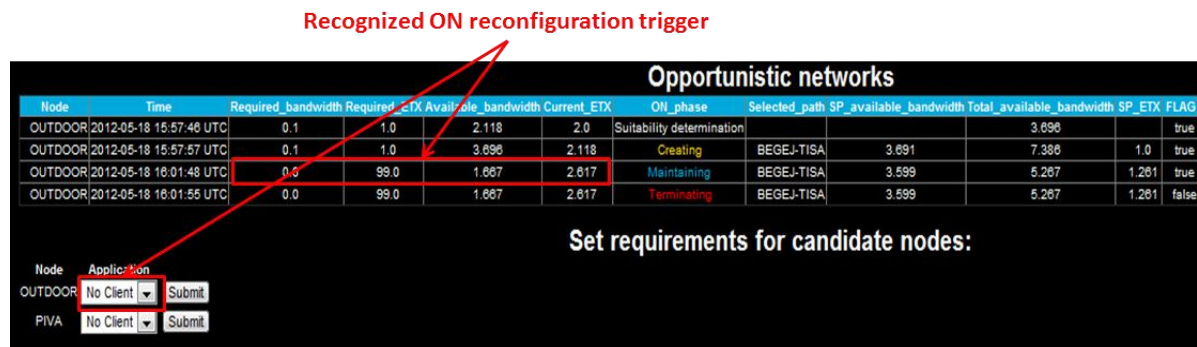


Figure 60 – Recognized trigger for ON reconfiguration/termination

The LCI performs monitoring of two WMNs in Novi Sad. One WMN (9 stations) is located at the campus of the University of Novi Sad (UNS) and the other (35 stations) covers various public areas all over the city (city squares, parks and walking areas). The UNS WMN offers internet access and several local services to university students and employees. The city WMN offers free internet access to everyone. These two WMNs have been monitored for the past 8 months. Gathered contextual parameters are stored into one contextual history database. By utilizing simple data analysis (calculation of average values, deviation and frequency of changes) and plotting various dependency diagrams, certain patterns in changes of monitored contextual parameters can be easily detected. Patterns in resource utilization are much easier detected in UNS WMN than in the city WMN. This can be explained by profile of location and users of every WMN. The UNS WMN is utilized by students and at a specific location (campus). There are many patterns in user's behaviour in this WMN (spatial and temporal request distribution for specific services). On the other hand, patterns in utilization of resources in the city WMN are far less obvious (accept basic temporal patterns, i.e. less users in early morning hours).

Regarding the UNS WMN the following statistics for the number of end users are collected:

- Total number of end users during the 8 month period (see Figure 61);
- The number of end users on one randomly selected AP during one randomly selected week (see Figure 62) and month (see Figure 63);

The total number of end users shows a typical distribution for a university campus. The summer time is the least active period (this period is used for maintenance of the UNS WMN which is shown with down time without end users during the first week of August). The randomly selected AP is located between faculty buildings (no student dorms in the vicinity). The number of end users also show typical pattern for this type of AP deployments. The number of end users drops during the weekends, holidays (12th of November is holiday in Serbia) and late night and early morning hours.

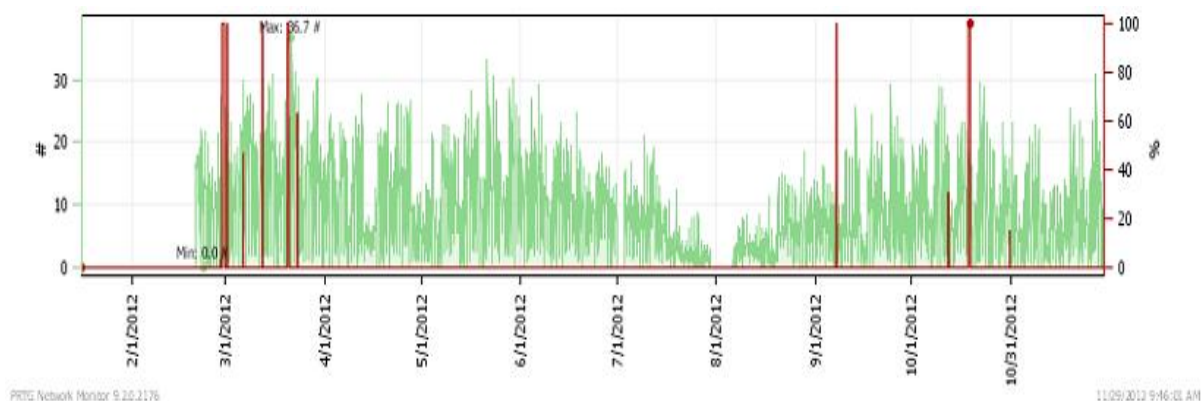


Figure 61 – Total number of end users at the UNS WMN during 8 month period

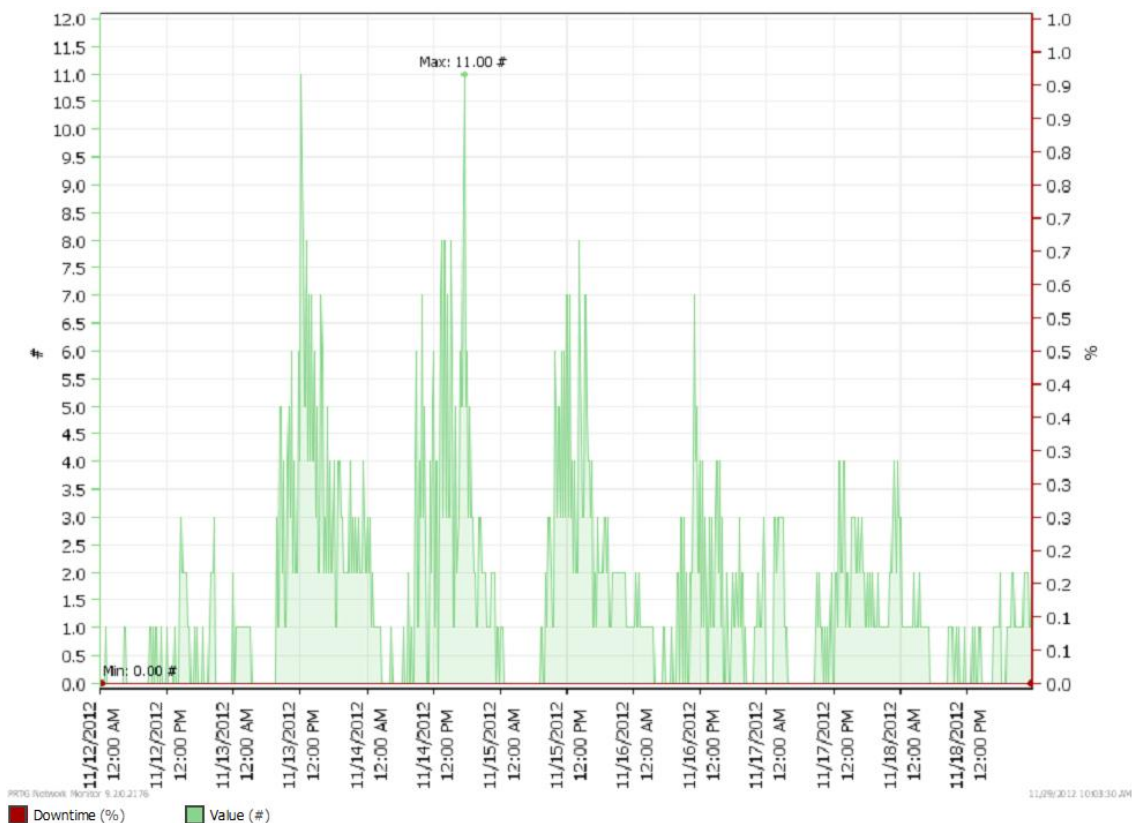


Figure 62 – The number of end users at one AP during one week

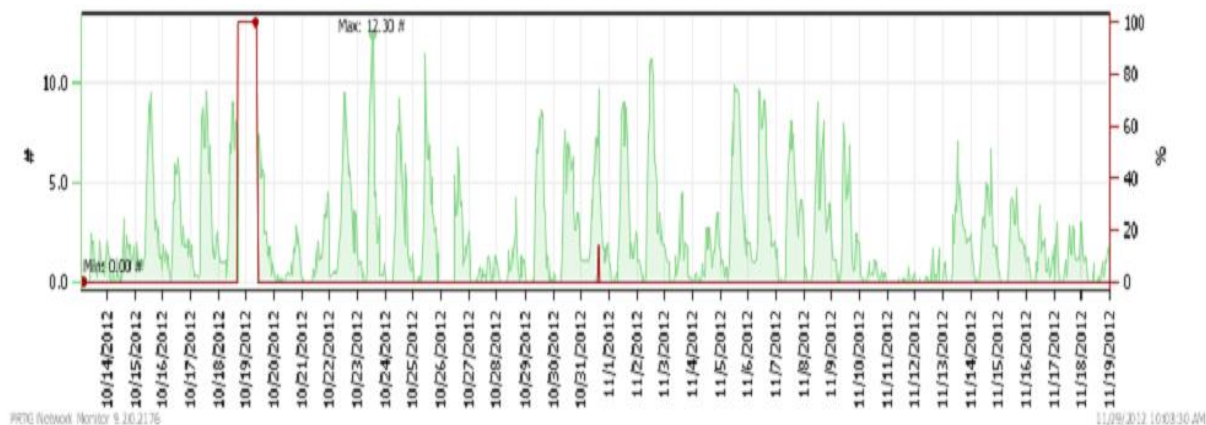


Figure 63 – The number of end users on one AP during one month

2.12.4.1 Impact of local data pre-processing on signalling load and performance of centralized algorithms

A detailed analysis of the signaling evaluation for the application cognitive multipath routing algorithm can be found in D3.3 [6]. There, the verification scenarios are described as well as the process of gathering of contextual parameters with the SNMP protocol.

The application cognitive multipath routing algorithm is a centralized algorithm, which means that all contextual parameters are gathered by the centralized management system from WMN nodes. This approach for context gathering and centralized decision making may suffer from issues which are specific for centralized management systems. Some of these issues are:

- Increased signaling load;
- Delayed detection of fast changes in values of monitored contextual parameters;

- Delayed reaction to changes in contextual parameters;
- Possibility that certain important changes in values of contextual parameters are missed;
- Single point of failure of the management process.

The signaling load related to gathering of contextual parameters (system monitoring) increases with the size of the managed system since every monitored device needs to send gathered contextual parameters to the centralized management system. In the case of multipath routing algorithm and WMNs the impact of the WMN size (number of APs) on total signaling load for system monitoring is presented in D3.3 [6]. Modern network monitoring systems gather batches of necessary contextual parameters in precisely defined time intervals (i.e. every 1 or 5 minutes). In this way the signaling overhead is controlled. However, this monitoring approach can result in delayed detection of crucial changes in values of monitored contextual parameters and, consequently, delayed response to the problematic situation. In some occasions, this monitoring approach can miss opportunity to detect fast changing parameter fluctuations.

However, the centralized context gathering, analysis and decision making has certain benefits as well. For example, centralized management system usually have enough processing power for performing complex data mining on the history of gathered contextual parameters. This data mining can result in detection of patterns in changes of contextual parameters and their interdependencies. Knowledge of these patterns can improve decision making process. Also, having all necessary information from all parts of a network allows centralized management systems to make better and more pervasive decision making. The broader knowledge about the whole network ensures that tasks of load balancing, spectrum selection and topology changes address more “global” picture of state of the network and its environment.

A hybrid network monitoring system, which combines distributed context pre-processing and centralized decision making, promise to be the best solution for cognitive and opportunistic resource management. Distributed pre-processing of monitored contextual parameters considers that WMN nodes (APs and GWs) perform local data gathering/sensing/measuring and provide initial analysis of the values of crucial contextual parameters. Two scenarios are possible:

- Contextual parameters are locally analyzed and when their value reaches pre-defined triggering level the “alarm” is sent to the centralized decision making system;
- The same as the first approach, but gathered values of contextual parameters are locally stored on WMN nodes and sent to the centralized management system periodically in order to be saved in the comprehensive database of context history.

Contextual parameters which are gathered with the network monitoring tools cover different aspects of the network and its environment. Some of these parameters (receiving and transmitting rate, number of lost packets, SINR, ETX of a path etc.) are important for resource management, some for diagnostic activities (power consumption, CPU load, interface status etc.) and some for identification of network elements and topology (IP and MAC addresses, utilization of radio channels, established links, GPS coordinates etc.). A detailed analysis of contextual parameters which can be obtained with the SNMP-based monitoring tools and form in which they are obtained is presented in the D3.3 [6].

If the first of the listed approaches for local data analysis is deployed in a WMN system, the network monitoring process will have no impact on control signaling overhead. This signaling overhead reduction may be important for networks with very limited capacity (i.e. wireless sensor networks). In this approach the centralized management system will utilize pre-defined resource management actions based on detected alarms. These alarms are sent by the WMN nodes when certain contextual parameters reach triggering values which require resource management and reconfigurations of the system in order to maintain the stability and performance. What this

approach is missing is the ability to proactively reconfigure the WMN system in order to meet certain predictable challenges which require specific resource management. The second local data analysis approach can store gathered contextual parameters in local memory of WMN nodes and send them periodically to the centralized management system. A custom script can be made so that the WMN node will send the batch of values of monitored contextual parameters when its interfaces (their capacity) are utilized beneath certain threshold (in this case the impact of the signaling overhead for gathering contextual data will be minimized), or if the internal memory dedicated to storing these data is full. These gathered contextual parameters form a central database of contextual history which is used for derivation of knowledge about utilization of network resources, patterns in changes of contextual parameters and their interdependencies. Derived knowledge gives provides the management system with ability to predict challenging situations before they become acute and proactively reconfigure the WMN system in order to meet requirements of specific challenge.

The stable version of the application cognitive multipath routing algorithm is currently implemented in completely centralized manner. However the second hybrid approach for context monitoring is under implementation in the Open platform WMN test-bed. This context gathering approach has impact on signaling overhead of the system, however through proper management of data gathering process the impact of the signaling overhead on performance of the system can be minimized. In addition, signaling overhead analysis provided in the D3.3 [6] shows that context gathering from WMN nodes has very small impact on performance of the system (signaling overhead is a very small percentage of total traffic load). Several proof of concept experiments are conducted with simple custom made scripts implemented in the OpenWRT system which is the operating system of the open platform WMN nodes.

The contextual parameters, which are currently used by the implemented algorithm for derivation of proper multipath set for backhaul BW aggregation/sharing, are:

- Link cost (ETX value);
- Signal to Interference plus Noise ratio (SINR);
- The number of transmitted and received packets - Tx and Rx;
- The number of lost packets on transmitting and receiving end;
- The number of clients;
- Used channel for backhaul communication;
- IP and MAC addresses;
- State of the interfaces and the node;
- End user's request for service/application;
- Type of the requested service/application.

The parameters: IP/MAC addresses, state of interfaces and nodes fall into the group of parameters which don't change often and therefore only changes in their values should be reported to the centralized management system. However, the address of the WMN node and its interfaces need to be sent to the centralized management system every time when the batch of gathered contextual values is sent in order to enable identification of the source of the received contextual data.

The link cost (Expected Transmission Count – ETX) is derived by the underlying OLSR single path routing protocol. This parameter is used by the multipath routing algorithm to decide whether or not a backhaul path in the WMN satisfies QoS requirements of a recognized service/application. Services such as video streaming and VoIP have specific requirements for allowed ETX values for their proper provision. The WMN nodes will report the ETX value of WMN backhaul paths/links

when it increases over a defined threshold (i.e. $ETX=2$). The ETX related alarm will be reported over the SNMP trap message which is sent to the centralized management system which knows that certain WMN backhaul path has ETX values which don't allow it to transmit certain class of packets (belonging to ETX sensitive services). ETX values are locally stored in a WMN node. The ETX value can be locally saved as many times as the local memory of WMN nodes allows. Also the gathered ETX values can be analyzed every 1 minute (this is the interval of context gathering in the current management system implementation) and simple statistics like median, standard deviation, min and max values can be derived and sent to the centralized management system, which gives a lot more information about the ETX contextual parameter than a single ETX value. In this way the signaling overhead is slightly increased. Since the ETX value is presented as a real number, its size in the memory is between 4 and 8 bytes. The OpenWRT system on open platform WMN APs uses real numbers with double precision which means that ETX values are 8 bytes in size. This leads to conclusion that local data pre-processing increases the amount of data which is sent every one minute from 8 bytes (single ETX value) to 32 bytes (median, standard deviation, min and max ETX values for the interval and in addition time stamps for the min and max values). For the purpose of experiments and plotting the results the current ETX value is also sent in the ETX batch. However, in practice this value will not be sent since the SNMP trap message (alarm) will be generated whenever the ETX threshold is reached.

The *Figure 64* shows how the centralized resource management system approximates changes in ETX values based on ETX values gathered every one minute during the period of 10 minutes. The interference in the WMN test-bed is intentionally increased so that ETX values are more likely to reach a defined threshold of $ETX=2$. Through gathering ETX values every one minute the resource management system has detected two triggering events based on increased ETX value. The *Figure 65* shows the ETX change approximation when min and max ETX values are sent for the 1 minute time interval when the current ETX value is measured. It can be seen how ETX value changes more erratically when compared with *Figure 64*. The important difference is that *Figure 65* shows that 2 additional triggering events are recognized. There are two new triggers in time intervals between the 7th and 8th minute and the 8th and 9th minute. These triggers are identified through local data pre-processing and triggers are reported on time through SNMP trap messages. By inclusion of median and standard deviation for every 1 minute slot, the plot diagram of the ETX dependency in time would be more precise and more accurate triggering moments would be presented.

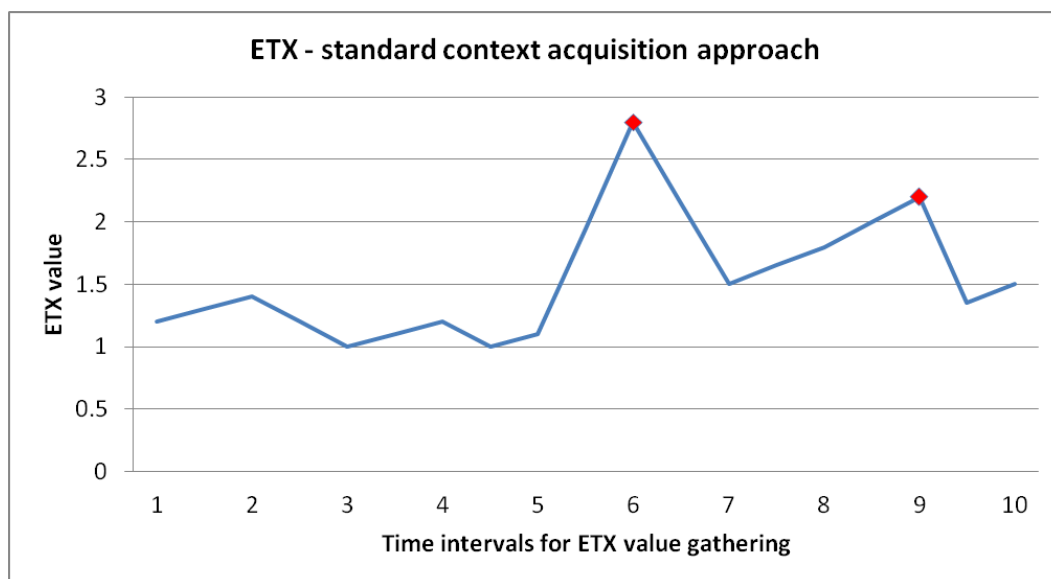


Figure 64 – Standard approach for ETX value gathering with detected threshold breaches

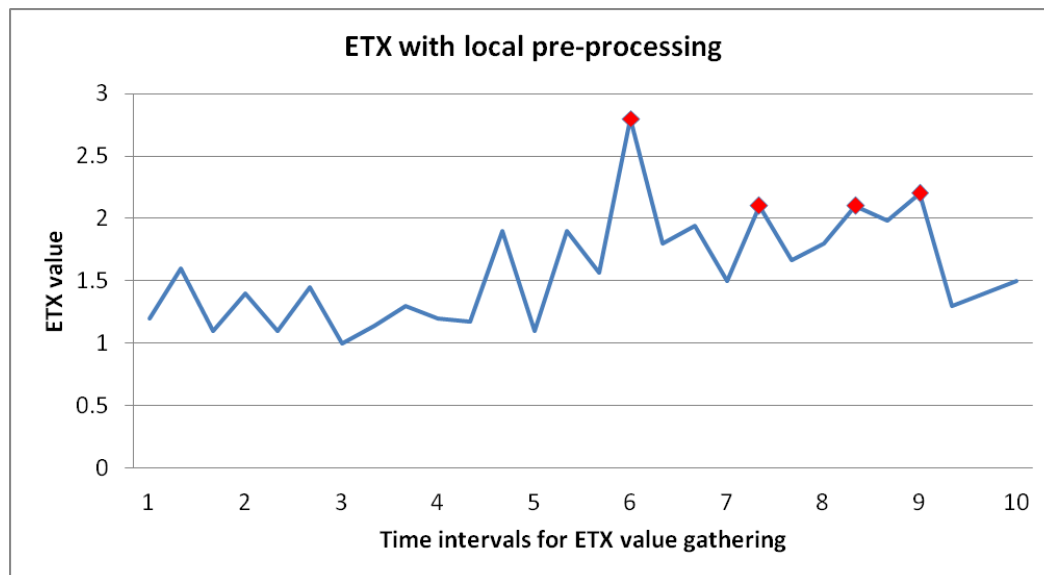


Figure 65 – Impact of local data pre-processing on accuracy of ETX awareness

The similar analysis conducted for the ETX parameter is valid for other parameters (SINR, Tx/Rx, packet drop count and used channel) as well.

The SINR parameter is gathered from every wireless interface of a WMN node. It is a standard parameter used for rate adaptation performed by the 802.11 MAC and PHY protocols. It is used for deciding whether the increased packet drop rate is a consequence of a congested link, or interference (link quality degradation). It can be gathered and pre-processed similarly as the ETX parameter. However, when the increased packet drop rate is detected (the corresponding alarm is triggered) the packet drop value is sent together with the current SINR value in order to recognize the source of the packet drop. The SINR has several triggering thresholds. A certain SINR value corresponds to a certain maximal bit rate achievable over a link utilizing a specific 802.11 protocol (for listing of SINR and corresponding bit-rate values please address the 802.11 standard). SINR fluctuations between these defined thresholds trigger alarms which indicate to the centralized decision making system that the supported capacity of certain WMN backhaul link has changed.

The Tx and Rx parameters indicate the load over interfaces/links of a WMN node. By measuring their values the resource management system can assess the load over every backhaul link and the level of their congestion (available capacity for transmitting new data flows). Triggering thresholds are set to values slightly less than maximal capacity of a link, which is derived from the SINR of a link and 802.11 protocol specifications. Therefore the thresholds are modified whenever the maximal capacity of a link changes due to interference or link quality degradation (appearance of a temporal obstacle). Local (on WMN nodes) assessment of the load over a backhaul interface provides better responsiveness with respect to detection of congested backhaul links.

The contextual parameters related to the end users (the number of clients per WMN node, clients request for particular service/application provision) are very important for proper decision making and for providing adaptivity to the resource management. Total number of active end users per WMN node (AP) and changes in value of this parameter are assessed through combination of centralized and distributed context acquisition. The Centralized monitoring system detects every end user's request since every request needs to be routed through the WiFi controller/management system. When end user disconnects from the network without proper logout, the WMN node detects that the link towards the end user is no longer active and sends SNMP trap message to the centralized management which updates the user distribution. User's request for service/application can be detected and recognized by the WMN node through implementation of layer 7 based packet inspection, which is previously described. However, this packet inspection is very stressful for a CPU

of a device which performs it. Centralized resource management server has enough processing power to perform packet inspection and detect patterns which are used for identification of requested service/application. However, a WMN node uses most of its processing power for performing networking tasks and layer 7 packet inspection of every request and new data flow could lead to major degradation of WMN's networking performance. User's requests are routed to the central part of the network (at least to a 3rd layer switch or WiFi controller). This fact is used for recognition of user's request through utilization of packet inspection techniques. Certain WMN deployments are offering limited number of services to their clients (i.e. university WiFi network can provide access to the video streaming server containing recorded lectures, access to a local database and access to the internet). In these networks the internet traffic is treated as best effort service while i.e. video streaming of lectures is treated as a "premium" service which needs to be provided with required level of QoS. Therefore, the university WMN can be managed in a way which provides backhaul BW aggregation only for services of lecture video streaming. In these examples, where services are easily distinguished, the packet inspection can be performed on WMN nodes without significant impact on the performance of the networking process. For this purpose the packet inspection process can be optimized so that it detects only the packet pattern which corresponds to the "premium" service. In this way the request recognition process can be distributed, which offloads the burden of packet inspection from a centralized management system to the WMN "cloud".

Obtained results from performed proof of concept experiment provide conclusions that the hybrid context acquisition approach promises the best tradeoff between signaling overhead and responsiveness of the decision making process. Local data pre-processing on WMN nodes will be further developed and implemented as part of the context aware, opportunistic and knowledge based resource management for WiFi networks which is in development by LCI.

2.12 UE-to-UE Trusted Direct Path

2.12.1 Problem formulation and algorithm concept

A WLAN using Wi-Fi technology can be obtained by interconnecting a set of " N " candidate stations (STA or STations) through an AP (Access Point). For Wi-Fi communications, the AP can perform a channel selection in order to choose the best channel to communicate with all STAs. This is usually performed by measuring directly the quality of the transmitted signals (received power, interference, BER (Bit Error Rate), PER (Packet Error Rate) etc.).

As described by prior art ([29] and references therein) the AP is fixed and cannot be dynamically changed. In 802.11 WLAN for instance, there is a clear distinction between an AP and a STA. Both are different physical equipment and have different capabilities. Also, the best channel choice is generally done only after the measurements are performed in order to indicate the quality of the channel; these measurements are done while the WLAN is already established. The prior art (e.g. [29]) already investigates the WLAN channel selection without communications; however these works assume that the AP is fixed and already known in advance. This aspect will be different from the algorithm proposed in this section, where the AP function and the best communication channel can both be dynamically allocated within the group to form the WLAN.

Here, the goal is to establish a WLAN (Wireless Local Area Network) between N devices (System 1) by using communications allowed by other technology (System 2) (see Figure 66). The N devices or users have double communication capabilities, since they can communicate using two different technologies (the one used by System 1 and the other used by System 2). System 1 can use Wi-Fi (i.e. terminal devices are stations or STA) and System 2 can use LTE (i.e. terminal devices are User Equipment or UEs). In the following, the notation UE/STA is used to emphasize the double functionality of a device: in the same time the device is a UE (belonging to LTE system) and a STA (belonging to the WLAN).

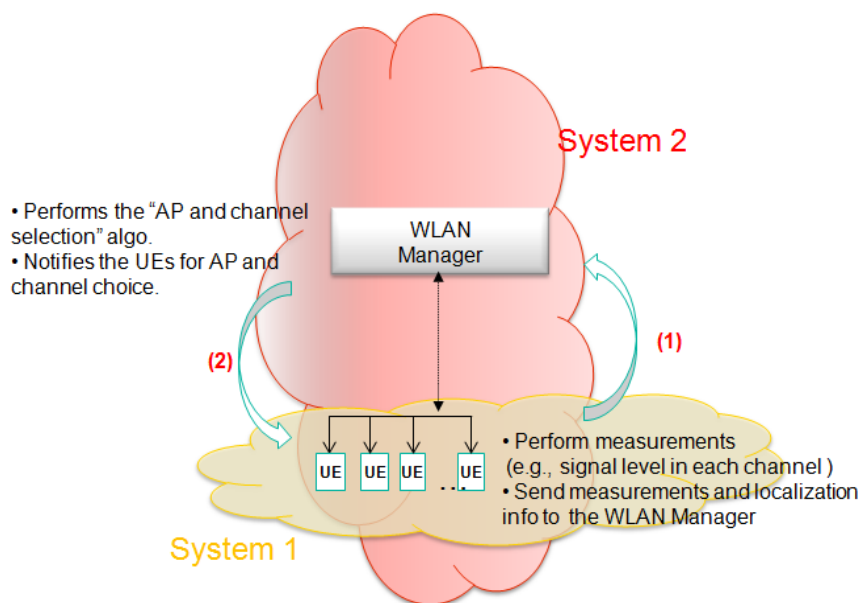


Figure 66: AP and Channel Selection during the creation of a WLAN

At the beginning, we consider that System 1 is not yet communicating, so our goal is to predict or to estimate the best channel and AP before the actual communication takes place, instead of the classical approach of directly measuring communication quality from System 1. Here, the approach is different from prior art for the following reasons:

- The choice of the combination best channel / best AP is based on prediction and not on measurements from the established WLAN. The only measurements needed here are made before the WLAN communications (before System 1 starts communication).
- Any station belonging to set N can be chosen to be the AP, and there is a set of M available channels that may be used. This means that both AP and communication channels can be dynamically set. A priori, any of these M channels may be used by other networks or systems nearby, so System 1 would have to share one of these channels with other incumbent systems.
- The selection of the best AP and channel used by System 1 is based on the information received through System 2.

Moreover, the WLAN establishment (i.e. System 1) is under the following criteria:

- 3GPP QCI (QoS Class Identifier) requirements for all the WLAN members. Particularly, we use Per-link (between two STA) PER requirements (3GPP QCI provides PER requirements).
- Mean transmit power minimization requirements in the WLAN.

More specifically, the problem can be formulated as follows: given a set of M opportunistic channels and a set of N candidate stations, how to choose an AP and an operating channel that allow fulfilling the aforementioned criteria?

The underlying scenario is depicted in Figure 66, we face the problem of a WLAN creation controlled by a WLAN Manager. The WLAN Manager is aware of the localization information of all the candidate stations via a cellular network for instance.

Please note that the WLAN Manager could be, for instance, an entity connected to MME (3GPP entity) and has access to the localization information from a network entity called E-SMLC (Evolved-SMLC). In this case, the communication between the N candidate stations (UE/STA) is performed

over LTE since it is considered that the N UE/STA have double communication technology. The WLAN Manager assigns the role of the AP to the selected terminal (and/or STA to unselected terminals) and triggers operation of the WLAN on the selected channel. In our proposal the latter function is done using LTE again. After this WLAN network is formed, the candidate stations will then start communicating through the AP and the channel proposed by the WLAN Manager. Therefore, the WLAN (System 1) is formed with a help of another system or technology (System 2) without a prior communication between the WLAN members.

Furthermore, the WLAN Manager is in charge of computing the AP and channel selection algorithm that we are proposing. The UE/STA essentially make measurements in the available channels and report to the WLAN Manager. After running the AP and channel selection algorithm, the WLAN Manager informs the UE/STA with the AP and channel choice for WLAN establishment.

It is worth noting that the objective here is only the initial selection of the AP and an operating channel. When the network changes (association of new STA, disconnection of existing stations or variations of channel conditions), another dynamic channel selection algorithms should be used to maintain the network service demand.

2.12.2 Algorithm specification

The proposed algorithm relies on the 3GPP LTE QCI, [30], to predict the most power efficient couple (i.e. AP, Wi-Fi channel) that allows reaching given QCI requirements in a per-link basis (per-link QCI requirement). More specifically:

- We use the PER requirements from the various 3GPP QCI, and the modulation characteristics of the given WLAN technology (e.g., IEEE 802.11b uses DBPSK (Differential Binary Phase Shift keying), DQPSK (Differential Quadrature Phase Shift Keying) and CCK (Complementary Code Keying) modulations) in order to derive the links' required SNR (Signal-to-Noise Ratio) noted by SNR_{req} .
- By considering the fact that, to reach the per-link PER the mean SINRs (Signal-to-Interference-plus-Noise Ratio) of all the links must be greater than the aforementioned SNR_{req} , we derive the minimum required mean transmit powers that the various transmitters should use to communicate to the various receivers in the WLAN.
- By considering that all the N STA can be used as AP, and that all the M opportunistic channels can be used for operating, we go through the set of all possible (AP, Wi-Fi channel) couples to find in each case the maximum of the aforementioned minimum required mean transmit powers. The most efficient couple (AP, Wi-Fi channel) is the configuration with the lowest of the "maximum of the required mean transmit powers".
- At the end, the first proposed solution allows finding, if it exists, a configuration that meets the per-link PER requirements. This algorithm does not change the transmit power of the UE/STA.
- The second proposed solution returns the most power efficient couple (AP, Wi-Fi channel) and the required mean transmit powers for all the links in the WLAN. This is useful when the transmit power of the UE/STA can be changed accordingly to meet a particular QoS.

For the two algorithms we will propose further, we compute the required SNR, SNR_{req} , as follows:

- For given QCI, we retrieve the corresponding PER.
- Using the error independence assumption, we derive the BER from the PER as follows:

$$\text{BER} = 1 - (1 - \text{PER})^{\frac{1}{L_{\max}}}$$

Where $L_{\max}$ is the maximum length of the packet (maximum number of bits contained in a packet).

- Considering the modulation type of the WLAN technology (e.g., IEEE 802.11b uses DBPSK (Differential Binary Phase Shift keying), DQPSK (Differential Quadrature Phase Shift Keying) and CCK (Complementary Code Keying) modulations) we retrieve the BER-SNR characteristic of the WLAN technology under AWGN (Additive White Gaussian Noise) assumption. The BER-SNR characteristic of given modulation in AWGN environment is the curve of BER versus E_b/N_0 , where E_b is the energy per bit, and N_0 is the noise power spectral density (noise power within a 1 Hz bandwidth). For most of the current modulations, the BER-SNR characteristic in AWGN environment is well known and is widely given in the literature (see [31]-[33]) Using the BER-SNR characteristic and given the BER value, we derive the value of E_b/N_0 .
- The required SNR, SNR_{req} , is then given by the following formula:

$$\text{SNR}_{\text{req}} = \frac{R_b}{W} \times \frac{E_b}{N_0}$$

Where R_b is the bit rate of the WLAN technology and W is the communication bandwidth of the WLAN technology.

The flowchart describing the SNR_{req} derivation is depicted in *Figure 67*.

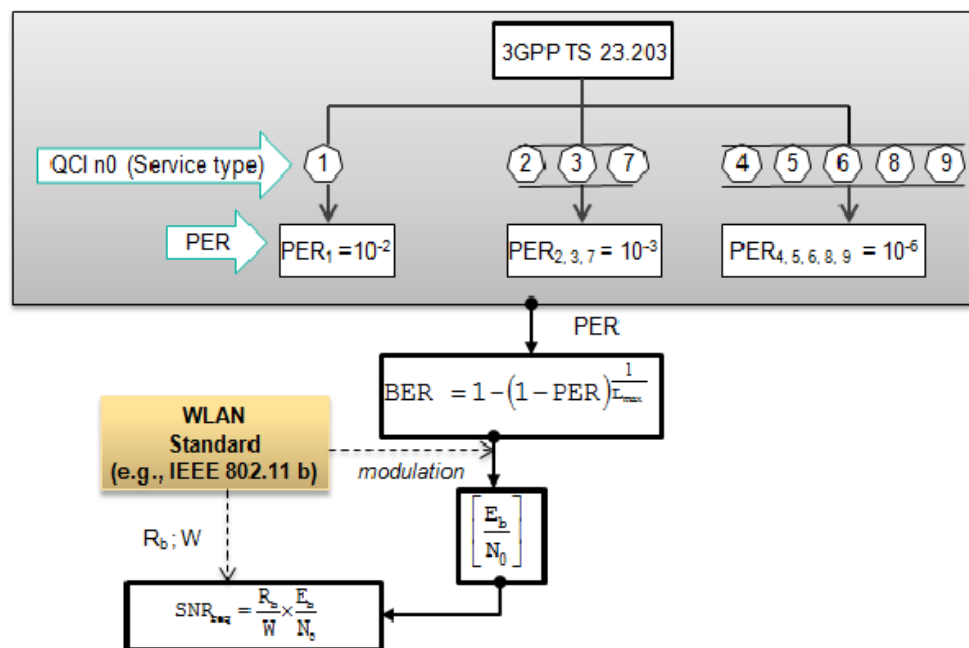


Figure 67: Flowchart of SNR_{req} derivation

Based on the above required SNR to reach a particular 3GPP QCI, the next two sections show two novel algorithms for WLAN configuration which jointly select an AP and allocate an operating channel for WLAN communications.

2.12.2.1 QoS-Constrained Joint Access Point Selection and Channel Allocation for WLAN Configuration

In this section, we present a novel joint AP selection and channel allocation for a WLAN configuration. The proposed solution is given under 3GPP QCI requirements. We consider a set of M opportunistic channels and a set of N candidate stations attempting to set a WLAN for communications. The goal is to select a couple (AP, Wi-Fi channel) that meets a particular QoS requirement. Moreover, each candidate station can be chosen as the AP, and each channel can be chosen as the operating channel. The QoS requirement consists in meeting the per-link PER requirements associated to a particular 3GPP QCI. Therefore, the proposed algorithm runs to find a (AP, Wi-Fi channel chn) configuration where all the per-link PER requirements are fulfilled. An example of configuration is given in Figure 87.

Given a particular link, we check the PER achievability as follows:

- We estimate the mean SINR of the link from UE/STA i to UE/STA j , operating in channel chn , as follows:

$$\text{SINR}(i, j, chn) = \frac{P_{Tx}(i, j) \times G_{i,j} \times H_{i,j} \times d_{i,j}^{-\alpha}}{N_0 \times W + I(j, chn)}.$$

Where $P_{Tx}(i, j)$ is the transmit power of UE/STA i to UE/STA j , and $I(j, chn)$ is the interference level in channel chn measured at receiver j . The term $N_0 \times W + I(j, chn)$ represents the harmful signal power in channel chn at received j ; this is to be measured by UE/STA j (before WLAN establishment) and sent to the WLAN manager. The term $G_{i,j}$ represents the antenna gain accounting for the antenna gain of both the transmitter i and the receiver j . $H_{i,j}$ represents the mean gain of the communication channel between transmitter i and receiver j . The distance between UE/STA i and UE/STA j is represented by $d_{i,j}$, and the path-loss exponent of the communication environment is noted α .

- If $\text{SINR}(i, j, chn) \geq \text{SNR}_{req}$ then PER requirement for link UE/STA i to UE/STA j in channel chn , can be achieved. Otherwise, the PER requirement cannot be achieved.

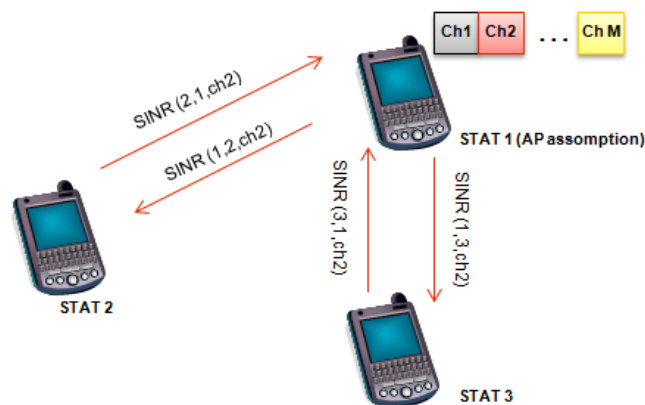


Figure 68: Links mean SINR for given AP and channel assumptions.

The flowchart for the per-link PER achievability checking is represented in Figure 69.

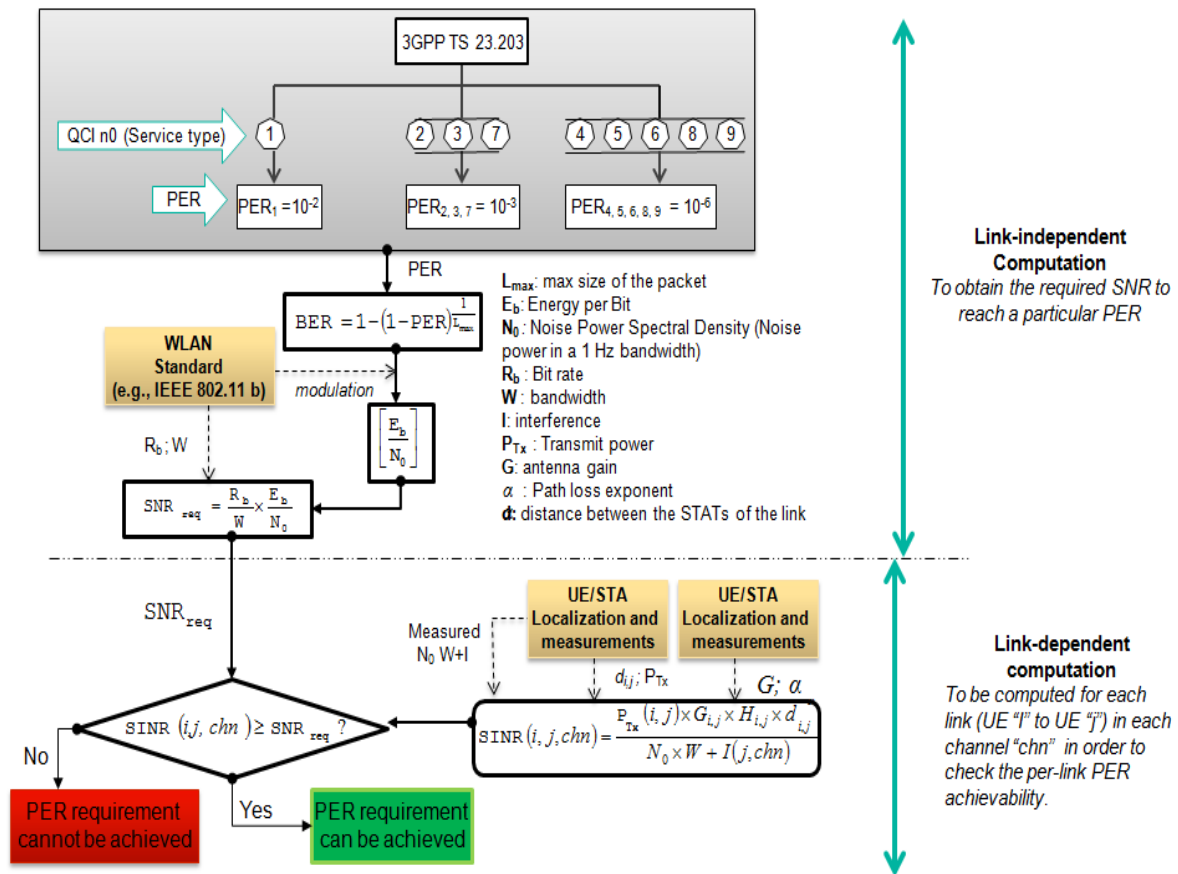


Figure 69: Flowchart of SNR_{req} derivation, and of the per-link PER achievability checking

As mentioned before the goal of our algorithm is, on the one hand to choose a UE/STA as AP, on the other hand to select an operating channel. The joint AP selection and channel allocation must meet the per-link PER requirements. Therefore:

- A given (AP, Wi-Fi channel) configuration is compliant if and only if the per-link PER requirement is achieved for all the links (uplinks and downlinks).
- It may be that there is no (AP, Wi-Fi channel) configuration where the above criterion is verified. In this case, the WLAN cannot be configured to meet the QCI requirements.

The general flowchart of the proposed WLAN configuration algorithm is depicted in Figure 70.

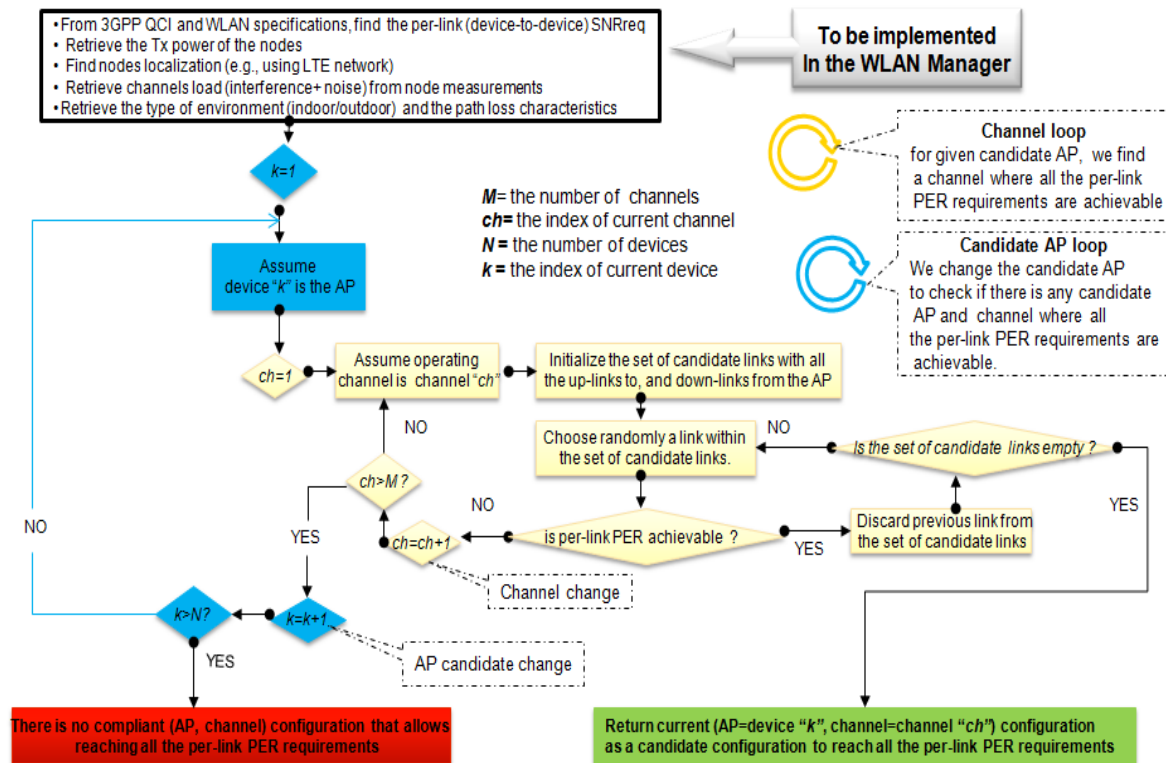


Figure 70: Flowchart of the joint AP and Channel selection for WLAN establishment: 1er algorithm. This algorithm allows finding one compliant couple (AP, Wi-Fi channel) that meets the per-link PER requirements.

2.12.2.2 QoS-Constrained Power Efficient Joint Access Point selection and Channel Allocation for WLAN Configuration

In this section, we present another novel joint AP selection and channel allocation for WLAN configuration. The proposed solution is given under 3GPP QCI requirements.

The current approach assume that the transmit power of the N UE/STA can be changed accordingly to meet given QoS. The algorithm runs to find the most power efficient configuration that meets the per-link PER requirements associated to a particular 3GPP QCI, it also returns the minimum mean transmit powers for all the links to meet the requirements.

Given a particular link, we compute the minimum required transmit power, to meet the per-link PER requirement, as follows:

- We ensure that the mean SINR of a particular link from UE/STA i to UE/STA j , operating in channel chn , is higher than the required SNR, SNR_{req} :

$$\text{SINR}(i,j,chn) = \frac{P_{\text{Rx}}(i,j,chn)}{N_0 \times W + I(j,chn)}$$

$$\text{SINR}(i,j,chn) \geq \text{SNR}_{\text{req}}$$

$$\Rightarrow P_{\text{Rx,min}}(i, j, \text{chn}) = \text{SNR}_{\text{req}} \times [N_0 \times W + I(i, j, \text{chn})].$$

Where $P_{rx}(i, j, chn)$ is the mean received power from UE/STA i to UE/STA j in channel chn , and $I(j, chn)$ is the interference level in channel chn at receiver j . The term $N_0 \times W + I(j, chn)$ represents the harmful signal power in channel chn at received j , this is to be measured by UE/STA j and sent to the WLAN manager. The term

$P_{Rx,min}(i, j, chn)$ is the minimum required mean received power (from UE/STA i to UE/STA j in channel chn) in order to reach the PER requirement.

- The minimum required mean transmit power, $P_{Tx,min}(i, j, chn)$, is given from the minimum required mean received power $P_{Rx,min}$ as follows:

$$P_{Rx,min}(i, j, chn) = P_{Tx,min}(i, j, chn) \times G_{i,j} \times H_{i,j} \times d_{i,j}^{-\alpha}$$

$$\Rightarrow P_{Tx,min}(i, j, chn) = \frac{P_{Rx,min}(i, j, chn)}{G_{i,j} \times H_{i,j} \times d_{i,j}^{-\alpha}}.$$

The flowchart for the minimum required transmit power derivation is given in Figure 71.

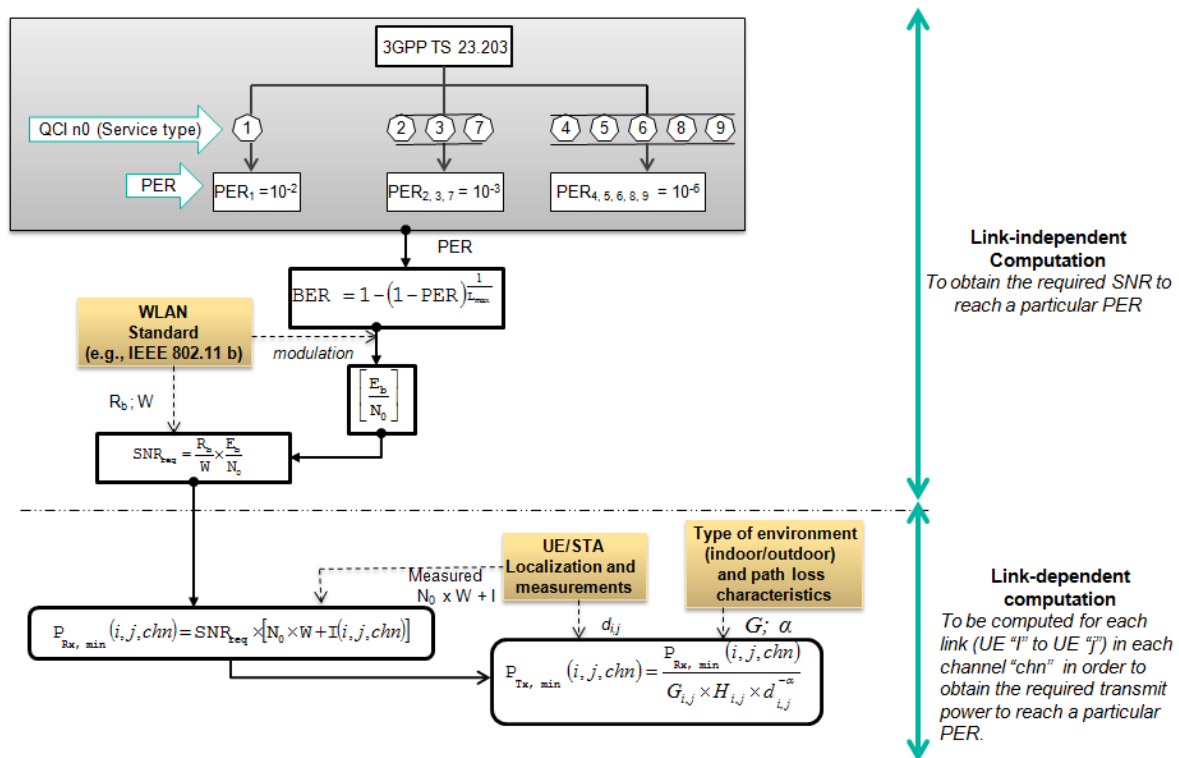


Figure 71: Flowchart of minimum required transmit power derivation

Here, the joint AP selection and channel allocation is the most efficient configuration compared to the others. We find the most power efficient couple (AP, Wi-Fi channel) as follows:

- For each AP assumption, let UE/STA k , we find for each operating channel assumption, let channel chn , the maximum of the set of links' minimum required mean transmit powers:

$$\text{Max}_{Tx,min}^{k,chn} = \text{Max}_{(i,j) \in \{(k,j); j \neq k\} \cup \{(i,k); i \neq k\}} \{P_{Tx,min}^{k,chn}(i, j, n)\}.$$

Where $\text{Max}_{Tx,min}^{k,chn}$ represents the maximum of the set of links' minimum required mean transmit powers, UE/STA k being assumed to be AP and channel chn being assumed to be the operating channel.

- We obtain the most power efficient channel, for UE/STA k assumed to be the AP, by finding the channel that exhibits the minimum of all the $\text{MaxP}_{\text{Tx, min}}^{k, \text{chn}}$ with $\text{chn}=1,2,\dots,M$, as follows:

$$\left[\text{MinMaxP}_{\text{Tx, min}}^{k, \text{Ch}_{\text{opt}}(k)}, \text{Ch}_{\text{opt}}(k) \right] = \text{Min}_{\text{chn}=1,2,\dots,M} \left\{ \text{MaxP}_{\text{Tx, min}}^{k, \text{chn}} \right\}$$

where $\text{MinMaxP}_{\text{Tx, min}}^{k, \text{Ch}_{\text{opt}}(k)}$ represents the minimum of the $\text{MaxP}_{\text{Tx, min}}^{k, \text{chn}}$ with $\text{chn}=1,2,\dots,M$. The term $\text{Ch}_{\text{opt}}(k)$ represents the index of the channel which exhibits minimum $\text{MinMaxP}_{\text{Tx, min}}^{k, \text{Ch}_{\text{opt}}(k)}$.

- The best AP assumption is obtained by finding, among all AP assumptions, the AP assumption which exhibits the minimum of $\text{MinMaxP}_{\text{Tx, min}}^{k, \text{Ch}_{\text{opt}}(k)}$ for $k=1,2,\dots,N$.

$$\left[\text{MinMinMaxP}_{\text{Tx, min}}^{k_{\text{opt}}, \text{Ch}_{\text{opt}}(k_{\text{opt}})}, k_{\text{opt}} \right] = \text{Min}_{k=1,2,\dots,N} \left\{ \text{MinMaxP}_{\text{Tx, min}}^{k, \text{Ch}_{\text{opt}}(k)} \right\}$$

where $\text{MinMinMaxP}_{\text{Tx, min}}^{k_{\text{opt}}, \text{Ch}_{\text{opt}}(k_{\text{opt}})}$ represents the minimum of $\text{MinMaxP}_{\text{Tx, min}}^{k, \text{Ch}_{\text{opt}}(k)}$ for $k=1,2,\dots,N$. The term k_{opt} is the index of the best AP assumption.

- Finally, the couple $(k_{\text{opt}}, \text{Ch}_{\text{opt}}(k_{\text{opt}}))$, where k_{opt} is the AP index and $\text{Ch}_{\text{opt}}(k_{\text{opt}})$ is the channel index, is the most power efficient (AP, Wi-Fi channel) configuration that allows reaching the required per-link PER.

The general flowchart of the second algorithm is depicted in Figure 72.

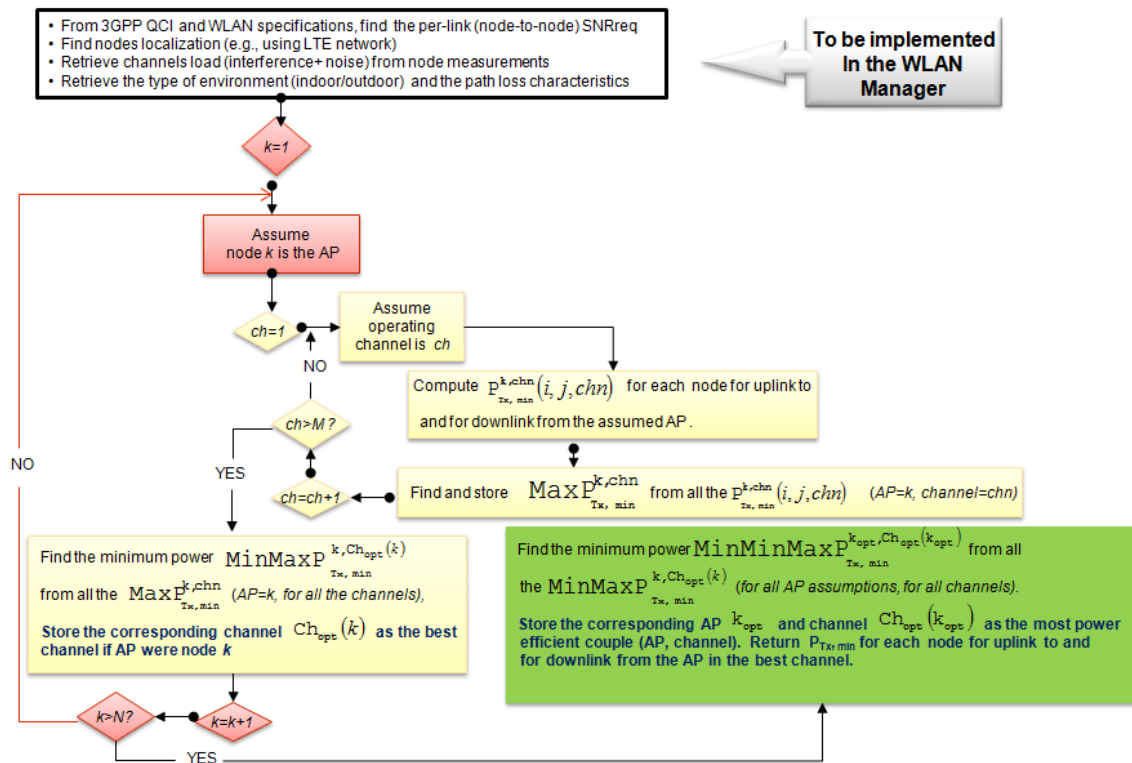


Figure 72: Flowchart of the joint AP and Channel selection for WLAN establishment: 2nd algorithm. This algorithm allows finding the most power efficient couple (AP, Wi-Fi channel) that meets the per-link PER requirements associated to a particular 3GPP QCI.

2.12.3 Integration in OneFIT architecture

The proposed WLAN establishment algorithms can be mapped to both CMON and CSCI management systems in the infrastructure nodes (WLAN Manager and the candidate stations). The reasons of the WLAN creation (suitability determination), the service demands, and the WLAN requirements should be mapped to the CSCI management system. The WLAN creation and maintenance are mapped in the CMON system.

2.12.4 Further performance evaluation results

Detailed performance evaluation results for the proposed algorithms can be found in D5.3. For the evaluation scenario, we considered two Wi-Fi channels and three smart-phones having both Wi-Fi capability and Access Point (AP) functionality. The smart-phones were located in different places within an office. We use Chanalyzer Pro of Metageek to visualize the 2.4 GHz band and to make measurements in the different Wi-Fi channels. In order to estimate the PER, we use CommView for Wi-Fi of TamoSoft. The evaluation consisted in estimating the PER for each link in each possible configuration (channel, AP assumption), then we have compared the different configurations to find the best one and to compare the best one with the configuration returned by the 2 algorithms. Please, confer to the OneFIT D5.3 for more detail.

2.13 Content conditioning and distributed storage virtualization/aggregation for context driven media delivery

2.13.1 Problem formulation and algorithm concept

The algorithm addresses the OneFIT scenario 5 “Opportunistic resource aggregation in the backhaul network”. More precisely, the use case related to aggregation of backhaul storage resources of WMNs is addressed. The algorithm’s task is to, based on contextual data gathered from the WMN environment and end users, provide appropriate WMN node selection for multimedia content placement and distribution. The criteria for node selection is based on request distribution, popularity of multimedia content, status of caching storage of WMN nodes and user’s behaviour (mobility, viewing patterns and used devices).

Algorithm’s aspects are included into LCI’s patent application with EU patent office: *EP12181758.9 – A method and apparatus to dynamically and opportunistically create a wCDN through reconfiguration and management of the underlying wireless network’s resources.*

2.13.2 Algorithm specification

The algorithm for node selection for content placement on WMN APs is based on a mathematical model in form of mixed integer linear program (MILP). This model takes into account distribution of users’ requests, available storage capacity of WMN nodes, and capacity of the links in WMN backhaul. Contextual data (database in Figure 73) regarding history of requests for content, user’s mobility, content profiles, status of WMN backhaul links and status of storage area in WMN nodes will be used for derivation of knowledge regarding:

- Temporal and spatial distribution of requests for particular content.
- User’s mobility patterns.
- Traffic patterns in WMN backhaul links.

This knowledge will be used for triggering content distribution to WMN nodes and among them and for decision making mechanism of the algorithm. Proactive caching provides initial network access node selection for placement of video files. To be able to achieve cognitive placement of files, proactive caching mechanism must take into account the following contextual data:

- Video content popularity (on local and more broad level);

- Spatial and time distribution of user requests;
- Profile of end users (their mobility patterns, viewing patterns, equipment capabilities, QoS requirements);
- Status of backhaul nodes (available storage in APs, popularity of currently cached data and available bandwidth for streaming);
- Backhaul traffic patterns.

By processing above mentioned contextual data, proactive caching mechanism selects appropriate WMN APs and places certain number of copies of multimedia content (video file's chunks) into their storage space. Reactive caching mechanism relies on monitoring system to detect changes in status of the WMN APs and users' requests which will trigger redistribution of content among WMN APs (trigger levels are defined by the service provider). Contextual data that is taken into account for this mechanism are:

- Changes in backhaul traffic (in order to avoid congestion on some backhaul links);
- Changes in spatial user request distribution (more requesting users move from one AP to another – changes not covered by the detected mobility patterns);
- Changes in status of WMN nodes (node down);
- Popularity of files ready for proactive caching (if a very popular file has to be cached in access points then storage space for it has to be provided).

Collected contextual data is stored in centralized database and used for derivation of cognition about network environment and system performance. By using this knowledge, the algorithm is able to derive an appropriate response to previously encountered situations and to predict the best response for newly encountered problems.

When the algorithm detects a user's request for file f , which contains N chunks (k), and if that particular file is not locally cached in WMN access points, then the algorithm has to decide whether or not to cache this file or to proceed with streaming from source server (see Figure 73). If decision is made to cache new file in WMN APs, then candidate nodes have to be detected and selected. During the streaming process to the end user, chunks of video file f are cached on selected APs and, at the same time, contextual data about system performance and resource utilisation is collected.

If requested video file is stored on WMN access points, algorithm needs to determine candidate nodes for streaming chunks of the video file to the requesting user. Previously collected contextual data and derived knowledge is used for selection of candidate nodes and derivation of streaming schedule that provides the best system performance and network resource utilisation. During the streaming process the algorithm gathers contextual data from database (filled by monitoring system) and if changes in system environment reach some threshold (event triggered change) re-caching possibilities are examined. If re-caching is not possible or is not expected to solve the current problem, then streaming is continued from a source server. During streaming session from the source server, system is monitored and streaming is continued from APs when system environment status allows it. If re-caching is selected as the solution for the encountered problem, then appropriate APs are selected and some chunks of the video file are moved among APs. When the content redistribution process finishes, the new streaming schedule is derived and streaming session continues. Contextual data about system performance, resource utilisation and QoS levels achieved on end user side are gathered during algorithm cycle and saved in contextual database for further knowledge derivation.

For every user request, ON will be created among selected WMN APs to support algorithm execution which will provide aggregation of storage resources in order to locally stream as many chunks of the requested video file as possible. ON will provide support for reactive caching mechanism by enabling gathering and exchange of contextual data which will enable detection of any changes in the access network environment. When values of gathered contextual data change above a certain threshold, reactive caching mechanism will start the re-caching process over the created ON.

The algorithm covers all ON management phases as shown in Figure 73. Detailed description of the mathematical model for node selection can be found in D4.2 [7].

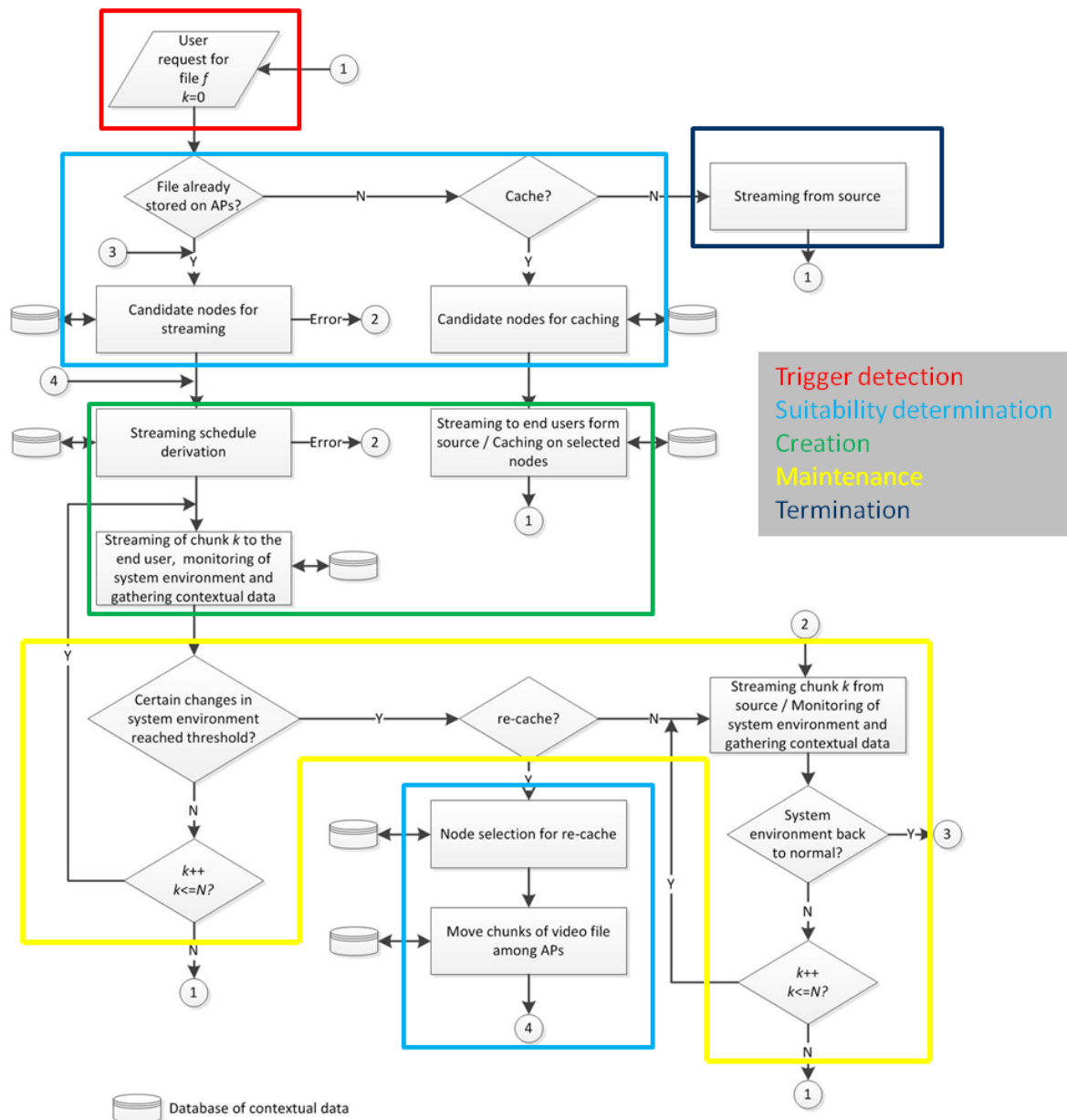


Figure 73: Mapping of the algorithm onto ON management phases

Research presented in [19] shows that content placement on WMN nodes (with cooperation between WMN nodes enabled – through ONs) has the same impact on total content delivery capacity (with respect to placed content) of the underlying WMN as introduction of new GWs into the WMN topology. As the local popularity of content placed on WMN nodes increases, the total streaming capacity of the WMN increases more significantly. A mathematical model for node identification and selection for content placement is also used for the derivation of optimal GW placement within the WMN topology with respect to contextual parameters mentioned above. This is used for practical validation of the algorithm's node selection capabilities. The algorithm's variant is used on the open platform WMN test bed which is described in [12]. This variant of the algorithm performs GW selection within the given WMN topology. The goal is to keep the number of WMN GWs at a minimum (predefined value) while maintaining network's capacity necessary for addressing the QoS requirements of currently connected users.

2.13.3 Integration in OneFIT architecture

Integration of the algorithm into the OneFIT architecture is described in D4.2 [7]. The role of the algorithm in the OneFIT scenario 5 is described in the section 4.5 of this document.

2.13.4 Further performance evaluation results

Detailed performance evaluation results for the algorithm can be found in D4.2. These results are obtained through extensive experimental campaigns performed in the custom built MatLab based simulator. The mathematical MILP model is validated in the simulator environment.

Next, the open platform WMN test-bed, which is described in the [12], is used for practical validation of the MILP model's node selection capabilities. The test-bed is configured so that in any moment only two WMN GWs can exist in the WMN topology. The position of WMN nodes is fixed and power levels of their interfaces are configured in a way which ensures that the complete WMN topology is known (all possible wireless backhaul links are known). This WMN topology is presented as a network graph where edges represent backhaul links while vertices represent WMN nodes. User request distribution is established with real mobile devices (laptops, smart-phones and tablets) and with dummy traffic sinks at WMN nodes. The UDP traffic generator is located in the core part of the test-bed. The SNMP based monitoring system gathers information about where mobile end users are connected (on which WMN APs) every 20 seconds. The initial request distribution is established with test-bed configuration presented on Figure 74. This distribution of end users is used, together with the derived network graph representing the test-bed topology, as input for the MILP model for node selection. The mathematical model is executed with the limitation that only two nodes need to be selected and QoS related constraints (minimization of the number of wireless hops, BW capacity of backhaul links and QoS requirements of detected application) need to be satisfied. Nodes WMN1 and 3 are selected and they are configured as WMN GWs by bringing up their Ethernet interfaces. It is important to note that all WMN nodes are connected via cable to the layer 2 switch, however only WMN GWs have active Ethernet port and WMN APs are connected to them over 802.11a based backhaul links.

As the next step, certain mobile users walk between WMN APs thus changing the request distribution among WMN nodes. As the number of end users connected to different WMN nodes changes, the MILP model is executed with different request distribution. The new set of WMN nodes are selected and these nodes are reconfigured into the new WMN GWs while the old WMN GWs are reconfigured into WMN APs and connected to corresponding GWs via wireless backhaul links. This procedure is shown in Figure 75. Figure 74 and Figure 75 show the course of the experiment which is used for practical validation of the algorithm's ability to properly select WMN nodes.

The performance evaluation experiments are focused on the impact of GW reconfiguration on QoS indicators on end user devices. For simplicity we will focus on portion of the test-bed network composed of the WMN 3, 4 and 5 and user nodes U4-U7. The mobile user U7 is walking from one WMN node to the other while other end users are static. All end users in the test-bed are requesting the same service. In the initial request distribution setup nodes WMN 1 and 3 are configured as WMN GWs. As the U7 walk from WMN 3 to the WMN 4, the request distribution changes. Handover of the U7 from the WMN 3 to the WMN 4 is detected by the SNMP protocol. This triggers execution of the MILP model with new input in form of the new request distribution. The MILP model calculated the new set of WMN GWs and the WMN 4 becomes the new GW while the WMN 3 becomes the WMN AP and its Ethernet interface is turned down. The WMN 3 connects to the WMN 4 over the new backhaul link established among available wireless backhaul interfaces of the WMN nodes. When the U7 finally reaches the WMN 5, the MILP model will select this WMN node to be the new GW and the WMN topology presented in Figure 75b) will be established.

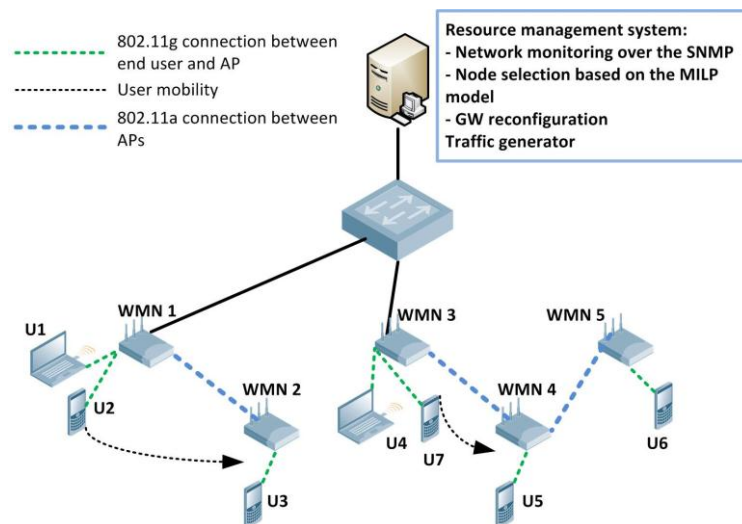


Figure 74 – Set up of the open platform WMN test-bed for validation of the GW selection algorithm

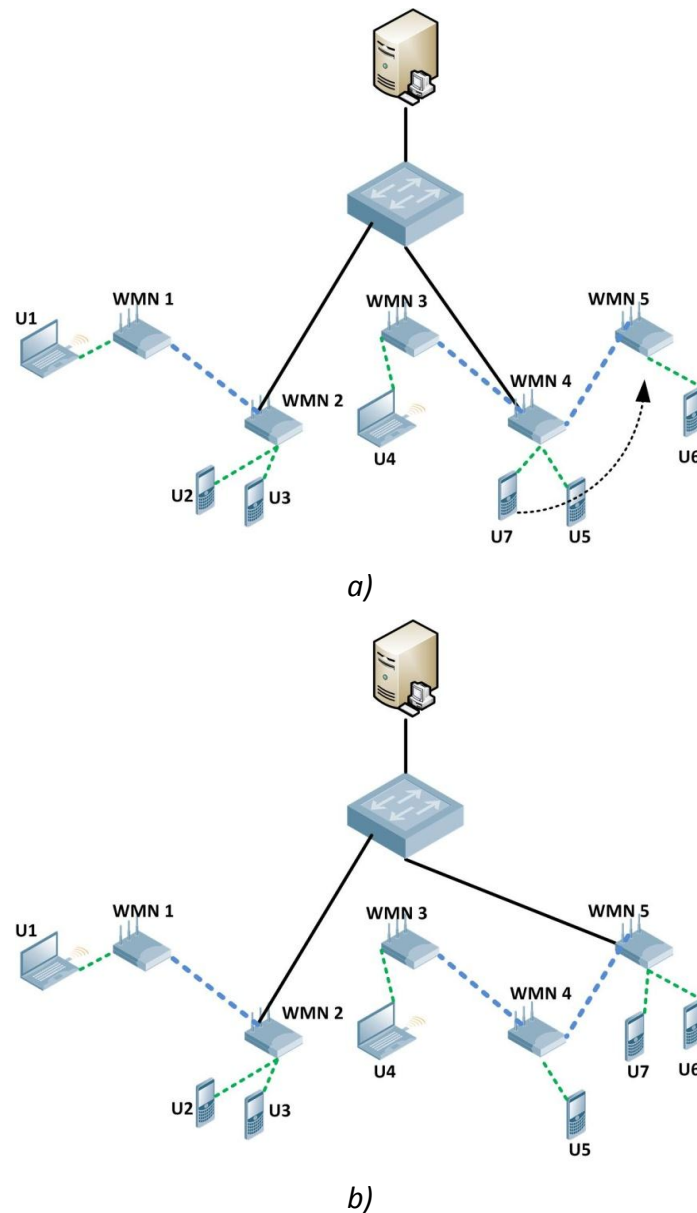


Figure 75 – User mobility results in change in request distribution and GW selection

The experiments are organized as follows:

- Three services are considered: VoIP (155Kbit/s), streaming of low quality video (Video 1 - 1Mbit/s) and streaming of high quality video (Video 3 - 3Mbit/s);
- Experiments are started when U7 starts walking towards the WMN 4;
- Three different service durations are tested while the U7 goes from the WMN 3 to the WMN 5: 1000, 2000 and 3000 seconds;
- Every experimental setup (service + service duration) is repeated 10 times.

The Iperf tool is used on nodes U4-U6 in order to check the impact of the GW reconfiguration on QoS indicators. The following QoS parameters are monitored on mentioned mobile users: achieved throughput, jitter level and packet loss percentage. The QoS on the U7 is not considered since this node experiences handover while moving between WMN nodes and its QoS would be affected whether or not the GWs are reconfigured.

Figure 76 depicts the impact of GW reconfiguration on jitter levels at the side of the U4. GWs are reconfigured as result of U7's mobility. It is clear that jitter level increases as the WMN 3, to which the U4 is connected, becomes the first and the second AP on the backhaul path to the new GW. The higher number of backhaul paths (wireless hops) results in increased jitter and, consequently, decreased QoS. Figure 77 shows the impact of GW reconfiguration and service provision duration on packet loss percentage. During the GW reconfiguration phases the connections are lost for short time interval (varying between 8 and 15 seconds) until the new GW is selected and addressing tables are updated. Since experiments are using the UDP traffic generator, packets sent during the GW reconfiguration phase are lost. Therefore, total packet loss rate at the U4 side increases in case where GW reconfiguration is executed. Service provision duration also impacts the total packet loss percentage in case when the same number of GW reconfigurations needs to be performed. The first 10 measures shown in Figure 77 correspond to the service session which lasts for 1000 seconds, the next 10 measures correspond to 2000 seconds long service session and finally the last 10 measures correspond to the service provision duration of 3000 seconds. The packet loss percentage in case when there is no GW reconfiguration is relatively constant in controlled environment. Longer session duration can be seen as impact of user's mobility speed on the packet loss rate when GWs are reconfigured in order to follow the mobile user. The first 10 measures in Figure 77 can be seen as packet loss rate for very fast moving U7, while the last 10 results in this figure can be seen as packet loss rate in case where the U7 is moving slowly. The conclusion of this analysis is that the GW reconfiguration process should not address fast moving end users if there are static or slowly moving users in the same WMN. The GW reconfiguration should correspond to more general user mobility patterns and changes in request reconfigurations.

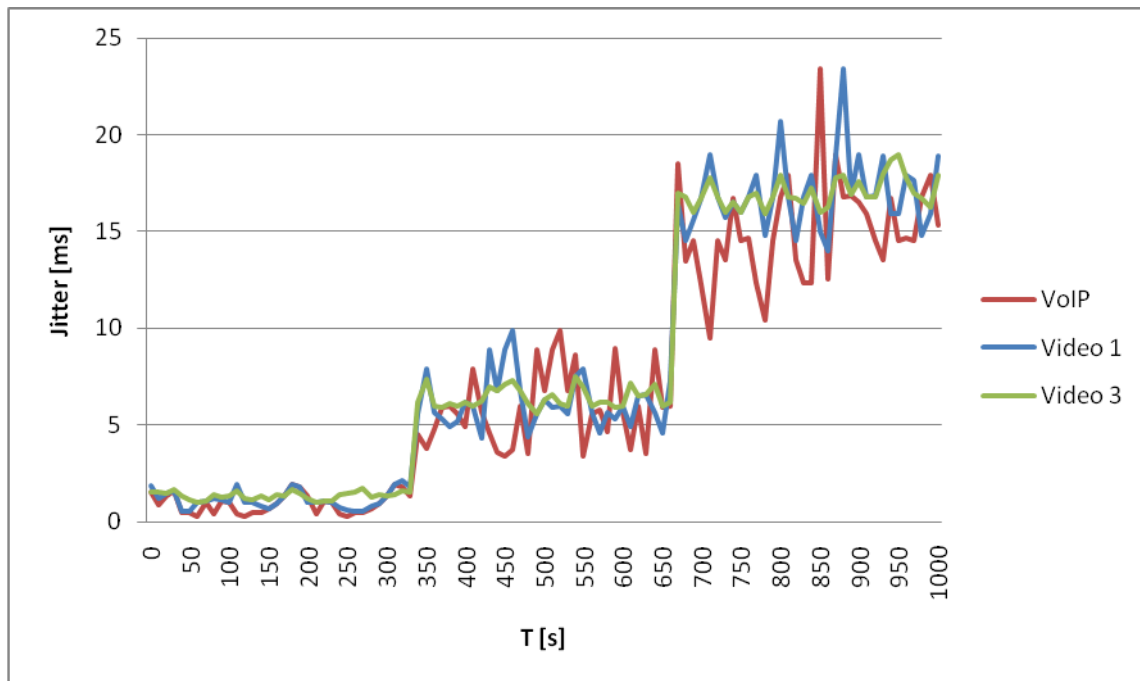


Figure 76 – Jitter levels measured at U4 side

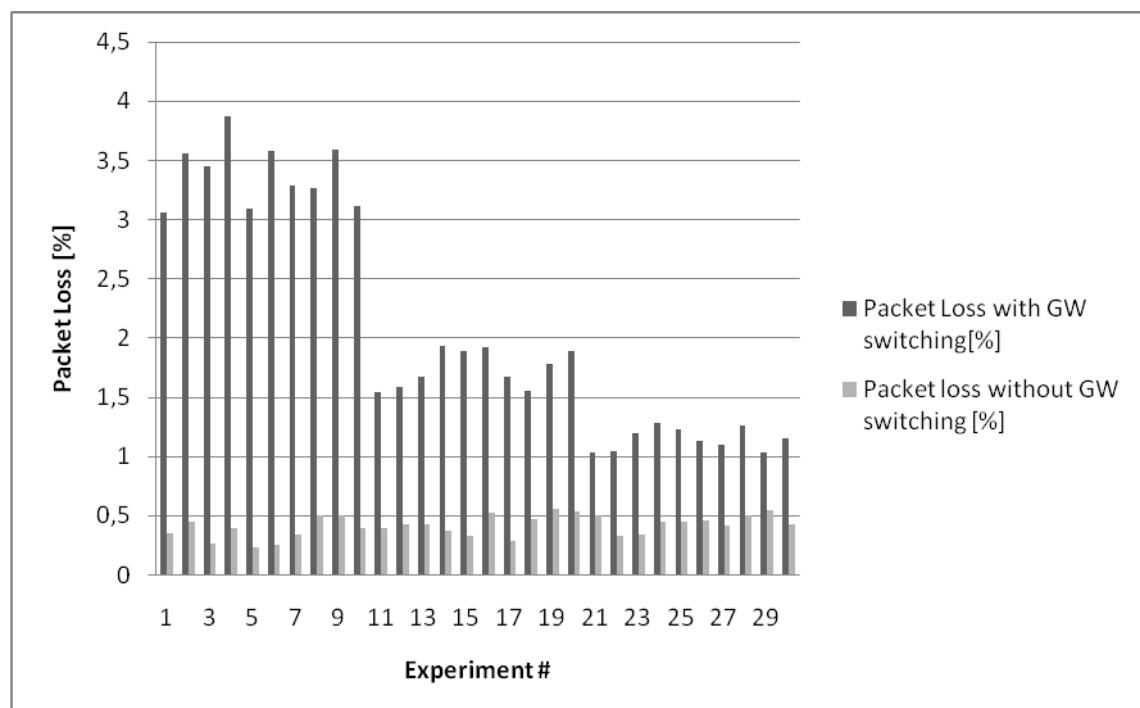


Figure 77 – Packet loss percentage measured at U4 side

The GW reconfiguration process has significant impact on QoS indicators at the side of end users which are not moving throughout the WMN. During the GW reconfiguration the connections is lost and, because the traffic generator sends UDP data streams, packets sent during this time are lost.

The conclusion of the research presented in [19] is that GW reconfiguration and content placement on WMN nodes has the same impact on the WMN capacity with respect to the locally stored content. The results presented above clearly indicate that GW reconfiguration has significant impact on QoS on end user side. Therefore, placing the popular content on WMN nodes will increase both the WMN capacity and the QoS achieved on end user devices. Further research and development will be focused on content placement on WMN nodes and practical validation of this approach.

2.14 Capacity Extension through Femto-cells

The problem considered here is the selection of the femtocells that will acquire traffic from the congested BS (i.e. the terminals that will participate in the ONs), as well as the minimum power allocation to the femtocells that is needed to cover the offloaded terminals.

2.14.1 Problem formulation and algorithm concept

For the proposed solution, a network like the one of Figure 78 is assumed. The network comprises a macro BS that faces congestion issues due to the traffic of its connected terminals. In addition, within the BS a set of femtocells is deployed. It is assumed that femtocells can operate at discrete power levels (i.e. a proportion of their maximum transmission power) which correspond to a transmission range. The target is to assign the most appropriate power-level to femtocells in order to serve terminals that are suitable to be handedover to the femtocells. The suitable terminals are depicted by the suitability determination phase.

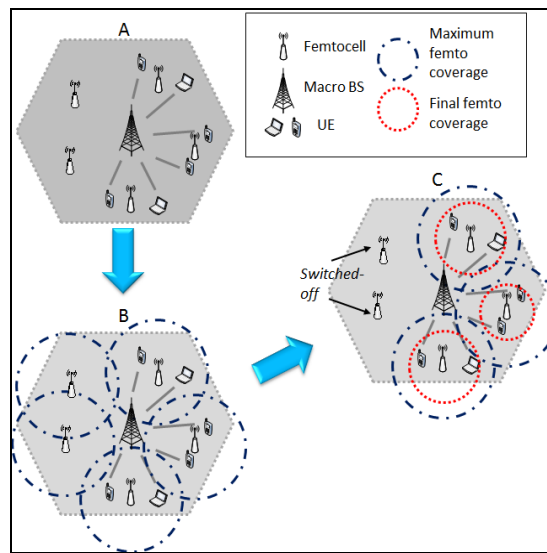


Figure 78: The concept of the Energy-efficient Resource Allocation

Let F be the set of femtocells and let B be the set of the macro BSs, while the set of terminals will be denoted with T . Moreover, PL will denote the set of power-levels to which the femtocells can be configured. In addition, the following decision variables are considered:

$$X_{ij} = \begin{cases} 1, & \text{if terminal } i \in T \text{ is assigned to} \\ & \text{infrastructure } j \in F \cup B \\ 0, & \text{otherwise} \end{cases} \quad (37)$$

$$Y_{kj} = \begin{cases} 1, & \text{if power-level } k \text{ is assigned to femtocell } j \in F \\ 0, & \text{otherwise} \end{cases} \quad (38)$$

In order to calculate the coverage of a femtocell according to its transmission power, the propagation model proposed by the authors in [20] is used. The allocation problem is an optimization problem where an objective function (OF) will be minimized, satisfying a number of constraints. Accordingly, the overall optimization problem can be formulated as follows:

$$\text{Minimize } OF = \sum_{j \in F} CP_j + \sum_{j \in F} \sum_{\substack{j' \neq j \\ j' \in F}} (\phi_j - \phi_{j'})^2 + \sum_{i \in T} \sum_{j \in F \cup B} X_{ij} \cdot d_{ij} \quad (39)$$

subject to:

$$\sum_{j \in F \cup B} X_{ij} = 1, \forall i \in T \quad (40)$$

$$\sum_{k \in PL} Y_{kj} = 1, \forall j \in F \quad (41)$$

$$\phi_j \leq cap_j, \forall j \in F \quad (42)$$

$$N_{ijk} \leq Y_{kj}, \forall i \in T, j \in F, k \in PL \quad (43)$$

$$X_{ij} \leq N_{ijk} \cdot Y_{kj}, \forall i \in T, j \in F, k \in PL \quad (44)$$

where

$$\phi_j = \sum_{i \in T} X_{ij}, \forall j \in F \quad (45)$$

$$CP_j = \sum_{k \in PL} (Y_{kj} \cdot l_k \cdot P_j) \quad (46)$$

$$N_{ijk} = \begin{cases} 1, & \text{if user } i \in T \text{ can be covered by femto } j \in F, \\ & \text{configured at power-level } k \in PL \\ 0, & \text{otherwise} \end{cases} \quad (47)$$

ϕ_j denotes the load of a femtocell j , i.e. the number of terminals that are served through the femtocell, while cap_j is the capacity of femtocell j , i.e. the maximum number of terminals that the femtocell can serve. In addition, d_{ij} is the Euclidean distance of terminal i from BS/femtocell j in metres, while P_j denotes the maximum transmission power of femtocell j in Watts and CP_j is the current transmission power of femtocell j . l_k depicts the proportion of the maximum transmission power to which a femtocell is configured to operate (e.g. $l_1=0.5$).

The OF in (39) tracks the power consumption of the femtocells, the load balancing factor among femtocells and the distance of the terminals from their serving infrastructure. More specifically, the first term of the function is the power consumption of the femtocells. The second term is used as a load balancing factor and is the square of the load difference between each femtocell. The usage of this term is to ensure that terminals will be equally distributed among the femtocells. Lastly, the third term expresses the distance among the terminals and their serving BS/femtocell. As far as the constraints are regarded, relation (40) depicts that each terminal is served by one and only one BS or femtocell, while relation (41) denotes that every femtocell can be configured to operate at only one power-level. Relation (42) illustrates the fact that the current load of a femtocell cannot exceed its capacity. Relation (43) expresses the fact that a terminal cannot be covered by a femtocell which operates at power-level k , if the power-level is not assigned to the femtocell. Relation (44) depicts that a terminal cannot be served by a femtocell operating at power-level k , if the terminal is not covered by the femtocell or if the power-level has not been assigned to the femtocell. Relation (45) denotes that the load of a femtocell is equal to the number of the terminals that are served through it, while relation (46) is utilized to calculate the current transmission power of a femtocell.

2.14.2 Algorithm specification

In order to solve the aforementioned optimization problem, an algorithm was developed, namely the Energy-efficient Resource Allocation (ERA) to femtocells.

As it was aforementioned, the lifecycle of an ON includes the suitability determination phase. This phase will provide the input to the ERA algorithm and specifically the terminals that should be

offloaded to the femtocells in order to solve the problematic situation. The set of these terminals will be denoted with $T' \subseteq T$. Apparently the mobility level of the femtocells will also be taken into account, due to the fact that terminals with high mobility level (e.g. ≥ 2 m/s) should not be rerouted to femtocells because they will need to proceed to a handover to a macro BS shortly. The ERA will also utilize the set of femtocells and their capabilities in terms of capacity and possible power-levels at which they can be configured to operate. The output of the algorithm will comprise the selection of the femtocells that are needed to serve the terminals provided by the suitability determination phase, as well as the minimum possible power allocation to the femtocells in order to cover these terminals.

The algorithm can be described in steps as follows. At *Step 1* the algorithm starts by forming the F' set which contains the femtocells with the smallest distance from the terminals of T' . Therefore, $F' \subseteq F$ can be expressed as follows:

$$F' = \left\{ f \in F : N_{ifk_{max}} \cdot d_{if} \leq N_{ijk_{max}} \cdot d_{ij}, (\forall i \in T) \wedge (\forall j \neq f) \right\} \quad (48)$$

Where $k_{max} \in PL$ is the maximum power-level at which the femtocell can be configured. At *Step 2* the algorithm configures all femtocells that belong to F' to operate at their maximum power level. It should be reminded that each power-level $k \in PL$ corresponds to a range of coverage cov_k . At *Step 3* the terminals are assigned to the femtocells for this power allocation and according to the aforementioned constraints. T_f will denote the set of terminals that were assigned to be served by femtocell $f \in F'$. Therefore, $T_f \subseteq T'$ can be expressed as follows:

$$T_f = \left\{ i \in T' : X_{if} = 1 \right\} \quad (49)$$

At *Step 4* the algorithm forms the F_d set which contains the femtocells that can reduce their power, i.e. they have not reached the minimum power-level $k_{min} \in PL$. Therefore, $F_d \subseteq F'$ can be expressed as follows:

$$F_d = \left\{ f \in F' : Y_{k_{min}f} = 0 \right\} \quad (50)$$

Step 5 checks if F_d is empty. If it is empty, the algorithm is terminated; otherwise it proceeds to *Step 6* where the next femtocell $f \in F_d$ is picked up. At *Step 7* the power-level of f is reduced to the next power-level and then *Step 8* checks if a terminal $i \in T_f$ exists that is no longer within the range of f due to the power-level reduction. If there is no such a terminal, the femtocell is able to reduce its power-level and the algorithm transits to *Step 7*. Otherwise, the algorithm must proceed to *Step 10* where the power-level of the femtocell must be set to the previous-value in order to continue covering the user. Then f is removed from F_d since it cannot decrease its power-level any more. Finally, the algorithm transits to *Step 5*. Figure 79 depicts the flow chart of the ERA algorithm.

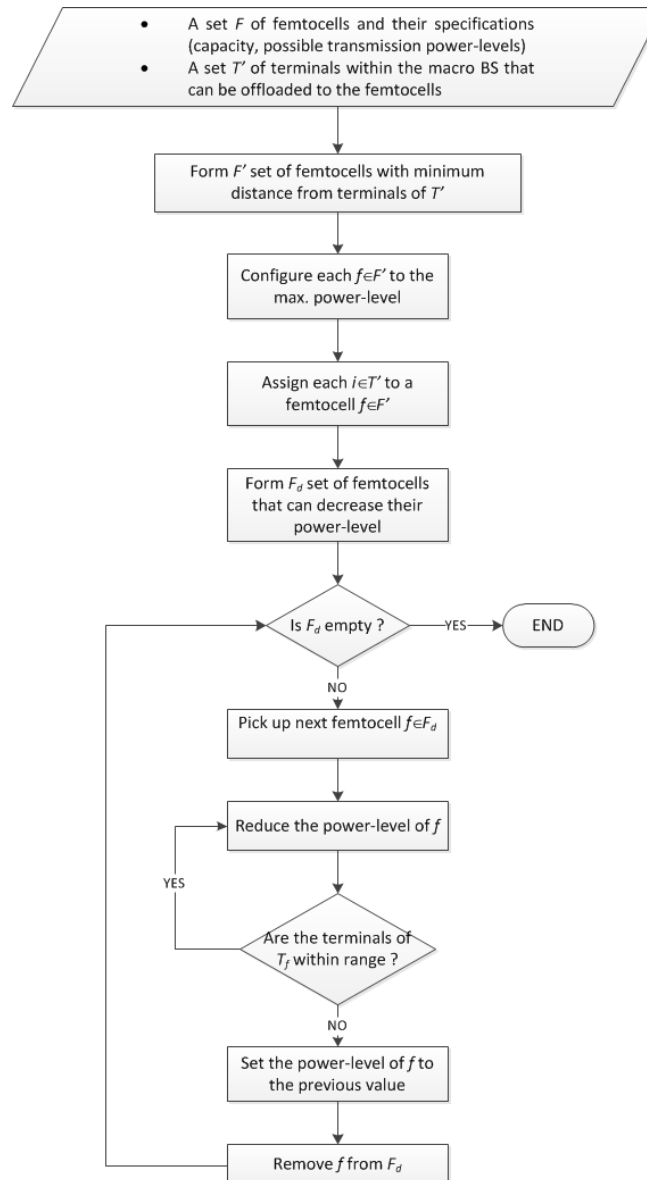


Figure 79: Flow-chart of the ERA algorithm

2.14.3 Integration in OneFIT architecture

Figure 80 provides a mapping of the functionalities of the ERA algorithm, which is associated with the ON creation functionality to the OneFIT functional architecture and the functionalities entities as defined in D2.2 [3].

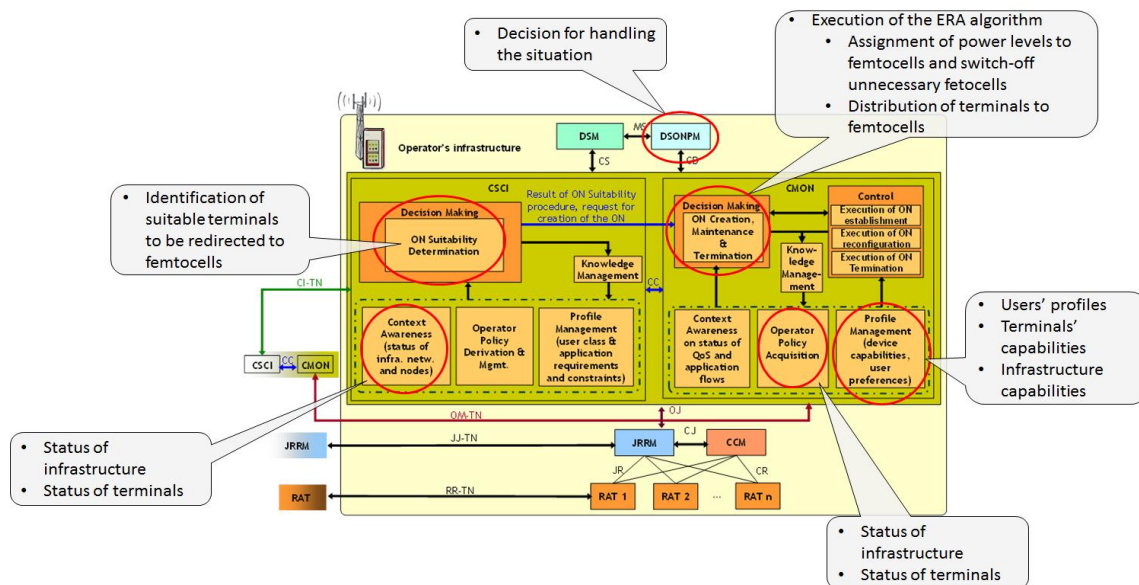


Figure 80: Mapping of the ERA concept functional entities to the OneFIT functional architecture

2.14.4 Further performance evaluation results

In this section, the ERA solution will be evaluated. To this respect, a macro BS is considered and within its area 9 femtocells are deployed (in a gridline) and 40 terminals. Figure 81 illustrates the aforementioned topology. As far as the traffic model of the terminals is regarded, message generation is a random variable, ranging from 3 to 7 secs uniformly distributed with mean of 5 secs. Furthermore, message size is a random variable in the range of 64 to 1024 Kbytes uniformly distributed. Also, the number of terminals that generate traffic in each interval is random.

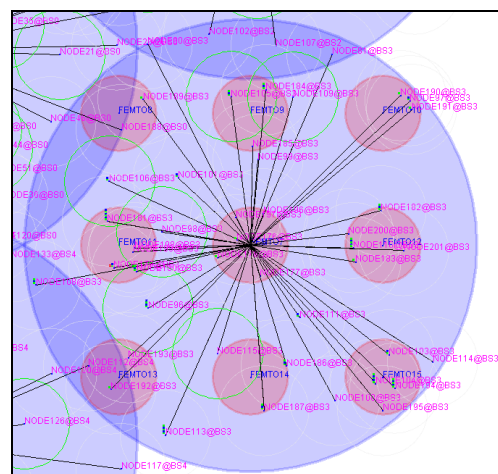


Figure 81: Considered topology

Moreover, the goal of the algorithm is to select the femtocells and assign them with the minimum possible power-level that is needed to cover the terminals that will be offloaded to the femtocells. As it was aforementioned in the previous sections, these terminals are provided as input to the ERA algorithm by the suitability determination phase. In order to proceed with the evaluation, 18 cases were considered. Each case comprises the number of terminals that will be rerouted to the femtocells, the mobility level of the terminals (in terms of m/s) and the distribution of the terminals within the area. Specific two distributions are taken into account. A “centralized” one which means that terminals are distributed at the center of the macro BS, and a “sparse” one which means that terminals are uniformly distributed within the area of the BS. Table 6 depicts the considered cases. It

should be noted, that due to the fact that the DRA algorithm does not accept input from the suitability determination phase, these cases cannot be examined for the DRA concept.

Table 6: Considered cases

Case	# terminals to be handedover to femtocells	Mobility level	Terminals distribution
1	6	0	Centralized
2	6	0	Sparse
3	6	1	Centralized
4	6	1	Sparse
5	6	2	Centralized
6	6	2	Sparse
7	12	0	Centralized
8	12	0	Sparse
9	12	1	Centralized
10	12	1	Sparse
11	12	2	Centralized
12	12	2	Sparse
13	18	0	Centralized
14	18	0	Sparse
15	18	1	Centralized
16	18	1	Sparse
17	18	2	Centralized
18	18	2	Sparse

Figure 82 illustrates the delivery probability of all 40 terminals in the congested BS before and after the solution enforcement for each one of the 18 investigated cases. The horizontal axis depicts the number of the case, while the vertical one presents the delivery probability. The general trend shows that as mobility level increases the delivery probability tends to drop. On the other hand, as more terminals are handed-over to femtocells, delivery probability tends to increase.

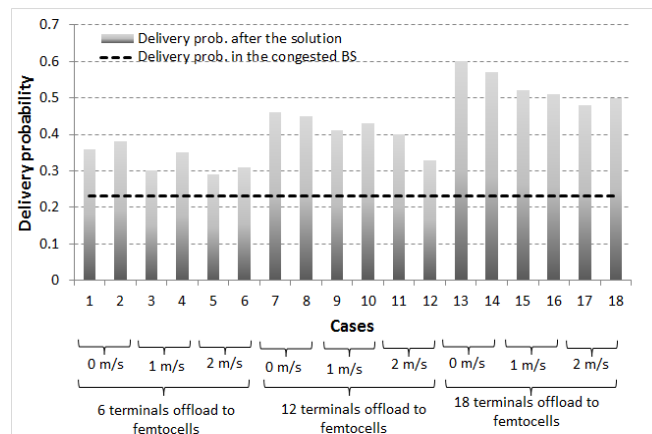


Figure 82: Delivery probability for each considered case

Figure 83 shows the average delay for delivered messages from all 40 terminals in the congested area before and after the solution for each one of the 18 cases. The horizontal axis depicts the number of the case, while the vertical one presents the average delay in seconds. In all cases, the proposed concept tends to perform better since the delay drops compared to the situation before the solution enforcement. Specifically, the decreasing delay is higher when more terminals are handed-over to femtocells. As the moving speed of the terminals increases, the average delay tends to increase as well.

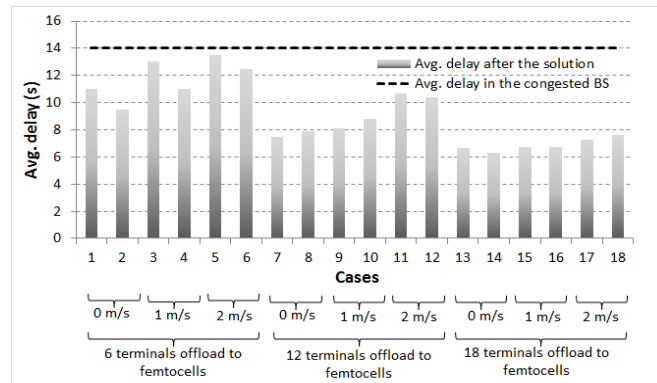


Figure 83: Average delay for each considered case

Furthermore, the output of the ERA algorithm, i.e. the power-level that was assigned to each femtocell is illustrated in Figure 84, as well as the number of terminals that each femtocell acquires is also counted in Figure 85.

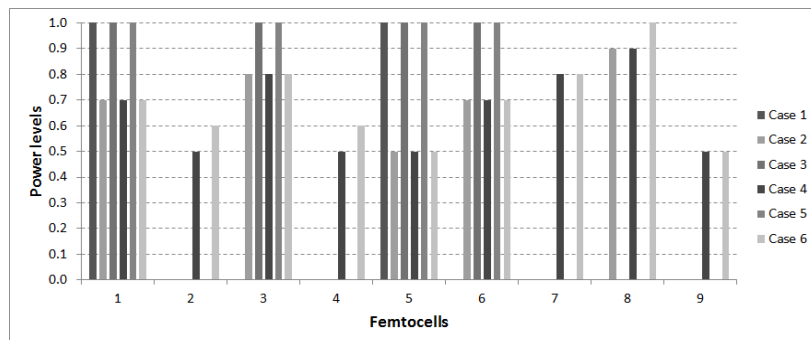


Figure 84: Power allocation to femtocells for each grouped case

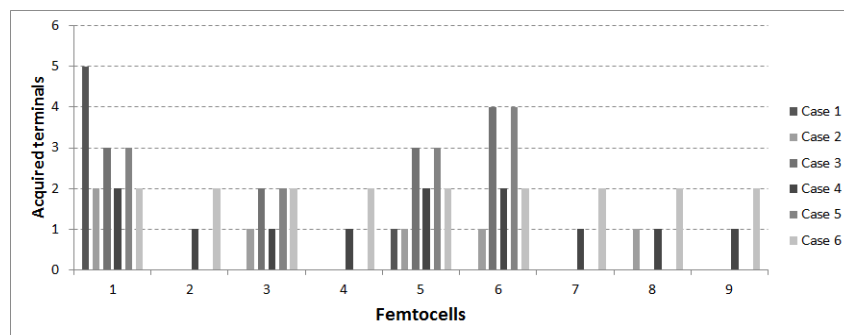


Figure 85: Number of terminals acquired by femtocells for each grouped case

In addition, the runtime of the ERA algorithm in msec is provided in Figure 86.

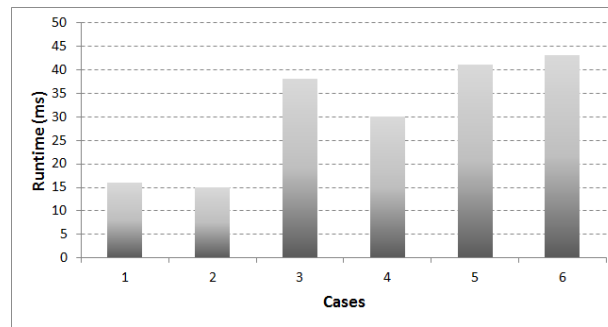


Figure 86: ERA runtime for each grouped case

Table 7 depicts the previously aforementioned cases grouped by the number of terminals offloaded to femtocells and according to the distribution of the terminals within the macro BS area, that correspond to the aforementioned results in Figure 84, Figure 85 and Figure 86. Apparently, in the “centralized” distribution the femtocells needed to be configured to high power-level in order to cover the increased number of terminals. On the other hand only the 25% of the femtocells had to be utilized, while the rest were turned off. On the other hand, in the “sparse” distribution the terminals were uniformly distributed so each femtocell served only 1-2 terminals. Thus, low power-levels were assigned to the femtocells. However, almost all femtocells needed to be utilized.

Table 7 : Grouped cases

Case	# terminals to be offloaded to femtocells	Terminals distribution
1	6	Centralized
2	6	Sparse
3	12	Centralized
4	12	Sparse
5	18	Centralized
6	18	Sparse

2.15 Support activity to validate ON algorithms on an offloading-oriented real-deployment testbed

2.15.1 Problem formulation and algorithm concept

The objective of this activity is to evaluate the performance of an ON suitability determination algorithm on a specific femtocell-populated urban environment, in order to elaborate some recommendations for the implementation of an ON-based macro-to-femto offloading mechanism in LTE.

This evaluation relies on the simulations conducted over a pre-defined LTE-2600 test scenario. An urban Manhattan-like area has been selected from real cartography of Barcelona city centre. It covers about 2 km², and there are 11 macro sites (i.e. 33 macro cells) are located in their actual positions. An additional layer of 291 femtocells has been set based on business estimations of real future femto deployments. To avoid artefacts in the results near the edges of the scenario, a smaller focus area (1 km²) has been defined, with only 15 macros (5 macro sites) and 126 femtos. Main results will be referred to this focus area.

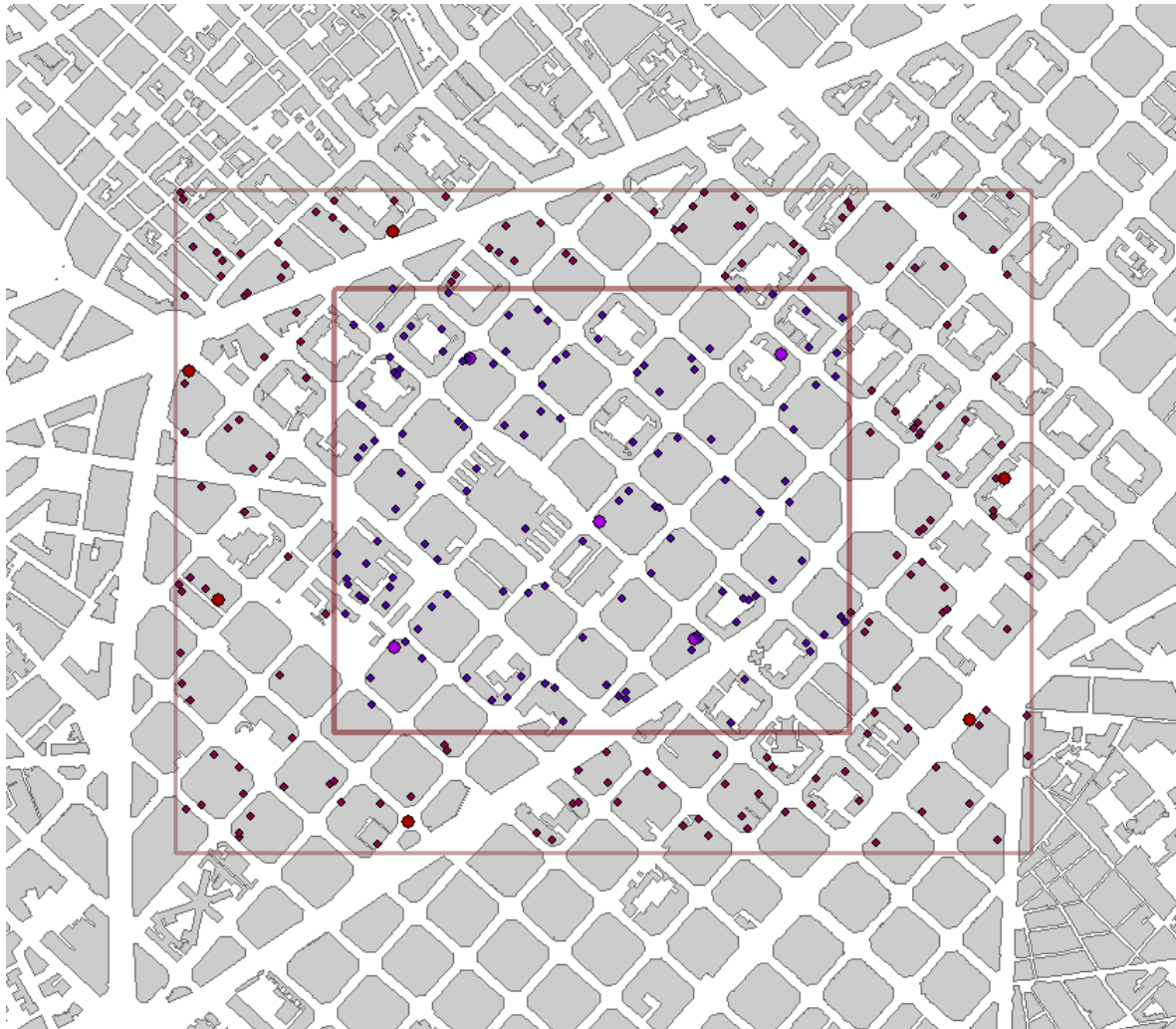


Figure 87: Base LTE test scenario and focus area

This base scenario has been properly adapted to obtain results for three different OneFIT scenarios, namely:

- *OneFIT Scenario 1 (Opportunistic Coverage Extension)*: Some of the macro sectors of the base scenario are turned off to create coverage gaps. Users immersed in these gaps will be able to connect to surrounding femtos by means of opportunistic networking.
- *OneFIT Scenario 2 (Opportunistic Capacity Extension)*: The number of simulated users will be increased in order to force the apparition of hotspots. Congested cells will be alleviated by transferring rejected and underachieving users to surrounding femtos on an opportunistic basis.
- *OneFIT Scenario 3 (Ad-hoc opportunistic network)*: A target throughput threshold will be set at different values in order to evaluate the suitability of the ON to support different bitrate-hungry ad-hoc applications. Underachieving users will then be opportunistically offloaded to the femto layer in order to reach the threshold.

The suitability determination algorithm under test comprises the following stages: firstly it searches throughout the simulated users to find those that are eligible for offloading (i.e. those macro users which receive enough power from a neighbouring femto). Then, it selects the best offloading strategy to fulfil the coverage/throughput requirements derived from the specific scenario. Finally,

the optimal ON configuration is determined, and the necessary femto policy changes are established. The following section comprehensively describes these stages.

2.15.2 Algorithm specification

2.15.2.1 Candidate node identification

The femtos and UEs eligible to be part of the ON are identified. These initial candidates are selected following a radio environment criterion: those UEs that could potentially be successfully served by a femto-node, and, correspondingly, those femtos with UEs under their coverage area.

Therefore, all macro-served UEs are evaluated, and those having a femto as best server or second-best server (but are not served by them due to the CSG policies), are identified as ON candidates.

2.15.2.2 Optimisation problem definition

An objective function to choose the nodes that will form the ON has been designed. The general formulation of this OF is the following:

$$OF = \beta \left[\frac{1}{N_M + N_F} \sum_{j \in MU_F} \sum_{j' \in MU_F} \left(\sigma_j \frac{\varphi_j}{\varphi_{jMAX}} - \sigma_{j'} \frac{\varphi_{j'}}{\varphi_{j'MAX}} \right)^2 \right] + \delta \left[\frac{1}{N_U} \sum_{i \in U} \sum_{j \in MU_F} d_{ij} \right] - \tau \left[\frac{1}{N_M + N_F} \sum_{j \in MU_F} \varphi_j \right] + \alpha \left[\frac{1}{N_U} \sum_{i \in U} \max(\varphi_{TGT} - \varphi_i, 0) \right] \quad (51)$$

Where:

- $\beta, \delta, \tau, \alpha$ are the weights of the four blocks of the function.
- N_M, N_F, N_U are the number of macro cells, femto stations and users, respectively.
- M, F, U are the sets of all the macro cells, all the femto stations and all the users, respectively.
- σ_j are weights to adjust the contributions of femtos and macros to the load balance.
- φ_j and φ_{jMAX} are the current and the maximum achievable DL throughput of node j , respectively.
- d_{ij} is the distance from user i to node j .
- φ_i and φ_{TGT} are the current and the target DL throughput of user i , respectively.

The OF is divided into four independent blocks:

- **Load balance (β block):** addresses how total traffic is shared among nodes. The bigger the difference between any couple of nodes, the higher this block becomes. σ_j weights allow to balance the contributions from macro nodes against the ones from femtos. Minimizing this block helps balancing the total traffic in the operator's network.
- **Distance to nodes (δ block):** focuses on the distance between users and serving nodes. The farther the user is, the higher this block gets. Minimizing this block helps reducing the transmission power (in both nodes and UEs), thus leading to less interference, longer battery lives and a greener footprint.
- **Total throughput (τ block):** assesses the total traffic of the network. The higher the load grows, the higher this block becomes. Maximizing this block helps serving as much users as possible.
- **Target throughput achievement (α block):** evaluates how each user is able to attain his target throughput, according to the requirements of the running application. The farther the user is to achieving his target, the higher this block gets. Minimizing this block leads to attain QoS requirements and thus, to a better user satisfaction.

Each of these blocks is preceded by a weighting factor (namely, β , δ , τ and α), so the objective function can be adapted to the coverage/throughput requirements derived from the scenario to be simulated (e.g. the focus may be to keep the load balance, to conserve energy levels of the UEs, to minimize the femto transmission power, etc.)

Therefore, the aim of the optimisation problem is to find the best offloading strategy (i.e., the best combination of UEs that should be transferred from the macro to the femto layer) which minimizes the value of the objective function.

2.15.2.3 Optimisation solving

For a scenario with a high number of UEs, minimizing the OF might become a computationally heavy task and it may take too much time to get a solution. To avoid this, an initial pre-assignment of candidate UEs to femtos can be guessed, using part of the available information of the scenario.

In particular, the distance from UEs to serving nodes and to candidate femtos is known. So, those UEs closer to their candidate femtos are pre-assigned to them. For each candidate UE, this pre-assignment leads to a direct decrease in the δ block of the OF, but it also may or may not enhance the β or τ blocks. If the pre-assignment does not produce a significant reduction of the OF, then it is rolled back, and a new UE is evaluated.

After the pre-assignment is complete, we have an OF value that is better than (or, at least, equal to) the original one, but there is no guarantee that the optimal value has been found. To accomplish that, a numerical optimisation method is performed next. A genetic algorithm approach has been selected: starting from the pre-assignment, new “generations” of assignments are chosen, evaluated, mixed and evaluated again. Most successful branches of assignments are promoted and finally a local minimum of the OF is found.

This process of pre-assignment followed by numerical optimisation is quite efficient and can be performed quickly even for a high number of UEs.

2.15.2.4 ON configuration determination

Once the best assignment of UEs to neighbouring femtocells has been found, the CSG policies of the affected femtos have to be changed in order to accept incoming connection request from candidate users (i.e. the femto policies are still CSG rather than OSG, and candidate UEs acquire temporary privileges to access the CSGs).

The change on the CSG policies effectively forces the offloading and the UEs begin the handover procedure from the macro to the femto layer. These handovers provoke a reallocation of the radio resources both in origin macro nodes and in destination femtocells. From the viewpoint of this algorithm, the ON is considered created at this point and performance KPIs can be collected.

2.15.3 Integration in OneFIT architecture

The suitability determination algorithm tested in this activity can be easily mapped into the OneFIT architecture as shown in Figure 88:

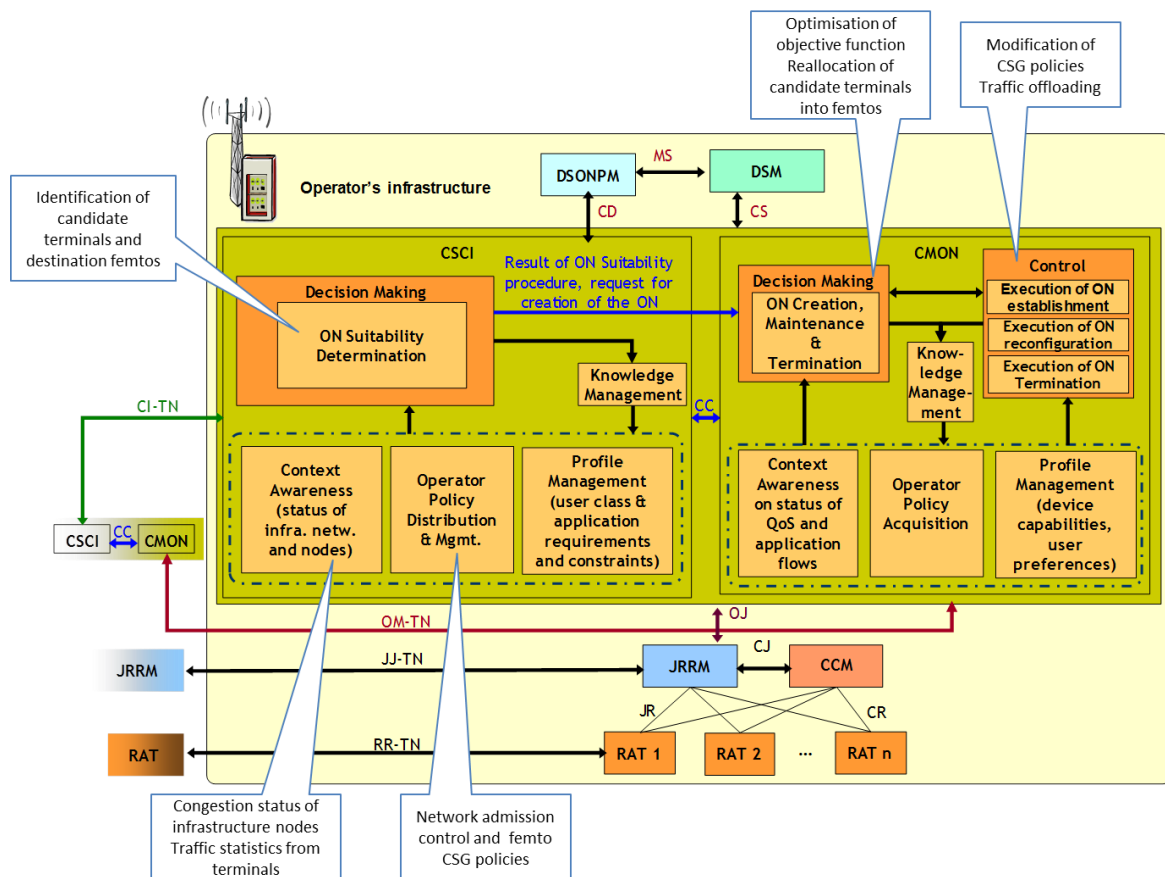


Figure 88 : Mapping of the algorithmic procedure to the OneFIT functional architecture

2.15.4 Further performance evaluation results

2.15.4.1 Baseline initial results

In order to assess the performance of the OF optimisation procedure, simulations have been carried out on the test scenario. The results from this test can be used as a baseline to benchmark the outcomes of applying the algorithm to the OneFIT scenarios.

In this test, 1200 UEs have been randomly placed throughout the scenario, where only 504 of them are inside the focus area. Further results will be referred to this area.

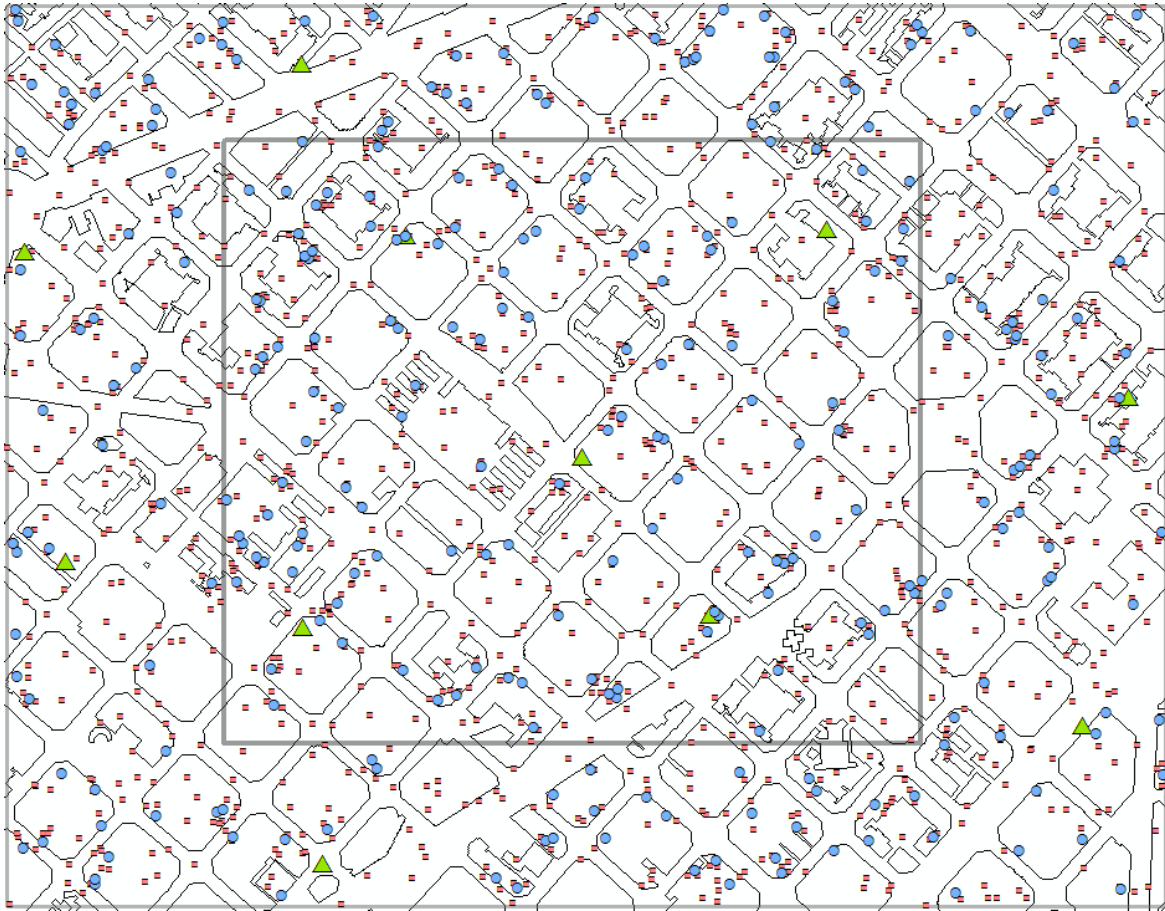


Figure 89: Baseline test scenario. Green: macros. Cyan: femtos. Red: UEs.

After running the node selection procedure, a total of 86 UEs are eligible to form ONs with neighboring femtos. For this generic scenario, the Objective Function has been configured using the following parameters:

- $\beta = 1$
- $\delta = 1$
- $\tau = 10$
- $\alpha = 0$
- $\sigma_{\text{MACRO}} = 2.5$
- $\sigma_{\text{FEMTO}} = 1$
- $\varphi_{\text{MAX_MACRO}} = 100$
- $\varphi_{\text{MAX_FEMTO}} = 20$

The minimization of the OF concludes that only 73 of the UEs should be derived to the femto layer. These UEs are extracted from 21 different macros (i.e. some of them were connected to macros outside the focus zone) and injected into 64 femtos.

The following table shows some KPIs measured before and after the ON is created:

Table 8: Baseline test results

		Without ON	With ON
Macro	Nr. of UEs	389	316
	Nr. of disconnected UEs	103	33
	Mean DL throughput per UE	0.9 Mb/s	1.2 Mb/s
	Total DL traffic in all nodes	257 Mb/s	332 Mb/s
Femto	Nr. of UEs	115	188
	Nr. of disconnected UEs	11	11
	Mean DL throughput per UE	14 Mb/s	11 Mb/s
	Total DL traffic in all nodes	1460 Mb/s	1961 Mb/s

Results show a noticeable performance enhancement: After the ON is created, the number of UEs that cannot connect to the macro layer is reduced by 68%, and the average per user DL throughput increases in 33%. The total traffic in the macro layer grows a 29% even when the number of users is a 19% lesser. At the femto layer, the average throughput reduces, as the limited resources have to be shared among existing and incoming UEs, but the total traffic is a 34% better.

The following figures shows the data rates experienced by final users before and after the ON is created. Bitrates in the macro layer tend to grow, whereas in the femto layer there is a significant increase, except for the maximum throughput.

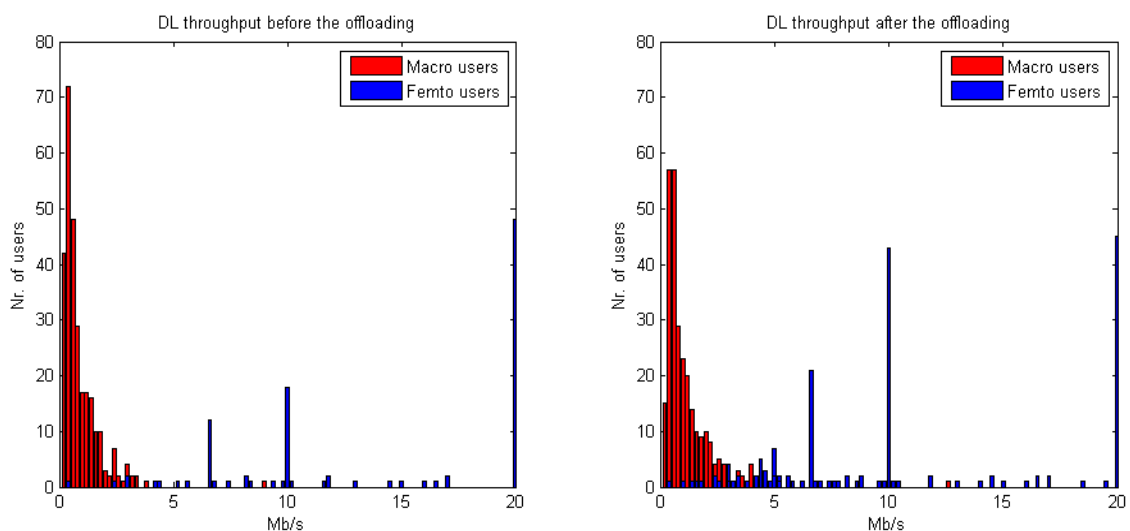


Figure 90: Per user DL throughput.

Those UEs that are part of the ON experience a minimum increase of their throughput (compared to the one they had in the macro layer) of +1 Mb/s, being the average value of +11 Mb/s. Additionally, out of the 73 offloaded users, 70 of them were disconnected when in macro, and then regained the connection thanks to the ON.

Next figure depicts how the total traffic handled by the macro and femto nodes is increased due to the presence of the ON. The effect is more evident in the macro layer, where there is an important increase in the maximum load.

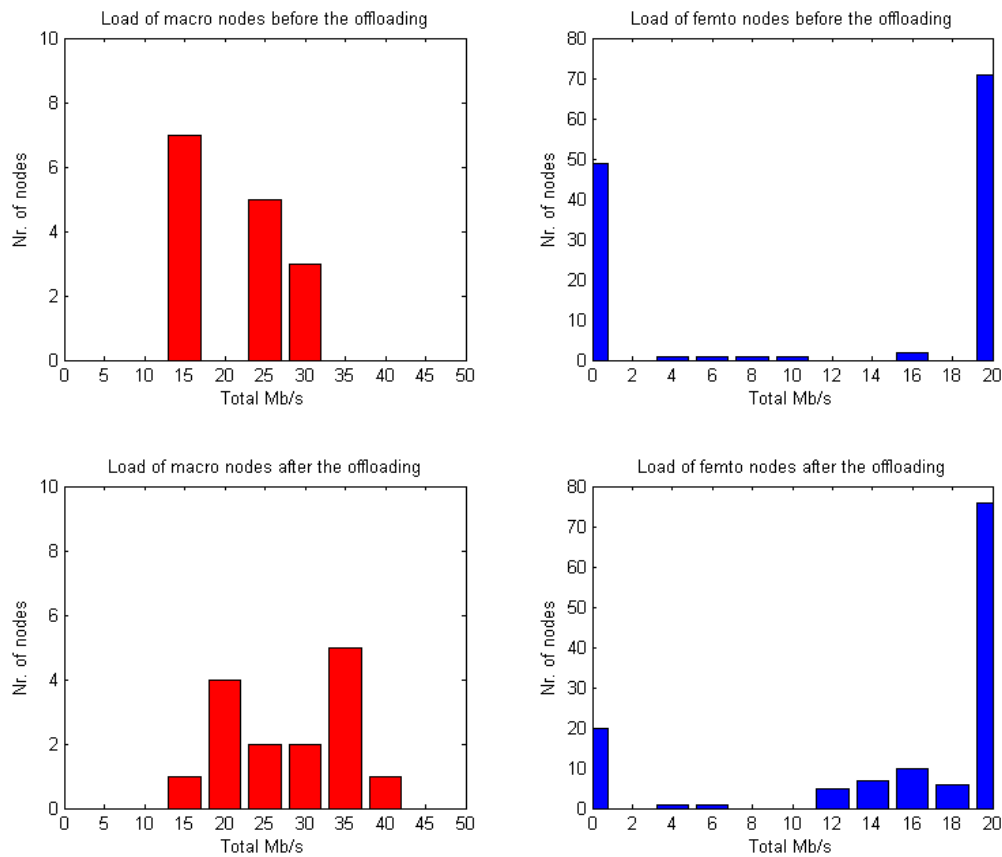


Figure 91: Per user DL throughput.

These figures confirm that the OF is well-designed and that the optimization method leads to a situation much better than the original.

2.15.4.2 Results for Scenario 1

This section describes the performance of the algorithm for solving the challenges of OneFIT Scenario 1 (see D2.1 [2] Section 4.1). In particular, a situation similar to use case 3 ("Coverage extension via an access point") is faced.

In our test scenario, the coverage of the macro layer is artificially degraded and the traffic will be partially derived to femtos by creating the corresponding ONs. The performance of the ON suitability algorithm will be assessed by direct comparison of the network KPIs before and after the ON creation.

Test case 1

In the first test case, a single sector of one of the macro sites fails and its users are derived to neighbouring cells.

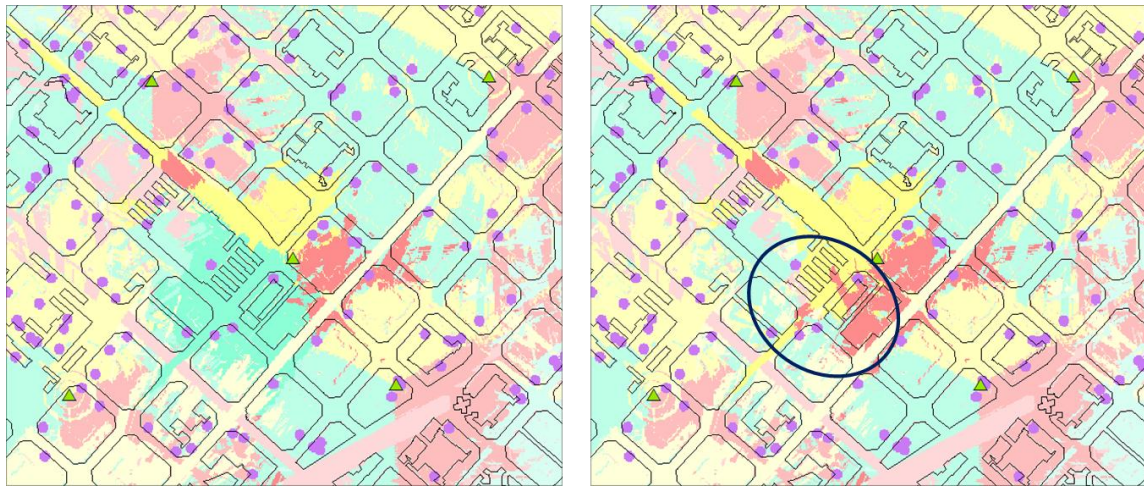


Figure 92: Best server map before and after a macro cell is turned off.

The node selection procedure selects 86 UEs as candidates for the ONs. Using the same configuration parameters for the OF than in the baseline simulation, the optimization finally selects 69 UEs and 62 femtos. Measured KPIs are the following:

Table 9: Test case 1 for Scenario 1 results

		Without ON	With ON
Macro	Nr. of UEs	389	320
	Nr. of disconnected UEs	105	38
	Mean DL throughput per UE	0.9 Mb/s	1.1 Mb/s
	Total DL traffic in all nodes	247 Mb/s	314 Mb/s
Femto	Nr. of UEs	115	184
	Nr. of disconnected UEs	11	11
	Mean DL throughput per UE	14.1 Mb/s	11.4 Mb/s
	Total DL traffic in all nodes	1466 Mb/s	1971 Mb/s

The average increase in the throughput of the offloaded users is +11.7 Mb/s (and the minimum, +1 Mb/s). The number of disconnected macro users that regain connection when joining the femto layer is 67 (out of 69 total offloaded UEs).

Results are quite similar to the ones of the baseline scenario, a bit worse for macro users and a bit better for femto ones, as expected (as not all macro users losing coverage can be successfully handovered to femto).

Test case 2

In the second test case, a whole macro sites fails and its users are derived to neighbouring cells.

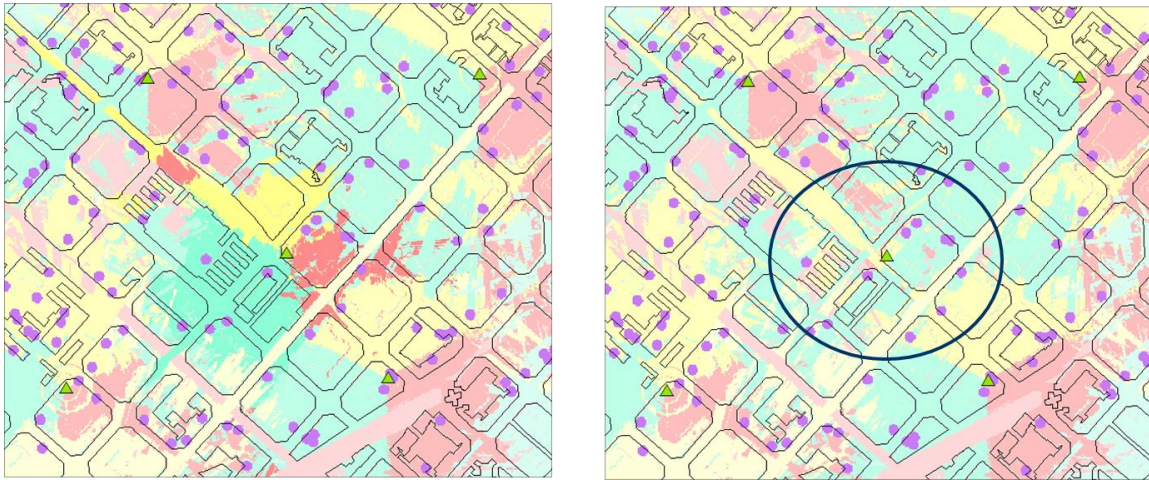


Figure 93: Best server map before and after a macro site is turned off.

The node selection procedure selects 90 UEs as candidates for the ONs. Using the same configuration parameters for the OF than in the baseline simulation, the optimization finally selects 73 UEs and 60 femtos. Measured KPIs are the following:

Table 10: Test case 2 for Scenario 1 results

		Without ON	With ON
Macro	Nr. of UEs	389	316
	Nr. of disconnected UEs	125	54
	Mean DL throughput per UE	0.8 Mb/s	1 Mb/s
	Total DL traffic in all nodes	206 Mb/s	267 Mb/s
Femto	Nr. of UEs	115	188
	Nr. of disconnected UEs	9	9
	Mean DL throughput per UE	13.9 Mb/s	10.9 Mb/s
	Total DL traffic in all nodes	1477 Mb/s	1958 Mb/s

The average increase in the throughput of the offloaded users is +10.7 Mb/s (and the minimum, +0.9 Mb/s). The number of disconnected macro users that regain connection when joining the femto layer is 71 (out of 73 total offloaded UEs).

Results are similar to the previous ones, though the performance is degrading for both macro and femto layers, as the number of users that cannot join an ON grows.

Test case 3

In the third test case, a three adjacent macro cells, focused in the same area fail simultaneously and its users are derived to neighbouring cells.

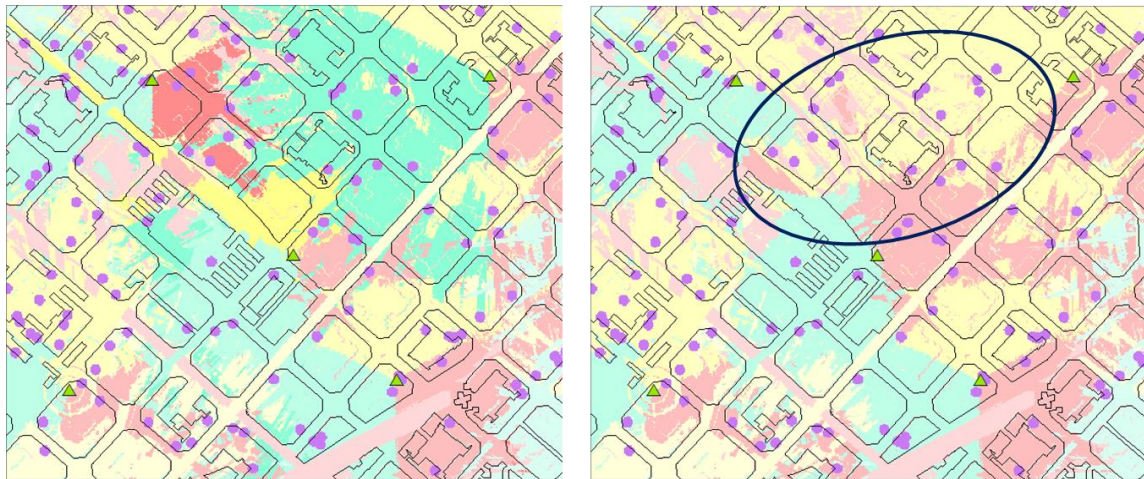


Figure 94: Best server map before and after 3 macro cells are turned off.

The node selection procedure selects 86 UEs as candidates for the ONs. Using the same configuration parameters for the OF than in the baseline simulation, the optimization finally selects 66 UEs and 57 femtos. Measured KPIs are the following:

Table 11: Test case 3 for Scenario 1 results

		Without ON	With ON
Macro	Nr. of UEs	389	323
	Nr. of disconnected UEs	129	65
	Mean DL throughput per UE	0.8 Mb/s	1.1 Mb/s
	Total DL traffic in all nodes	219 Mb/s	282 Mb/s
Femto	Nr. of UEs	115	181
	Nr. of disconnected UEs	10	10
	Mean DL throughput per UE	14 Mb/s	11.6 Mb/s
	Total DL traffic in all nodes	1465 Mb/s	1983 Mb/s

The average increase in the throughput of the offloaded users is +11.7 Mb/s (and the minimum, +1.3 Mb/s). The number of disconnected macro users that regain connection when joining the femto layer is 64 (out of 66 total offloaded UEs).

Results are slightly better than the ones of test case 2: even though the number of failing macro cells is the same, the number of affected UEs is better and they are geographically concentrated, so it is easier to manage them.

Test case 4

The fourth test case is a stress test, designed to evaluate ability of ONs to enhance the resilience of the operator's network when facing a generalized node failure. In this case, all five macro sites of the focus area are progressively disconnected, increasing the need for macro UEs to offload to the femto layer.

The following figures show the coverage maps of the focus zone for the different macro failure subcases:

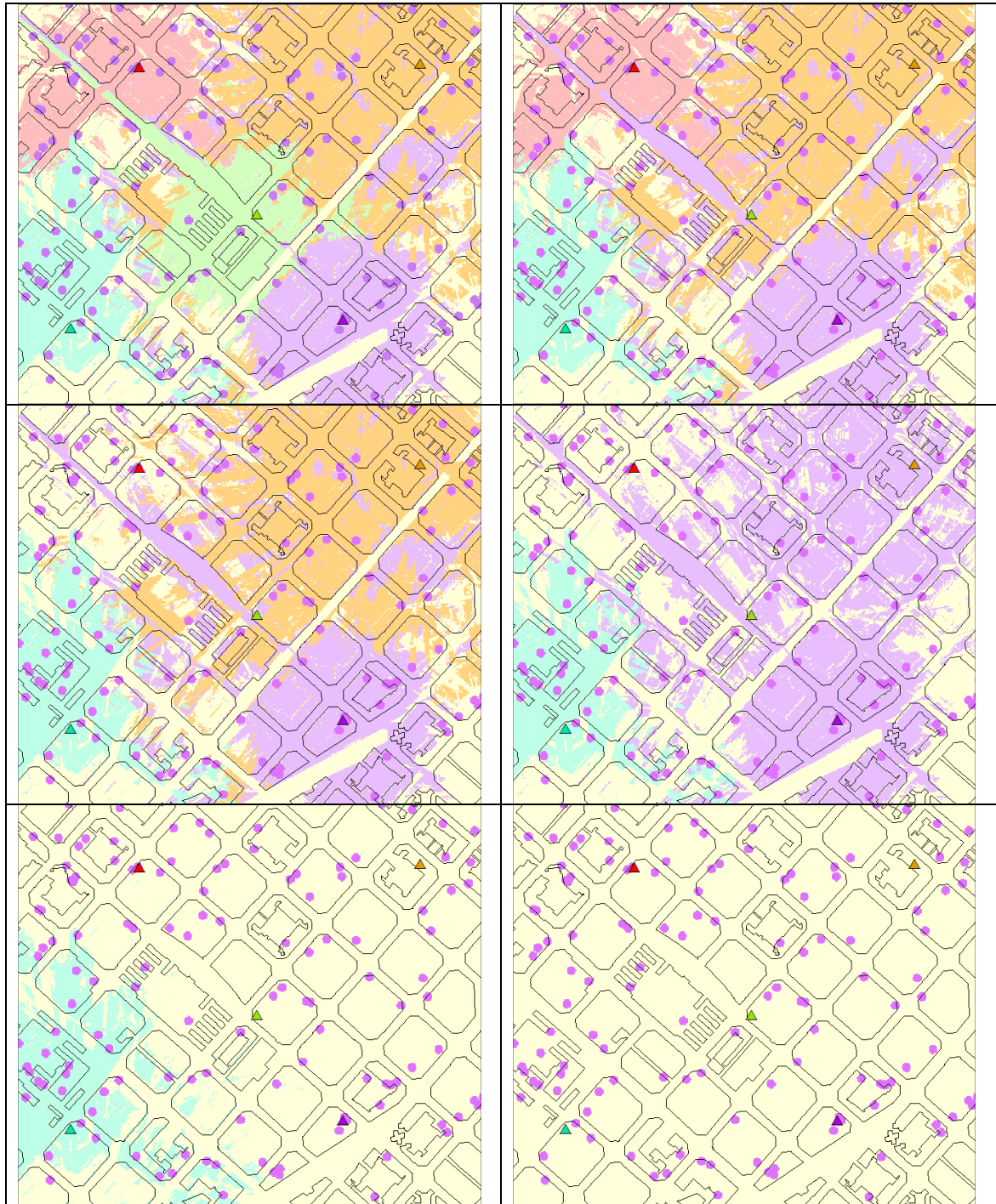


Figure 95: Best server maps as all 5 macro cells are subsequently turned off.

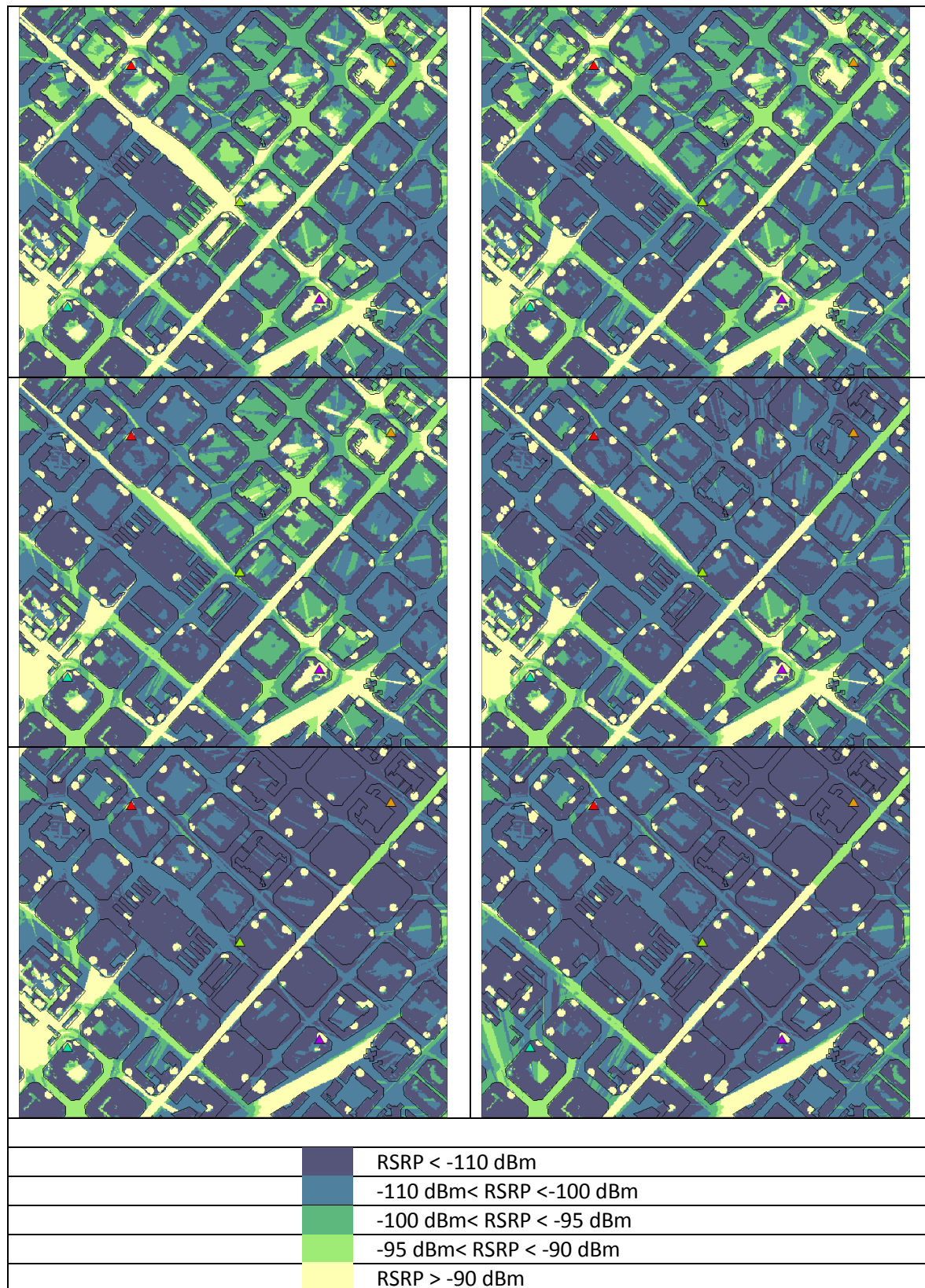


Figure 96: RSRP maps as all 5 macro cells are subsequently turned off.

Even with all five macro sites turned off, there is still enough coverage of the macro layer (from macros outside the focus zone) to absorb some of the macro users.

Measured KPIs for the stress test are presented in the following figures showing their evolution as the number of failing macro sites grows.

In the first place, the number of disconnected macro users grows almost linearly as the sites are turned off. The number of candidate UEs also grow, but with a less steep slope (not all UEs are close enough to a femto to be eligible), as do the number of finally selected nodes. Even so, the reduction on the number of disconnected users after the ONs are created is quite noticeable, being over 40% in most cases.

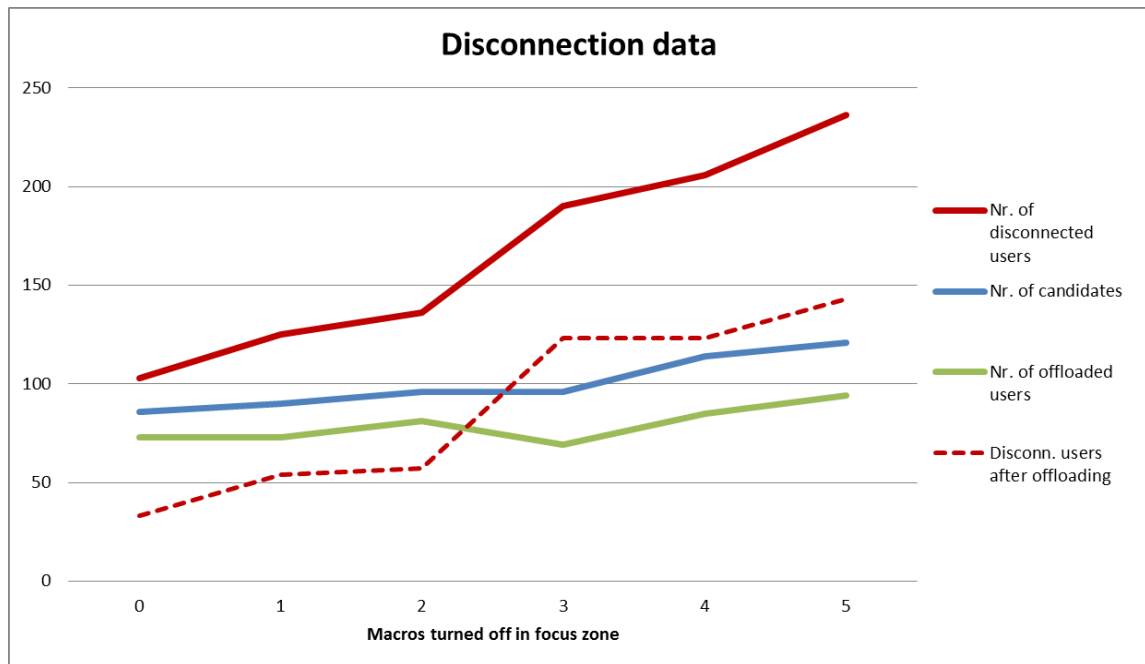


Figure 97: Number of disconnected macro users.

In fact, when most of the macro sites are turned off, there are more UEs connected to femto nodes than to the macro layer:

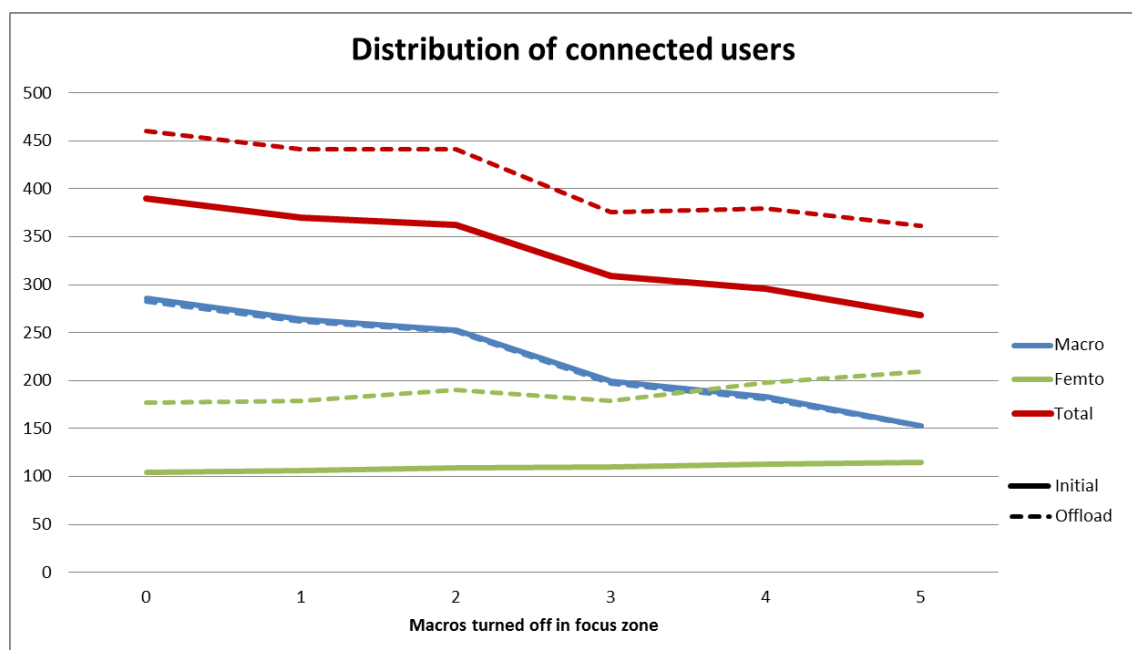


Figure 98: Distribution of connected users.

Regarding the DL throughput, the value for the average macro user tends to slightly increase after the ON is created (although, in percentage, the increase is over 30%), and for the average femto user, the decrease is noticeable (in percentage, over 20%).

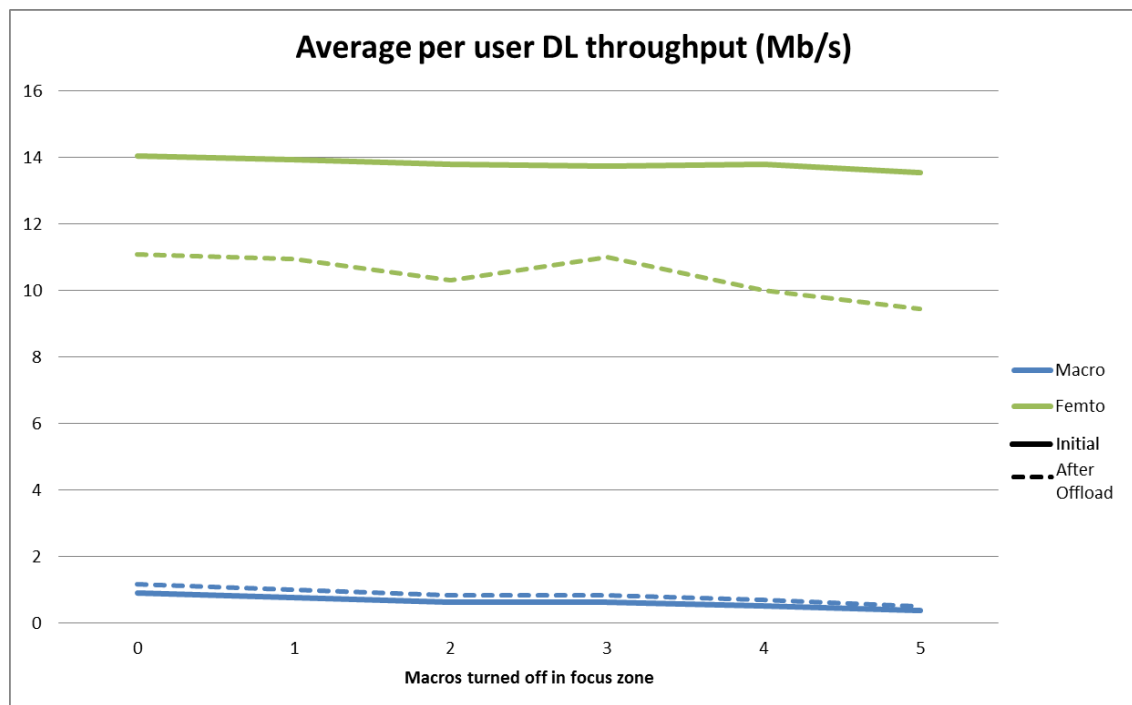


Figure 99: Average DL throughput.

The total DL traffic generated by macro UEs decreases as the sites turn off, because the number of UEs that can connect to macros outside the focus zone is low. However, the presence of the ON enhances these figures by nearly 30%. In the femto case, the traffic initially handled remains constant, but the enhancement is about the same magnitude than in the macro case.

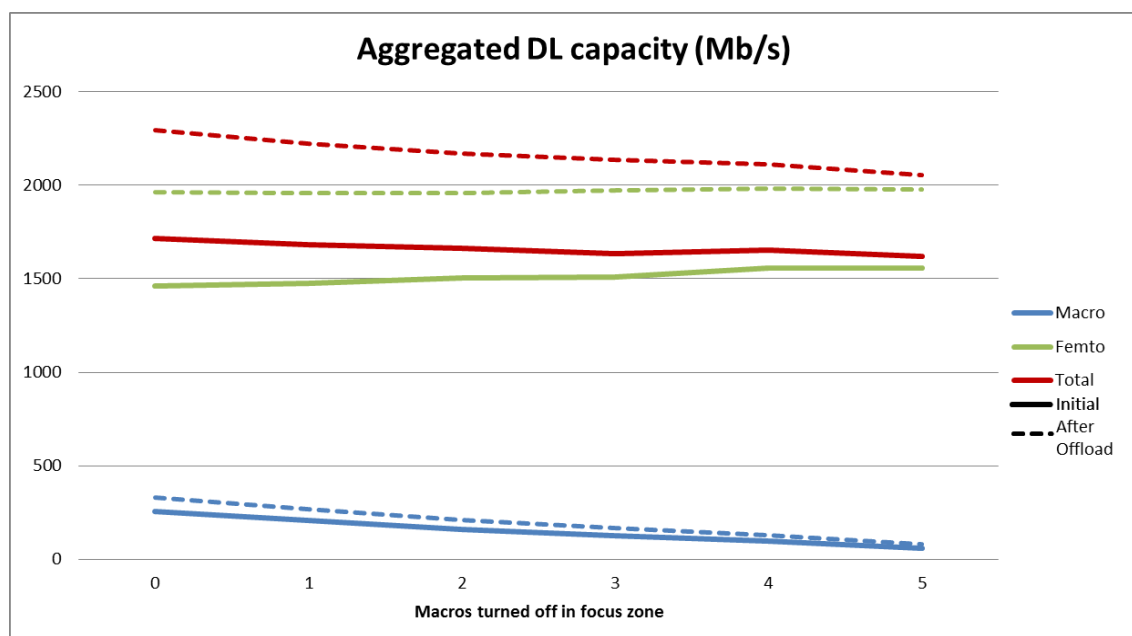


Figure 100: Total DL traffic.

For the UEs that join the ON, an important enhancement in their DL throughput is reported. That increase is better than 1 Mb/s, and usually over 10 Mb/s (although it decreases as the macros are turned off). This is due to the fact that almost all offloaded users were disconnected before the ON is created, so they get a great enhancement.

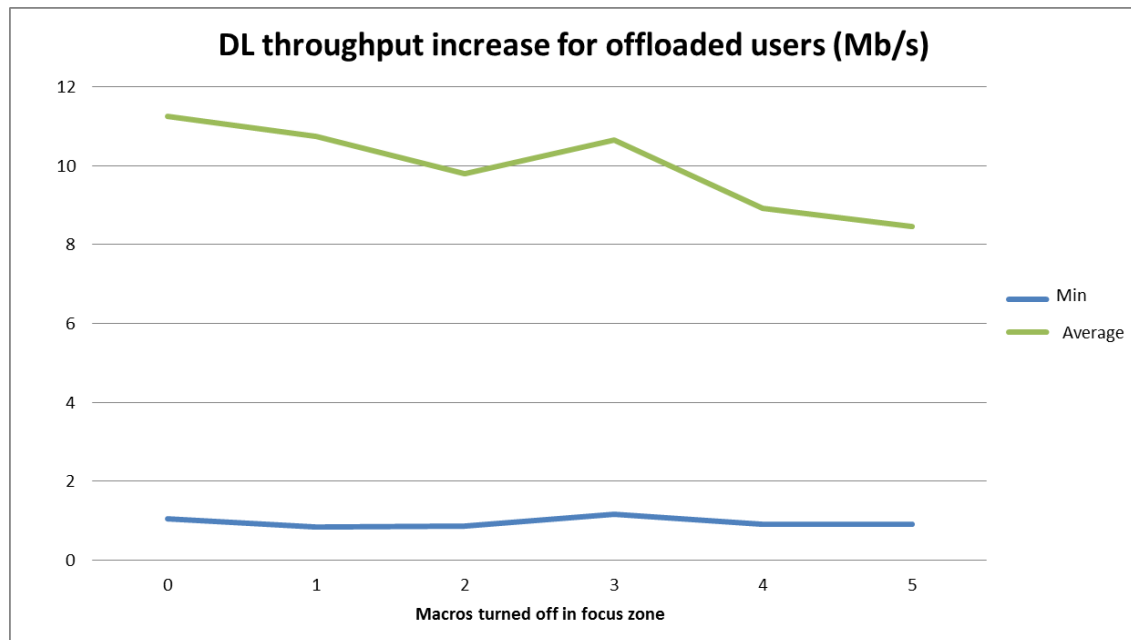


Figure 101: Throughput enhancement for ON users.

Conclusions for Scenario 1

The use of a macro-to-femto ON approach shows good performance results for addressing coverage problems of the infrastructure network. A lot of users with connection problems can easily be reconnected via ON, as long as the neighbouring femtos are distributed in a proper way to absorb incoming traffic.

The entrance of external users to the femto layer leads to a noticeable decrease in the per-user DL throughput, as expected: the resources of a single femto are quite limited and they have to be shared among existing and new UEs. Nevertheless, the total throughput handled by all the nodes in the scenario increases, due to the reduction in the disconnection rate.

2.15.4.3 Results for Scenario 2

This section describes the performance of the algorithm for solving the challenges of OneFIT Scenario 2 (see D2.1 [2] Section 4.2). In particular, a situation similar to use cases 1 and 2 ("Congestion solving for cell-edge users" and "Macro/femto management") is shown.

In our test scenario, the number of UEs of the macro layer is artificially increased, so the cells become saturated. In this situation, macro cells react by disconnecting low performance users and by lowering the throughput of the remaining ones. Introducing femto ONs, part of these underperforming users can be transferred to the femto layer, thus both enhancing the performance of those UEs and alleviating the situation of the macros.

Test case 1

In the first test case, the number of UEs in the vicinity of the central macro site of the scenario is progressively increased: from the initial 1200 scenario in +100 users steps. The objective is to get that macro into saturation and evaluate if the ON approach is able to turn it back to a normal state. Measured KPIs for the test case are presented in the following figures showing their evolution as the number of UEs grows.

In the first place, the number of disconnected macro users grows quickly as the number of users increases. However, the number of candidates grows very slowly, as the number of femtos available in the vicinity of the central site is limited. Therefore, the reduction in the disconnection rate is almost constant in number (and decreasing in percentage).

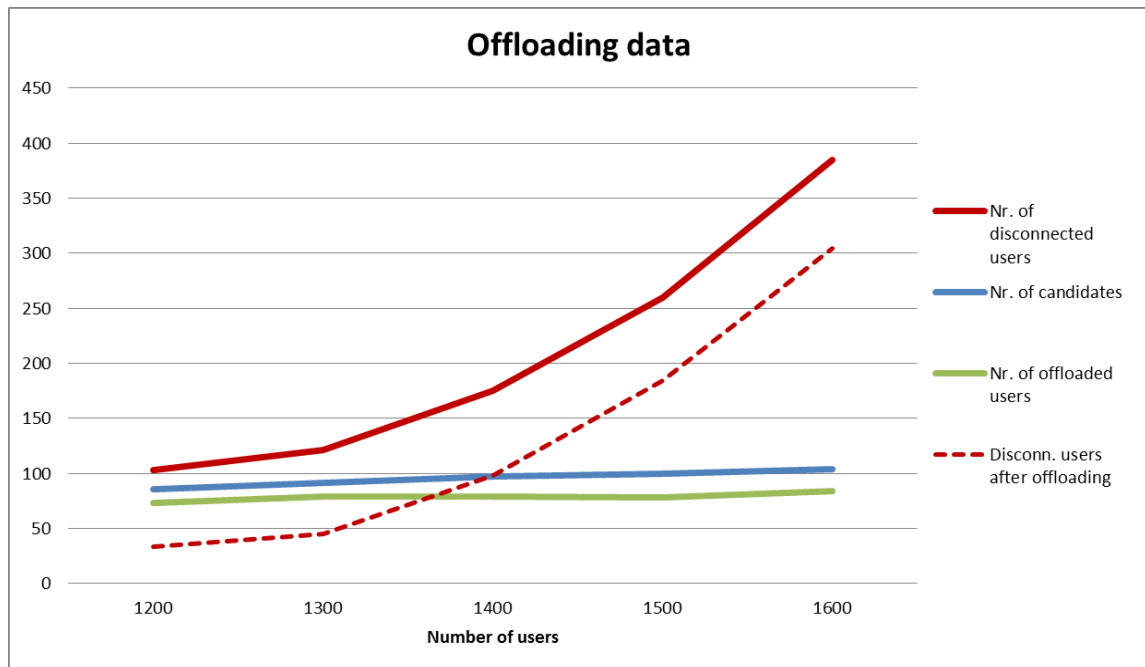


Figure 102: Number of disconnected macro users.

In this case, the number of users connected to macros is always strictly superior to the femto connections, but there is a noticeable saturation on the number of users that macros can handle.

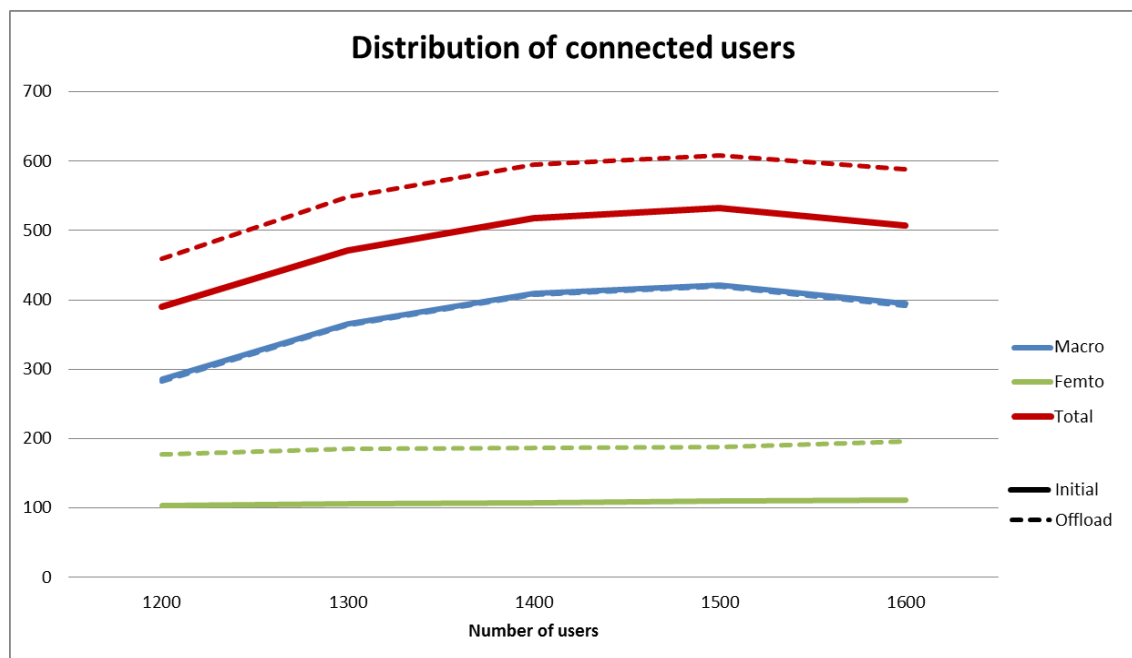


Figure 103: Distribution of connected users.

If we focus on what happens at the saturated cell, the situation is even worse: the number of disconnected users grows very fast and the neighboring femtos are not enough to attract them. Thus, the disconnection enhancement becomes negligible as the number of users increases:

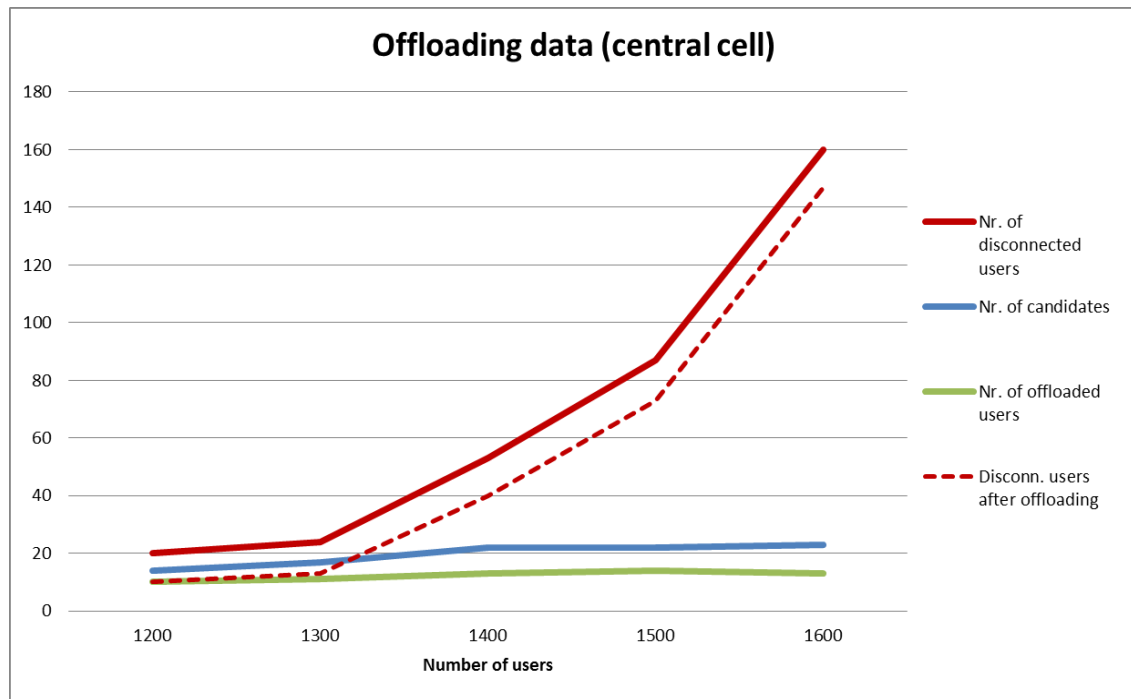


Figure 104: Number of disconnected macro users.

Regarding the DL throughput, the value for the average macro user tends to slightly increase after the ON is created (in percentage, the increase is over 20%, but decreasing with the number of UEs), and for the average femto user, the decrease is noticeable (over 20%, almost stable).

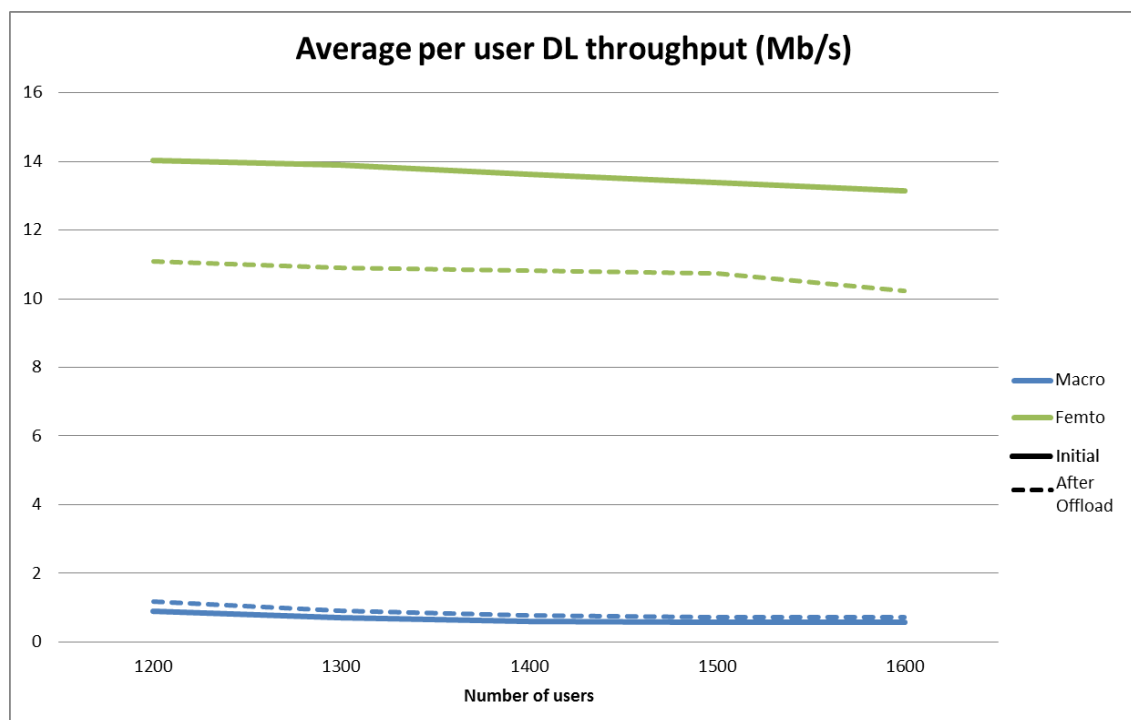


Figure 105: Average DL throughput.

In the central saturated cell, figures are lower than average (and worsening with the number of users):

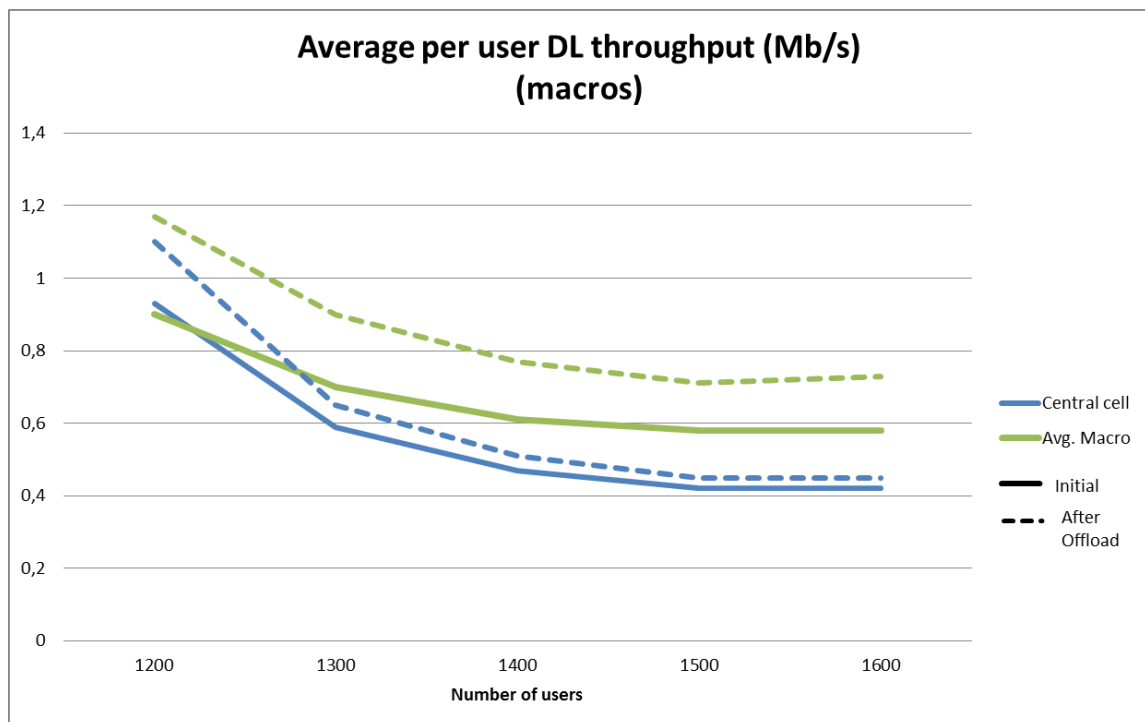


Figure 106: Macro DL throughput.

The total DL traffic handled by macro and femto nodes remains almost stable as the number of users grows because the available resources are constant. When the ON is up, some free resources are allocated to offloaded users (leading to a capacity increase of 35% in the femto layer and over 20% in the macro). However, the lack of enough femtos to absorb the traffic, makes this increase almost stable when the number of users grows.

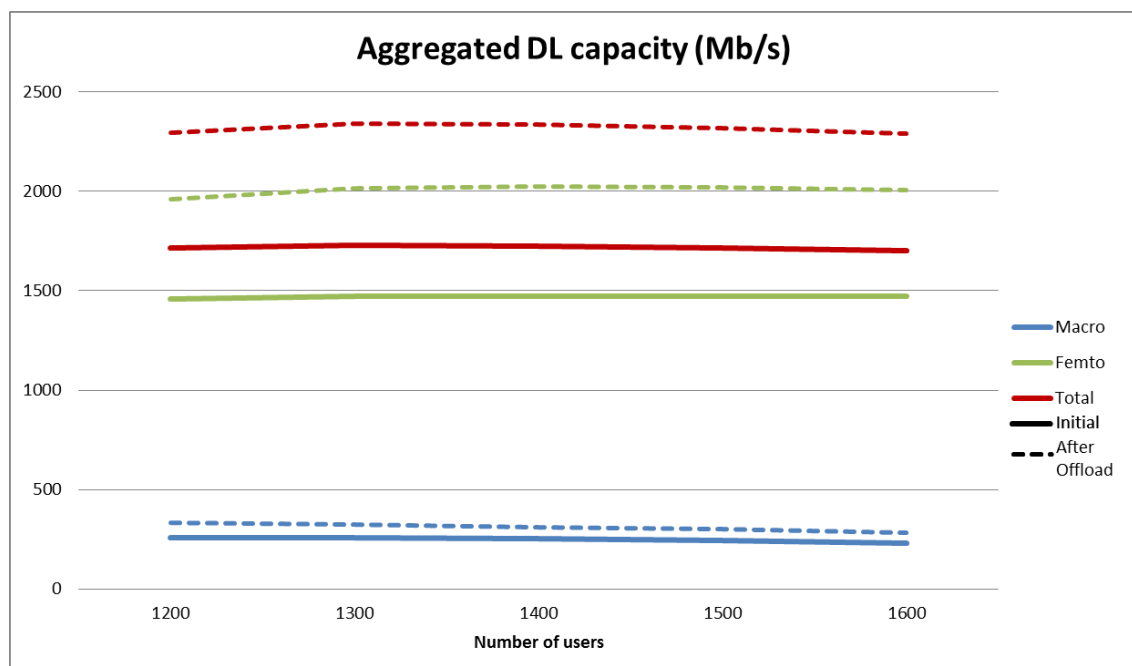


Figure 107: Total DL traffic.

The total traffic on the saturated cell, however, noticeably decreases in the test, and the enhancement due to the ON is very limited:

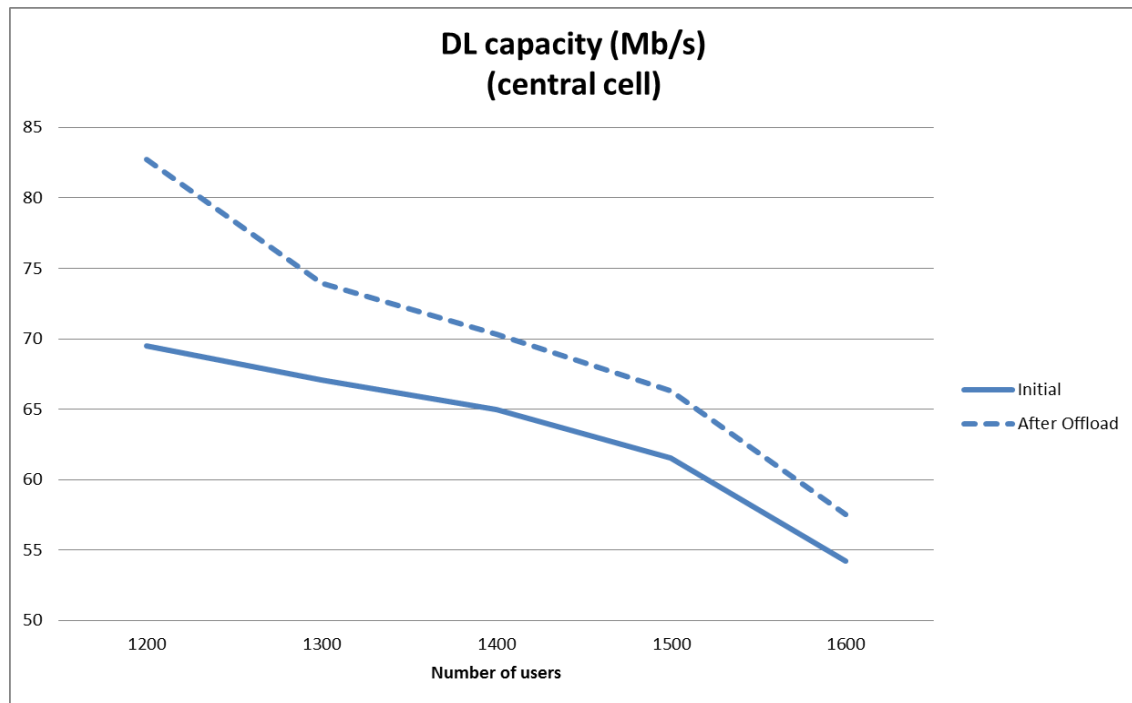


Figure 108: Total DL traffic.

For the UEs that join the ON, an enhancement in their DL throughput is measured. That increase is better than 0.5 Mb/s, and usually over 10 Mb/s (although it slightly decreases as the number of users grows). This is due to the fact that almost all offloaded users were disconnected before the ON is created, so they get a great enhancement.

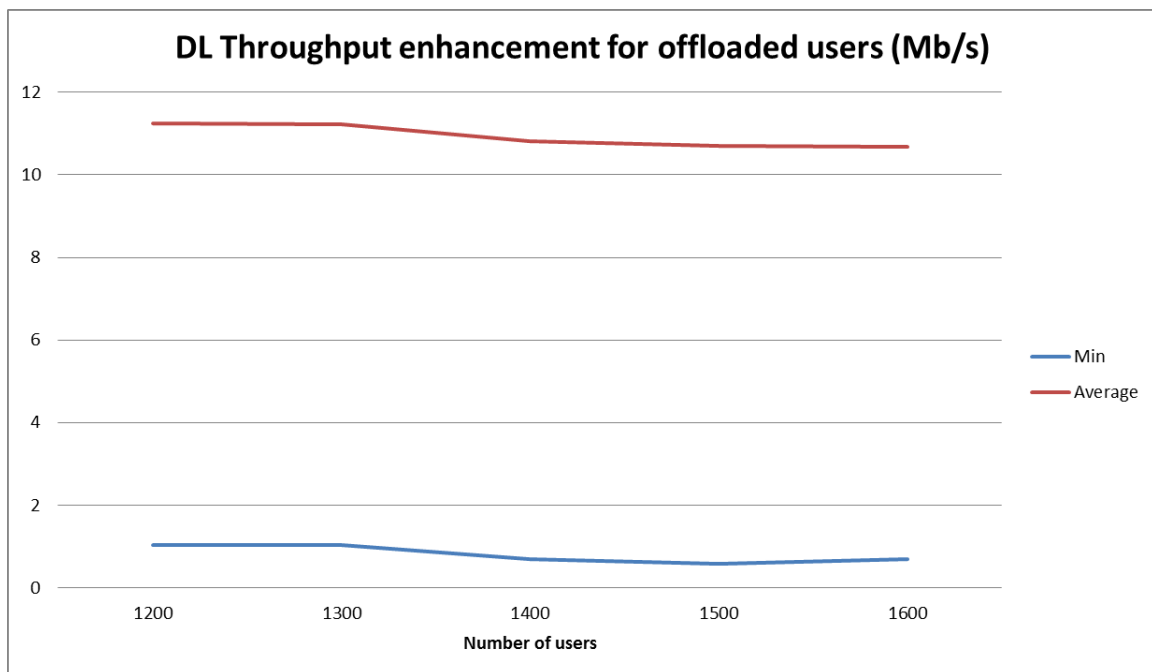


Figure 109: Throughput enhancement for ON users.

Test case 2

In the second test case, the number of UEs in the focus are of the scenario is progressively increased: from the initial 1200 scenario in +500 users steps. The objective is to get all five macro sites into saturation and evaluate if the ON approach is able to turn them back to a

normal state. Measured KPIs for the test case are presented in the following figures showing their evolution as the number of UEs grows.

In the first place, the number of disconnected macro users grows linearly as the number of users increases. The number of candidate nodes grows also linearly, but much slower, as the number of femtos of the focus area is limited. Therefore, the reduction in the disconnection rate decreases in percentage as the total number of UEs grows.

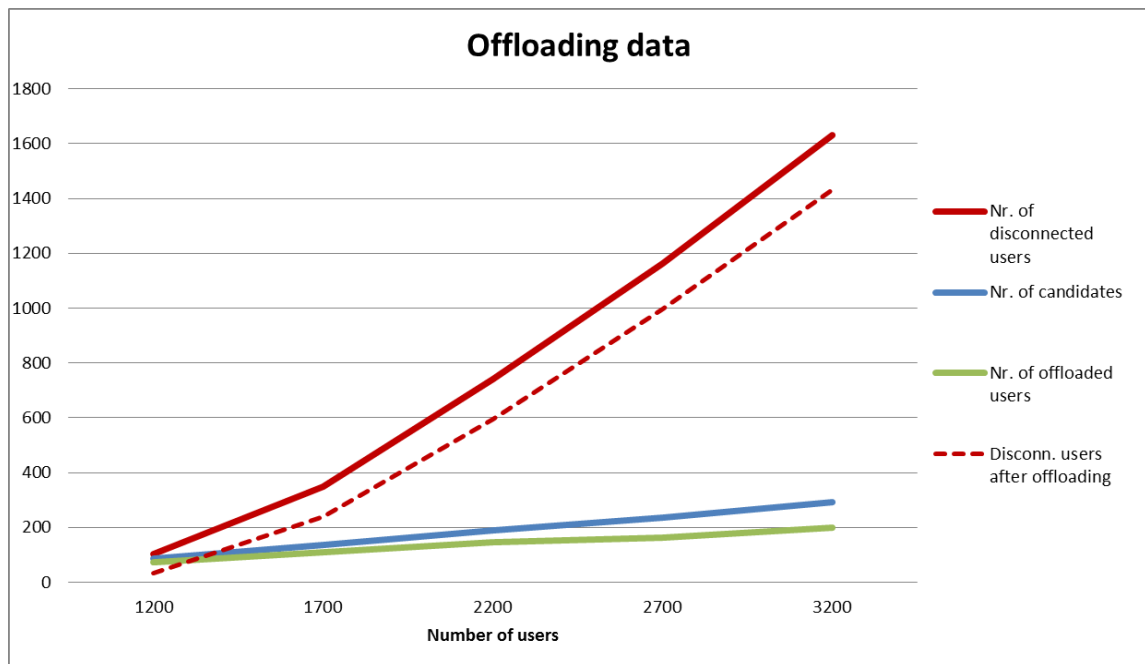


Figure 110: Number of disconnected macro users.

In this case, the number of users connected to macros is always strictly superior to the femto connections, due to the huge amount of UEs. The difference in the macro connection rate before and after the ON is created is almost negligible, due to the fact that almost all UEs in the ON were taken from the disconnected pool.

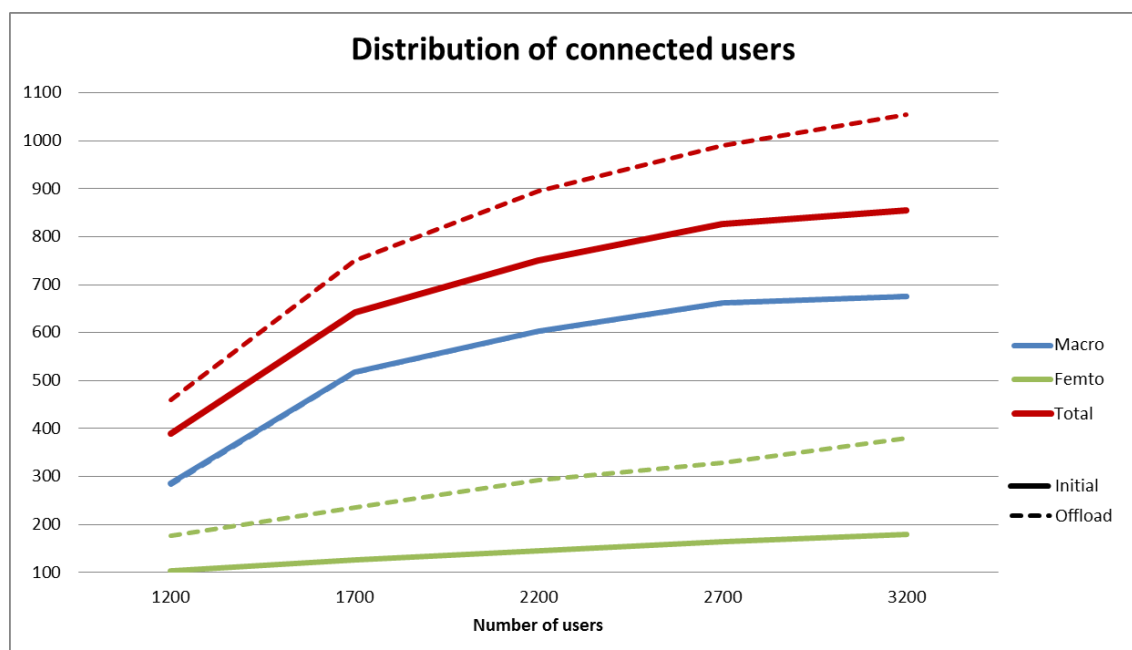


Figure 111: Distribution of connected users.

Regarding the DL throughput, the value for the average macro user tends to slightly increase after the ON is created (in percentage, the increase is over 15%, but decreasing with the number of UEs). The DL throughput of the average femto user decreases linearly hitting a 70% of the initial value in the most populated scenario, and the extra decrease due to the ON is also quite noticeable (ranging from 20% to 45% as the number of users grows).

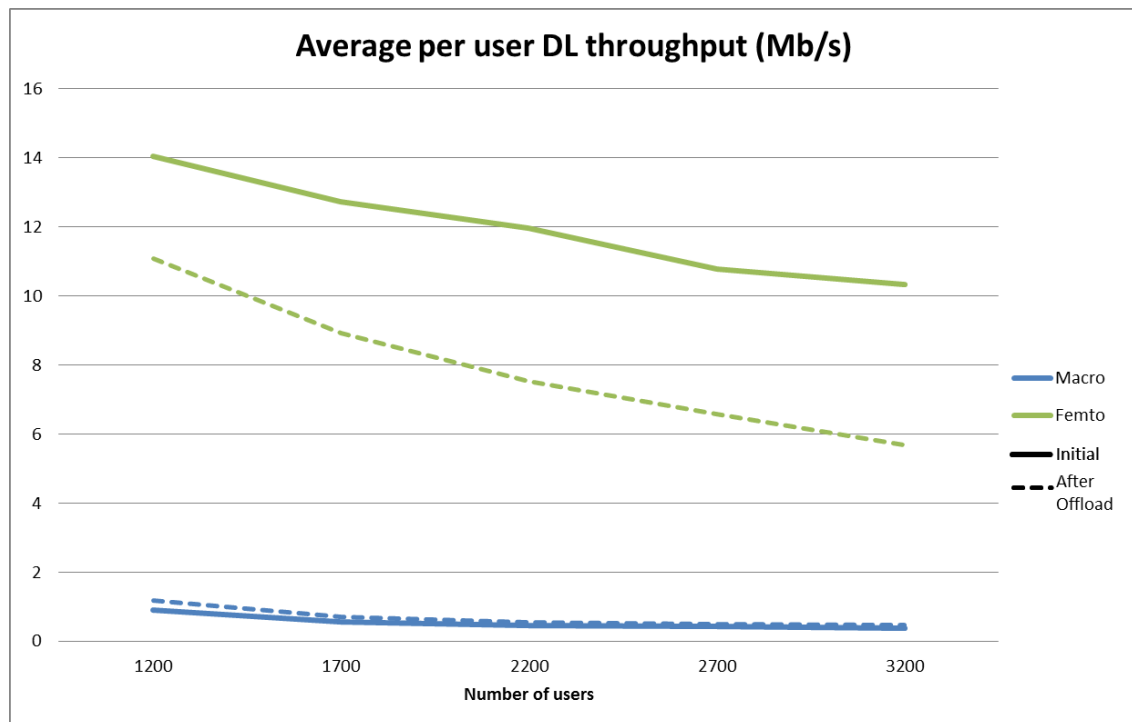


Figure 112: Average DL throughput.

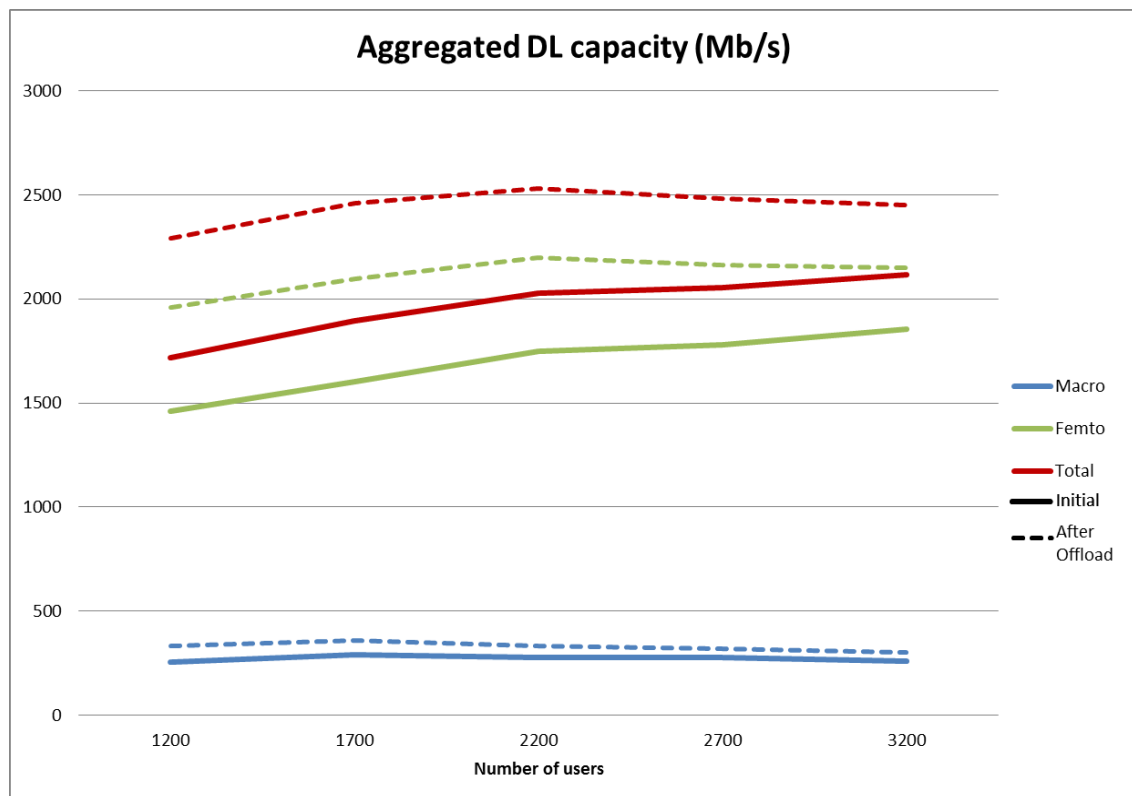


Figure 113: Total DL traffic.

The total DL traffic handled by macro nodes remains almost stable as the number of users grows because the used resources are nearly constant. On the other hand, the traffic in the femto nodes increases steadily as more and more UEs are absorbed by the macro layer. When the ON is up, free resources are allocated to initially disconnected users, leading to an important total capacity increase (ranging from 35% to 15% depending on the number of users), mainly coming from the capacity handled by the femto layer. This capacity increase is reducing due to the lack of enough femtos to absorb the incoming traffic.

For the UEs that join the ON, an enhancement in their DL throughput is measured. The minimum value of the increase can be very low for the most populated scenario, but in average ranges from over 11 Mb/s to 5 Mb/s (decreasing in a steadily manner). This is due to the fact that almost all offloaded users were disconnected before the ON is created, but as the number of UEs grow, destination femtos become also quite saturated (averaging 4 UEs per femto for the most populated case).

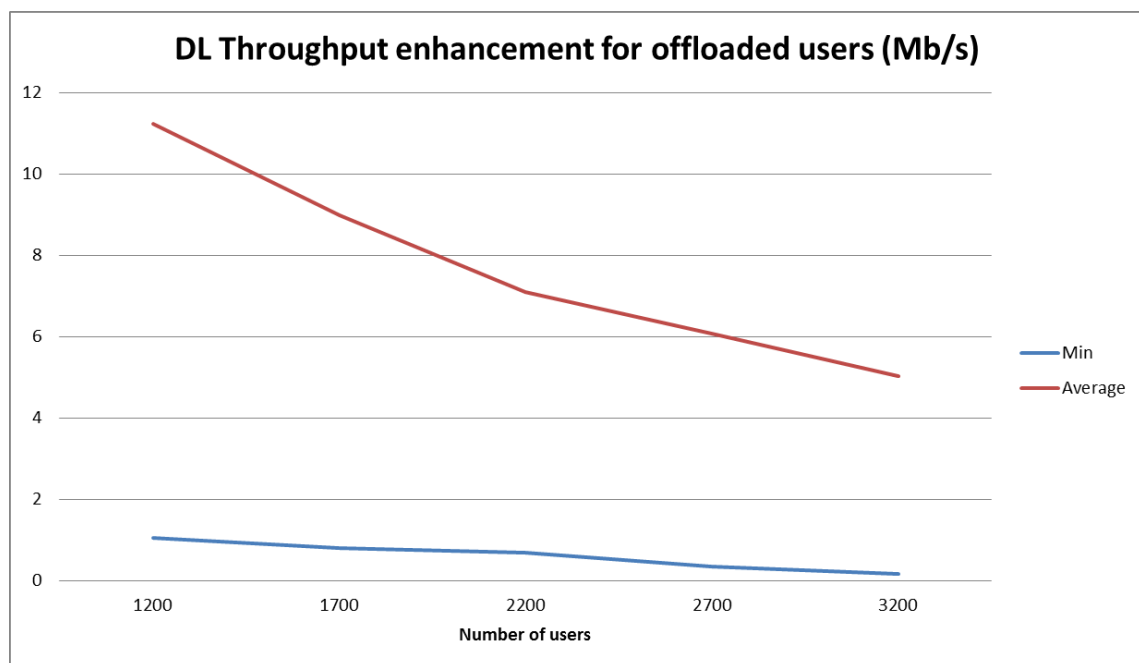


Figure 114: Throughput enhancement for ON users.

Conclusions for Scenario2

The use of a macro-to-femto ON approach shows interesting performance results for addressing capacity problems of the infrastructure network, although only limited congestion levels can be alleviated. When the number of UEs in a limited geographical area grows over expected values, the disconnection rate boosts, due to the lack of macro resources to absorb all of them. However, these disconnections help maintaining a constant capacity level at the macro layer.

The algorithm presented here, selects the candidates for the ON mainly among those disconnected UEs (in order to maximize the final throughput). Unfortunately, the number of candidates is usually low, as the number of available femtos is limited. The higher the femto density, the higher the congestion level this algorithm is able to handle.

Those UEs that are transferred to the femto layer get an important throughput increase, which leads to a better utilisation of the resources of the network, even in high congestion situations.

2.15.4.4 Results for Scenario 3

This section describes the performance of the algorithm for solving the challenges of OneFIT Scenario 3 (see D2.1 [2] Section 4.3). In particular, a situation similar to use cases 1 and 3 (“Infrastructure offload” and “ONs as platforms for location-specific services”) is addressed.

In our test scenario, the number of UEs of the macro layer is artificially increased (not as much as in Scenario 2), so the cells become partially saturated. In this situation, some of the UEs cannot achieve the target throughput demanded by the running applications, so the service they need cannot be provided with enough QoS. The creation of an ON to offload part of the traffic to the femto layer may help these underachieving UEs to reach the proper QoS level, thus improving their user experience.

In the following tests, the same configuration parameters for the Objective Function presented in 2.15.4.1 are used, with the exception of $\alpha = 50$ and φ_{TGT} variable depending on the service to be evaluated (0.5 Mb/s for voice, text/IM and background data services, 1 Mb/s for web surfing and low-definition video services and 3 Mb/s for high-definition video services).

Test case 1

In the first test case, the number of UEs in the focus zone of the scenario is progressively increased (from the initial 1200 scenario in +300 users steps) and the three different throughput targets are considered. The objective is to evaluate the ability of the ON approach to help as much UEs as possible achieving the target in low and medium load situations. Measured KPIs for the test case are presented in the following figures showing their evolution as the number of UEs grows.

In the first place, the number of disconnected macro users after the ON is created is almost independent from the target, as those UEs that are not suitable for offloading are the same in all cases:

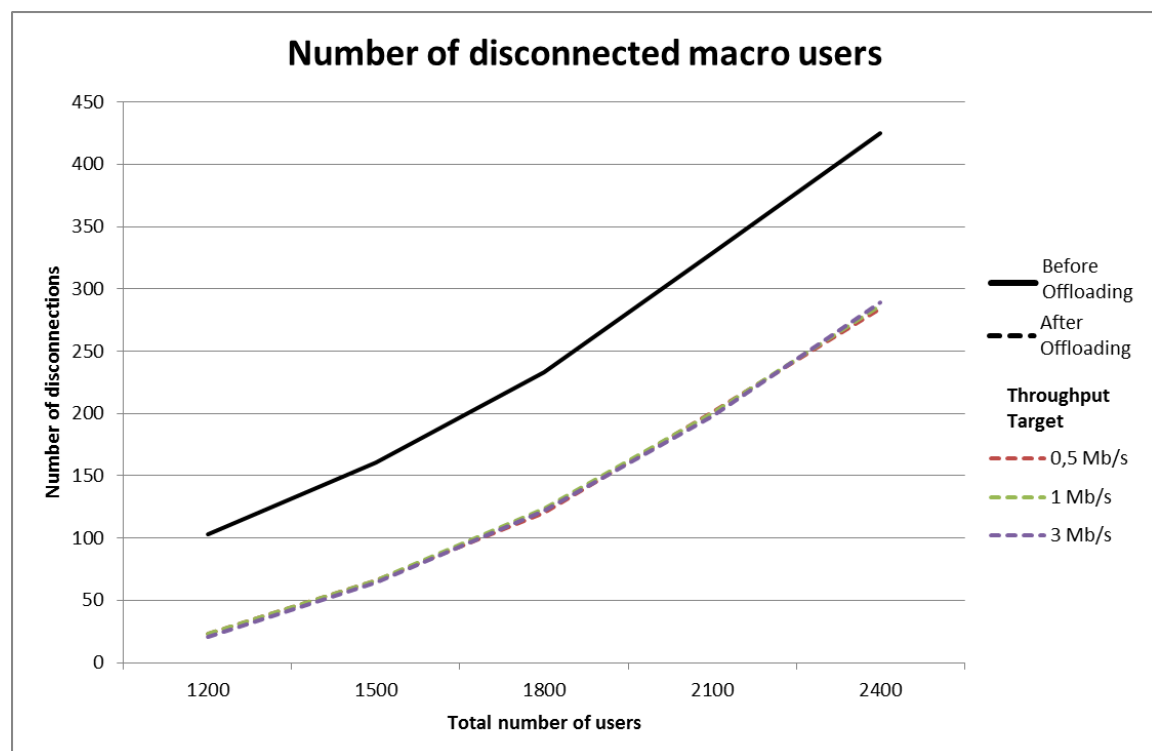


Figure 115: Number of disconnected macro users.

The users that finally get to be part of the ON are also quite similar in the three cases:

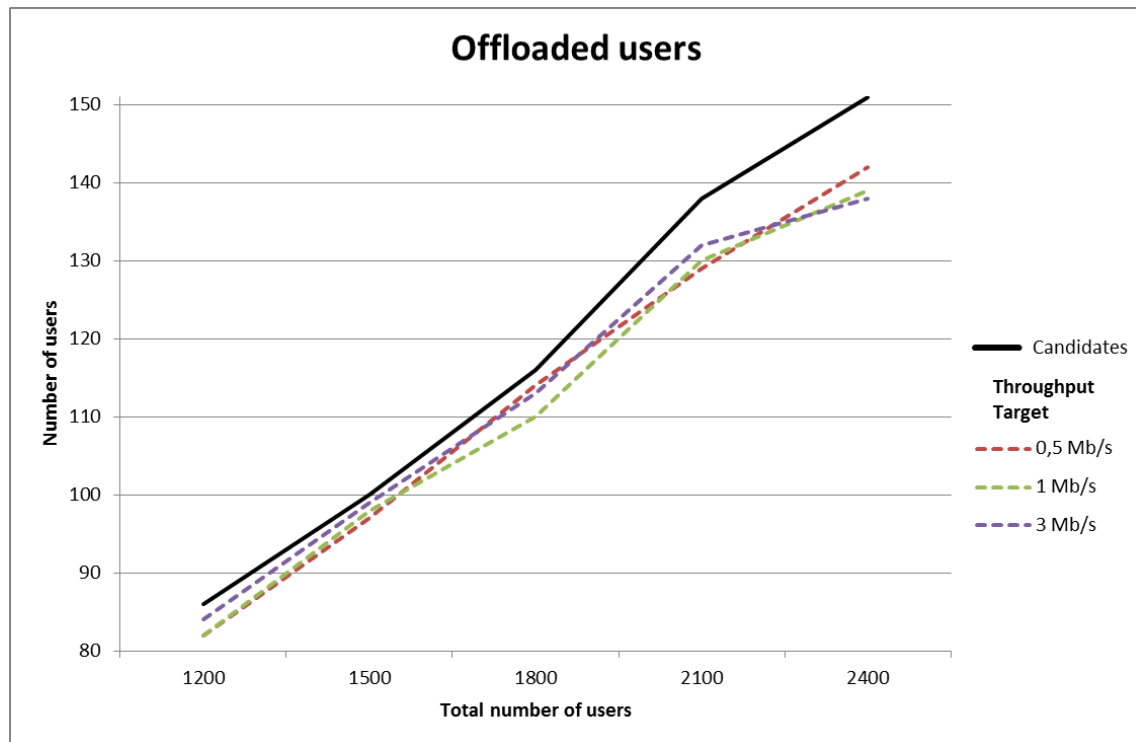


Figure 116: Number of users transferred to the ON.

Regarding the final DL throughput, the effect is similar, as it is unaffected by the target. For users that remaining in macro, there is an increase on the per user throughput over 20%, while for those that end in the femto layer, the decrease is higher than 30%.

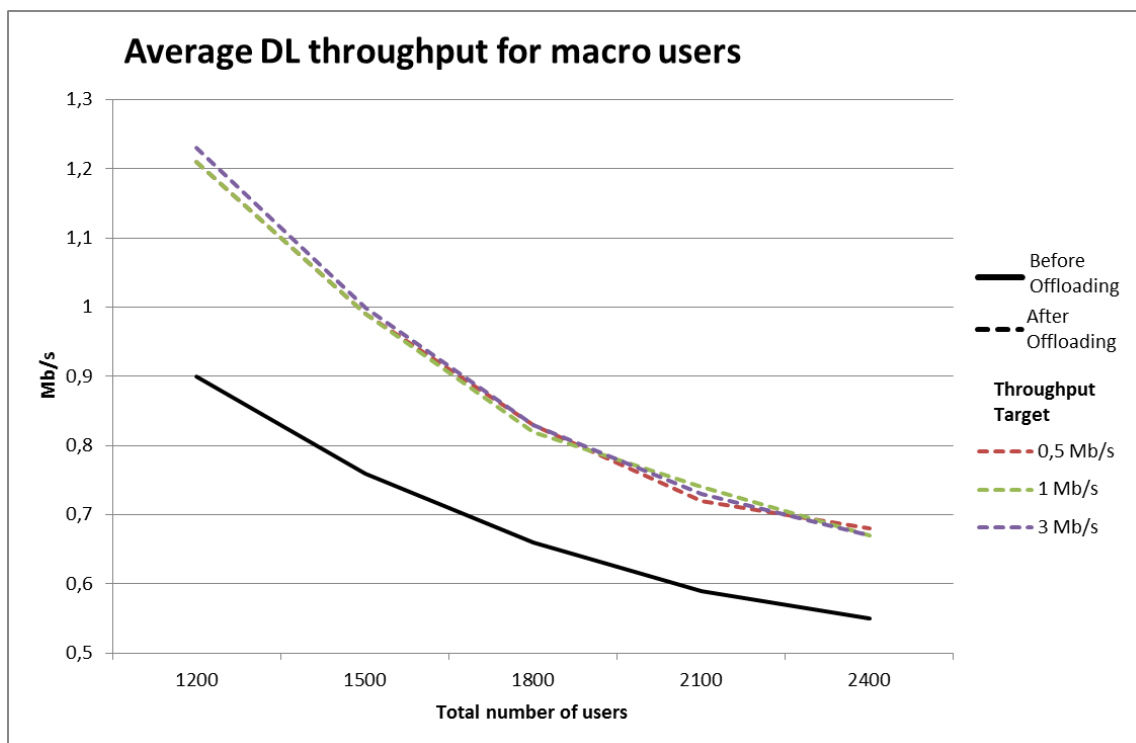


Figure 117: Average DL throughput in the macro layer.

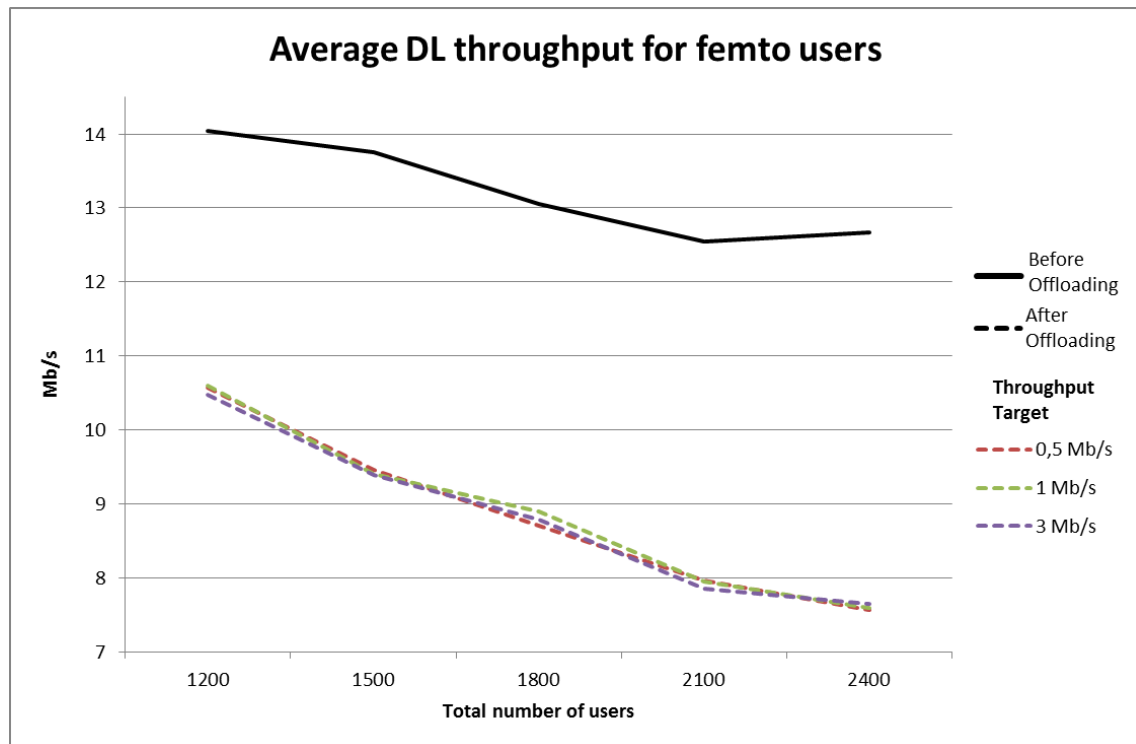


Figure 118: Average DL throughput in the femto layer.

And, obviously, for the increase in the total traffic handled by the nodes in the scenario, the result is also independent from the target and well over 20% in all cases:

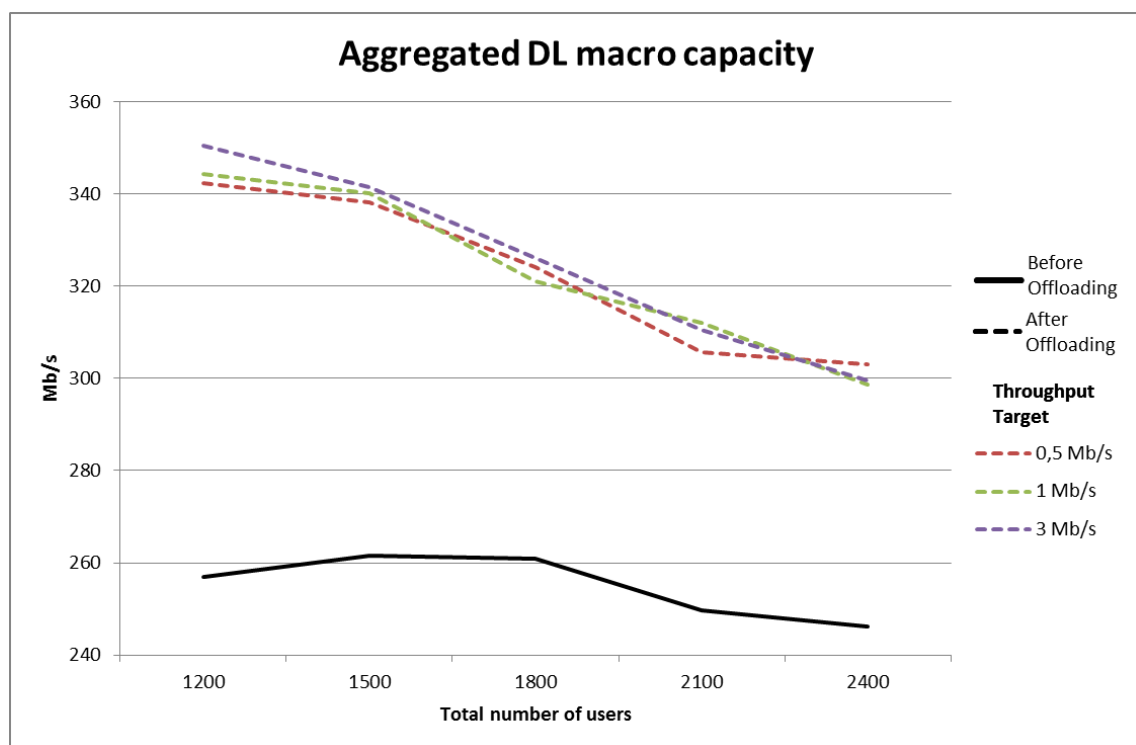


Figure 119: Total DL macro traffic.

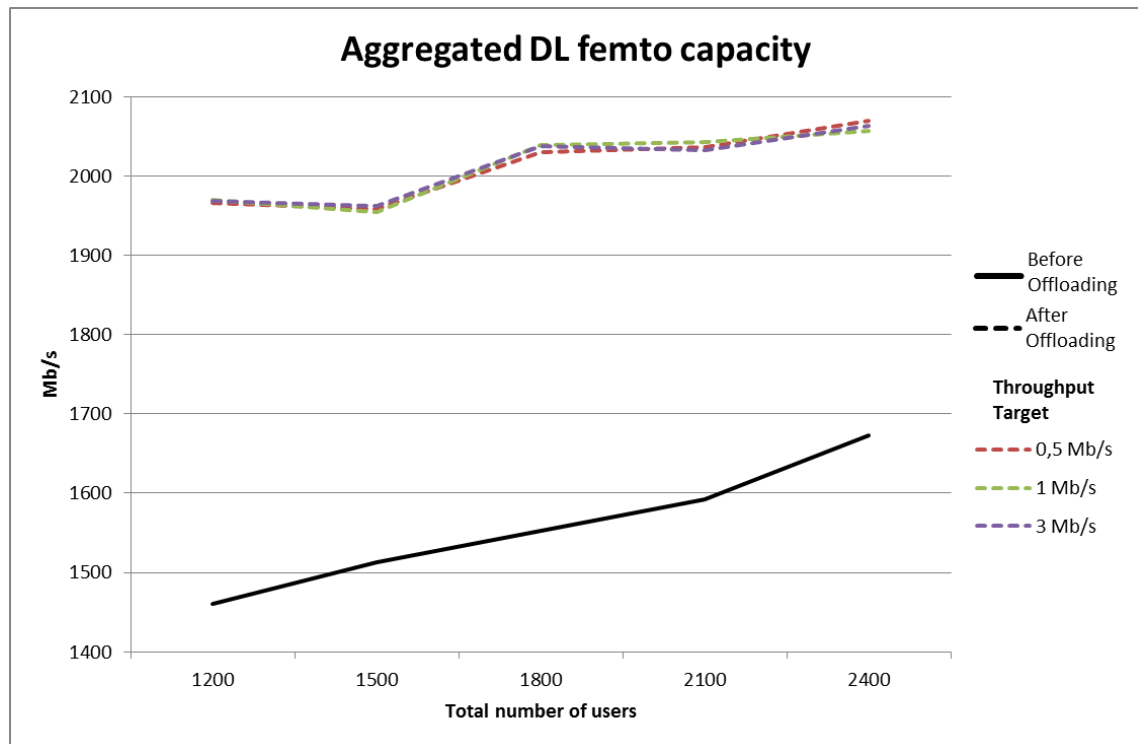


Figure 120: Total DL femto traffic.

Finally, focusing on how UEs achieve their throughput targets, there are clear differences for each subcase. For macro users, the increase in the achievement rate is quite noticeable for the lower target, even when the saturation grows. This increase reduces for higher targets, being almost negligible in the highest case.

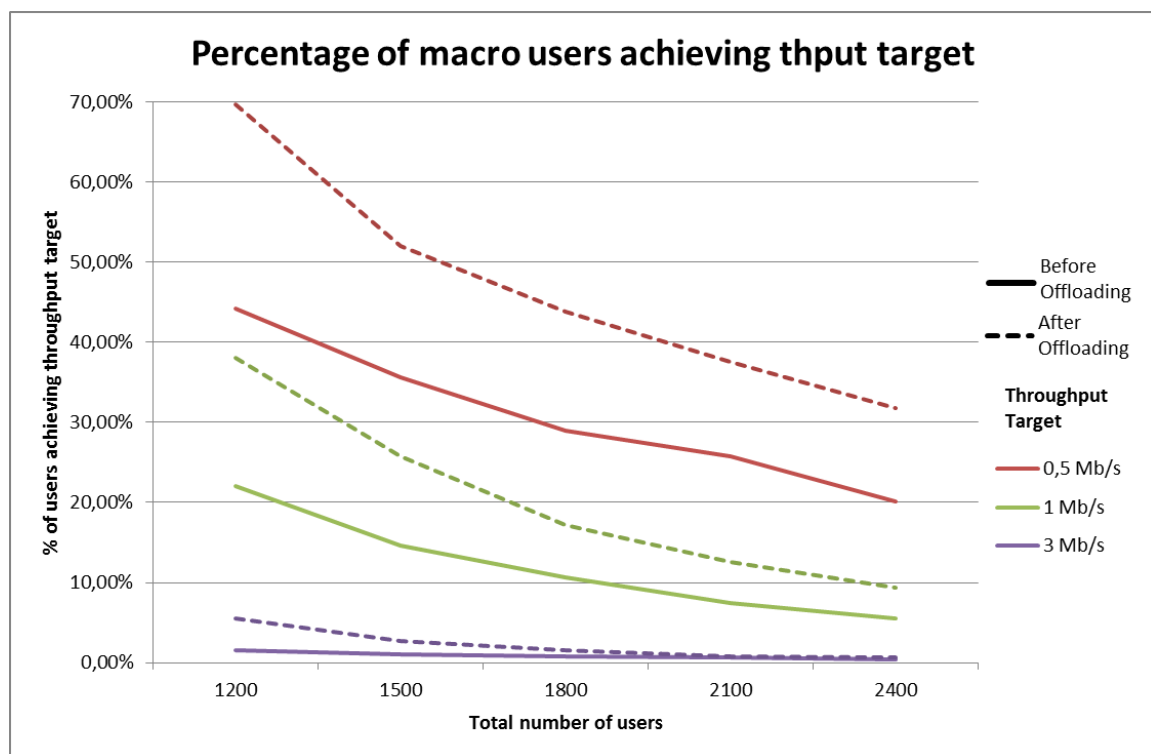


Figure 121: Achievement rate for macro UEs.

For femto users, the achievement rate is very high (it is even before the ON is created). It increases a bit for the lower targets after the offloading, due to the entrance of new achieving

users, but it slightly decreases for the highest target, because the resources have to be shared among a growing number of users and they may not be enough in highly-populated femtos.

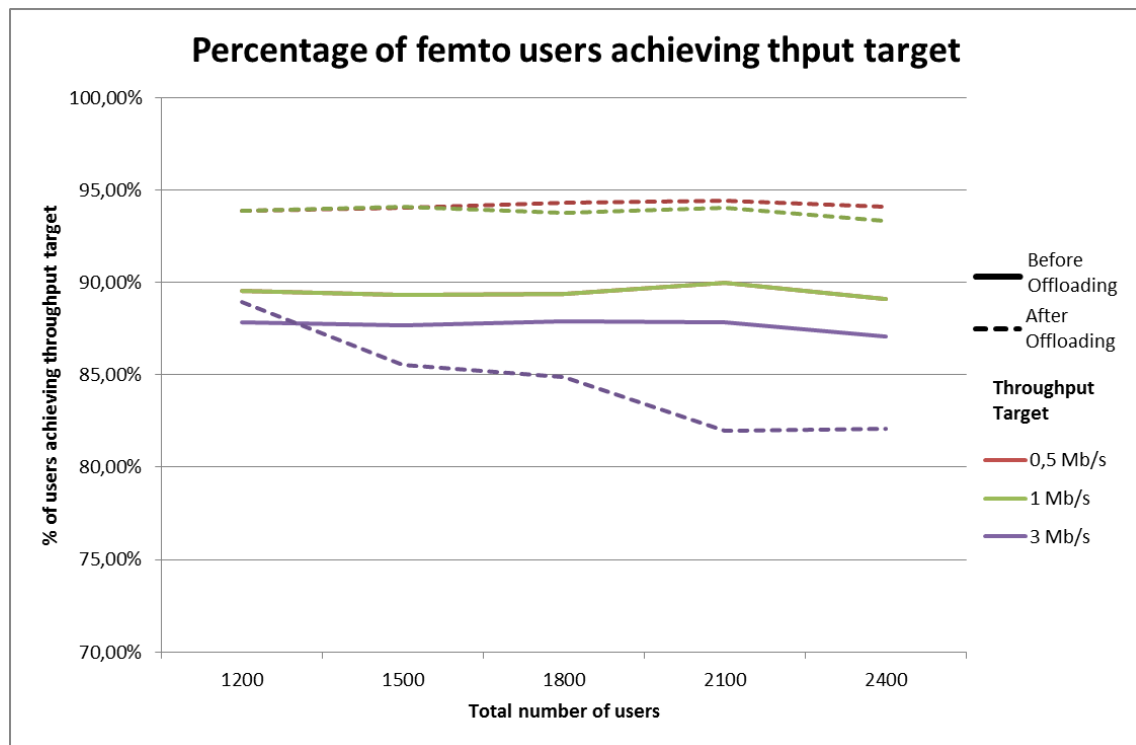


Figure 122: Achievement rate for femto UEs.

In fact, most of the macro users that become part of the ON, achieve their targets, although the rate in the highest target case is noticeably lower:

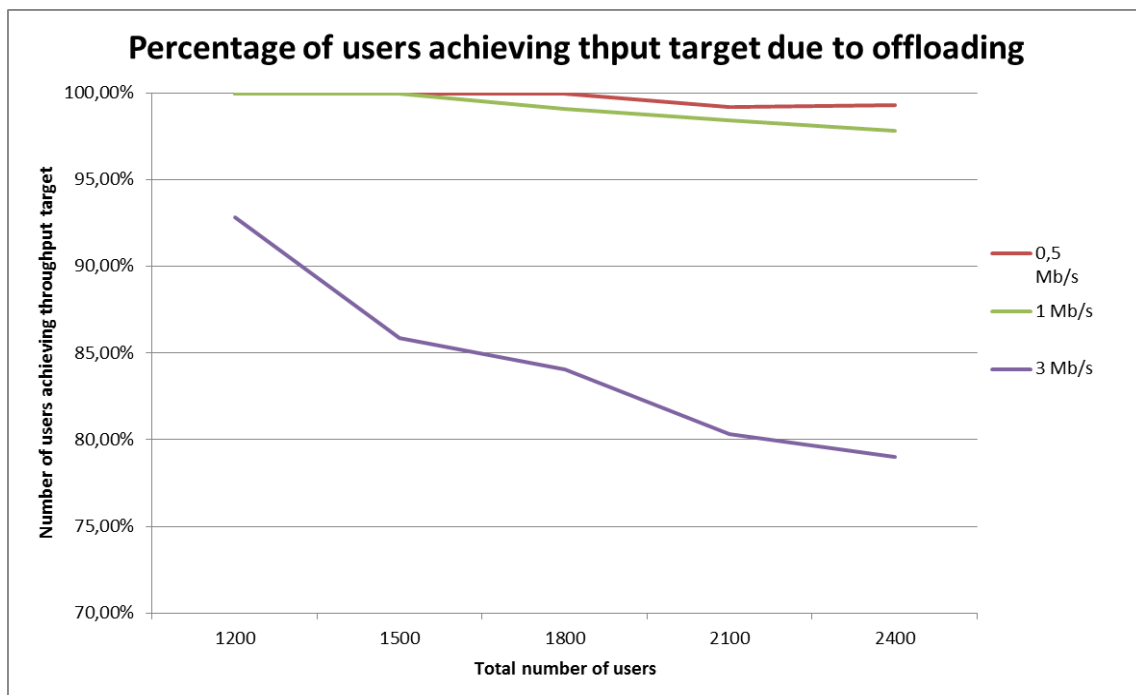


Figure 123: Achievement rate for UEs in the ON.

Test case 2

In the second test case, the number of UEs in the focus zone of the scenario is progressively increased (from the initial 1200 scenario in +300 users steps) and three different traffic mixes are considered:

- Traffic mix A: 70% of UEs demand a 0.5 Mb/s service, 20% demand 1 Mb/s and 10% demand 3 Mb/s.
- Traffic mix B: 50% demand 0.5 Mb/s, 30% demand 1 Mb/s and 20% demand 3 Mb/s.
- Traffic mix C: 30% demand 0.5 Mb/s, 40% demand 1 Mb/s and 30% demand 3 Mb/s.

The objective is to evaluate the ability of the ON approach to help as much UEs as possible achieving the target in low and medium load situations. Measured KPIs for the test case are presented in the following figures showing their evolution as the number of UEs grows.

In the first place, the number of disconnected macro users after the ON is created is almost independent from the traffic mix, as those UEs that are not suitable for offloading are the same in all cases:

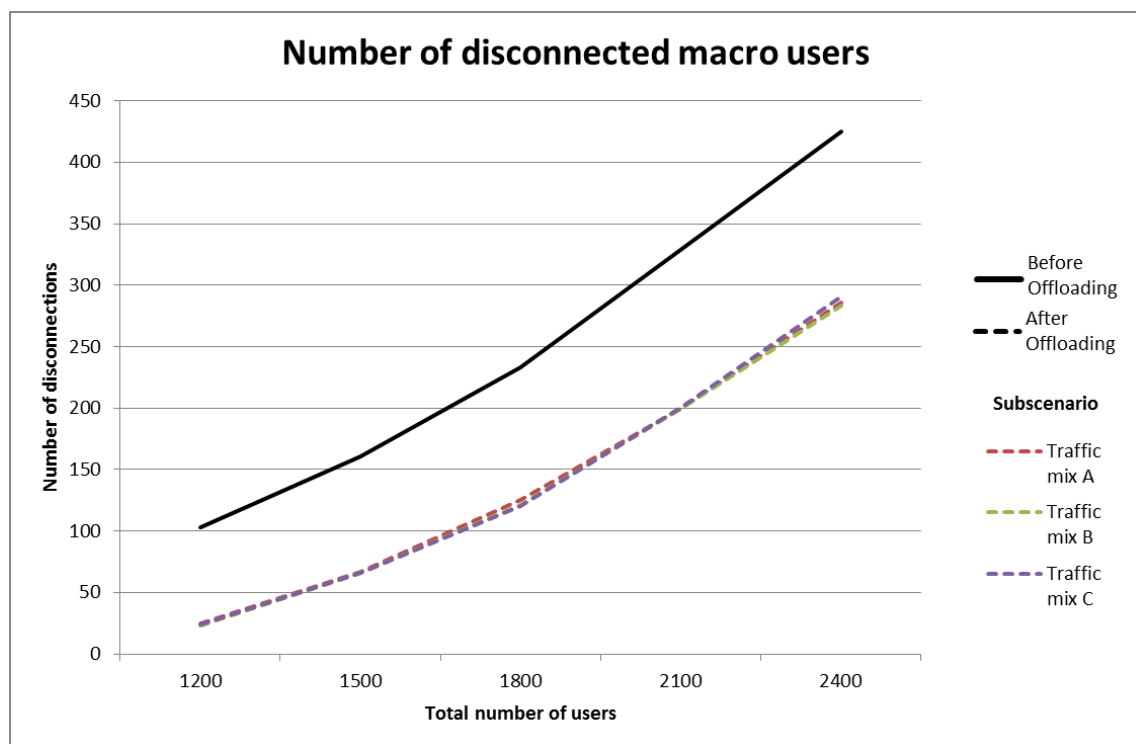


Figure 124: Number of disconnected macro users.

The users that finally get to be part of the ON are also quite similar in the three cases:

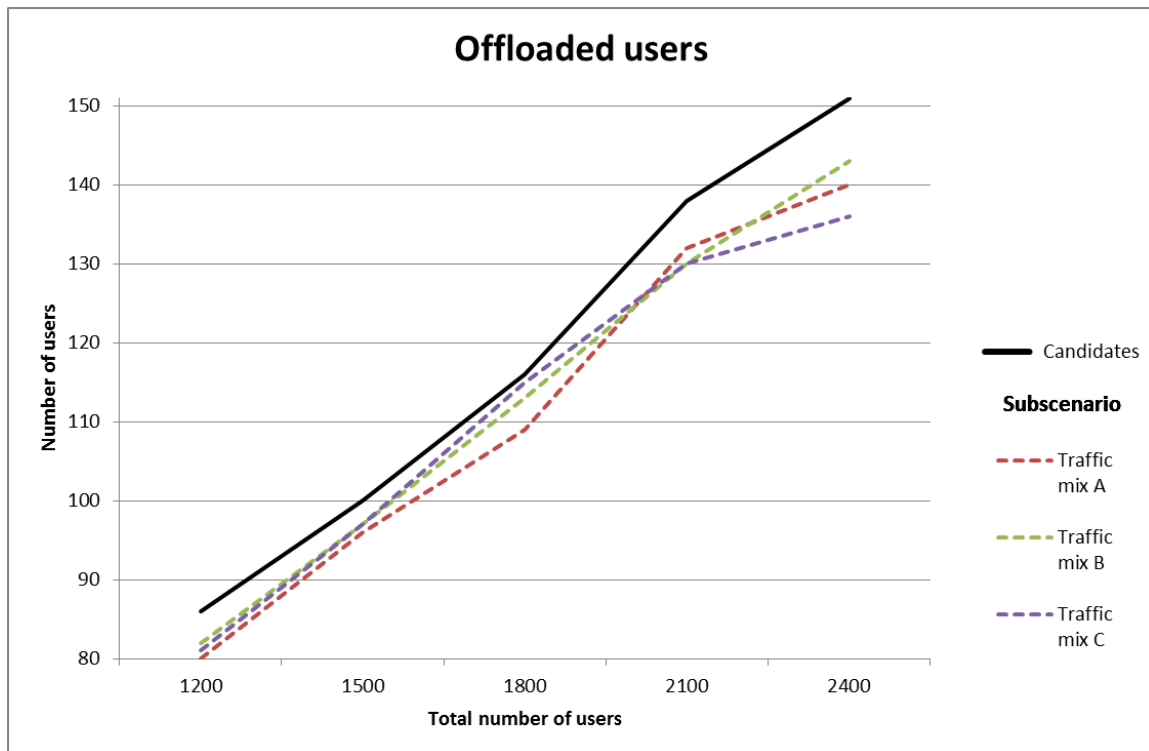


Figure 125: Number of users transferred to the ON.

Regarding the final DL throughput, the effect is similar, as it is unaffected by the traffic mix. For users that remaining in macro, there is an increase on the per user throughput over 20%, while for those that end in the femto layer, the decrease is higher than 30%.

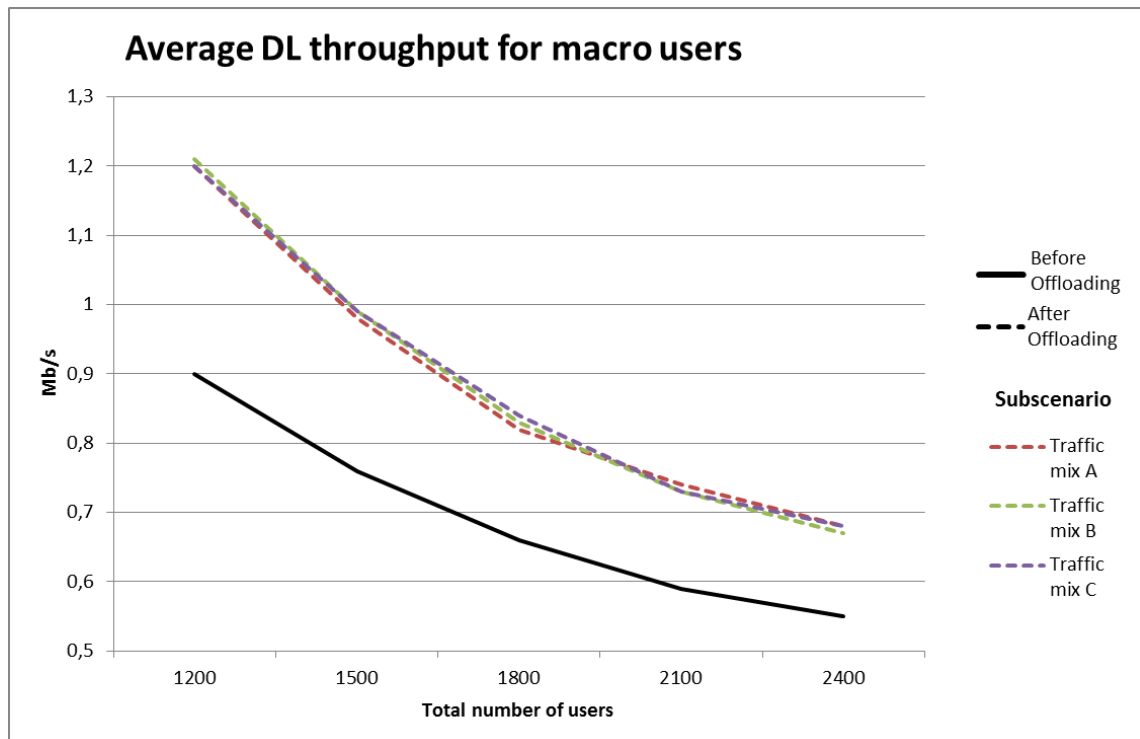


Figure 126: Average DL throughput in the macro layer.

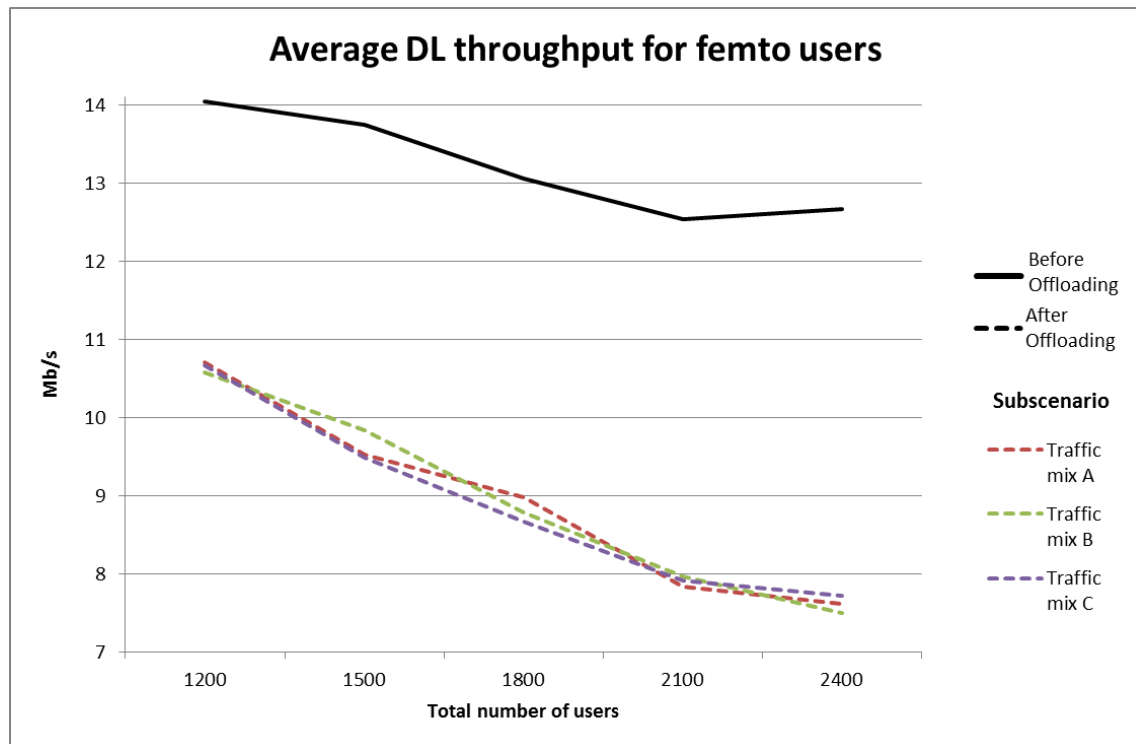


Figure 127: Average DL throughput in the femto layer.

The conclusion is that the increase in the total traffic handled by the nodes in the scenario is also independent from the target and well over 20% in all cases:

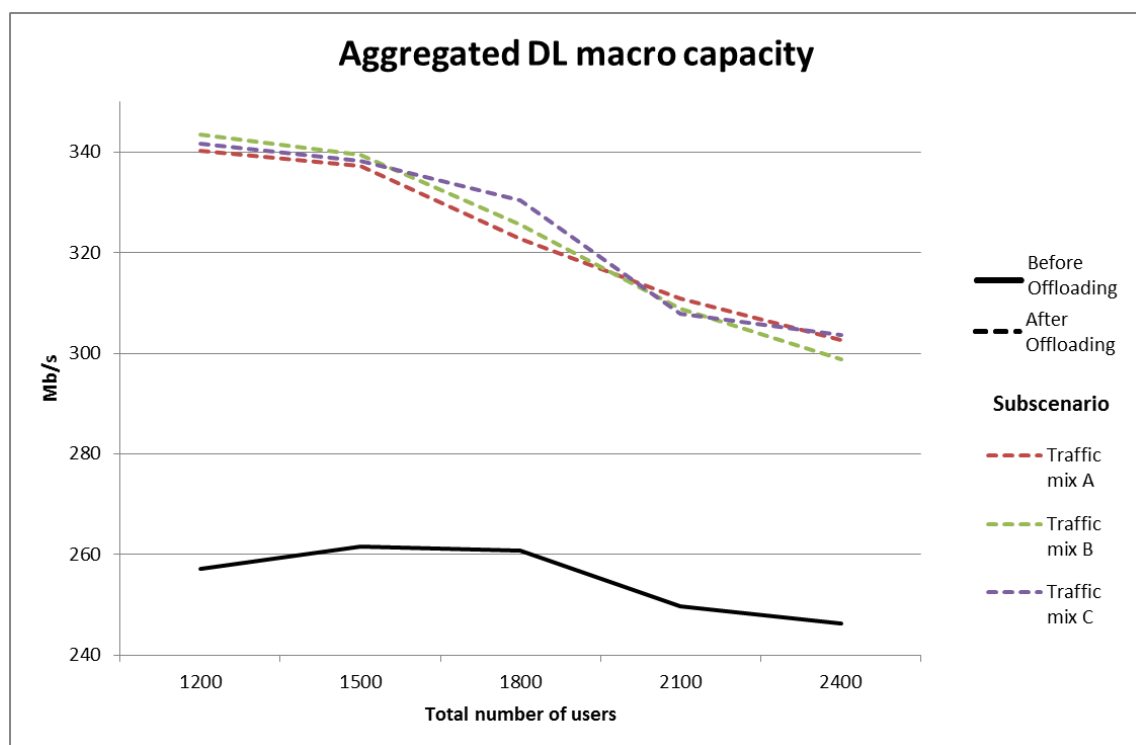


Figure 128: Total DL macro traffic.

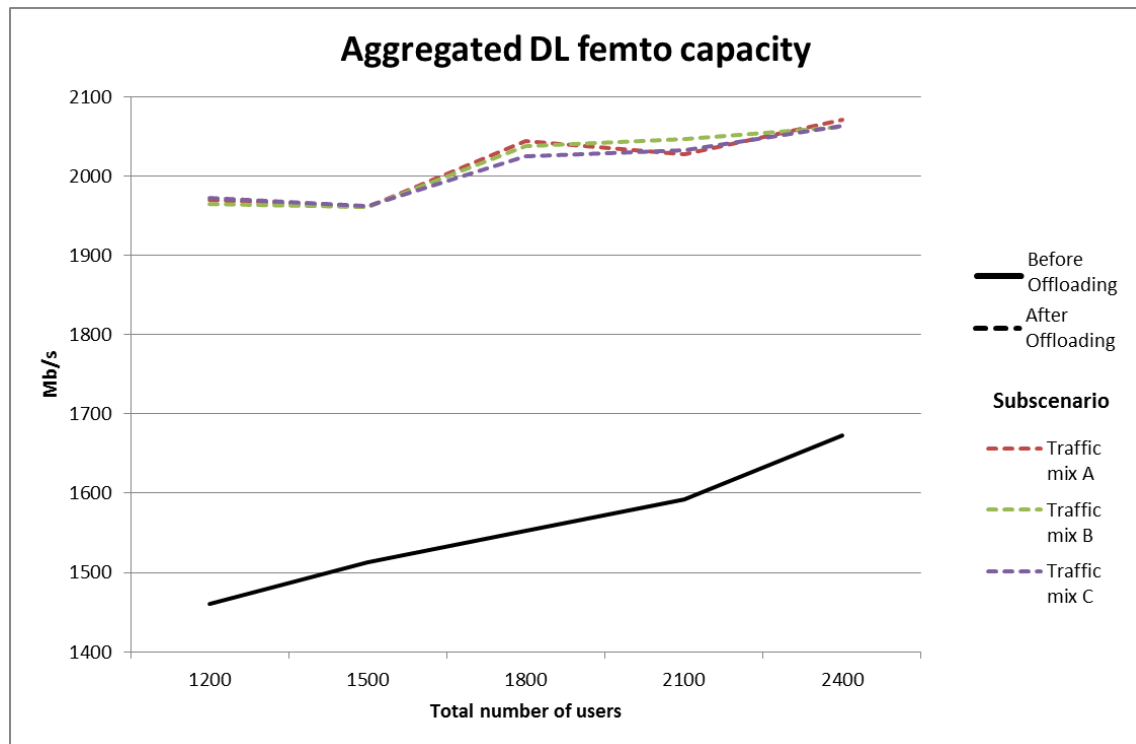


Figure 129: Total DL femto traffic.

Finally, focusing on how UEs achieve their throughput targets, there are clear differences for each traffic mix. For macro users, the increase in the achievement rate is quite noticeable, being higher in traffic mix A and lower in mix C.

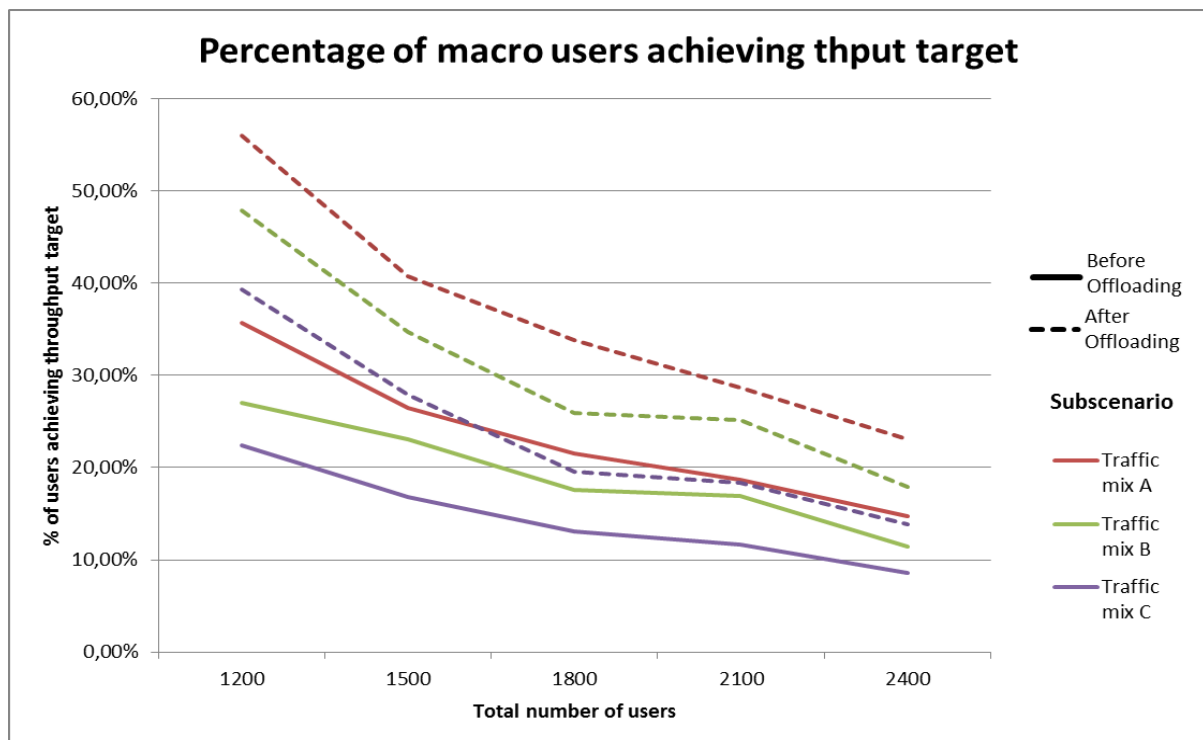


Figure 130: Achievement rate for macro UEs.

For all traffic mixes, the highest achievement rate is obtained by those users with lower targets, and is negligible for the 3 Mb/s target:

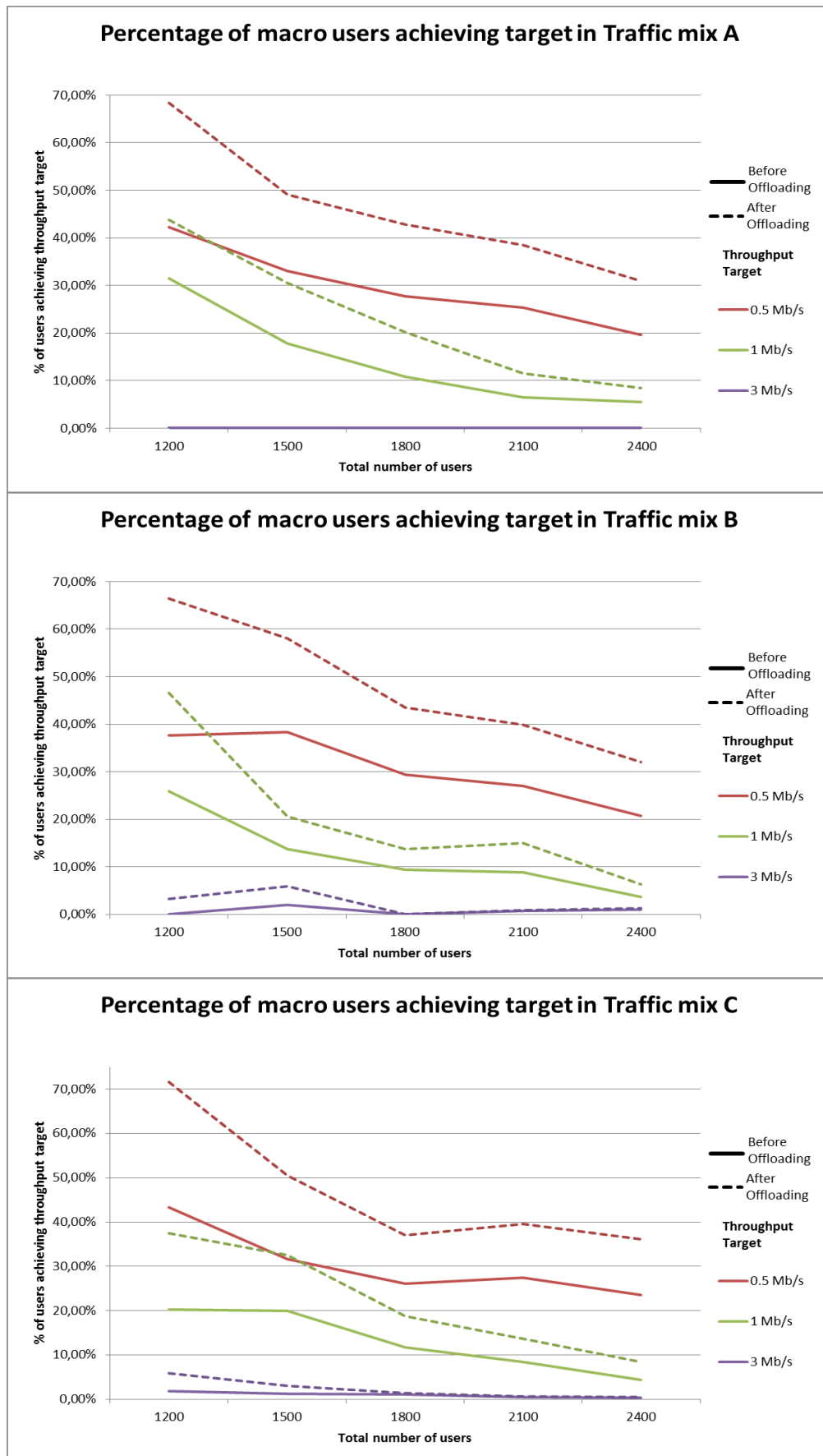


Figure 131: Achievement rate for macro UEs in traffic mixes A, B and C.

For femto users, the achievement rate is very high (it is even before the ON is created) and it increases a bit in all mixes after the offloading. The increase remains almost stable with the growth of users in the scenario.

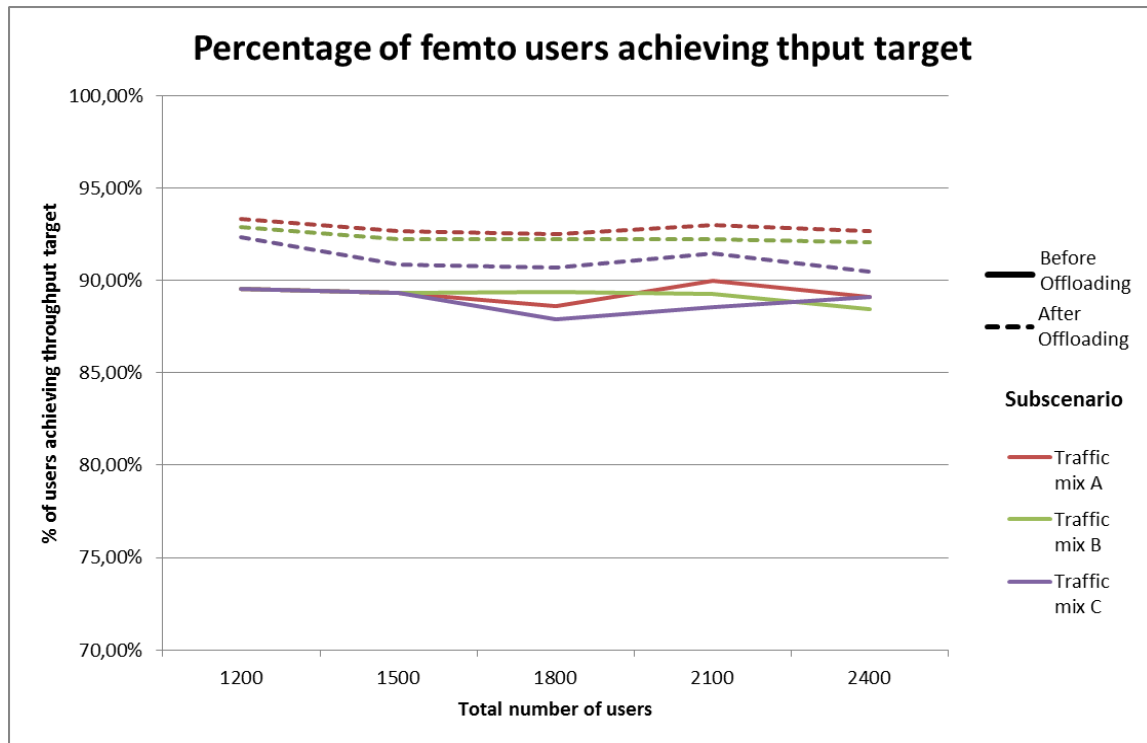
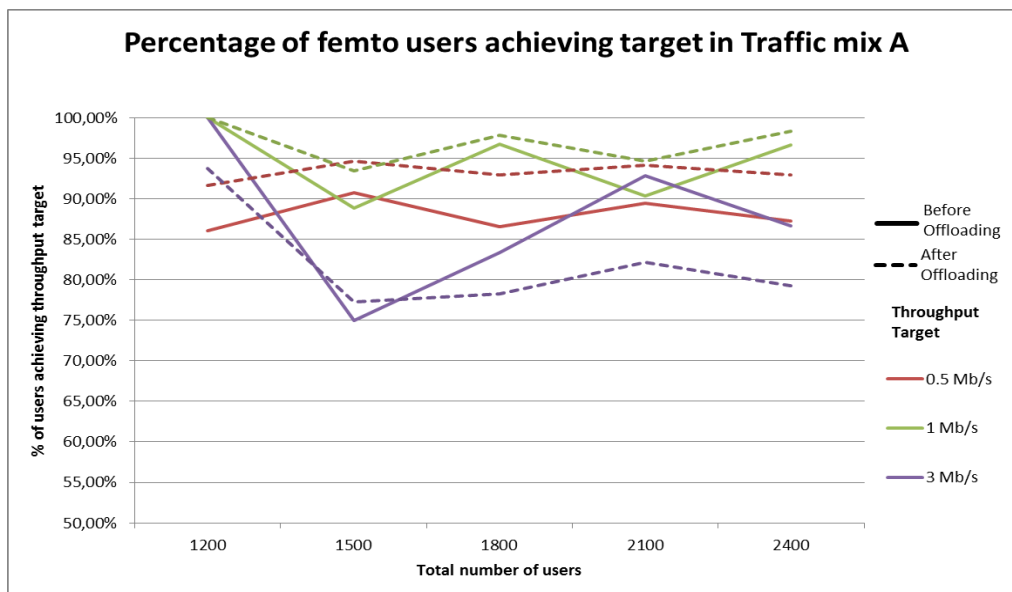


Figure 132: Achievement rate for femto UEs.

The behaviour considering the throughput targets is a bit chaotic. In general, the achievement rate tends to grow after the ON is created for the low and medium targets, but there is a slight decrease in the highest target case:



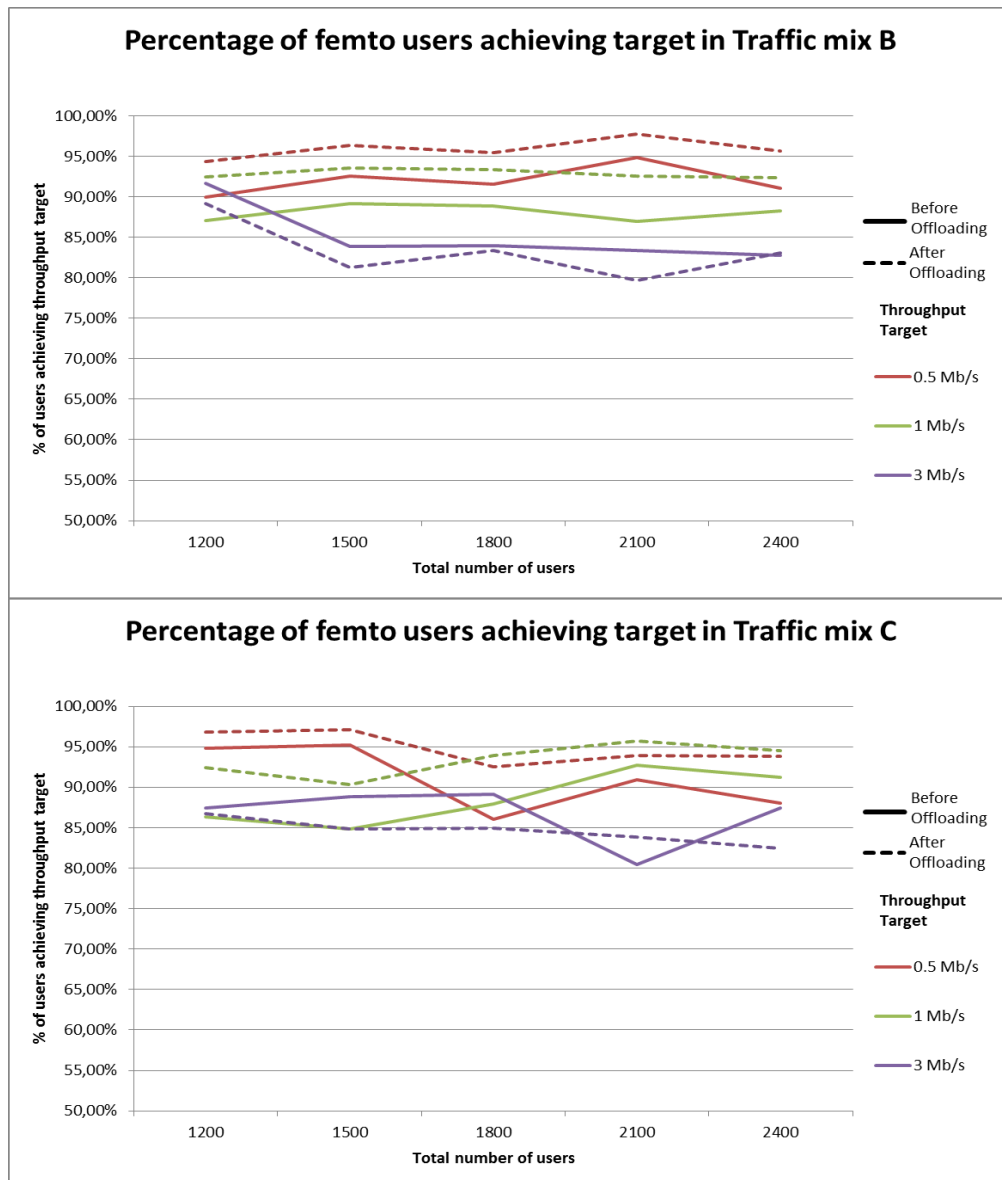


Figure 133: Achievement rate for femto UEs in traffic mixes A, B and C.

The achievement rate for those macro users that become part of the ON, is very high for all mixes:

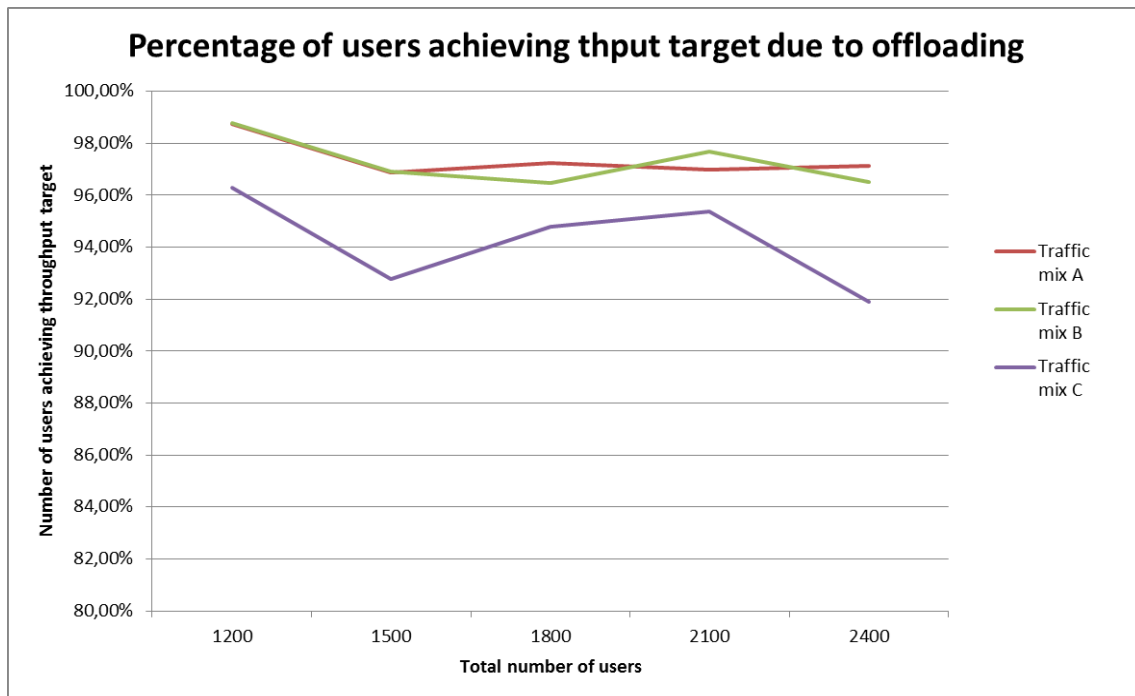
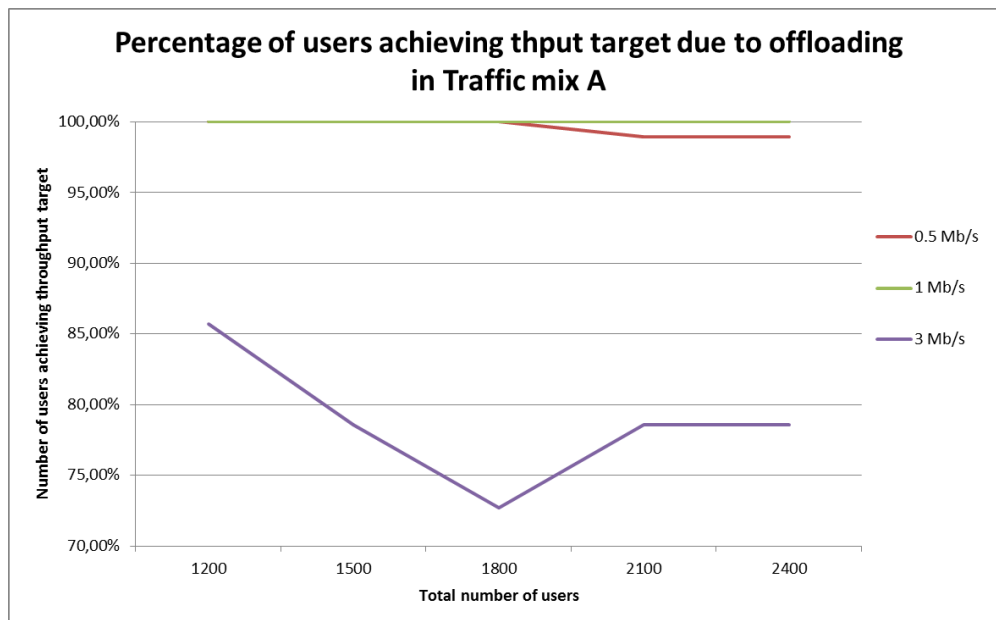


Figure 134: Achievement rate for UEs in the ON.

Focusing on the throughput targets, all users in the ON achieve their targets for the 0.5 Mb/s and 1 Mb/s cases, and most of them for the 3 Mb/s target case, almost independently of the load level:



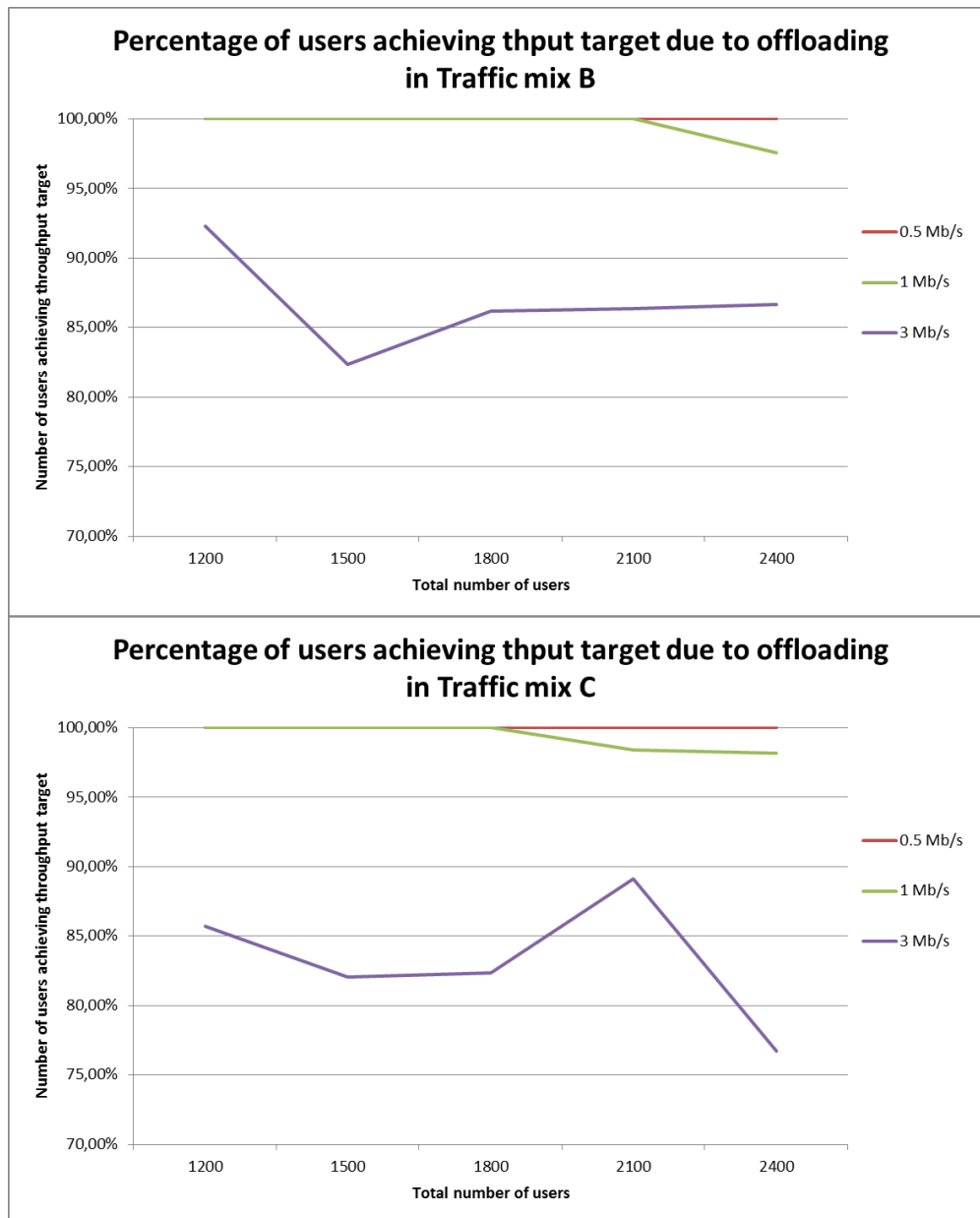


Figure 135: Achievement rate for UEs in the ON for traffic mixes A, B and C.

Conclusions for Scenario3

The performance of the Objective Function in a QoS-oriented scenario is as good as in coverage and capacity ones, leading to noticeable enhancements for users inside and outside the ON.

Macro users are mainly only able to achieve low and medium throughput targets, and the achievement rate shows an increase after ON users leave the macro layer. For high data rates, the achievement is very low and the enhancement is even lower.

For the femto users, the situation is reversed: most of the initial users are able to reach demanded bitrates, and when the new femto users join the ON, the achievement rate increases a bit, except for the highest targets, that experience a substantial reduction.

The users that originally belonged to the macro layer and then joined the ON obtain a very high achievement rate after the offloading, a bit lower in the high target case.

The congestion situation of the involved macros greatly affects the performance and the achievement rate of the users. A high congestion may decrease the enhancements achieved by the ON, especially for bitrate-hungry applications.

In summary, the ON offloading strategy leads to a higher user satisfaction, as the number of UEs achieving their demanded QoS levels increases.

3. Functional Components of the comprehensive OneFIT Solution

3.1 OneFIT solution for Spectrum selection

3.1.1 Description

The general cognitive management solution for spectrum selection considered in OneFIT is shown in Figure 136. It is based on the interaction between a decision making entity and a knowledge management functional block. Both elements are residing in the infrastructure of the operator that is governing the ONs.

The knowledge management functional block includes on the one hand the current knowledge on spectrum use indicating the status (e.g. idle/busy) of the available spectrum portions as well as different features of each portion (e.g. measured noise and interference, etc.). This information can be processed and stored in a database in the form of different statistics reflecting the experience of past situations. Such database can be used by learning methods to support the decision making processes. Moreover, mechanisms to cope with non-stationary environments are needed, so that it can be detected if some relevant changes in the environment have occurred that require to regenerate totally or partially the statistics in the knowledge database of historical information. This is the case of the so-called Reliability Tester (RT) that has been presented in section 2.4.2.1.

The decision making entity contains two main elements. The first one is the spectrum selection, that decides which spectrum portion is to be assigned to each communication link in the network. The spectrum selection strategy is based on an objective function such as the fitness factor and it can include spectrum aggregation capabilities as well as it can integrate the RAT selection together with the spectrum selection. In turn, the decision on the method to obtain knowledge on spectrum use will select the most adequate strategy and configuration for acquiring knowledge about the status of the different spectrum portions (e.g. sensing method, control channel, etc.).

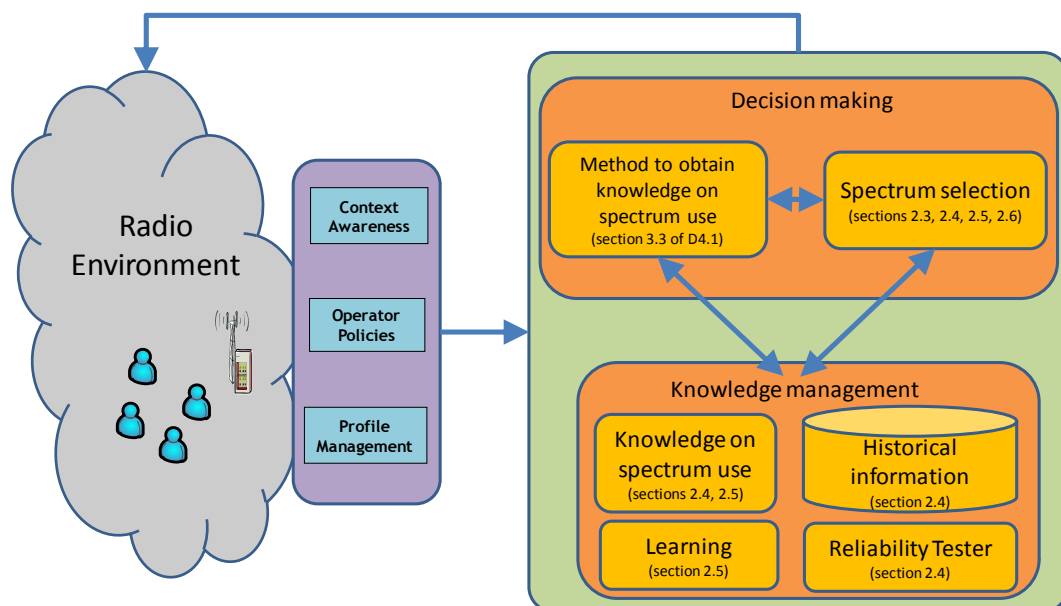


Figure 136: General framework for spectrum selection

Both decision making and knowledge management blocks use the information captured from the radio environment where the network operates. This information is categorized in terms of context awareness, operator policies and user/application profiles, as detailed in the following:

- **Policy related information:** This information contains knowledge about frequency bands that are permitted to be used for ON purposes, transmission power constraints in each band, and allowed bandwidths. Policies may indicate also the method to obtain knowledge on spectrum use in specific bands and, in case of sensing, they can define different sensing parameters such as probability of detection, sensing threshold, and minimum time required for spectrum sensing. Finally, this type of information also includes the preferences with respect to the use of one or another spectrum band.
- **Profile related parameters:** Each mobile device involved with ON creation needs to exchange information about its own parameters and capabilities. This includes device capabilities such as spectrum aggregation capability, spectrum sensing capabilities, and network interface capabilities (e.g. supported bit rates and bandwidths of each network interface). This category also includes information about the application requirements such as minimum bit rate, latency, application duration, etc. used to guarantee Quality of Service (QoS) for different applications. Finally, other aspects such as terminal location or terminal speed are also considered.
- **Context awareness information:** This refers to information about how the spectrum is used in the different bands, including spectrum occupancy for the specific time/place where the ON operates. The available spectrum is organized in pools characterized by their central frequencies and bandwidths. This category of inputs also includes the measurements to monitor the degree of QoS of the applications in the ON, in terms of the actual bit rate that is achieved by a given link in the assigned pool.

3.1.2 Evaluation

3.1.2.1 Benefits of using utility functions

In section 4.2.2 of deliverable D4.2 [7] the benefits of using utility functions in a scenario under dynamic variations in the interference of the different spectrum pools was analysed. There, the effect of two utility functions in terms of the fittingness factor definition was studied and compared against a random selection scheme. It was obtained that both functions achieved a better performance in terms of dissatisfaction probability than the random scheme that does not make use of utility distributions. In addition, the proper selection of the utility function can also lead to improved performance from a system perspective.

3.1.2.2 Benefits of applying knowledge management to historic information

In the following some results are presented to illustrate the importance of the Knowledge Management in the framework shown in Figure 136. This entity extracts the relevant information for the decision making based on the historical information regarding previous use of the different spectrum pools.

The considered scenario is an indoor DH (Digital Home) environment consisting of a single floor of dimensions 16.8m x 30.4m organized in six different rooms where $L=2$ radio links need to be established between DH nodes.

The possible candidate spectrum pieces for establishing the ON links are the following $P=3$ candidate spectrum pools:

- Pool #1: $BW_1=20$ MHz bandwidth in the 2.4 GHz ISM license-exempt band;
- Pool #2: $BW_2=16$ MHz bandwidth in the 600 MHz TV White Space (TVWS) band that can be operated opportunistically;
- Pool #3: $BW_3=20$ MHz bandwidth in a 2.6 GHz licensed band of the Mobile Network Operator (MNO) serving as the DH management service provider.

Radio propagation losses experienced at DH receivers are modelled using the COST 231 model given by:

$$L(dB) = L_o + 20\log f(MHz) + 10\alpha\log d(m) + N_w L_w \quad (52)$$

where $L_o = -27.55dB$, α is the propagation coefficient with the distance $d(m)$, N_w is the number of traversed walls between transmitter and receiver and L_w is the attenuation of one wall dependant on the material and width. Based on a measurement campaign, the estimated values of the propagation model in the considered indoor scenario resulted to be $\alpha=2.6$ and $L_w=5.1dB$. An additional shadowing loss with standard deviation 3 dB and decorrelation distance 1 m has been also considered to model the effect of objects in the building.

The transmitted power is assumed to be 20 dBm in all three pools.

Just as an example, Figure 137 presents the achievable bit rates (based on Shannon's capacity) in the different positions of the considered environment for the three considered pools, and for a transmitter located in the position shown in Figure 137(a). Figure 137(b) shows the effect of an interference source of 20 dBm arisen from a neighboring building, located 13 m on the right side from the reference floor.

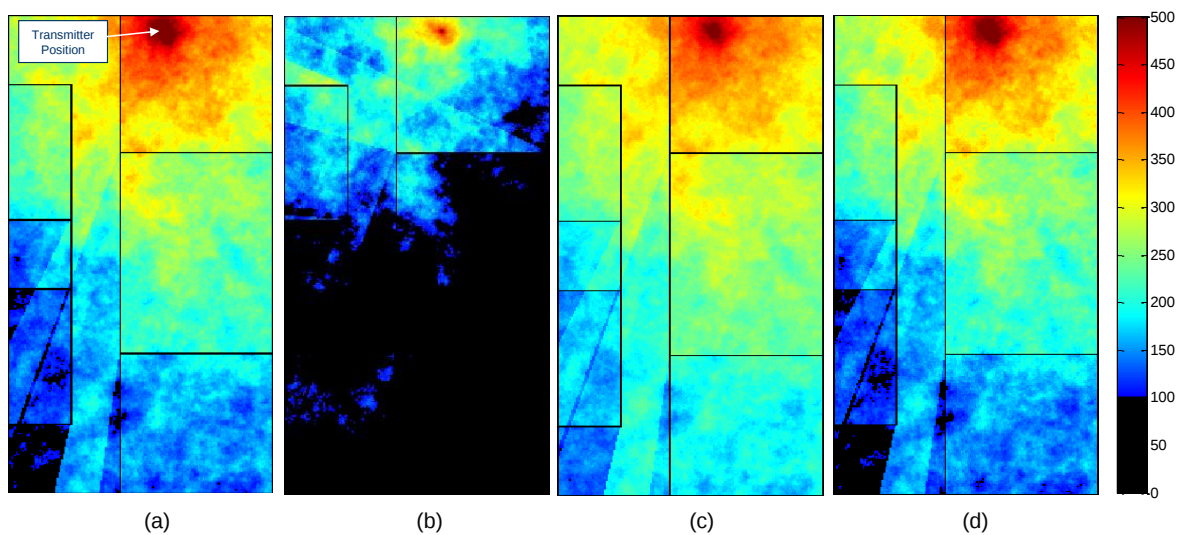


Figure 137: Achievable bit rates (Mbit/s) with the different configurations (a) Pool #1, (b) Pool #1 in the presence of an external interference, (c) Pool #2, (d) Pool #3

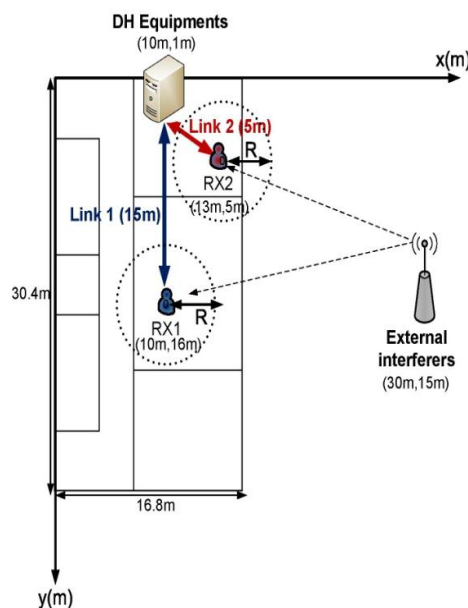


Figure 138: Location of the receivers for the two considered links

In Figure 138 the specific locations of the receivers for the two links considered in this study are plot. In order to assess the sensitivity of the proposed framework to the accuracy of generated statistics, the KD data is filled out assuming the fixed positions of the DH nodes while actual link sessions are uniformly placed in a radius R around these positions as described in Figure 138.

Link #1 is associated to low-data-rate sessions ($R_{req,1}=20Mbps$, $T_{req,1}=2min$) while link #2 is associated to high-data-rate sessions ($R_{req,2}=200Mbps$, $T_{req,2}=20min$). The l -th link generates sessions with constant session duration $T_{req,l}$ and session inactivity periods exponentially distributed with average $1/\lambda_l$.

External interferers are located at coordinates $(30m, 15m)$ as shown in Figure 138 and, when they are active, their transmit power is $P=20dBm$. Due to the dynamics of these external interferers' activity, total noise and interference power spectral density $I(p)$ experienced in each pool $p \in \{1..P\}$ is modelled with a two-state discrete time Markov chain jumping between a state of low interference $I_0(p)$ and a state of high interference $I_1(p)$ with transition probabilities $P_{01}(p)$ (i.e. probability of moving from state I_0 to I_1 in a simulation step of 1s) and $P_{10}(p)$ (i.e. probability of moving from state I_1 to I_0). Based on these probabilities, the average durations of the $I_0(p)$ and $I_1(p)$ states are respectively given by:

$$\overline{I_0(p)} = \frac{1}{P_{01}(p)} \quad (53)$$

$$\overline{I_1(p)} = \frac{1}{P_{10}(p)} \quad (54)$$

In our specific case, pools #1 and #2 alternate between $I_0(p)$ and $I_1(p)$ randomly with transition probabilities for pool #1 $P_{10}=55.5 \cdot 10^{-5}$ and $P_{01}=3.7 \cdot 10^{-5}$ and for pool #2 $P_{10}=9.25 \cdot 10^{-5}$ and $P_{01}=1.32 \cdot 10^{-5}$. Based on these probabilities, the average durations of the high interference state are, respectively, $\overline{I_1(1)}=30min$ and $\overline{I_1(2)}=2min$ while the average durations of the low interference state are, respectively, $\overline{I_0(1)}=7.5h$ and $\overline{I_0(2)}=2.5h$. On the contrary, pool #3, relying on the interference control mechanisms existing in the operator licensed band, is assumed to be always in state $I_0(p)$.

In order to smoothen the short-term variability of interference conditions, $R(I,p)$ values experienced within each of the $I_0(p)$ and $I_1(p)$ states are averaged. The average achievable bit-rates by one link l in the different pools are given in Table 12.

Table 12: Average achievable bit rates (Mbit/s) in the different pools

	Pool #1		Pool #2		Pool #3
	State I_0	State I_1	State I_0	State I_1	State I_0
Link #1	148.5	41.2	187.8	33.3	143.9
Link #2	243.2	119.6	256.5	96.1	229.6

As for the operator preferences, assuming that DH service provider aims at keeping the MNO band as much as possible available for other services and only use it here whenever the other pools are not able to provide the bit-rate requirements, ISM and TVWS bands are considered to have a higher preference factor for both links, that is $(\lambda_{l,1} = \lambda_{l,2} = 0.9; \lambda_{l,3} = 0.1)$, $l \in \{1, 2\}$.

To assess the influence of the different components of the proposed spectrum selection framework, the following variants will be compared:

- Spectrum Selection only (SS): This is the proposed fittingness factor-based spectrum selection without the support of the KM (i.e. using only the last measured value of the fittingness factor) and without spectrum handover support.
- Spectrum selection supported by Knowledge Manager (SS+KM): This is the proposed fittingness factor-based spectrum selection supported by the KM block but without spectrum handover support.
- Spectrum selection supported by both Knowledge Manager and spectrum mobility (SS+KM+SM): This is the proposed complete fittingness factor-based spectrum selection solution supported by both the KM block and the Spectrum Mobility (SM) algorithm that checks the convenience of executing spectrum handovers (SpHOs) either after variations in the interference of some spectrum pools or when a given link is released.
- *Rand*: This implements only the spectrum selection at ON creation and performs a random selection among available pools.

In addition to the metrics that were considered in section 3.5.1.1 of deliverable D4.2 [7], namely the dissatisfaction probability and the spectrum HO rate, some additional performance metrics have been included in this study:

- Regret level of link l in using pool p ($Regret(l,p)$): It is defined as the conditional probability of not respecting the pool priority constraints given that end-users are satisfied. Equivalently, it is the probability that there exists another pool able to provide the desired bit rate with higher preference than the allocated one, that is:

$$Regret(l, p) = Prob \left[\exists p' \in Av_Pools, (R(l, p') - R_{req,l}) \times (R(l, p) - R_{req,l}) > 0; \lambda_{l,p'} > \lambda_{l,p} \right] \quad (55)$$

- The fraction of time in which pool p is used by link l sessions ($Usage(l,p)$)
- An efficiency measure of link l operation: It evaluates, for each link l , the likelihood that the end-user is satisfied while respecting the preference constraints in using the different pools. It is then defined as:

$$Eff(l) = \sum_{p \in \{1..P\}} Usage(l, p) \times (1 - Dissf(l, p)) \times (1 - Regret(l, p)) \quad (56)$$

where $Dissf(l,p)$ is the dissatisfaction probability of link l when using the pool p .

- The global efficiency defined as a weighted sum of $Eff(l)$ where weights are the fractions of the different link traffic loads with respect to the total traffic load:

$$Global_Eff = \frac{\sum_{l \in \{1..L\}} \lambda_l \times T_{req,l} \times Eff(l)}{\sum_{l \in \{1..L\}} \lambda_l \times T_{req,l}} \quad (57)$$

Figure 139 plots the global efficiency as a function of the total offered traffic load $\lambda_1 T_{req,1} R_{req,1} + \lambda_2 T_{req,2} R_{req,2}$. Note that this metric, as defined in (56) and (57), integrates the effect of the dissatisfaction probability, the regret and the usage metrics. Focusing in this sub-section on the comparison between Rand, SS and SS+KM, the results show that the introduction of KM leads to a very important increase of the global efficiency level (see Figure 139) with respect to both the random and the SS approaches. The reason is that, whenever interference increases in pools #1 and #2 (i.e. they move to state I_1), the corresponding measured value of $F_{l,p}$ will be LOW. As a result, strategy SS that just keeps this last measured value of $F_{l,p}$ will decide in the future to allocate only pool #3. Then, the network will not be able to realize the situation when pools #1 and #2 move again to the low-interference state I_0 and become adequate for the link #2. This forces SS to unnecessarily assign pool #3 with the corresponding increase in the regret level. On the contrary, the use of the KM component considers the temporal properties of the $F_{l,p}$ statistics to disregard the last measured value and use an estimated value instead whenever a certain amount of time has passed since this last measure was taken. Correspondingly, sometime after the interference increase, the network will

allocate again pools #1 and #2 to link #2, thus being able to identify if they have re-entered in the low-interference state.

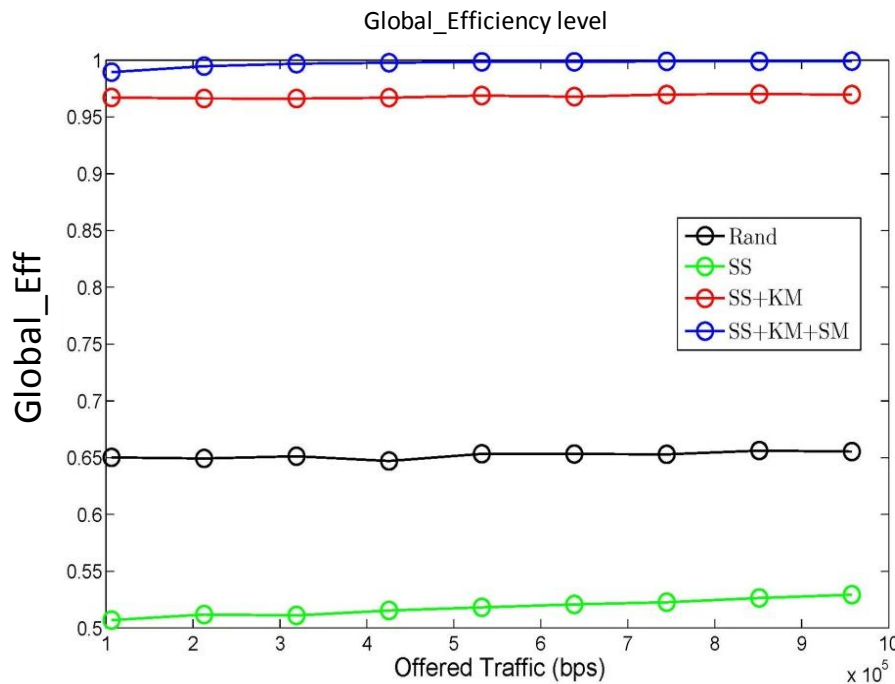


Figure 139: Global efficiency level for the considered approaches

3.1.2.3 Benefits of including operator preferences in the spectrum selection

In order to account for the benefits of including operator preferences in the spectrum selection process, the same scenario described in section 3.1.2.2 is considered, in which the three available pools have different preference factors. In particular, the ISM and TVWS bands are considered to have a higher preference factor for both links than the MNO band, that is $(\lambda_{i,1} = \lambda_{i,2} = 0.9; \lambda_{i,3} = 0.1), i \in \{1, 2\}$.

Figure 140 plots the corresponding regret level of link #2 in using pool #3 ($Regret(2,3)$), which is the pool that the operator prefers not to allocate, for the different strategies considered in section 3.1.2.2, namely the random selection, the use of only SS, the use of SS together with KM and the use of SS+KM+SM. Note that Pools #1 and #2 are not considered in the results since the considered preferences always make $Regret(2,1)=Regret(2,2)=0$. Results reveal that the inclusion of the operator preferences in the framework, either in SS+KM or SS+KM+SM, allow a significant reduction of the regret level than a strategy like the random case that does not take into account these references. Moreover, the improvement in regret level is increased with the progressive consideration of the knowledge management (KM) and spectrum mobility (SM), because in this case the system is able to capture better the actual interference conditions existing in the environment and thus to adapt accordingly, particularly by avoiding the allocation of the pool #3 unless it is strictly necessary.

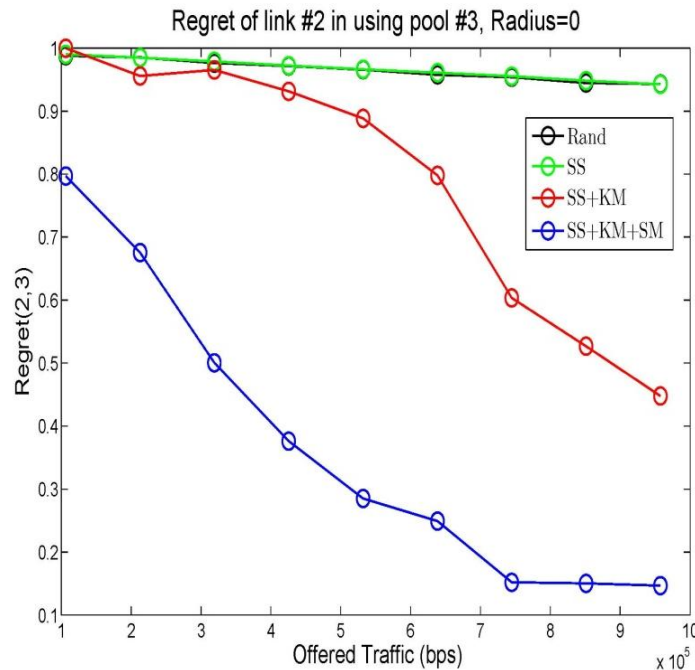


Figure 140: Regret of link #2 in using pool #3 for the considered approaches

3.1.2.4 Benefits of applying learning mechanisms

As indicated in section 2.5.4, learning based models and machine-learning techniques can offer practical solutions to challenge of spectrum opportunity identification and for configuration /optimization of spectrum selection for FCs. The proposed model enables the FCs to autonomously sense their environment and to tune their parameters accordingly, in order to operate under restrictions of avoiding/causing interference in heterogeneous networks deployments. We also studied the convergence of the played game when active node (FCs) adopted different learning strategies, to mimic the real radio environment where nodes have different capabilities and objectives. Nodes work autonomously so no cooperation or information exchange between multiple nodes is required. The total experienced interference and the reconfiguration cost has been also studied in order to provide a reasonable comparison between different learning strategies and to highlight the cost associated with the learning process. Specifically we show that the proposed strategies can identify unused resource blocks with high degree of accuracy, co/cross-layer interference can be reduced significantly, thus resulting in average cell throughput increases of 30% and satisfaction probabilities increased by 47%, in the considered scenarios.

3.1.2.5 Benefits of providing adaptability to algorithmic solutions

3.1.2.5.1 Adaptability provided by spectrum Handover

To analyse the benefits introduced by Spectrum Mobility capabilities, let consider the same conditions as in section 3.1.2.2 and let focus now on the comparison between strategy SS+KM+SM, which provides the capability of executing spectrum handovers, and strategy SS+KM. It can be observed in Figure 139 that the introduction of the SM functionality results in further improving the global efficiency level at all traffic loads with respect to SS+KM. This is due to the fact that performed reconfigurations by SM allow reconfiguring the allocated spectrum pool whenever external interferers for pools #1 or #2 show up. Similarly, the execution of SM will lead to release pool #3 whenever a better pool among pools #1 or #2 can be found, which significantly reduces the regret level in using pool #3.

To further assess the proposed strategy in terms of the cost associated to the reconfigurations performed by the SM functionality, *Figure 141* plots the average number of SpHOs per session experienced by *SS+KM+SM* for each of the two possible triggers of the spectrum mobility SM, namely the change in the $F_{l,p}$ value of the currently assigned pool or the release of another link. Results indicate that the total amount of SpHO signaling per session is below 0.2 for all considered traffic loads. The analysis of the fractions associated to each trigger reveals that most of SpHOs are triggered by changes in $F_{l,p}$ values. This is because, on the one hand, in the considered DH environment, there is at most one active session for each link which makes not likely to experience a SpHO due to a link release due to the low traffic load. On the other hand, the average durations of the high-interference states ($\overline{I_1(1)} = 30\text{min}$ and $\overline{I_1(2)} = 2\text{min}$) are comparable to the link session durations ($T_{req,1} = 2\text{min}$ and $T_{req,2} = 20\text{min}$), which makes likely to experience a change in $F_{l,p}$ values during link session.

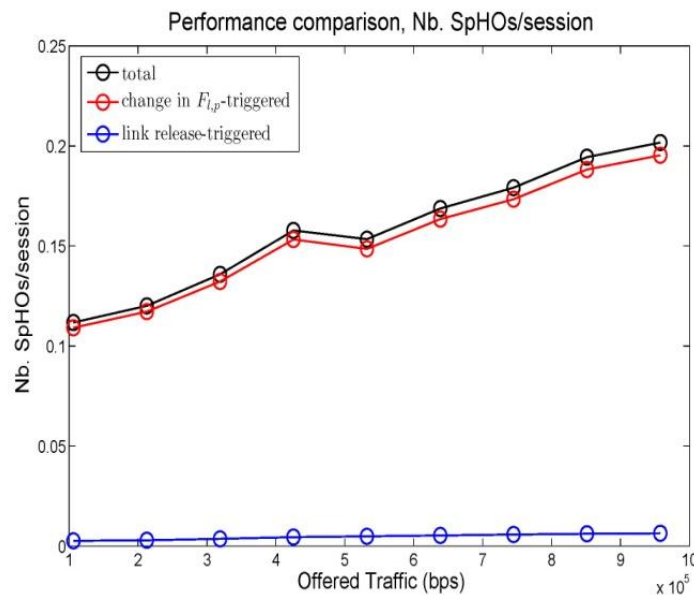


Figure 141: Spectrum HO rate

To further assess the robustness of the proposed strategy, let consider now that there is some uncertainty in the position of the receivers with respect to the location where the database statistics were obtained. This is modelled assuming that the actual location of the receivers is in a certain radius R around the position where the statistics were obtained. Focusing on the *SS+KM+SM* approach, *Figure 142(a)* plots the satisfaction probability when different values of R are considered. *Figure 142 (b)* and *Figure 142 (c)* plot the corresponding regret level in using pool #3 and the SpHO rate per session, respectively. Results show that as R increases, the global satisfaction level decreases a bit while still keeping a sufficiently satisfactory level in the order of 98% for $R=8\text{m}$ (see *Figure 142 (a)*). The regret level in using pool #3 is degraded when R increases as seen in *Figure 142 (b)*. This is because the estimation of $F_{l,p}$ values assuming the initial fixed position is likely to be invalid as DH nodes are placed further from the initial positions. This results in some erroneous spectrum selection decisions which particularly tend to unnecessary use pool #3 thus degrading the regret level. The trigger of the SM functionality after measurements of $R(l,p)$ are taken in the new positions and valid $F_{l,p}$ values become available will result in frequent SpHOs events to the erroneously selected pools. This turns into the corresponding increase in the spectrum HO rate that can be seen in *Figure 142 (c)*.

Consequently, it can be observed that a proper SpHO algorithm is useful to overcome the inaccuracies in the database information associated to the uncertainty in the position of the receivers which can lead to some erroneous decisions during initial spectrum selection. This capability can be used to cope to some extent with small variations in the environment and keep

adequate levels of satisfaction probability. For more significant variations associated to non-stationary environments, additional techniques such as the reliability tester (RT) described in section 2.4.2.1 will be required.

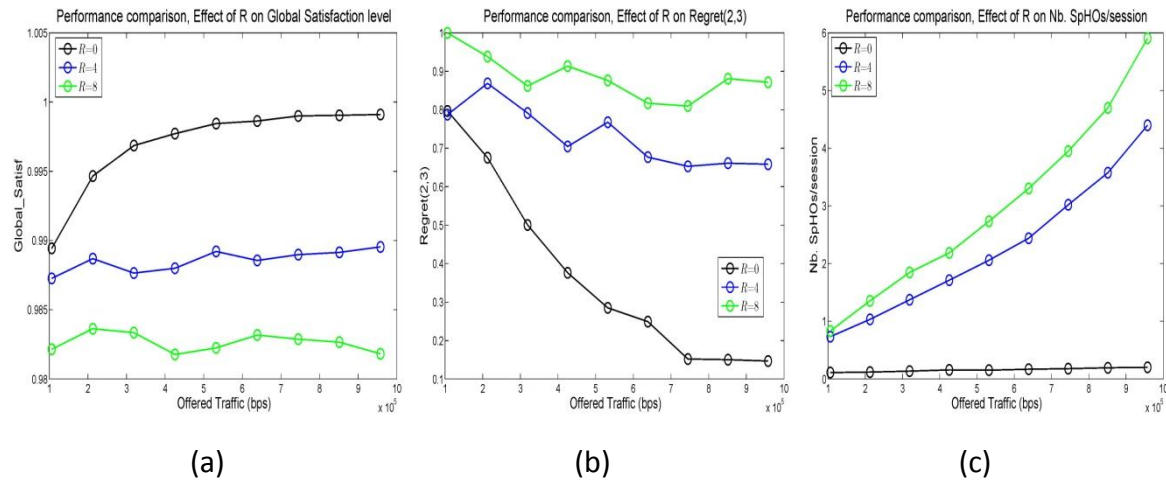


Figure 142: Sensitivity to R in terms of (a) Global satisfaction level, (b) Regret of link #2 in using pool #3, (c) Spectrum HO rate per session

3.1.2.5.2 Adaptability provided by utility-based spectrum aggregation

The proposed utility-based spectrum aggregation and allocation algorithm considered the complexity of spectrum aggregation as well as the channel switching and the total achievable throughput. These three objectives are integrated into a weighted sum utility function and the learning mechanism is utilized to decide the optimal weights setting which can depend on the (possibly periodic) variations in metric of interest. In the scenario that the system has pre-defined thresholds for each performance metric based on system level Key Performance Metrics (KPIs), the spectrum aggregation and allocation algorithm aims to maintain the performance of each objective remain close as possible to the pre-defined thresholds. When system detects the degradations in any of the performance metrics, the Q-learning process is triggered to find the optimal weights setting depending on the situation. The reason of degradation of the performance metrics could be the radio environmental changes such as the primary users' activity pattern's change as well as changes in the pre-defined thresholds. The weights setting obtained as the result of the learning are applied to the utility function of the spectrum aggregation and allocation algorithm. The proposed spectrum aggregation and allocation algorithm allows for the automated (versus manual setting) management of complex interactions and trade-offs between different metrics while the weight setting is adaptable and is modifiable depending on the application/environmental conditions encountered.

3.1.2.5.3 Adaptability provided by reciprocity based reinforcement learning

With the best response channel selection strategies, the SUs will consider whether their beliefs have any negative effects. Our belief model suggests that error exists in the belief. For such reason, we will assume that the dynamics of the cognitive radio networks will appear reasonably consistent to the SUs if the values of beliefs stabilize as the time passes. We consider a relatively simple scenario

where there are two SUs and two channels with idle probabilities 0.6 and 0.8. The belief factors $\delta_{i,n}$

are uniformly distributed between 2 and 7. The initial channel selection strategies $(p_1^1(1), p_2^1(1))$ are set to be (0.5, 0.5) and (0.7, 0.8). It can be shown from Figure 143 that the trajectory converges to the same optimal channel selection strategies.

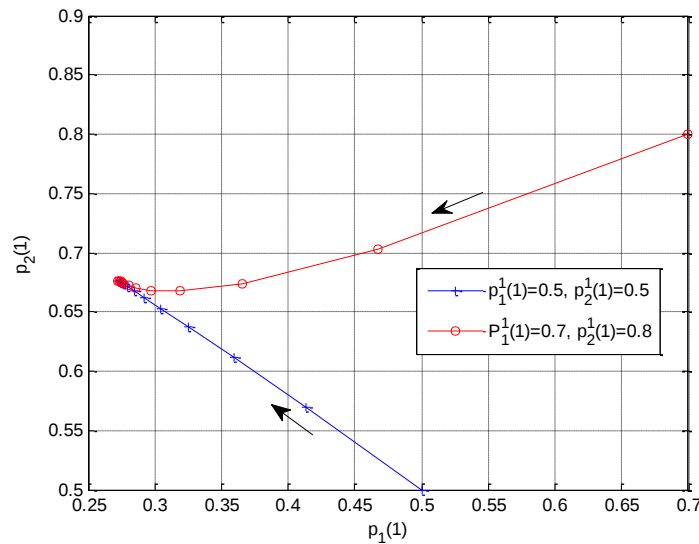


Figure 143: Strategy dynamics with different initial value of $p_1(1)$ and $p_2(1)$

3.1.2.6 Benefits of considering a reliability tester for non-stationary environments

In the following the robustness of the spectrum selection framework in front of non-stationary environments is analysed. Such robustness can be achieved thanks to the capability to detect that relevant changes have occur in the environment so that the statistics that are stored in the Knowledge Database need to be regenerated. In the proposed framework this detection is performed by the reliability tester specified in section 2.4.2.1.

The evaluation is performed in a variant of the Digital Home environment scenario described in section 3.1.2.2. It is depicted in Figure 144. It is assumed that the link #1 receiver is situated in a fixed position while the link #2 receiver may jump between the positions #0 and #1 indicated in the Figure 144.

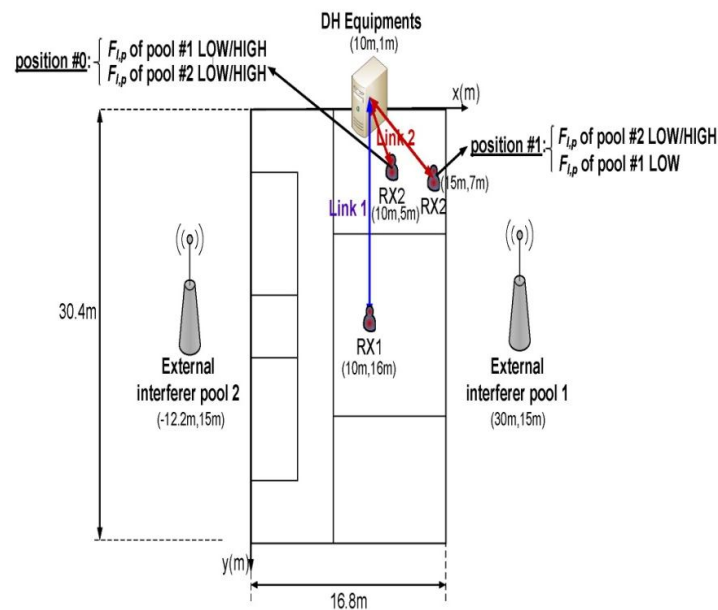


Figure 144: Considered DH environment for evaluating capability to adapt to non-stationary environments

The external interferers of spectrum pools #1 and #2 are located in two different positions as shown in Figure 144. Their transmit power is 20 dBm during the high interference state and 10 dBm during the low interference state. The transition probabilities for pool #1 are $P_{10}=55.5 \cdot 10^{-5}$ and $P_{01}=3.7 \cdot 10^{-5}$, while for pool #2 they are $P_{10}=833.33 \cdot 10^{-5}$ and $P_{01}=55.5 \cdot 10^{-5}$. Based on these probabilities, the average durations of the low-interference state are, respectively, $\overline{I_0(1)}=7.5h$ and $\overline{I_0(2)}=0.5h$. In turn, the average durations of the high-interference state are, respectively, $\overline{I_1(1)}=30min$ and $\overline{I_1(2)}=2min$. Pool #3, relying on the interference control mechanisms existing in the operator licensed band, is assumed to be free of interference all the time. Based on the interference model, the average bit rates achievable in each spectrum pool are indicated in Table 13.

The rest of parameters are the same as indicated in section 2.4.2.1.

Table 13: Average achievable bit rates (Mbit/s) in the different pools

	Pool #1		Pool #2		Pool #3
	State I_0	State I_1	State I_0	State I_1	State I_0
Link #1	88.8	32.5	106.6	55.45	229.3
Link #2	228	161.9	244.2	191.2	361.7

The evaluation is performed using the same indicators described in section 2.4.2.1. To see the influence of the reliability tester (RT), the following two variants of the proposed framework are compared:

- *SS+SM*: This approach makes use of both the SS and SM algorithms, but does not include the RT to detect changes and regenerate the KD statistics in case.
- *SS+SM+RT*: This is the complete approach that includes SS, SM and RT functionalities.

The detection of changes made by the RT is executed through the monitoring of the following KPIs:

- The average dissatisfaction probability, $Dissf(I)$.
- The average number of SpHOs performed per session $SpHO(I)$.
- The average fraction of using pool #3, $Usage(I,3)$.

In order to assess robustness to changes in radio and interference conditions of the different spectrum pools, the following two possible changes will be considered during system operation affecting the positions of the link #2 receiver as indicated in Figure 144:

- *Change 1*: after initially generating KD statistics in position #0, the link #2 receiver is placed in position #1. Correspondingly, the receiver moves from a position where the fittingness factor value for pool #1 alternates between LOW and HIGH values to a position where the fittingness factor is always LOW regardless of interference activity.
- *Change 2*: after initially generating KD statistics in position #1, the link #2 receiver is placed in position #0. In this case, the receiver moves from a position where the fittingness factor value of pool #1 is always LOW to a position where it alternates between LOW and HIGH depending on interference activity.

Figure 145 shows the online evolution of the dissatisfaction level of link #2 ($Dissf(2)$). For a better visualization, the time evolution on the x-axis is shown in terms of the number of established link #2 sessions considered for updating $Dissf(2)$ along the simulation. *Change 1* occurs after 9750 sessions. To analyse the impact of RT both variants *SS+SM* and *SS+SM+RT* are considered. Traffic load is

$\lambda_2 T_{req,2} R_{req,2} = 0.9Er$. Notice that only link #2 sessions are considered for a better analysis of system reactivity since link #1 is all the time satisfied. It is worth pointing out that whenever a change is detected by the RT, $SS+SM+RT$ continues to use the old KD statistics until the new KD statistics become available. Figure 146 plots the corresponding number of SpHOs per session. Moreover, considering that pool 3 is the less preferred from the operator's perspective, Figure 147 plots the relative regret level in using this pool (defined as the regret weighted by the usage of the pool #3 $Usage(2,3) \times Regret(2,3)$).

The results in Figure 145 show that, after *Change 1*, the use of RT functionality ($SS+SM+RT$) results in a significant decrease in $Dissf(2)$. The reason is that, after the position change, the strategy $SS+SM$ without any support from the RT still relies on the out-of-date KD statistics previously generated in position #0, so it may decide in some cases to assign pool #1 to link #2 sessions. However, this turns out to be a wrong allocation because in the new position #1, pool #1 has always a LOW value of fittingness factor regardless of the interference conditions. Correspondingly, this degrades the dissatisfaction probability (see Figure 145). In addition, much more frequent SpHOs (Figure 165) are required to change pool #1 whenever it is assigned. On the contrary, when RT is used ($SS+SM+RT$), new KD statistics are generated in position #1 after detecting *Change 1*. As a consequence, the estimated fittingness factor value $\hat{F}_{l,p}$ of pool #1 is always set to a LOW value and pool #1 is never assigned in the future to link #2. This results in a significant gain in both the dissatisfaction level (Figure 145) and the number of performed SpHOs (Figure 165).

In Figure 147 it can be observed that the relative regret level of using pool #3 performs similarly for both strategies before and after the change. Before the change either pool #1 or #2 can be allocated because both alternate between LOW and HIGH fittingness factor values. This marginalizes the risk of unnecessarily using the less preferred pool #3, since it is not likely to simultaneously have a wrong estimation of LOW value for the $\hat{F}_{l,p}$ of both pools which would lead to an unnecessary assignment of pool #3. Then, after the change the use of pool #1 is excluded in both strategies (at first assignment for $SS+SM+RT$ or after initially assigning it and then executing a SpHO for $SS+SM$). As a consequence, whenever the estimate $\hat{F}_{l,p}$ of pool #2 is wrongly set to a LOW value, pool #3 will be unnecessarily assigned with the corresponding increase in the relative regret level that can be observed in Figure 147 after the change.

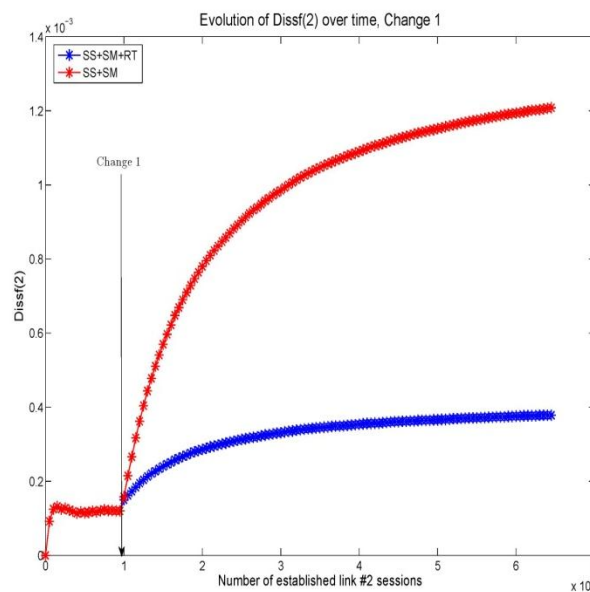


Figure 145: Evolution of dissatisfaction probability in the presence of Change 1

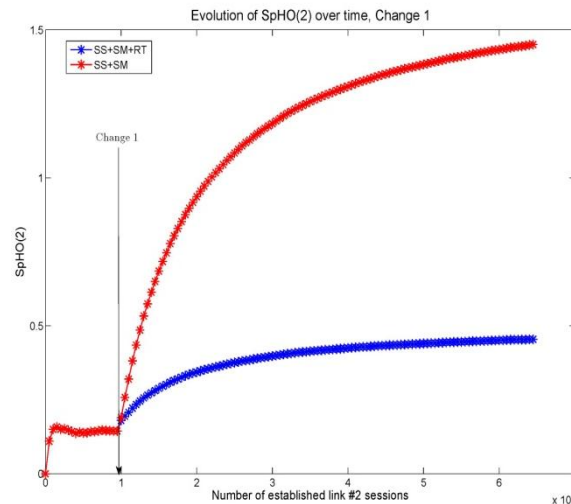


Figure 146: Evolution of the SpHO rate in the presence of Change 1

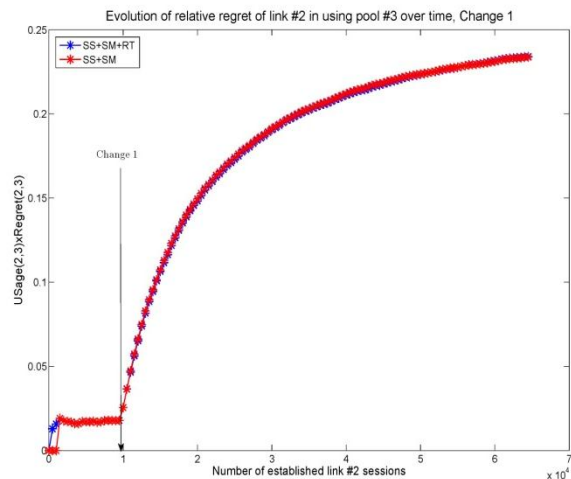


Figure 147: Evolution of the relative regret $Usage(2,3) \times Regret(2,3)$ for pool #3 in the presence of Change 1

Focusing now on *Change 2*, Figure 148 shows the online evolution of the dissatisfaction level of link #2. The position change occurs after 9750 sessions. Figure 149 plots the corresponding number of SpHOs per session. Figure 150 plots the relative regret level in using pool #3.

The results show that after *Change 2*, the use of RT functionality results in a significant reduction of both the dissatisfaction level (see Figure 148) and the number of performed SpHOs (see Figure 149). This is because pool #1, initially not used in position #1, becomes often the most suitable pool for link #2 after the position change has been detected and the new KD statistics are regenerated for the new position. On the contrary, when the RT is not supported (*SS+SM*), *Change 2* has no impact on the observed dissatisfaction level (Figure 148) and number of SpHOs (Figure 149), because link #2 continues to exclude pool #1 based on the old KD statistics. In contrast to *SS+SM+RT*, this results in penalizing the regret level of using pool #3 whenever this pool is assigned instead of pool #1 (see Figure 150).

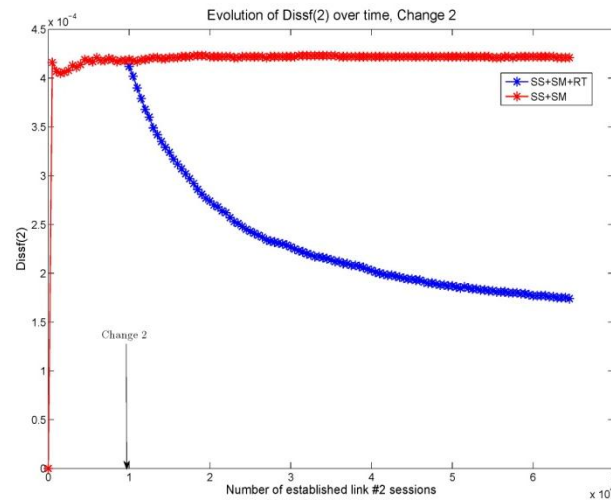


Figure 148: Evolution of dissatisfaction probability in the presence of Change 2

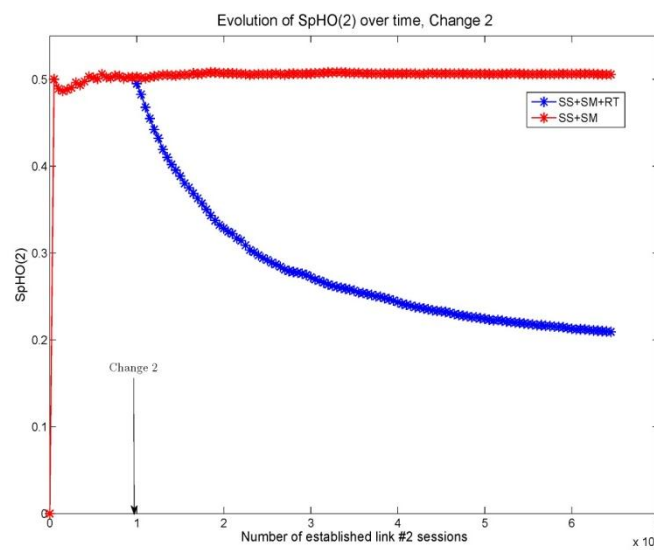


Figure 149: Evolution of the SpHO rate in the presence of Change 2

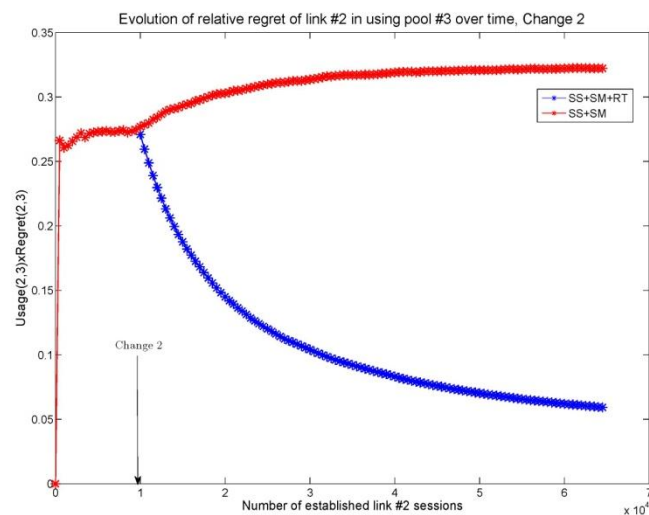


Figure 150: Evolution of the relative regret $Usage(2,3) \times Regret(2,3)$ for pool #3 in the presence of Change 2

3.1.2.7 Benefits of the joint consideration of the RAT and spectrum selection

This section is a continuation of work started in D4.2 [7] section 4.2.2.6. This section will introduce further results when channel selection has been included to algorithm introduced in D4.2 [7] section 2.3. The used simulation scenario is same as in section 2.3. In Figure 151 represent selected RAT for each traffic type during the maintenance phase. Browsing users mainly utilize 802.11a and 802.15.3c. LTE is mainly reserved for streaming user to meet their high data rate requirements. It is worth of notice that on simulations we have used LTE also on TV bands. Thus on Figure 151 LTE users from two different bands are counted together. Also for voice users LTE is selected to reduce delay which is more severe on ISM bands due to other users' appearance to same bands causing possible band switching. On the other hand 802.15.3c provides high data rates whenever its usage is sufficient and we can see that all data types utilize the band. The usage drops on 802.15.3c refer to situations where transmission range is high causing severe path loss and 60 GHz band usage is not sufficient. Since browsing users have least demanding data rate and delay requirements their usage is directed to ISM and TV bands more often to offload traffic from operator own band (IMT).

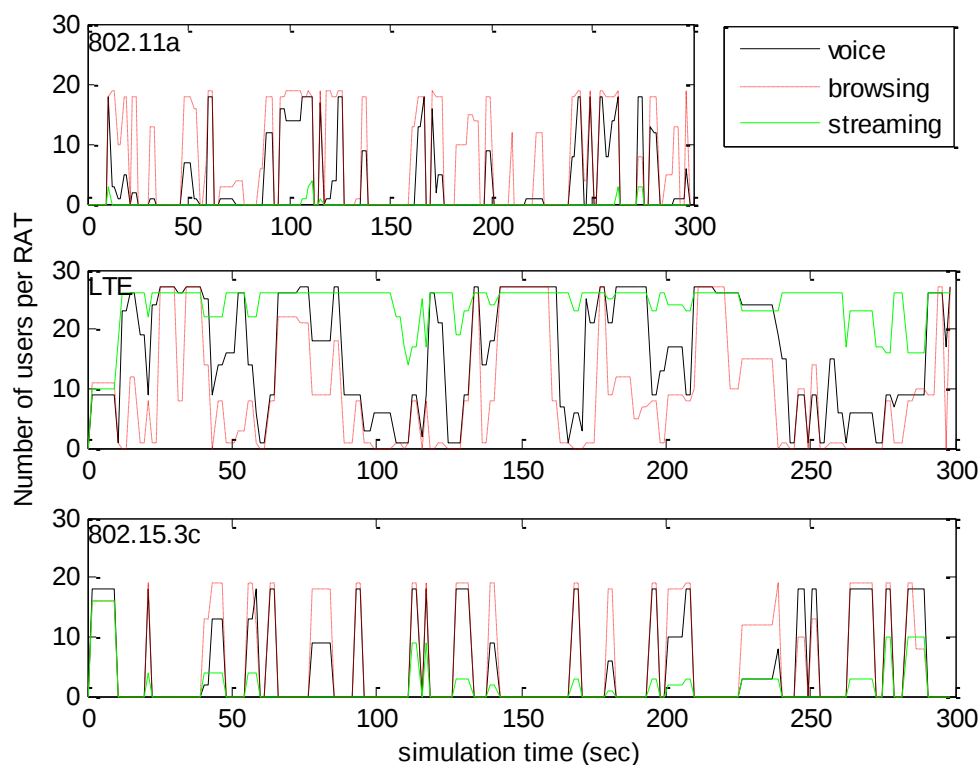


Figure 151: RAT selection for different data types.

3.1.2.8 Benefits of including reciprocity based reinforcement learning in stochastic channel selection

In this section, to show the performance gain achieved by the proposed algorithm, we consider a general case, where there are 6 channels with idle probabilities characterized by Bernoulli distributions with evenly spaced parameters ranging from 0.2 to 0.9. The belief factors $\delta_{i,n}$ are randomized in accordance with the number of SUs. As following, we numerically compare the overall network performance of the proposed algorithm in terms of accumulated utilities with two existing channel selection protocols: (1) adaptive random channel selection scheme [23]; and (2) no learning scheme, under which each SU i selects channel n with probability $p_i(n) = \frac{\theta_n}{\sum_{l \in N} \theta_l}$.

We can see from Figure 152 that the proposed algorithm and the adaptive random channel selection scheme can achieve better system performance than the no learning scheme. In addition, the curves show that the accumulated utilities of the channel selection schemes increase versus the number of SUs, but decrease when the number of SUs exceeds the number of channels. The reason is simple: when $I \leq N$, with more SUs, the utilization of the licensed channels is better exploited; however, when the $I > N$, the collisions among different SUs cannot be avoided, causing the reduction in the overall system performance.

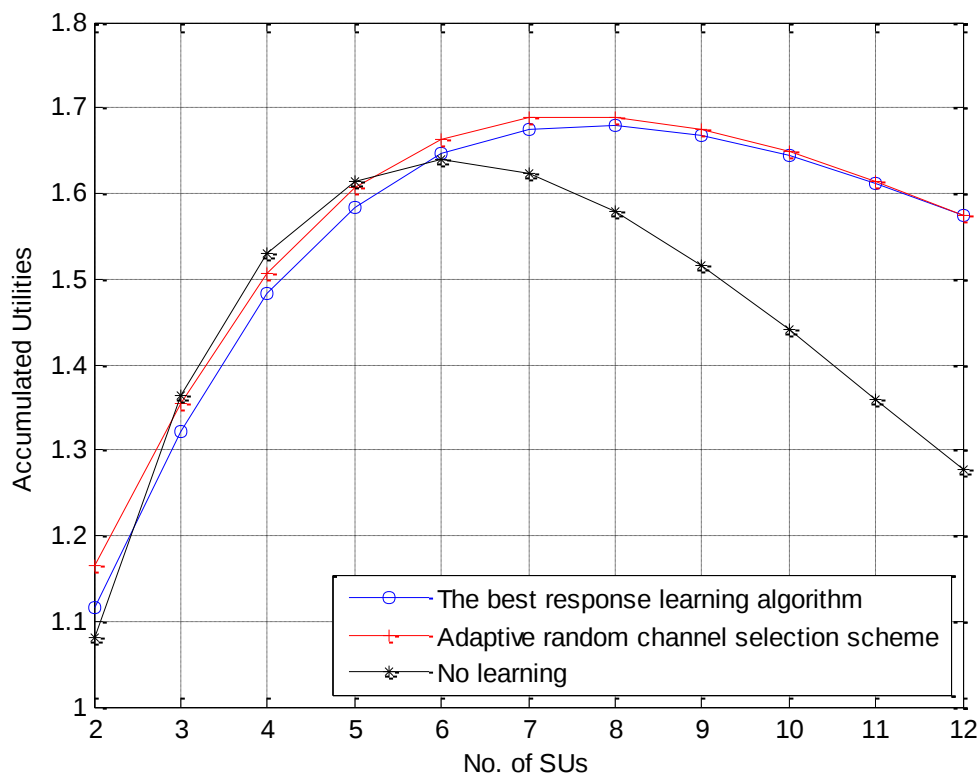


Figure 152: Comparison of the accumulated utilities corresponding to different channel selection schemes.

3.2 OneFIT solution for Selection of Nodes and Routes

3.2.1 Description

During the OneFIT project, specific algorithms have been developed in order to cover multiple facets of the selection of nodes and routes for the efficient creation of ONs. In the D4.2 document [7], there has been specific description of the general framework for the synergic operations of nodes and routes selection. Specifically, in [7] it is analysed that the identification and selection of nodes and routes are conducted both locally and in a centralized manner. Locally, decisions are made within network nodes (terminals and infrastructure nodes) with help from centralized operator's management (policy acquisition, profile assignment, broader knowledge). Also, decisions can be made in a centralized management system which monitors the whole operator's network and has more processing power and broader knowledge for pervasive decision making.

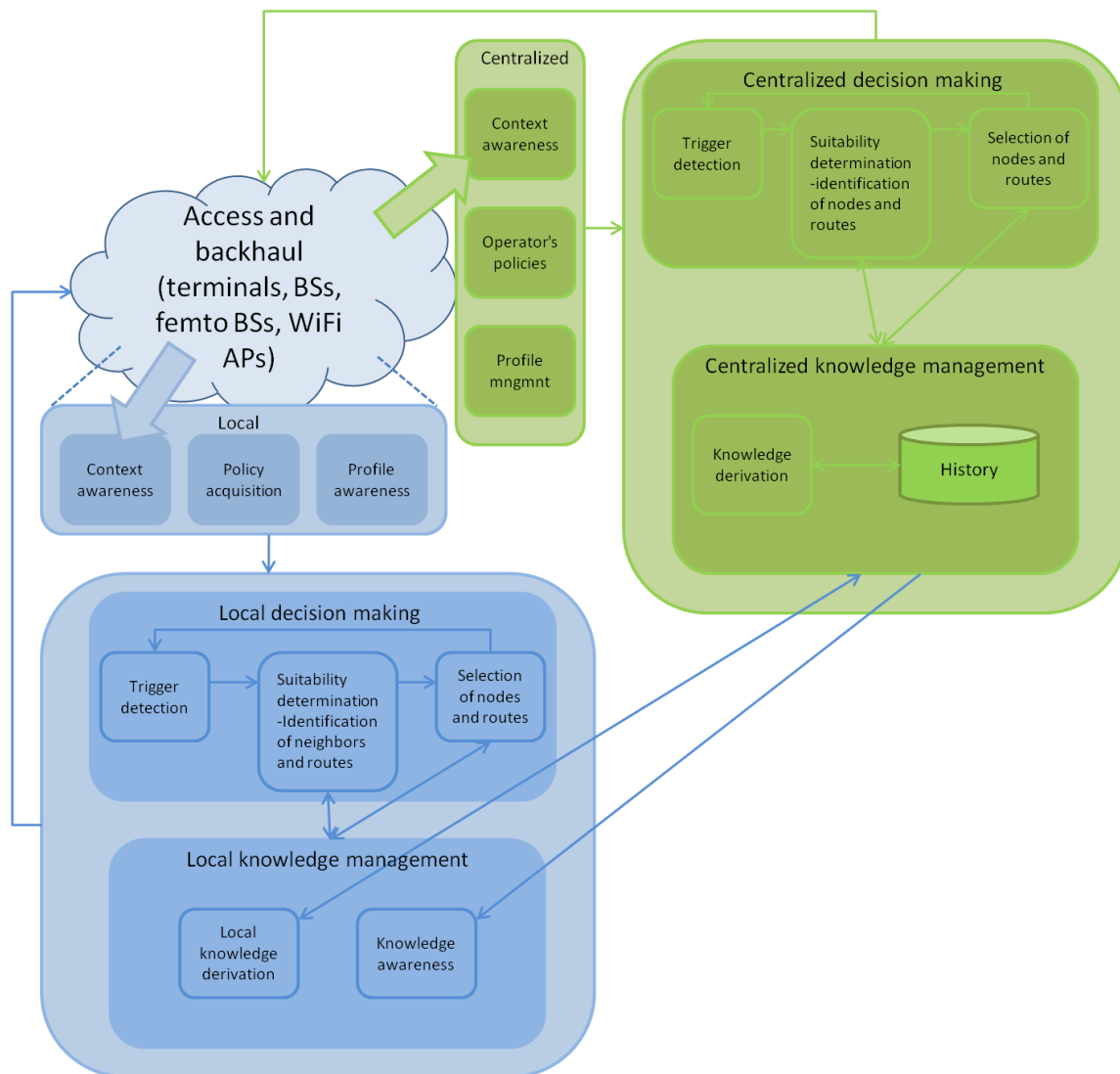


Figure 153: Mapping of “selection of nodes and routes” relevant algorithms to ON phases

In D4.3 steps forward have been made in order to clarify the complementarity of various algorithms by identifying potential interconnections. Figure 154 illustrates the identified algorithms and the mapping of these algorithms to ON phases. Moreover algorithms have been categorized to algorithms related to the wireless access and algorithms related to backhaul segments. The identified algorithms related to the wireless access are:

- Knowledge-based suitability determination and selection of nodes and routes
- Dynamic/ Energy-efficient resource allocation (capacity extension through femtocells)
- Direct D2D communication (UE-to-UE direct path)
- Multi flow routes co-determination
- Route pattern selection

The following algorithms are related to the management of backhaul segments:

- Application cognitive multipath routing in wireless mesh networks
- Content conditioning and distributed storage virtualization / aggregation for context driven media delivery

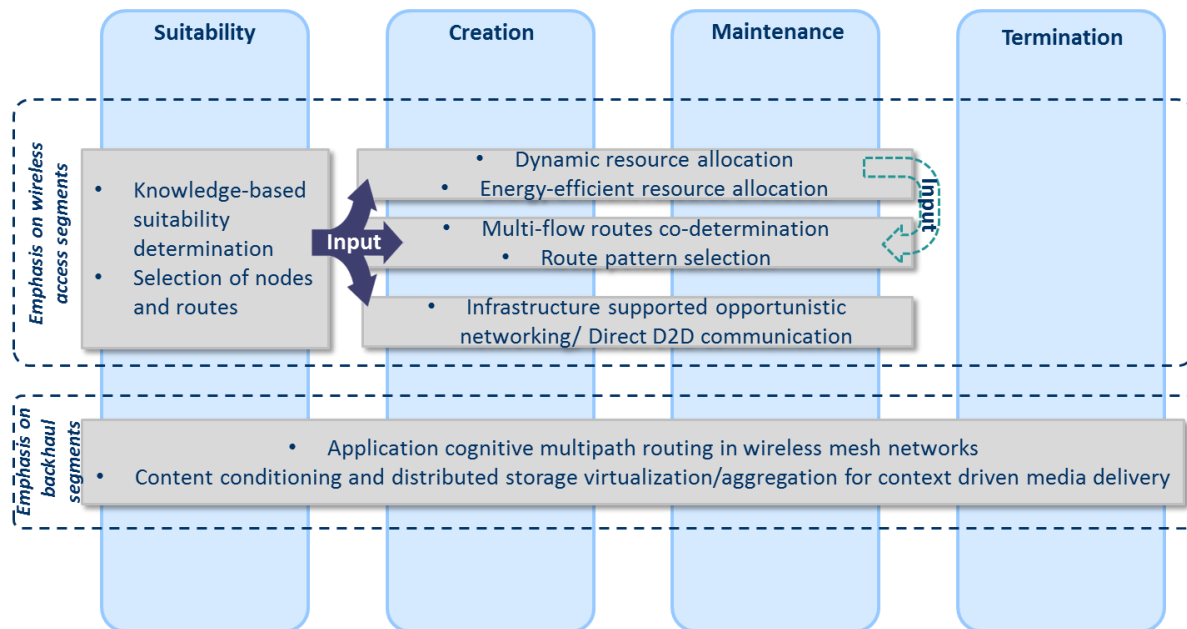


Figure 154: Mapping of “selection of nodes and routes” relevant algorithms to ON phases

Moreover, Figure 155 illustrates a synergic flowchart which shows potential sequence of execution of algorithms related to the management of the wireless access.

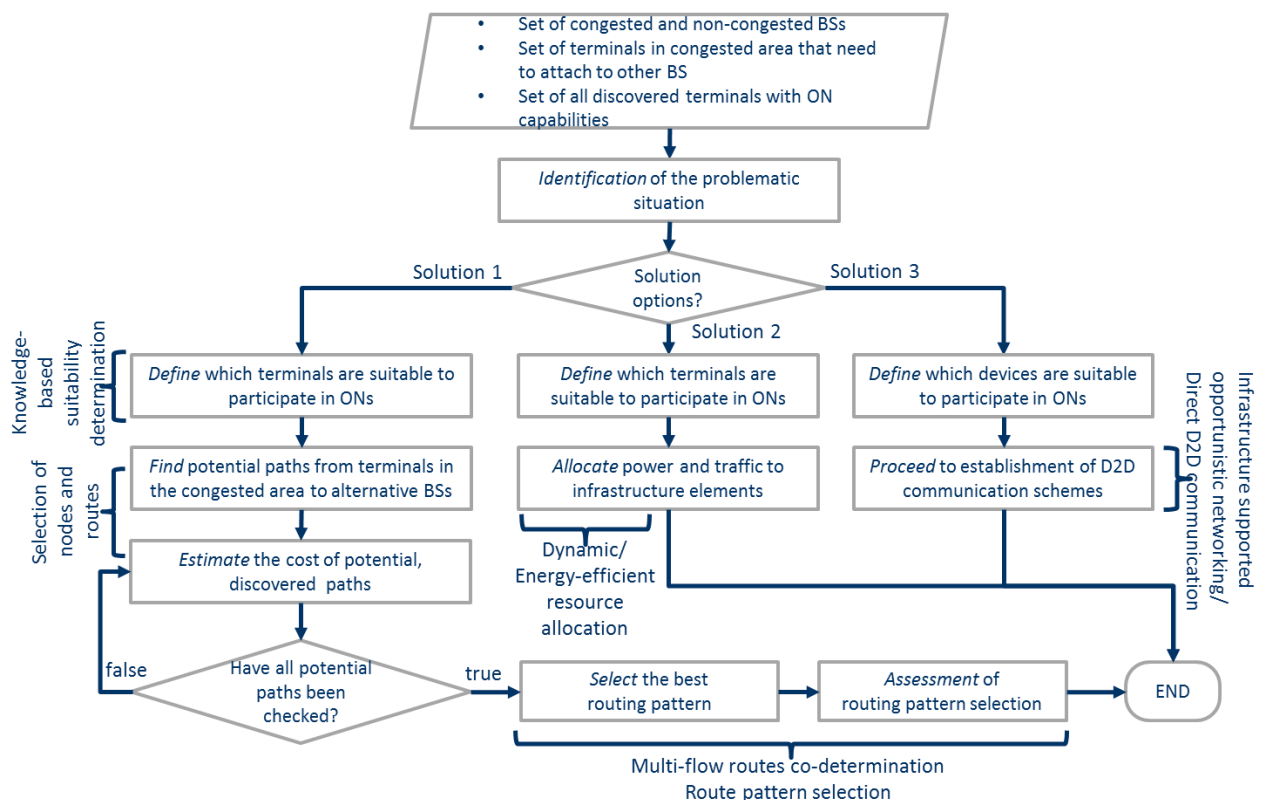


Figure 155: Synergic flowchart of “selection of nodes and routes” relevant algorithms

The flowchart uses as input a set of congested and non-congested BSs, a set of terminals in a congested area that need to attach to alternative BSs and a set of all discovered terminals with ON capabilities. After the identification of the problematic situation, 3 potential solution options are provided, which reflect the proposed algorithms. The potentials solutions are separated to the knowledge-based suitability determination and selection of nodes and routes which then leads to

the multi-flow routes co-determination/ route pattern selection; the dynamic/energy-efficient resource allocation (capacity extension through femtocells); and the direct D2D communication.

Furthermore, there have been specific efforts to combine inputs and outputs of various proposed algorithms in order to be able to identify the actual intersection points of the different solutions focusing on the wireless access. To that respect, Figure 156 shows that the output of knowledge-based suitability determination shall be used as input for the selection of nodes and routes, the dynamic/energy-efficient resource allocation and the direct D2D communication. Also, the output of the selection of nodes and routes shall provide useful input to the multi flow routes co-determination algorithm.

According to this approach, the interaction of proposed algorithmic solutions related to the management of the wireless access becomes clearer. Potential benefits from such an approach are also analysed in the section 3.2.2 that follows.

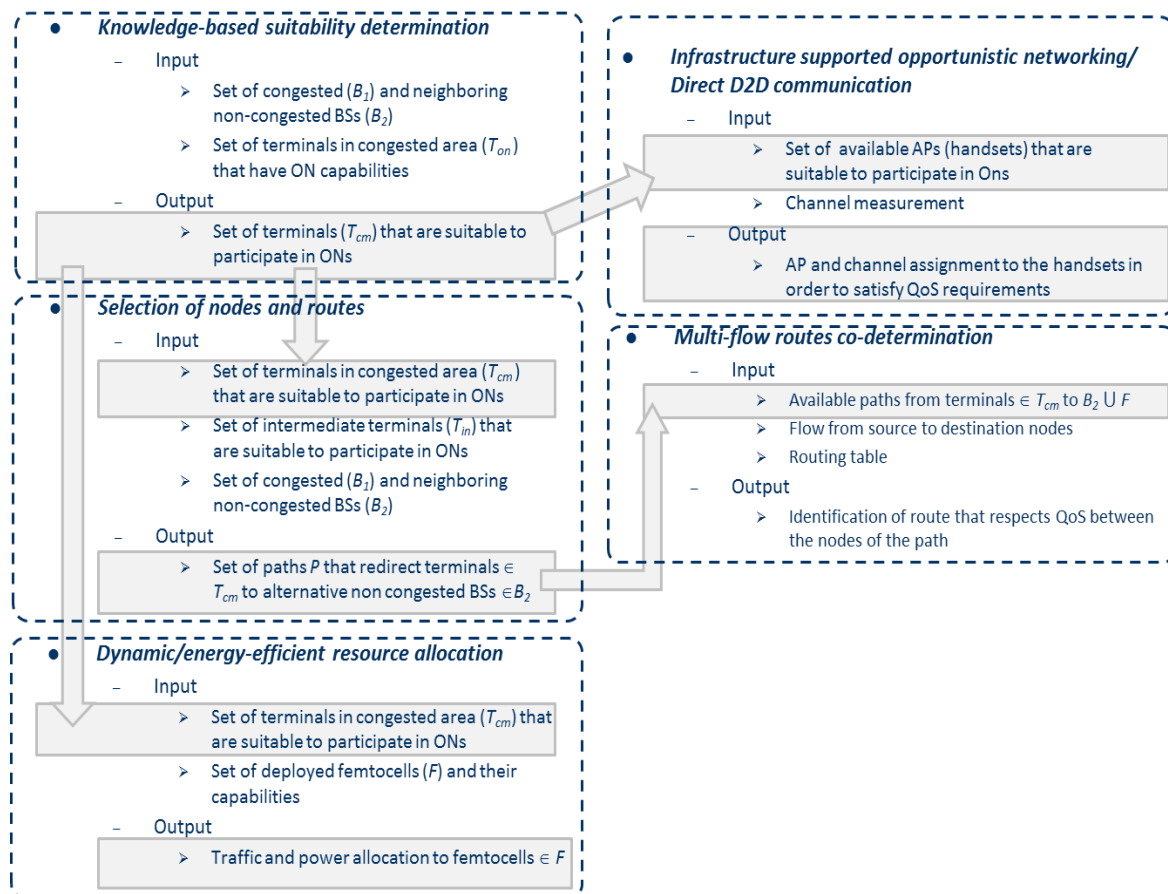


Figure 156: Synergetic inputs/outputs of “selection of nodes and routes” relevant algorithms

The synergetic solution for the backhaul resource management and related node and route selection is described in the section 4.5 of this document. Three of the OneFIT algorithms are used in the OneFIT scenario 5 which tackles the challenge of opportunistic backhaul resource aggregation. Algorithms described in sections 2.12 and 2.14 of this deliverable are developed, implemented and tested specifically for the scenario 5 and therefore their focus is on node and route identification and selection for management of backhaul resources. All performance evaluation and validation results presented for these algorithms can be mapped onto synergetic OneFIT solution for backhaul resource management.

3.2.2 Evaluation

This section intends to provide advantages of the various approaches regarding the selection of nodes and routes in order to present specific, consolidated solutions. To that respect, specific benefits have been identified in order to be able to show the importance of each metric to the ON performance. Initial versions of the identified benefits are provided in [7], so here enhanced versions are given.

3.2.2.1 Benefits in energy consumption of the infrastructure

This subsection analyses the benefits in energy consumption of the infrastructure. A first result set has been provided since D4.2 [7]. In D4.3 an enhanced result set is provided by taking into account extra studied cases as shown in the following table (Table 14) and as has been already described in section 2.7.

Table 14: Studied cases

Case	Moving terminals of each BS	Speed (m/s)
1	0	0
2	6	1
3	6	2
4	6	4
5	3	1
6	3	2
7	3	4

According to the studied cases as far as the consumption in the infrastructure is concerned, it is suggested that after the creation of an ON, the terminals that are located in more distant locations from the BS and having the capability to participate in an ON, will be these which will be offloaded to the neighboring BSs. Therefore, the remaining terminals that are closer to the BS in terms of distance shall need decreased transmission power in order to communicate with the BS. To that respect, Figure 157 illustrates the aforementioned impact to the infrastructure before and after the creation of ONs.

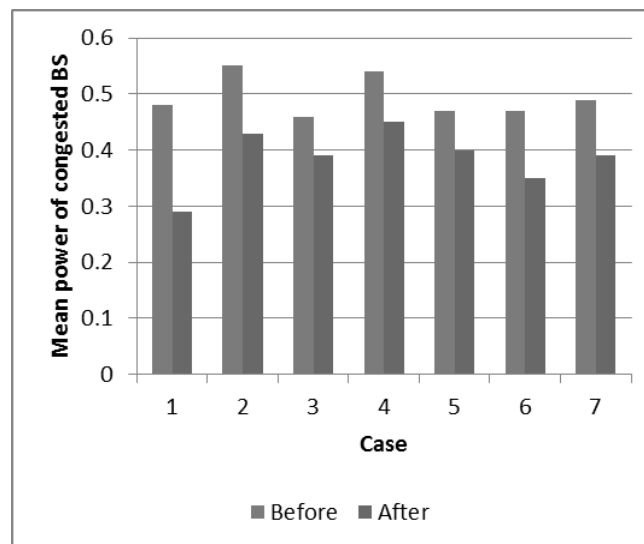


Figure 157: Mean power of congested BS

Moreover, regarding the impact of mobility to the energy consumption of the non-congested BSs, it has been observed that for low mean velocity values (which in turn lead to more stable ONs), the mean power increases after the solution due to the fact that the non-congested BSs acquire a proportion of the traffic of the congested BS. For higher speed values, the ONs (after some time) are more likely to be destroyed (due to mobility), hence the traffic acquired by non-congested BSs is lower. The impact is depicted on Figure 158.

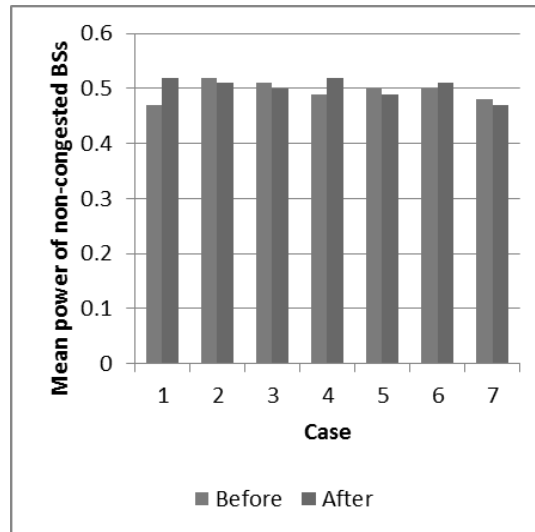


Figure 158: Mean power of non-congested BSs

In addition, opportunistic content placement on WMN nodes and aggregation of backhaul storage resources of WMNs results in significant reduction of energy consumption in content delivery services as presented in [28] and [7].

3.2.2.2 Benefits in energy consumption of the terminals

This subsection analyses the benefits in energy consumption of the terminals. Specifically, according to the additional studied cases presented in Table 14, benefits derive for edge terminals (i.e. terminals that are located near the BS coverage border and can switch to ON). As Figure 159 suggests, after the creation of the ON, the edge terminals switch to ON and this means that they are connected to closer neighbors compared to their previous connection to a distant BS. Therefore, the power that is needed for the communication among the terminals is lower than the one that is needed for the communication with the distant BS.

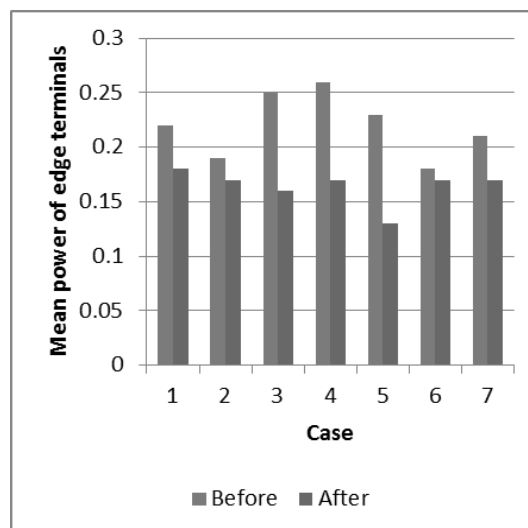


Figure 159: Mean power of edge terminals (terminals that switch to ON)

3.2.2.3 Communication benefits through the exploitation of ONs

Another identified benefit through the proposed solutions on selection of nodes and routes are the benefits related to communication through the exploitation of ONs. Specifically, a network is considered as in section 2.14.4 that comprises a macro BS with 9 deployed femtocells within its coverage area and 40 terminals.

Figure 160 shows the average delay for delivered messages from all 40 terminals before and after the solution with respect to the number of terminals that are offloaded to femtocells, their mobility level and the spatial distribution. In all cases, the proposed concept tends to perform better since the delay drops compared to the situation before the solution enforcement. Specifically, the decreasing delay is higher when more terminals are offloaded to femtocells, while as the moving speed of the terminals increases, the average delay tends to increase as well. In general, the delay is reduced by around 15% in average after the exploitation of the neighbouring femtocells.

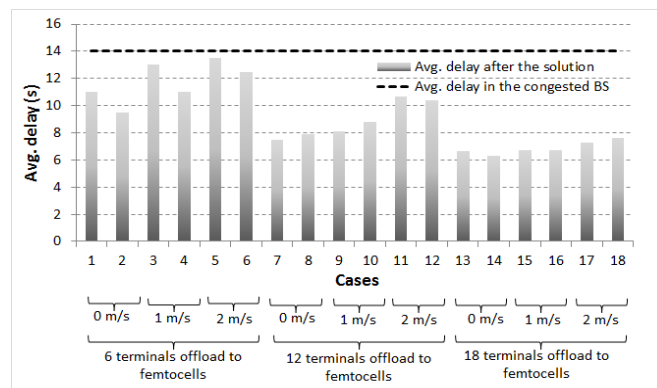


Figure 160: Network average delay before and after solution enforcement

3.2.2.4 Resource utilization benefits in terms of increased capacity of the underlying network

Nodes and routes for creating the ON can be selected with the goal to increase the resource utilization in underlying networks. A detailed analysis of impact of opportunistic resource management on capacity of underlying networks can be found in [7] and in description of algorithms presented in sections 2.12 and 2.14 of this document.

4. Dynamic Operation of the ON

This section provides a description of the overall dynamic operation of the opportunistic networks, based on the functional integration and synergic operation among the different algorithms. Evaluations of the integrated algorithms are presented for selected cases and scenarios.

Scenario + Phase (S: Suitability det., C: Creation, M: Maintenance) per Algorithm	Centralized / Distributed	SCE1 Opp. Coverage extension	SCE2 Opp. Capacity extension	SCE3 Infrastructur e supported Opp. Ad-Hoc Networking	SCE4 Opp. traffic aggregation in RAN	SCE5 Opp. Resource aggregation in the back- haul
1. Suitability for direct D2D communication (ALUD)	Cent.	-	-	S	-	-
2. Suitability for the coverage ext. scenario (ALUD)	Cent.	S	-	-	-	-
3. Mod. decision flow approach for selecting freq., bandwidth and RAT (VTT)	Cent.	SCM	-	SCM	-	-
4. Fittingness Factor based spectrum selection (UPC)	Cent.	CM	CM	CM	CM	CM
5. Machine Learning based Knowledge Acquisition on Spectrum Usage (UNIS)	Cent.	CM	CM	CM	CM	-
6. Techniques for Aggregation of Available Spectrum Bands/Fragments (UNIS)	Cent.	CM	-	-	CM	-
7. Knowledge-based suitability determination and selection of nodes and routes (UPRC)	Cent.	S	S	-	-	-
8. Route pattern selection in ad hoc network (TCS)	Dist.	SM	SM	SM	SM	-
9. Multi-flow routes co-determination (TCS)	Dist.	M	M	M	-	-
10. QoS and Spectrum-aware routing techniques (UNIS)	Dist.	CM	CM	CM	CM	-
11. Techniques for Network Reconfiguration – topology control (UNIS)	Cent.	CM	-	-	-	-
12. Application cognitive multi-path routing in wireless mesh networks (LCI)	Cent.	-	-	-	-	SCM
13. UE-to-UE Trusted Direct Path (NTUK)	Cent.	-	CM	CM	CM	CM
14. Content conditioning and distributed storage virtual aggregation for context driven media delivery (LCI)	Cent.	-	-	-	-	SCM
15. Capacity Extension through Femto-cells (UPRC)	Dist.	-	CM	-	-	-
16. Validation of ON algorithms on an offloading-oriented real-deployment testbed (TID)	Cent.	SC	SC	SC	-	-

Table 15: Mapping of Algorithms to Scenarios and Phases

4.1 Scenario 1 "Opportunistic Coverage Extension"

Scenario 1 "Opportunistic coverage extension" describes a situation in which a device (here: UE#1) cannot connect to the network operator's infrastructure, due to lack of coverage or a mismatch in the radio access technologies. In order to provide mobile access to UE#1, another node (or several nodes) must provide a relaying service towards the infrastructure. Therefore, an ON is created. In the example shown in Figure 161: Opportunistic Coverage Extension, the ON consists of 3 nodes: UE#1, UE#2 (which provides the relaying service) and BS#1.

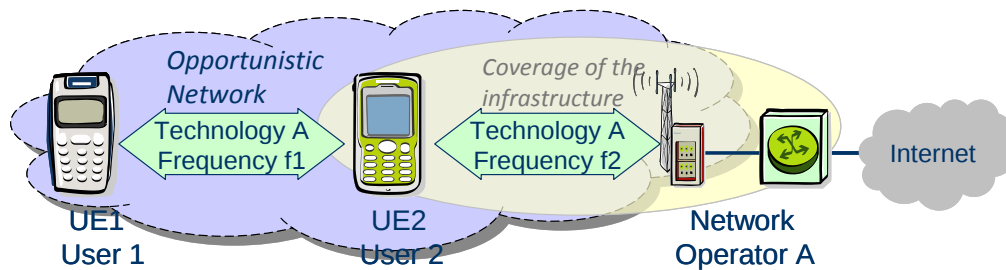


Figure 161: Opportunistic Coverage Extension

This scenario is described in D2.1 [2] Section 4.1 and message sequence charts are provided for a network initiated scenario as well as for a terminal initiated scenario in D2.2.2 [4] Section 5.1 where message sequence charts. Further details on the information flows can be found in D3.2 [5] Section 5.1. The algorithms involved in scenario 1 are listed in the Figure 162 below depicting at which phase each algorithm is processed.

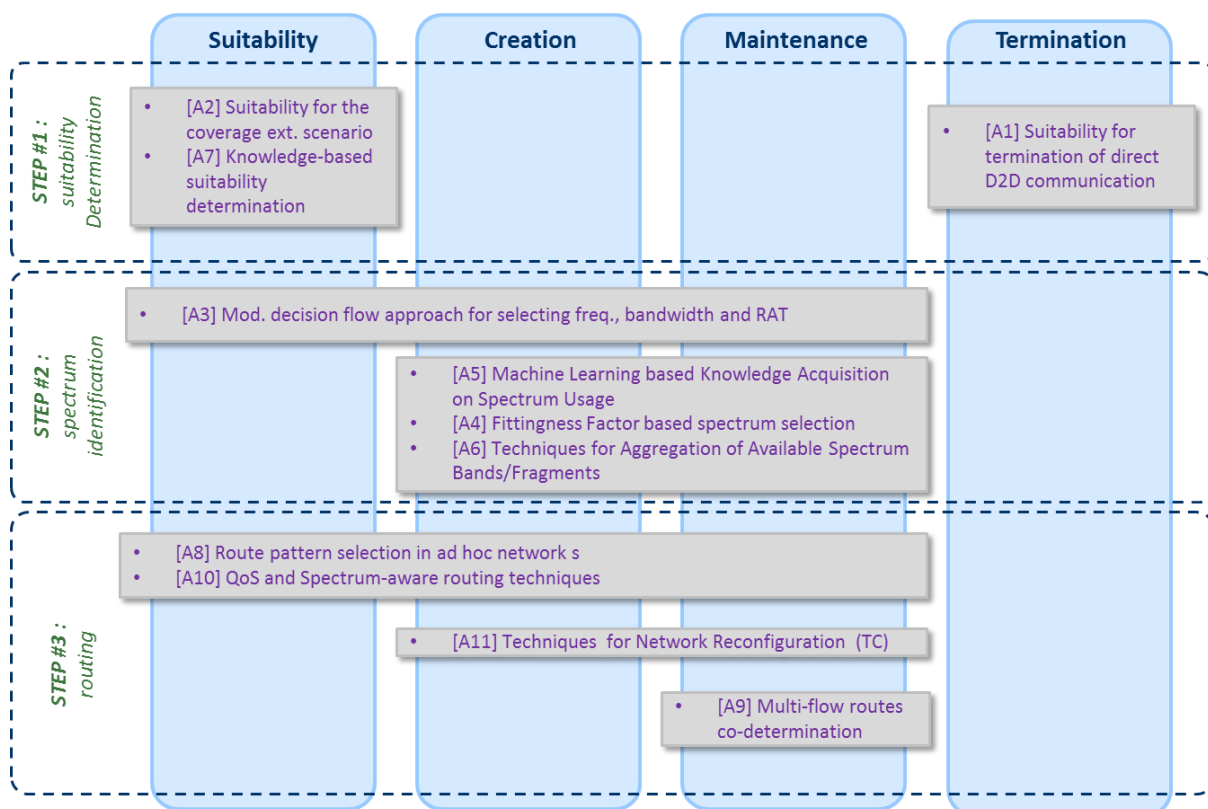


Figure 162: Mapping of the Algorithms to the different phases for Scenario 1

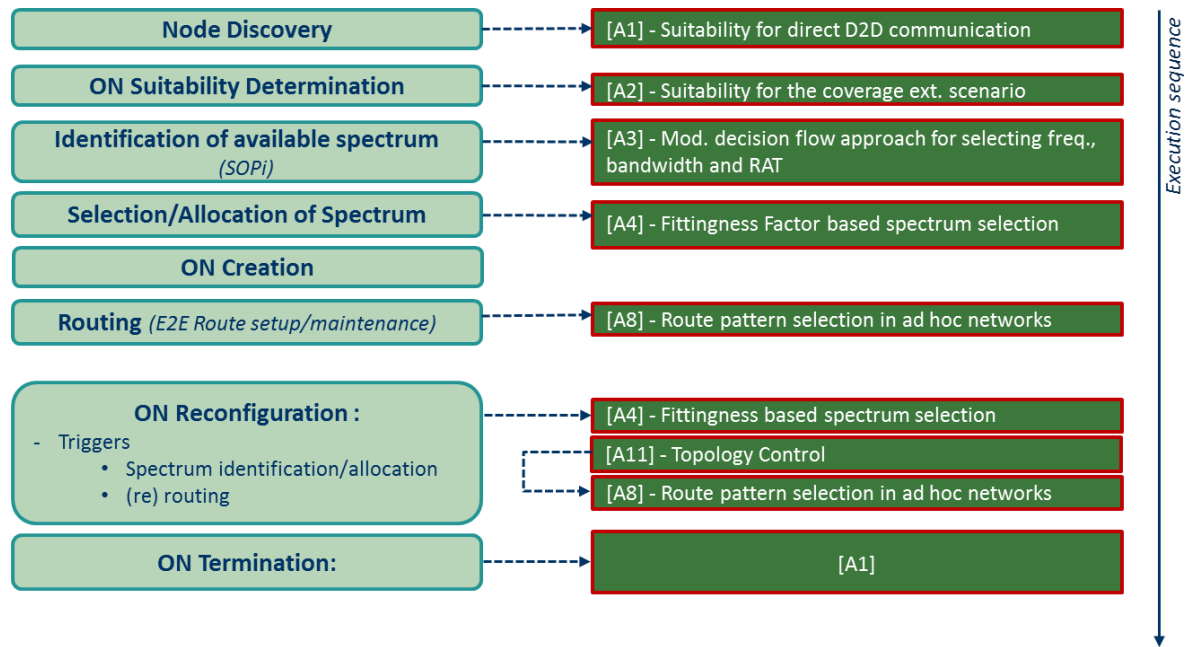


Figure 163 – Example algorithm execution sequence for the OneFIT scenario 1

4.1.1 ON Suitability

Upon loss of coverage or when a device needs to start a communication without any direct access to a BS, it initiates the suitability procedure. This procedure is depicted on the figure below:

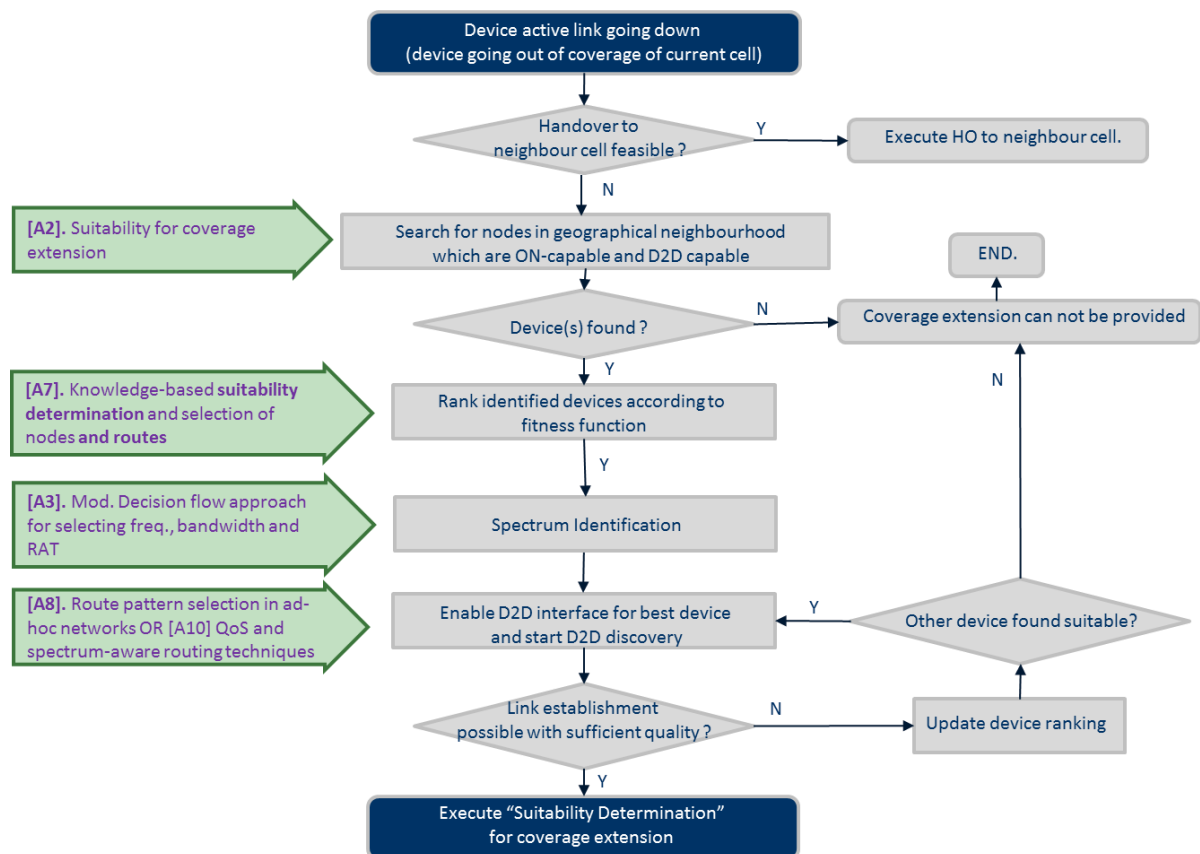


Figure 164: Scenario 1 - Mapping of algorithms for ON Suitability phase

This phase is mainly directed by the algorithm [A2]. When the BS detects a loss of coverage and if it is impossible to perform any handover, the infrastructure selects a set of nodes according to the fitness function, by processing the algorithm [7]. In order to communicate directly between the terminals, it is identified the available spectrum, by processing the algorithm [3]. Then it is established a route between the node in loss of coverage and the infrastructure, by processing the algorithms [8] or [10].

4.1.2 ON Creation

Once the ON suitability/determination has been identified, some algorithms are processed in order to prepare the ON creation, in order to establish properly the ON. The order in which algorithms depicted in Figure 162 are executed during the ON creation phase is presented in Figure 165.

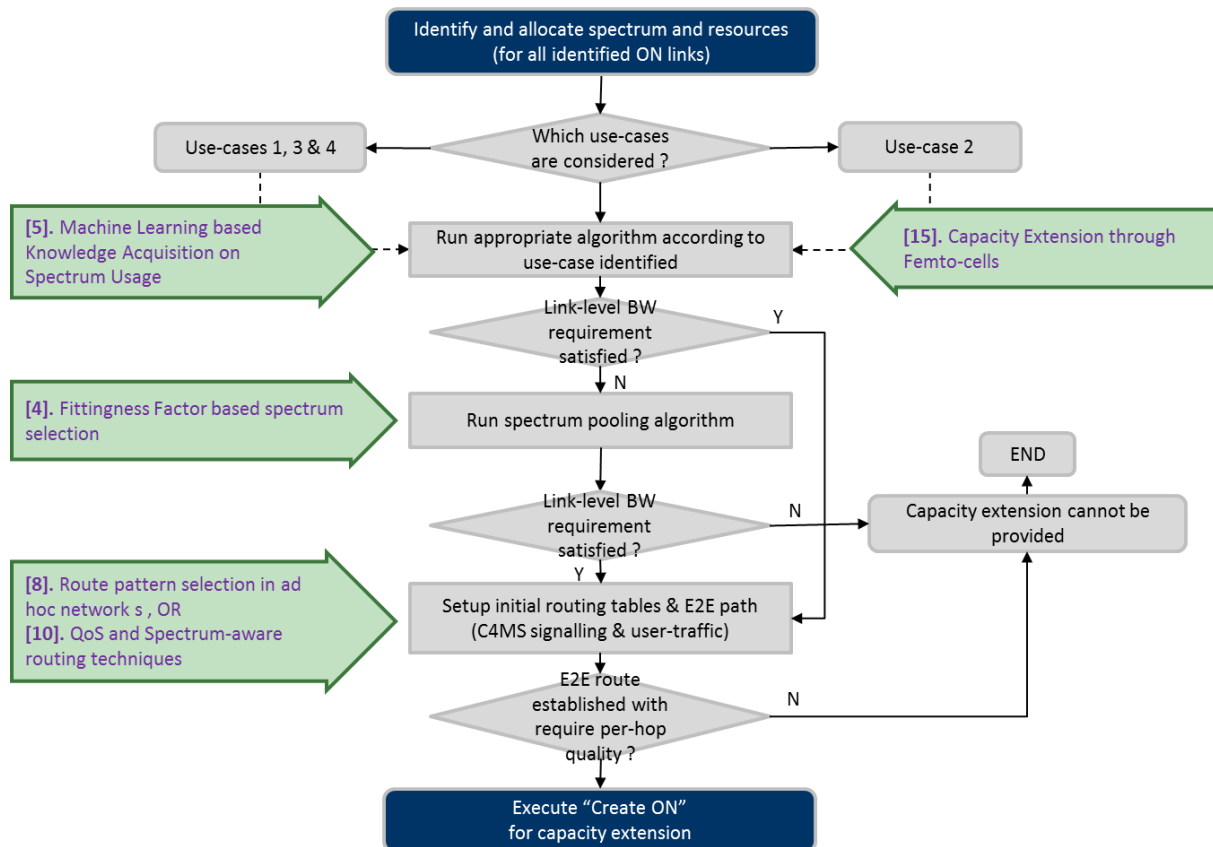


Figure 165: Scenario 1 - Mapping of algorithms for ON Creation phase

4.1.3 ON Maintenance

After the creation phase, the ON is operational and data flows can be exchanged through the ON.

The list of involved algorithm for the ON maintenance phase is depicted in Figure 166.

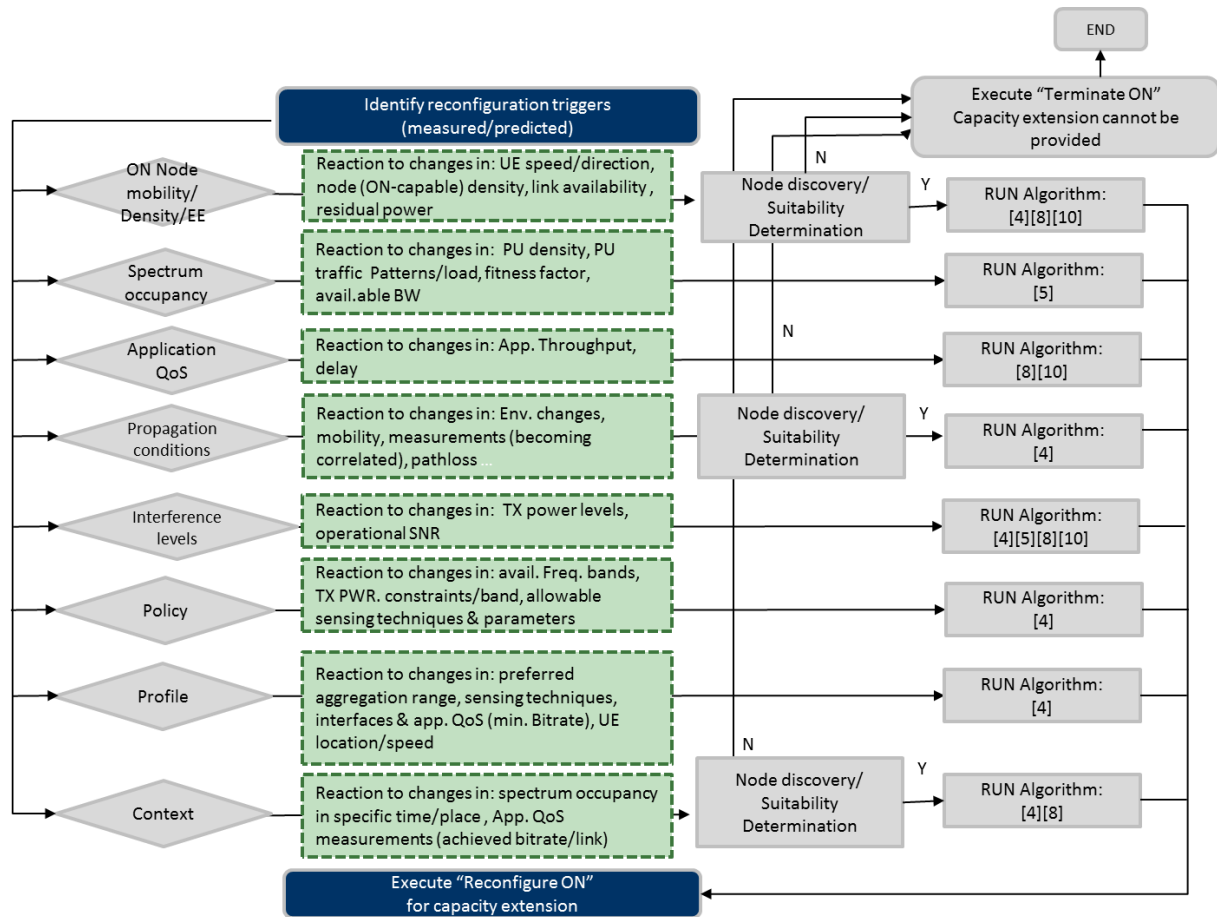


Figure 166: Scenario 1 - Mapping of algorithms for ON Maintenance phase

4.1.4 ON Termination

The termination phase is initiated when, during the maintenance phase, the monitoring of the concerned node has retrieved a BS coverage, either when the reason why the opportunistic network has been created is no longer available (for example at the end of data transfer).

4.1.5 Performance evaluation

4.1.5.1 Performance evaluation of spectrum selection

In the following, the impact of spectrum selection in the coverage extension scenario is evaluated. For that purpose, let assume a scenario with dimensions 10 km x 10 km, and the transmitter of the infrastructure located in the central point. For those nodes that are outside the coverage area of the infrastructure, the possibility to form an ON with another node having direct coverage is considered. The ON is composed then by 2 links: the direct link that connects the infrastructure with a terminal having coverage, and the ON link that connects the node without coverage with another terminal in the coverage area. It is assumed that the node with direct coverage can use its full capacity for relaying the ON link.

Given the different characteristics of the propagation for the direct link (i.e. from a mobile node to a base station) and for the ON link (i.e. from a mobile node to another mobile node), the following general propagation model is assumed, whose parameters are particularized for each case [26]:

$$L_p \text{ (dB)} = K + \beta \log f \text{ (GHz)} + \alpha \log d \text{ (km)} + S \quad (58)$$

S(dB) corresponds to the shadowing, following a Gaussian distribution with mean 0 and standard deviation σ dB. Shadowing is spatially correlated with exponential autocorrelation and de-correlation distance d_{corr} . The generation of 2D spatially correlated shadowing follows the methodology of [27] based on filtering a set of independent shadowing samples using a 2D filter defined from the Fourier transform of the exponential autocorrelation function. The shadowing of the direct and the ON links are independent.

Three possible spectrum pools are considered for the spectrum selection process for the ON link. The first option makes use of 20 MHz in the license-exempt ISM band at 2.4 GHz. The second option uses the TV White Space (TVWS) band at 600 MHz operated opportunistically so that transmission is allowed whenever no interference is generated to TV receivers. Finally the third option uses some MNO (Mobile Network Operator) licensed band owned by the operator at 1800 MHz.

Table 16 presents the considered scenario parameters for the evaluation, considering both the direct link and the ON link characteristics.

Table 16: Scenario parameters.

Direct link		
Frequency (f)		900 MHz
Bandwidth (BW)		3 MHz
Noise and interference spectral density (I_o)		-164 dBm/Hz
Transmit power (P_{Tmax})		40 dBm
Propagation model	α	37.6
	β	21
	K	122.1 dB
	σ	6 dB
	d_{corr}	100m
ON link		
Option 1: ISM band	f_1	2.4 GHz
	BW_1	20 MHz
	$I_{o,1}$	-164 dBm/Hz
	$P_{T,1}$	17 dBm
Option 2: TVWS band	f_2	600 MHz
	BW_2	8 MHz
	$I_{o,2}$	-164 dBm/Hz
	$P_{T,2}$	20 dBm
Option 3: MNO band	f_3	1800 MHz
	BW_3	5 MHz
	$I_{o,3}$	-164 dBm/Hz
	$P_{T,3}$	21 dBm
Propagation model	α	40
	β	30
	K	141.7 dB
	σ	6 dB
	d_{corr}	40 m

A minimum bit rate of $R_{req}=5$ Mb/s (measured at physical layer from the Shannon bound) is needed by the application. This requirement poses the limits for the coverage of both the direct and the ON

links. More specifically, in Figure 167(a) the SNR (Signal to Noise Ratio) is plotted in the different points of the scenario. Based on this SNR and the Shannon bound, Figure 167 plots the area where the direct communication at R_{req} is possible.

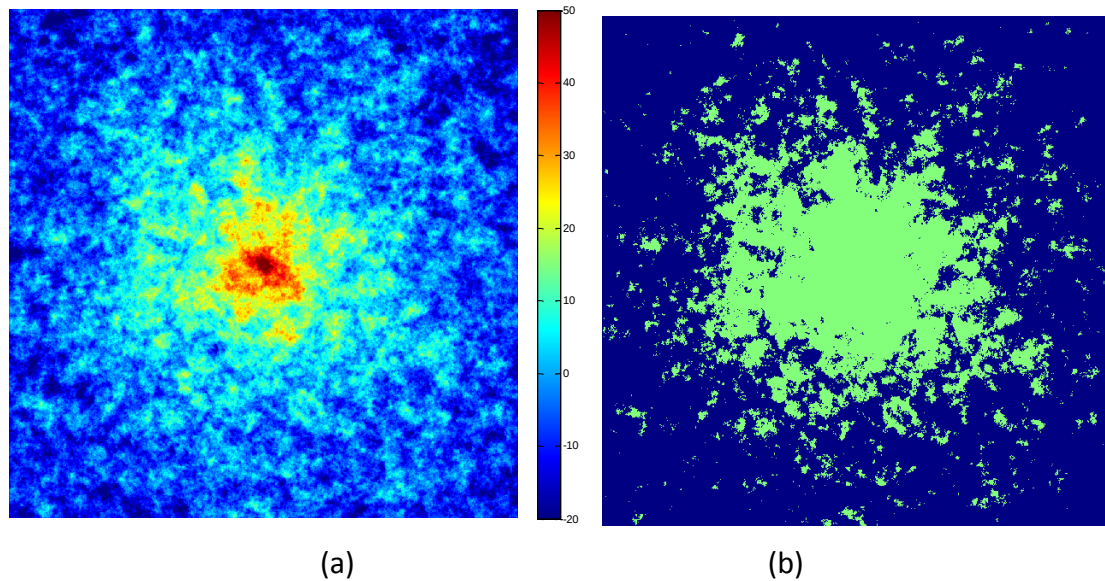


Figure 167: (a) SNR (dB) with the direct link, (b) area (in green) where the direct communication can be established

Based on the coverage map depicted in Figure 167 for the direct link, the relevant KPI considered for the analysis is the probability of having coverage either direct or through the ON, for a terminal located at distance D from the infrastructure transmitter. This probability will be analysed depending on the decisions made by the spectrum selection algorithm. Results for each distance D are averaged for a total of 359 positions with a separation of 1° .

As a first result, Figure 168 plots the probability of having direct coverage as a function of distance D . For comparison purposes, the reference case without shadowing is plotted. In this case, note that there is a distance limit around 2.4 km beyond which it is not possible to achieve the required 5 Mb/s with the direct link. In turn, when shadowing is considered, there is a gradual reduction in the probability as the distance D increases, and for 2.4km the probability of having direct coverage is 50%.

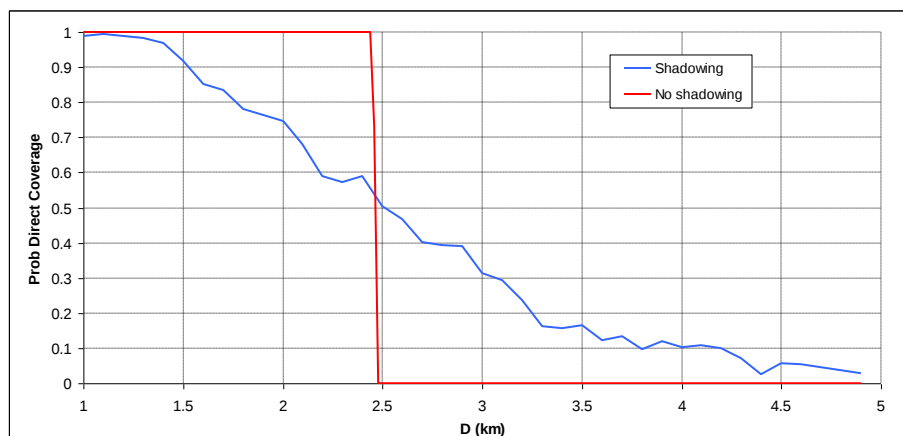


Figure 168: Probability of having direct coverage as a function of the distance D to the infrastructure transmitter with and without shadowing

Considering now the different choices that can be made by the spectrum selection algorithm for the ON link, Figure 169 plots the resulting total coverage probability for the three considered options for

the ON link to provide service to those nodes outside the direct coverage area. Computations assume in this case a node density $\eta=50$ nodes/km². Firstly, it can be noted how the selection of one or another spectrum pool has a high influence over the coverage probability. Clearly, the use of TVWS, operating at a lower frequency, allows having higher coverage probabilities than the use of the MNO or the ISM bands. Secondly, it is observed how the effect of shadowing is overall beneficial from the perspective of coverage probability. For instance, with TVWS for the ON link it is possible to have a significant coverage probability above 80% even at large distances such as 4 km. On the contrary, in the case without shadowing, the ON is only able to extend the coverage in roughly 400m beyond the limit of 2.4km when TVWS are used. To complement this result, Figure 170 plots the corresponding value of the surface S_i consisting in the points of the direct coverage area that can be reached by a node without coverage located at distance D . It can be seen that the area is larger with TVWS than with the rest of bands, and also that its value is larger when there is shadowing than when there is not. In the latter case, this area tends quickly to zero, while with the presence of shadowing the area is kept at non-zero values even at large distances. This fact, combined with the node density of $\eta=50$ nodes/km², allows having a significant probability of finding a node to form the ON even at large distances. Note that for distances below 2.4 km, S_i is not defined in the case without shadowing because all points have direct coverage.

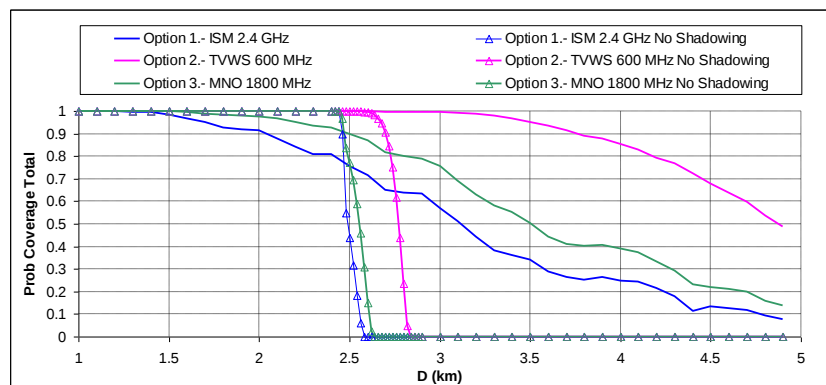


Figure 169: Coverage probability as a function of distance D for the different spectrum band alternatives, with and without shadowing.

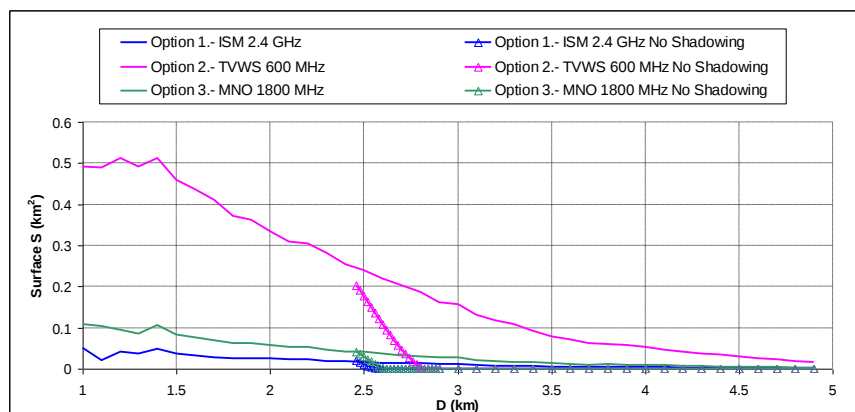


Figure 170: Average value of the surface S_i with direct coverage that can be reached through an ON link for the different options.

Figure 171 plots the variation of the coverage probability thanks to the ON formation as a function of node density η when terminals are at $D=4$ km in case of shadowing. The same trend with respect to the best behaviour of the TVWS band is observed like in previous results. It can also be noticed how the coverage probability is very sensitive to the node density. Focusing for instance on the TVWS, the coverage probability can range from around 15% for the case with the lowest node density considered here up to around 95% for the largest density. Note that for $D=4$ km, the case

without shadowing (not plotted in the figure) would yield a coverage probability of 0% regardless the node density, because at this distance the surface S_i with direct coverage reachable with the ON link is 0 for all the considered options as seen in Figure 170.

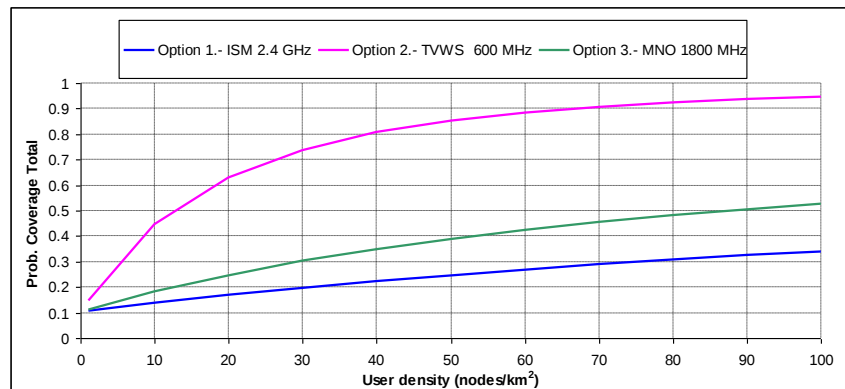


Figure 171: Coverage probability as a function of the node density η at distance $D=4\text{km}$ when shadowing is considered

To summarise the main observations, this study has reflected that a smart spectrum selection strategy is very relevant for establishing the ON link that provides the coverage extension in this scenario. It has been shown that the choice made by the spectrum selection algorithm can have an influence on the total coverage probability. Moreover, the node density assumed, also has a relevant impact because it affects the probability of being able to find and reach a terminal with direct coverage.

4.1.5.2 Evaluation of selection of nodes through a fitness value

For the evaluation of scenario 1 a specific algorithm has been developed, namely “selection of nodes through a fitness function”. Figure 172 illustrates the indicative duration of connections in static nodes (subject to their energy level). As the figure suggests, as more nodes are accepted in the ON, an increase of the ON’s duration may be observed.

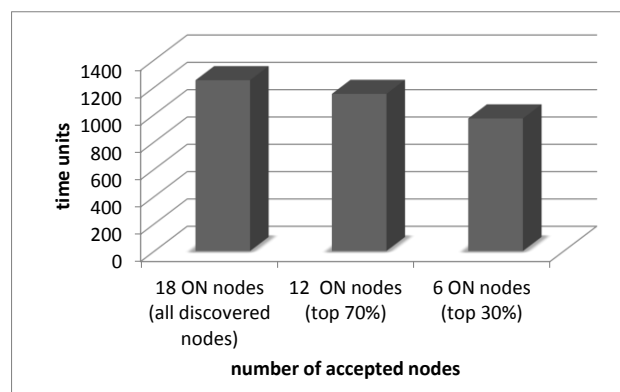


Figure 172: Indicative duration of connections by taking into account the number of accepted ON nodes.

Moreover, Figure 173 illustrates the percentage of aborted messages with respect to the number of accepted ON nodes. As the figure suggests, as more nodes are accepted in the ON, an increase in the percentage is observed since it is more likely to experience a failure in a link. Also, the number of aborted messages corresponds to the messages which are dropped due to a link drop at the moment of their transmission (so the message is not successfully transmitted to the receiving node).

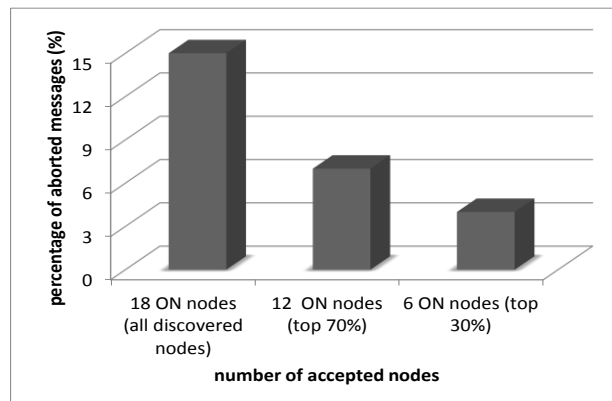


Figure 173: Percentage of aborted messages by taking into account the number of accepted ON nodes.

4.2 Scenario2 “Opportunistic Capacity Extension”

Scenario 2 “Opportunistic Capacity Extension” depicts the creation of an opportunistic network that derives the traffic from a localized hot-spot to non-congested access points. This scenario may also include cases such as dynamic spectrum management between macrocells and underlying micro-, pico- and femtocells, or 3G traffic offloading towards Wi-Fi APs.

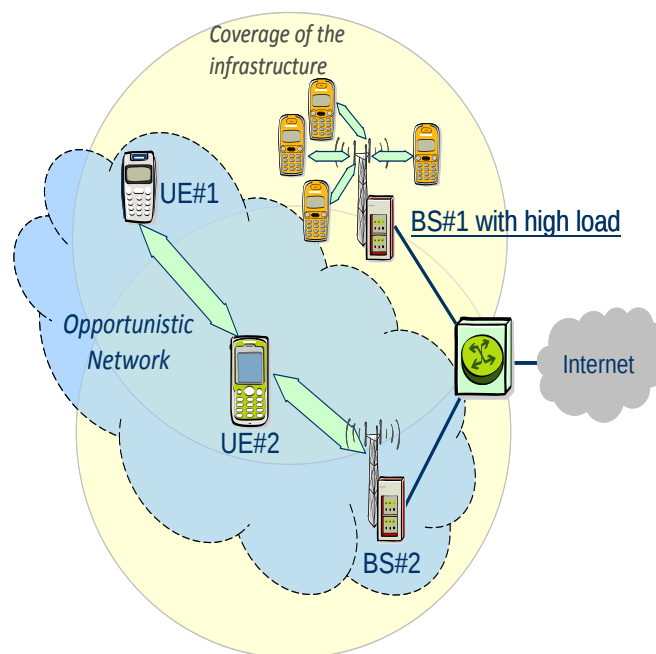


Figure 174: Opportunistic Capacity Extension

This scenario is thoroughly described in D2.1 [2] Section 4.2, including four different use cases. Some message sequence charts are shown in D2.2.2 [4] Section 5.2 (MSCs for both a network initiated scenario and a terminal initiated scenario are provided). Further details on the information flows among devices and infrastructure can be found in D3.2 [5] Section 5.2 and 5.3.

The algorithms involved in scenario 2 are shown in the following figure, including the phases at which each algorithm is triggered.

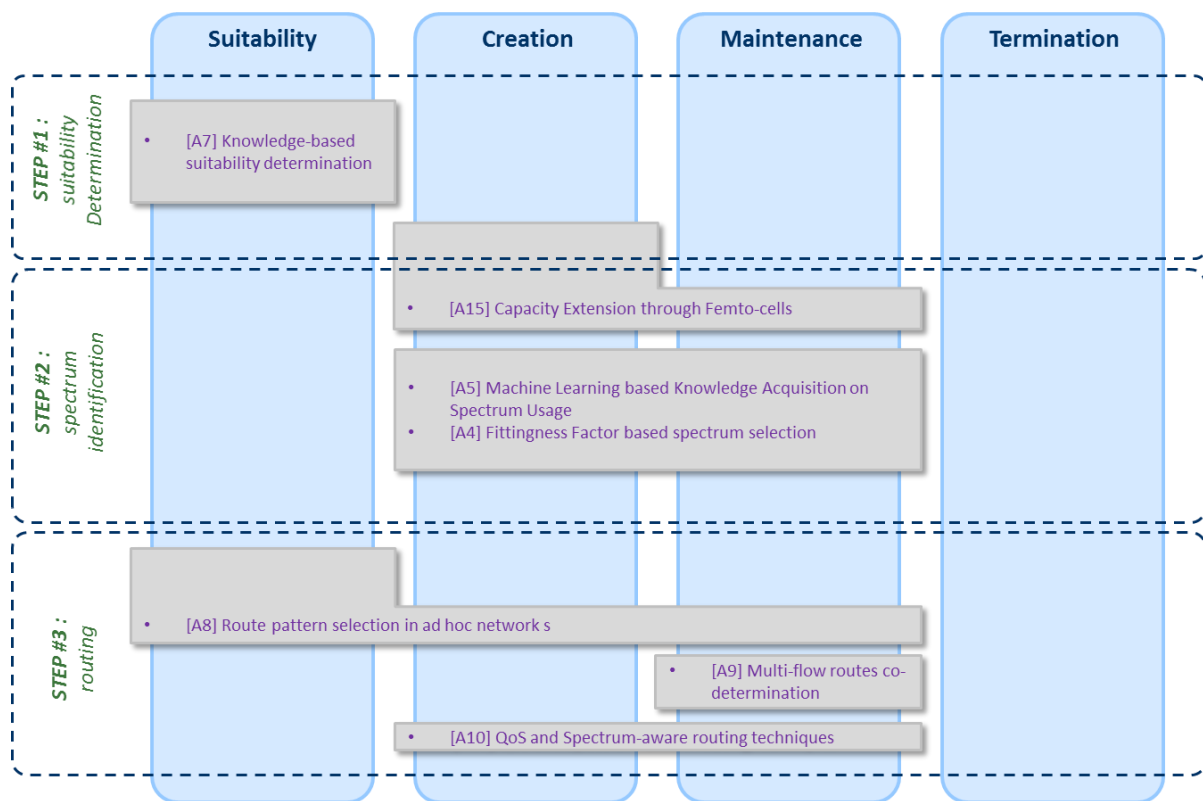


Figure 175: Mapping of algorithms to ON phases in Scenario 2

The following subsections describe the order in which events should happen for each ON management phase and which algorithms should be triggered according to them. But, in first place, in order for the first phase (ON Suitability Determination) to be initiated, two procedures must be running so that all the necessary information from the existing nodes and the environment is collected:

- **ON Information procedure:** Involves information concerning status and capabilities of neighbouring nodes, location and profiles of connected users, operator policies, etc. This information is retrieved from nodes and UEs and is collected at the infrastructure CSCI, by means of C4MS signalling (messages INR, INA, INI, as shown in D3.3 Part B [6], section 1.1.)
- **ON Node discovery procedure:** Identification of additional nodes (e.g. non-infrastructure APs) that may be useful for the creation of a future ON. This information is collected by infrastructure CSCI by means of beacon/broadcast or probe signals (a more detailed description can be found in D3.3 Part B [6], section 1.2.)

Additionally, the infrastructure CSCI entity must permanently monitor the utilisation of radio resources in order to detect traffic imbalances among nodes or to forecast the apparition of congestion situations. The QoS of certain infrastructure-UE links (e.g. for cell-edge users) must also be monitored so as throughput under achievements due to congestion can be detected.

Once the congestion / throughput underachievement / resource imbalance has been identified (or forecasted) in a located area, the infrastructure may decide to trigger the creation of an Opportunistic Network to solve (or avoid) such situation. In that case, the ON management phases described in subsections hereafter will occur:

4.2.1 ON Suitability

The following figure shows the flow diagram for this phase:

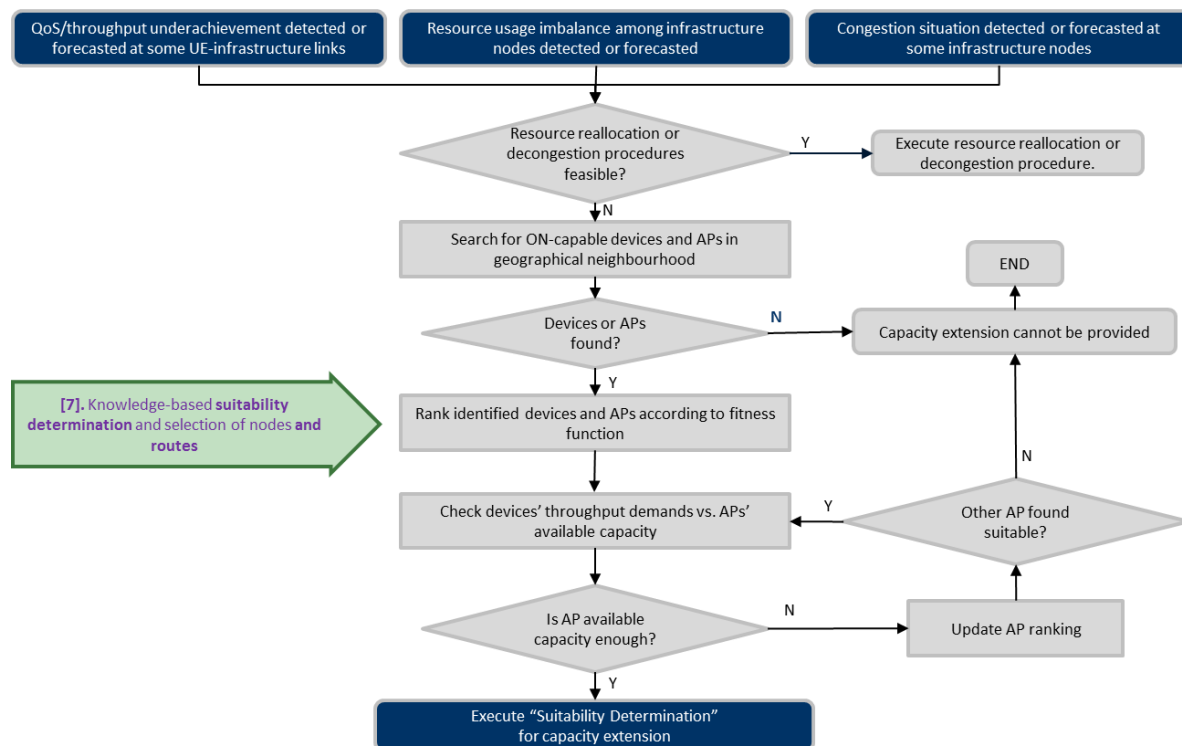


Figure 176: Scenario 2: Flow diagram for ON Suitability phase

In order to begin the suitability procedures, the CSCI entities of the problematic infrastructure nodes must in first place retrieve additional context information. Neighbouring candidate nodes (infrastructure or non-infrastructure) are contacted via DSONPM, while candidate UEs must be contacted through their serving infrastructure nodes.

Once all necessary information has been gathered, and other standard decongestion procedures have been discarded, the CSCI tries to find a set of configuration parameters that may allow the future ON to properly address or solve the congestion problem. To do so, the OneFIT comprehensive solutions (see chapter3) must be used:

- **OneFIT solution for spectrum:** Selects the most appropriate frequency band for the ON operation. This band may be a compact frequency block or a virtual block made up of several non-adjacent chunks, but it must be available, able to support the estimated offloaded throughput and comply with the interference requirements. Spectrum availability will be checked both in licensed and unlicensed band, according to the operator's policies.
- **OneFIT solution for nodes and routes:** Determines the best routing strategy to divert the traffic from the congested nodes towards the candidate nodes (or a subset of them). This strategy may include selecting APs from different RATs, change femtocells policies (e.g. from CSG to OSG), establishing UE-to-UE direct links, etc.

In the particular case of Scenario 2, both procedures can be completed by triggering the algorithm described in section 2.7.

Finally, the determination of the most suitable configuration triggers the following phase.

4.2.2 ON Creation

The creation of the ON is preceded by a negotiation procedure, so that congested nodes provide the selected candidates with additional information about the traffic to be offloaded, the profiles of the running applications and involved users, the operator's policies, etc. This additional information is used by destination nodes to evaluate (and negotiate if needed) the resource allocation necessities

for absorbing the diverted traffic. The negotiation procedure is performed by means of C4MS messages ONNR and ONNA, (see D3.3 Part B [6], and section 1.4.)

Finally, the negotiation leads to the definitive configuration of the ON, and its creation can be initiated. The creation procedure involves the interchange of additional signalling among the involved nodes (the C4MS messages to be used are ONCR and ONCA, see D3.3 Part B [6], and section 1.5.) and also the initiation of specific procedures to establish the new links as shown in the following flow chart:

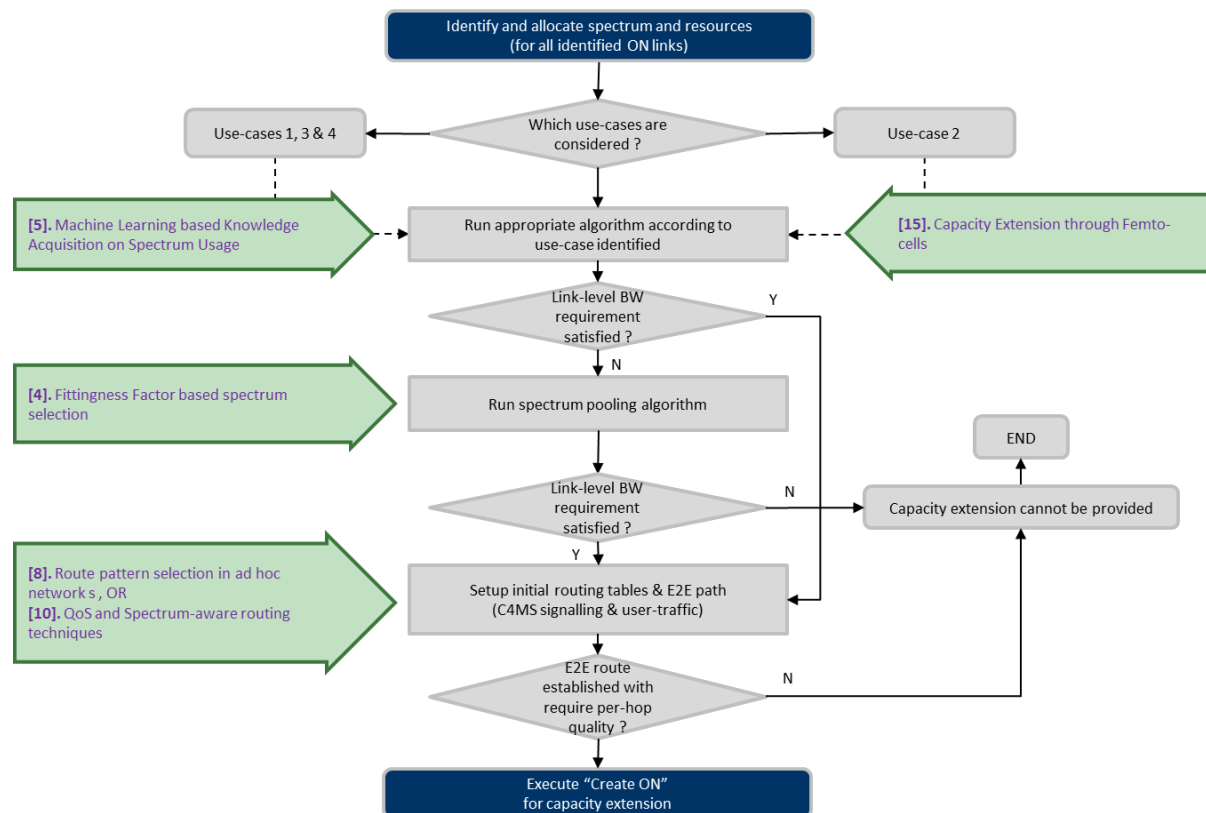


Figure 177: Scenario 2: Flow diagram for ON Creation phase

Such procedures, oriented to select the proper spectrum pool, to establish the initial routing tables or to choose the supporting femto node set, are executed by triggering the algorithms described in sections 2.4, 2.5, 2.8, 2.10 and 2.14.

4.2.3 ON Maintenance

The maintenance phase runs throughout the lifespan of the Opportunistic Network, and mainly consists of monitoring its performance and executing corrective actions when needed.

The ON performance is monitored by the estimation of the QoS on the newly established links to verify that profile requirements are fulfilled, as well as by the evaluation of the resource allocation in the involved nodes, to avoid further congestion situations.

Besides, the performance of the original (and now out-of-the-ON) nodes should also be monitored in order to verify if the situation that triggered the ON creation improves, worsens or no longer exists.

The following flow chart shows how different situations can trigger different corrective actions in order to keep the ON fine-tuned:

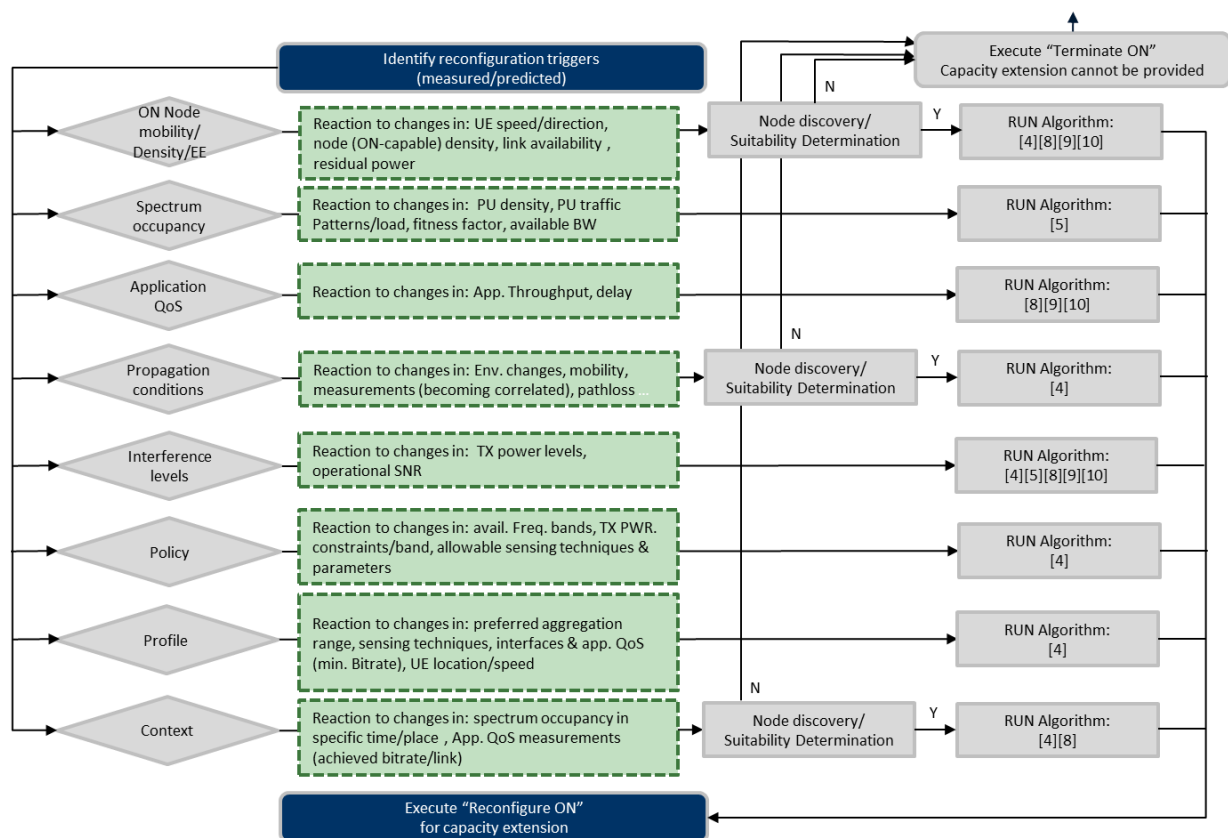


Figure 178: Scenario 2: Flow diagram for ON Maintenance phase

Therefore, corrective actions (i.e. modifications of the ON configuration and/or topology) should be executed when:

- Changes on the radio environment are detected. In these cases, new resource allocations or even a new band selection may be needed.
 - Changes on spectrum availability: may need the search of new bands, by executing of the algorithm described in 2.5
 - Changes on the propagation conditions: may need the search of new nodes and the execution of the algorithm described in 2.4
 - Change on the interference level: may need the search of new bands or new routes (as described in 2.4, 2.5, 2.8, 2.9 or 2.10)
- The topology of the ON changes (e.g., some of the participating UEs may move and become unreachable, or may decide leave the ON due to lack of battery life). In this case, new routes must be calculated and some resource allocation may be needed (using algorithms described in 2.4, 2.8, 2.9 or 2.10)
- There are changes on the context under which the ON was created, requiring the reallocation of resources and frequency bands.
 - Changes on the network policies (available bands, power, etc.) or on the devices' profiles (preferred bands, interfaces, etc.), may need new band selections (see 2.4)
 - Changes on the application context and requirements (e.g. target QoS, average bitrate, etc), may need to reallocate resources and routes (see 2.8, 2.9 or 2.10)
- New candidate nodes are detected (e.g. some new nodes or UEs appear close to the area of influence of the ON and estimations show that they may enhance the performance). In this

case, a negotiation procedure must be initiated to try to include the new nodes in the ON, followed by new resource allocations, re-routings and, perhaps, new frequency assignments.

- Congestion conditions in the original nodes out of the ON improve. In this case, some of the participating nodes and/or UEs may be extracted from the ON and diverted back to the infrastructure. To do so, some negotiations may be needed, along with resource/frequency allocations and re-routings.
- Congestion conditions in the original nodes out of the ON worsen. In this case, some additional nodes and/or UEs from outside the ON may need to join it in order to alleviate the congestion. To do so, some negotiations may be needed, along with resource/frequency allocations and re-routings.

The modifications of the ON are requested via C4MS messages ONMR and ONMA (see D3.3 Part B [6], section 1.6.), but if new negotiation or creation procedures are initiated, more complex signalling may follow.

4.2.4 ON Termination

The termination phase is initiated when, during the maintenance phase, the monitoring of the external nodes determines that the situation that triggered the ON no longer exists. This may happen when the resource occupation of the originally congested nodes falls below normal thresholds, or when the estimation of the QoS levels that previously-underachieving-users might reach outside the ON are acceptable.

Another situation that could trigger this phase may rise when some of the events listed in the previous subsection cannot be solved by the application of the stated corrective actions. In such case, the ON can no longer be maintained and termination is invoked:

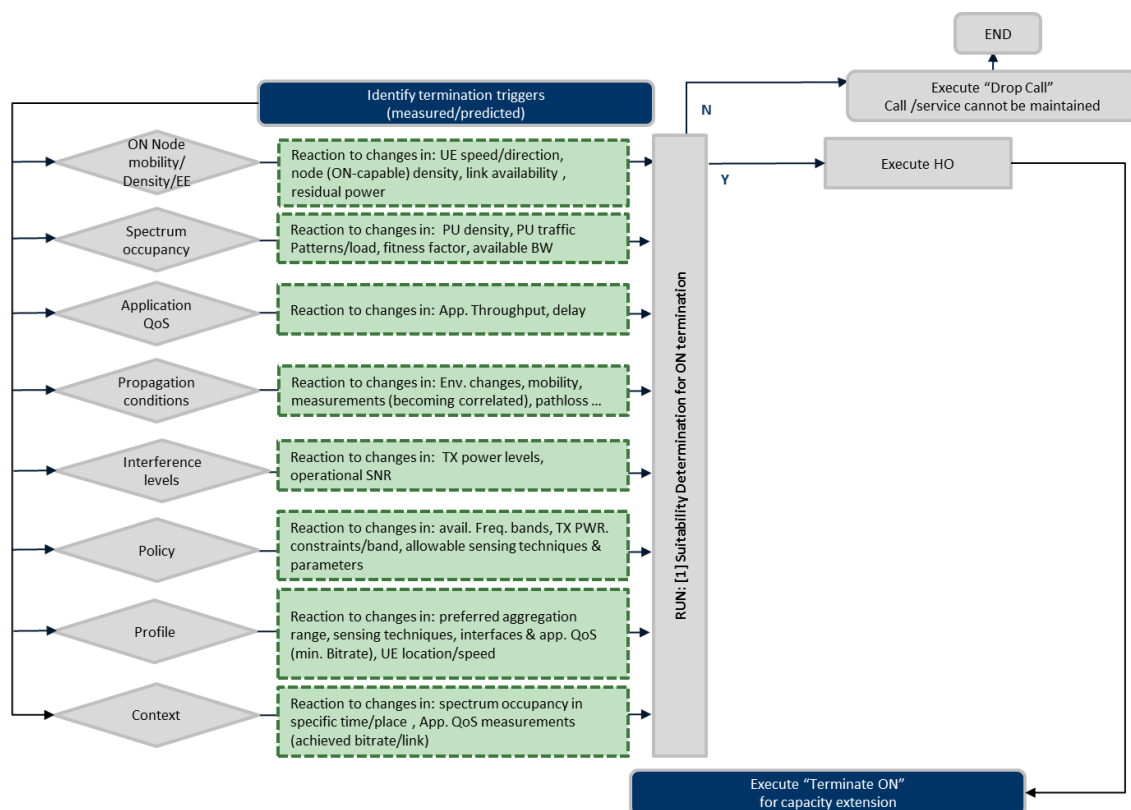


Figure 179: Scenario 2: Flow diagram for ON Termination phase

The termination procedure is initiated by means of C4MS signalling among the CMON entities of the involved nodes (messages ONRR and ONRA, see D3.3 PartB [6], section 1.7). This phase may need

the initiation of specific procedures to recover the old links (procedures such as femto policies modification, network attachments, intra- and inter-RAT handovers, etc.)

4.2.5 Performance evaluation

In this section indicative results are provided in order to evaluate the performance of the opportunistic capacity extension solution. Further results are presented at sections 2.7 and 2.14

4.2.5.1 Performance evaluation of resource allocation to femtocells with respect to spectrum selection

In this section, the impact of the fittingness factor-based spectrum selection mechanism on the energy-efficient resource allocation (ERA) to femtocells algorithm is studied.

A network is considered that comprises a macro BS with 9 femtocells uniformly deployed within its area that can be configured to operate at 6 discrete power levels (50-100% of their maximum transmission power). In addition, 40 users are also uniformly distributed. Two cases are considered. In the first case the available spectrum blocks are estimated to have high fittingness factor (i.e. greater than 0.7), while in the second case the available spectrum blocks are estimated to have low fittingness factor (i.e. lower than 0.5). Figure 180 illustrates the power allocation to the femtocells for these two cases. Due to the fact that in the first case the fittingness factor was high, the power level that was assigned was low since there was no interference in the communication (e.g. from the neighboring cells). On the other hand, when the fittingness factor was low, there was need to increase the power level of the small cells due to the increased interference. In general, the power consumption was decreased about 22% when the fittingness factor was high, in comparison with the low fittingness factor case.

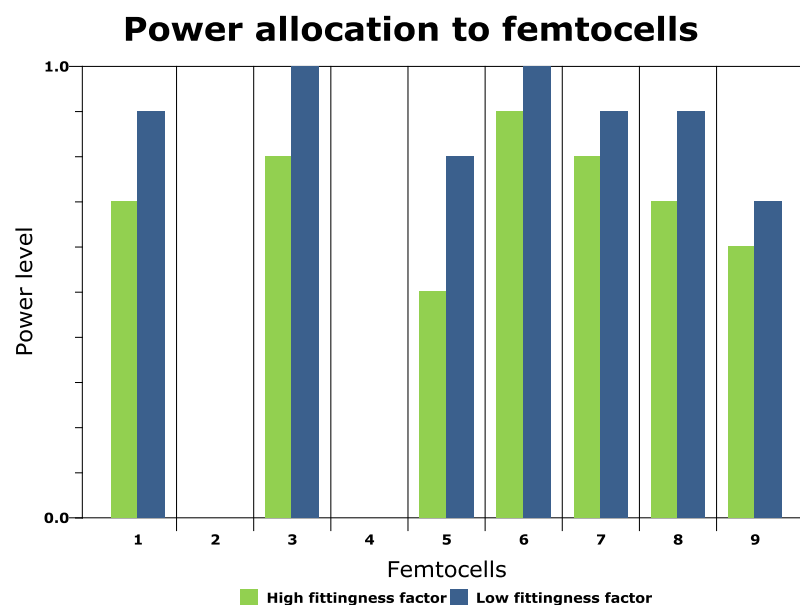


Figure 180: Power allocation to femtocells with respect to the fittingness factor of the available spectrum

4.3 Scenario 3 "Infrastructure supported D2D communication"

Scenario 3 "Infrastructure supported opportunistic device-to-device networking" shows the creation of an infrastructureless opportunistic network between two or more devices for the local exchange

of information (e.g. peer-to-peer communications, home networking, location-based service providing, public safety communication, etc.).

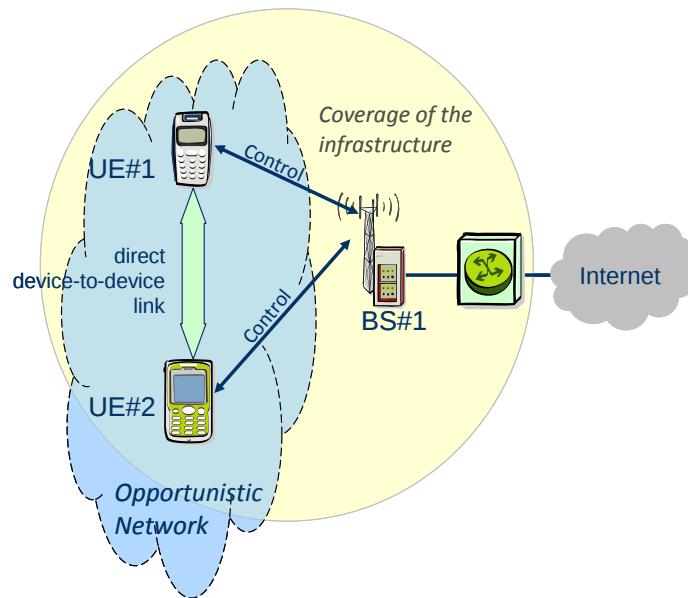


Figure 181: Infrastructure supported opportunistic device-to-device networking

This scenario is described in D2.1 [2] Section 4.3 and message sequence charts are provided for a network initiated scenario as well as for a terminal initiated scenario in D2.2.2 [4][4] Section 5.3 where message sequence charts. Further details on the information flows can be found in D3.2 [5] Section 5.4. Figure 182 gives an overview on which algorithm can be applied in which phase for scenario 3.

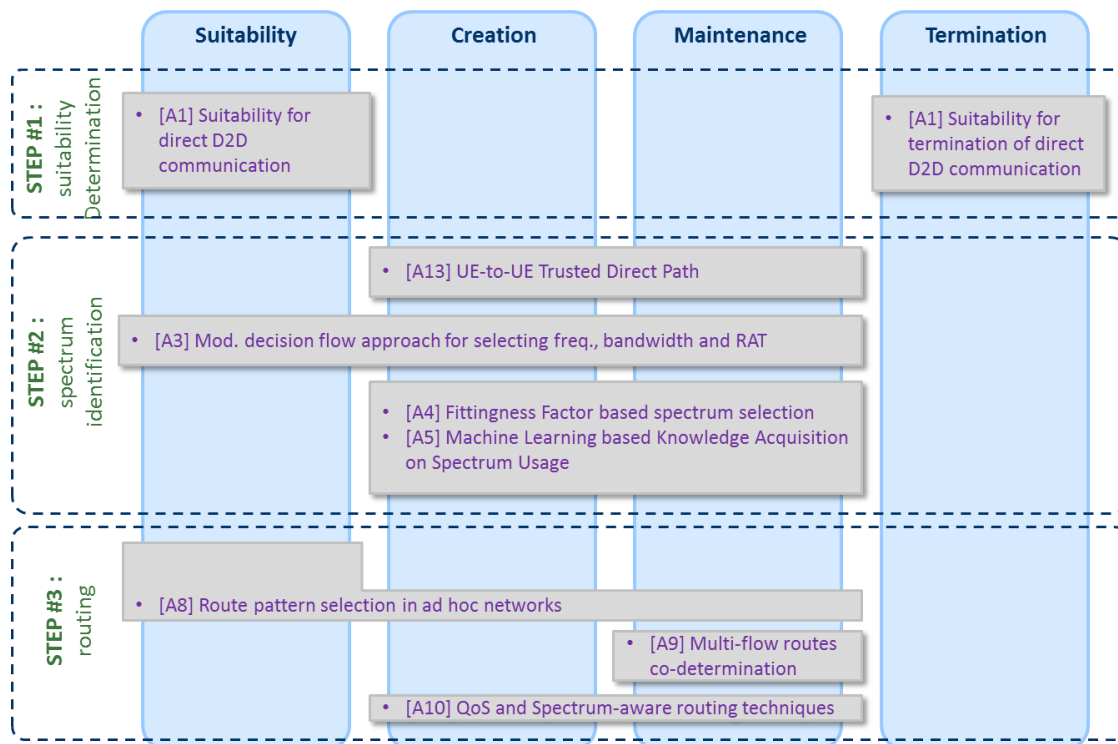


Figure 182: Mapping of the Algorithms to the different phases for Scenario 3

4.3.1 ON Suitability

For the infrastructure guided suitability determination, all involved devices must be registered to the network and context information including device capabilities must be provided to the network.

When a user is establishing a new session, the algorithm on “Suitability determination for direct device-to-device (D2D) communication” presented in section 2.1 describes how the network decides if a direct between the two devices is feasible or not.

Figure 183 shows the suitability determination procedure where after successful suitability determination the ON Creation takes place as described in the next sub-section.

4.3.2 ON Creation

During the ON creation, spectrum, bandwidth, RAT and routes must be selected. Some of these parameter may also already be determined during the suitability determination phase.

The algorithms [A3], [A4] and [A5] can be used for the selection of spectrum, bandwidth and RAT. Algorithm [A13] can be used to establish a UE-to-UE Trusted Direct Path.

In the case that several devices communicate and communication has to take place over several D2D hops, the route can be determined by the algorithms [A8], [A9] and [A10] as also shown in Figure 183.

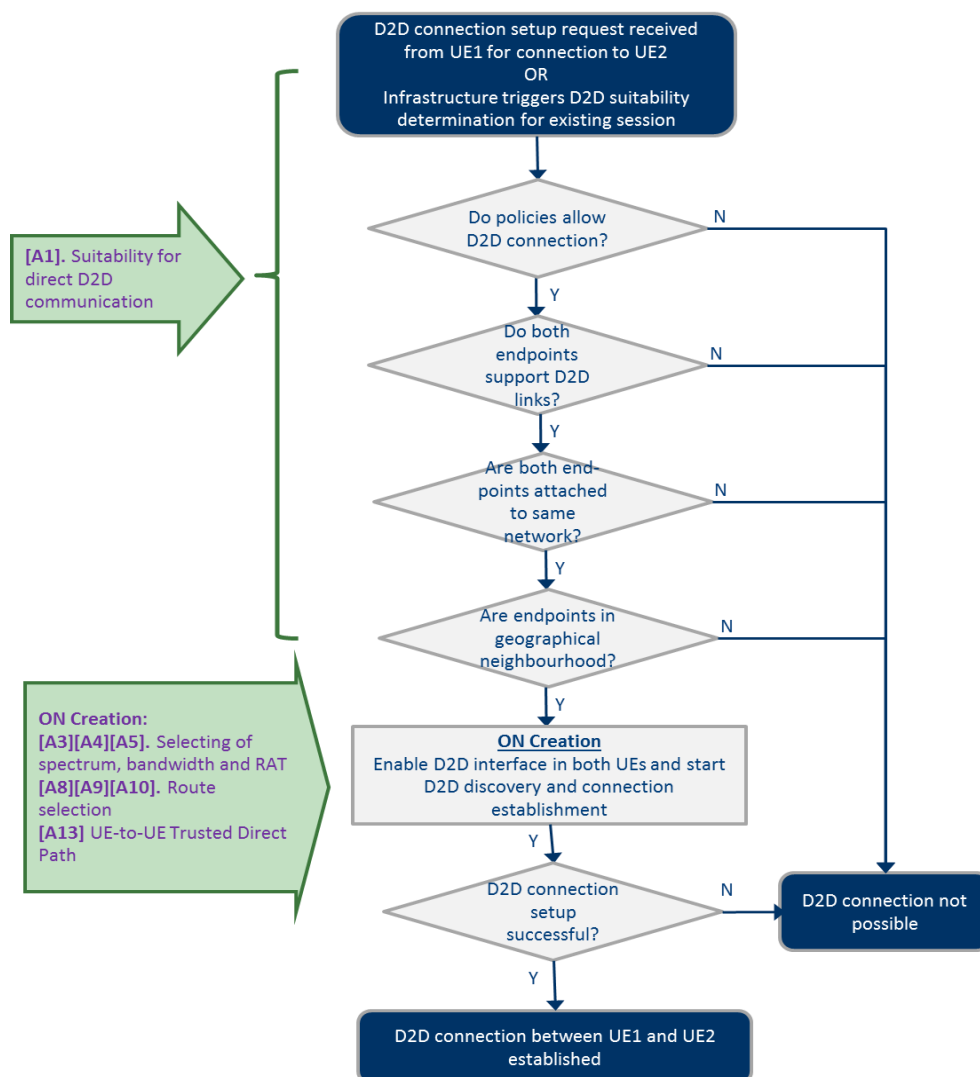


Figure 183: Scenario 3: Mapping of the Algorithms to the Suitability phase and to the Creation Phase

4.3.3 ON Maintenance

In the Maintenance Phase, the same algorithms as in the ON Creation phase can be used to modify the used spectrum, bandwidth, RAT or routing.

4.3.4 ON Termination

The ON can be terminated for different reasons. The standard case is that the ON will be terminated upon user request because the D2D communication session is no longer needed. Another case is that the ON can no longer be maintained because e.g. the devices are moving out of proximity of each other and the radio path will break. As described in section 2.1.2, a handover shall then be made to route the traffic via the infrastructure and the ON can then be terminated.

4.3.5 Performance evaluation

4.3.5.1 Performance evaluation of algorithm A5

In this section some performance evaluation results are presented on the performance of the spectrum selection strategy [A5] in this scenario, during ON creation and maintenance phases. Following the flow diagram of Figure 183, the evaluation considers the ON creation, ON maintenance and ON termination stages for establishing, maintaining and releasing the different ON links, assuming that the suitability determination has already been carried out. The algorithm is executed in the creation and maintenance stages.

More specifically, results consider the fitness factor-based spectrum selection algorithm described in section 2.4 in the Digital Home environment simulation framework considered in section 3.1.2.6. In this scenario, the infrastructure controls the spectrum allocated to the different links in an ON composed by different devices that are interconnected in a DH environment. The considered indoor DH (Digital Home) environment consists in a single floor of dimensions 16.8m x 30.4m organized in six different rooms. The simulation assumptions are essentially the same as in section 3.1.2.6, considering two different ON links having access to three different spectrum pools located in the ISM band (pool 1), in the TVWS band (pool 2) and in the mobile network operator band (pool 3). The operator preference with respect to the use of pool 3 is lower than for the use of the other two pools. The positions of the link transmitters and receivers are indicated in Figure 184, where also the positions of the external interferers are presented. These interferers follow the patterns described in section 3.1.2.6.

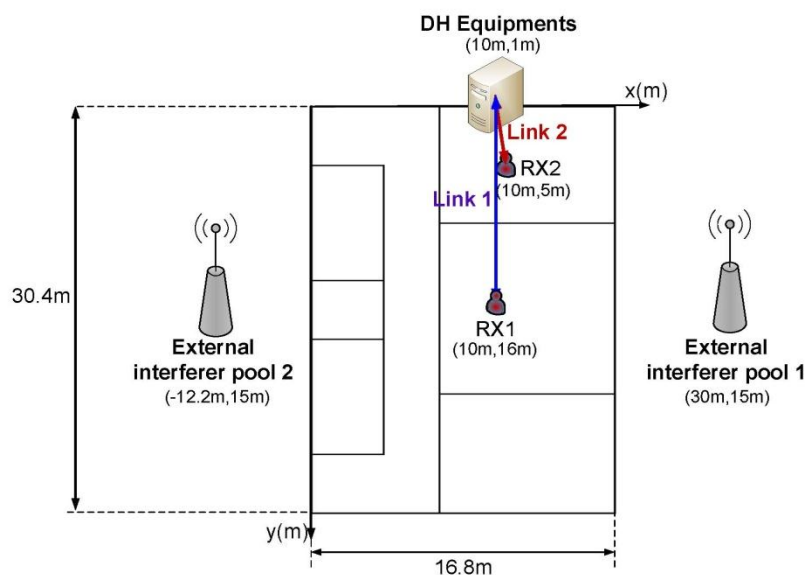


Figure 184: Location of the receivers for the two considered links

In order to assess the influence of the different components of the proposed framework, the following algorithms are considered as benchmark references:

- *Rand*: This implements only the spectrum selection at ON creation and performs a random selection among available pools.
- *Optim*: This scheme is an upper bound theoretical reference. In each simulation step, the procedure assigns the combinations of pools - active links that maximize the total instantaneous throughput at a given time instant k defined as:

$$\max \left(\sum_{l \in \text{Active_Links}(k)} \min(R_{\text{req},l}, R(l, p^*(l), k)) \right) \quad \text{s.t. maximizing} \quad \sum_{l \in \text{Active_Links}(k)} \lambda_{l, p^*(l)} \quad (59)$$

where $\text{Active_Links}(k)$ and $R(l, p^*(l), k)$ respectively denote the list of active links and the measured bit-rate $R(l, p^*(l))$ at time k .

The considered KPIs are those explained in section 3.1.2.2, namely the global efficiency level, the dissatisfaction probability, the relative regret when using the less preferred pool #3 (defined as the regret weighted by usage of pool #3, that is $\text{Usage}(2,3) \times \text{Regret}(2,3)$), and the spectrum HO rate during ON maintenance.

Figure 185 plots the global efficiency (Global_Eff) as a function of the total offered traffic load. For a better analysis of the impact of the different factors influencing Global_Eff , Figure 186 and Figure 187 plot, respectively, the dissatisfaction level of link #2 ($\text{Dissf}(2)$) and the relative regret level of link #2 in using pool #3. Results for link #1 are not presented since it is all the time satisfied (i.e., the bit-rate is always above the requirement of 20Mbps) and it never uses pool #3.

The results in Figure 185 show that the proposed strategy (denoted as SS+SM+RT) leads to a very important increase of the global efficiency compared to *Rand*. This is because, on the one hand, the estimated $F_{l,p}$ values by the KM together with the SM support help to assign the most suitable pools for link #2 sessions, which reduces the risk of being dissatisfied as reflected in Figure 186. On the other hand, the prioritization of pool #1 and #2 tends to avoid using pool #3 as much as possible which significantly improves $\text{Regret}(2,3)$ (see Figure 187).

Moreover, the proposed SS+SM+RT strategy performs very closely to the upper-bound optimal scheme in terms of global efficiency for most of the traffic loads, mainly thanks to the support of the KM and SM components. The small deviation observed for low traffic load is mainly due to the slight loss in $\text{Regret}(2,3)$ seen in Figure 187 motivated by the time spent before changes in interference conditions for unused pools can be captured in the fittingness factor values (note that updates in the fittingness factor values based on measurements are only performed when a pool is assigned, so for low traffic load conditions this can delay a bit these updates).

To further assess the proposed strategy in terms of the cost associated to the reconfigurations performed by the SM functionality, Figure 188 plots the average number of SpHOs per session experienced by SS+SM+RT for each of the two possible triggers of the spectrum mobility SM, namely the change in the $F_{l,p}$ value of the currently assigned pool or the release of another link. The results show that the total amount of SpHO signalling per session is below 0.06 for all considered traffic loads. The analysis of the fractions associated to each trigger reveals that, for all traffic loads, most of SpHOs are triggered by changes in $F_{l,p}$ values. This is because, on the one hand, in the considered DH environment there is at most one active session for each link, which makes not likely to experience a SpHO due to a link release. On the other hand, most of the average durations of the low- and high-interference states are comparable to the link session durations ($T_{\text{req},1}=2\text{min}$ and $T_{\text{req},2}=20\text{min}$) which makes very likely to experience a change in $F_{l,p}$ values during a given link session.

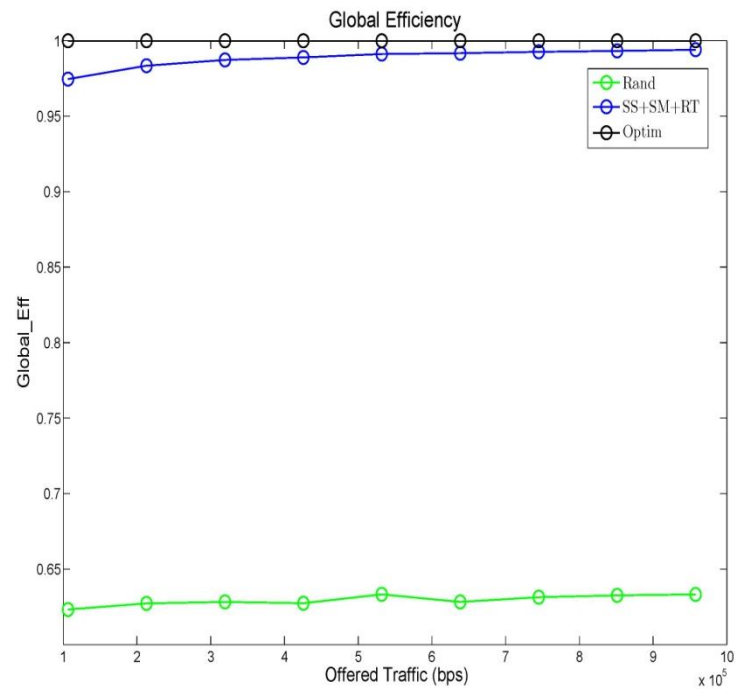


Figure 185: Global efficiency as a function of the offered traffic

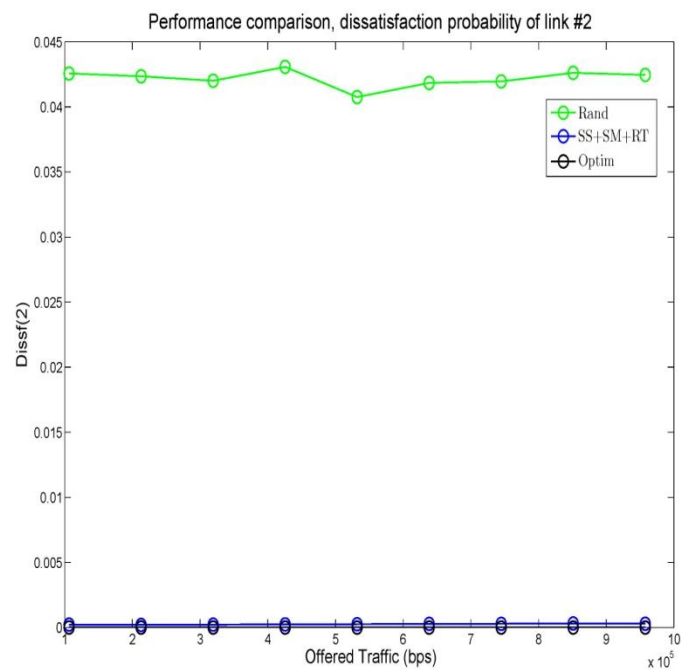


Figure 186: Dissatisfaction probability of link #2 as a function of the offered traffic

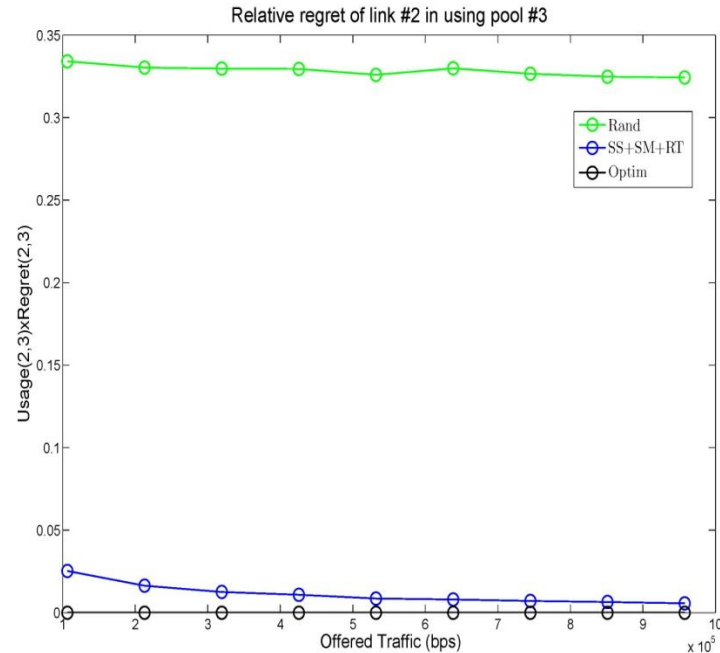


Figure 187: Relative regret of link#2 when using pool #3 as a function of the offered traffic

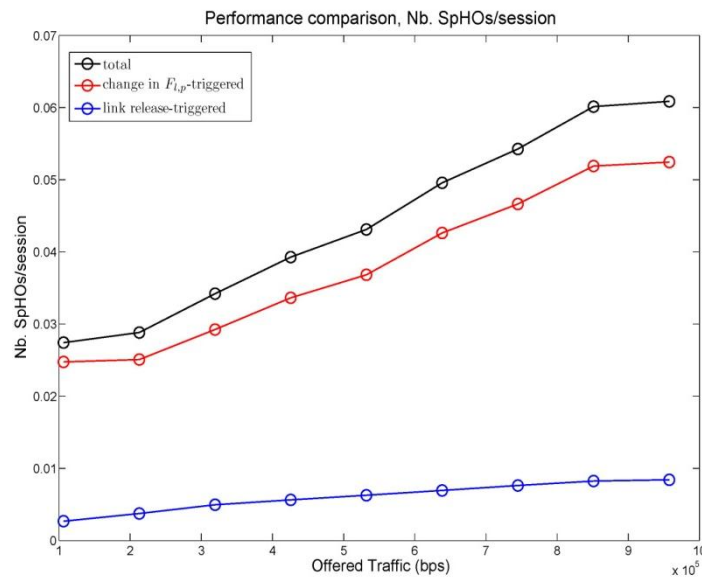


Figure 188: Spectrum HO rate

4.4 Scenario 5 “Opp. resource aggregation in the backhaul network”

The OneFIT Scenario 5 addresses opportunistic resource management in the wireless backhaul side of the network. A detailed description of the scenario and its use cases can be found in D2.1 [2].

Creation of wireless backhaul links between infrastructure elements requires additional radio interfaces for wireless stations (BSs, APs etc.). However, there are several examples of wireless networks where wireless stations communicate with each other over wireless backhaul links (i.e. WMNs and wireless bridge networks). Wireless mesh networks are selected for validation of the OneFIT scenario 5 since these networks have wireless backhaul links and mesh type topology which provides great possibilities for context informed resource management/aggregation/sharing. Also, these networks are identified as ones which can benefit the most from ON based resource management. Further, the OneFIT scenario 5 will be presented from the WMN point of view.

Context acquisition is important for enabling context informed decision making with respect to ON management. Since WMNs use TCP/IP protocol stack for networking, the SNMP protocol can be used for network monitoring and context acquisition. This protocol is supported by the majority of available equipment. The SNMP based C4MS variant ensures its usability in real WMN deployments without needs for modification on side of wireless equipment. Contextual information regarding users behaviour (the way in which users utilize WMNs and offered services) can be gathered by intercepting user's requests for network access and services (standard approaches for gathering user related contextual parameters). Gathered contextual parameters, if stored into a database, can form a database of contextual history which can be used for derivation/recognition of patterns in changes of certain contextual parameters and inter-dependencies between parameters. This knowledge can be used for detection/prediction of triggers for ON suitability determination (see section 2.12 of this document for detailed description of the trigger recognition functionality).

Two of the OneFIT algorithms (sections 2.12 and 2.14) address all ON management phases for Scenario 5 and are specifically developed for use cases of this scenario. These algorithms address nodes and routes selection/identification related technical challenges. Challenges of spectrum selection are covered by the algorithm presented in section 2.4 Spectrum identification (suitability determination phase) is not covered by any of the algorithms, however this feature is not required since channels which can be used for backhaul communication are well known within used RAT (i.e. widely used 802.11a for backhaul communication in 802.11 based WMNs). Mapping of these algorithms onto OneFIT technical challenges and ON management phases is shown in Figure 189. Sequence of execution of these algorithms throughout the ON management phases is shown in Figure 190.

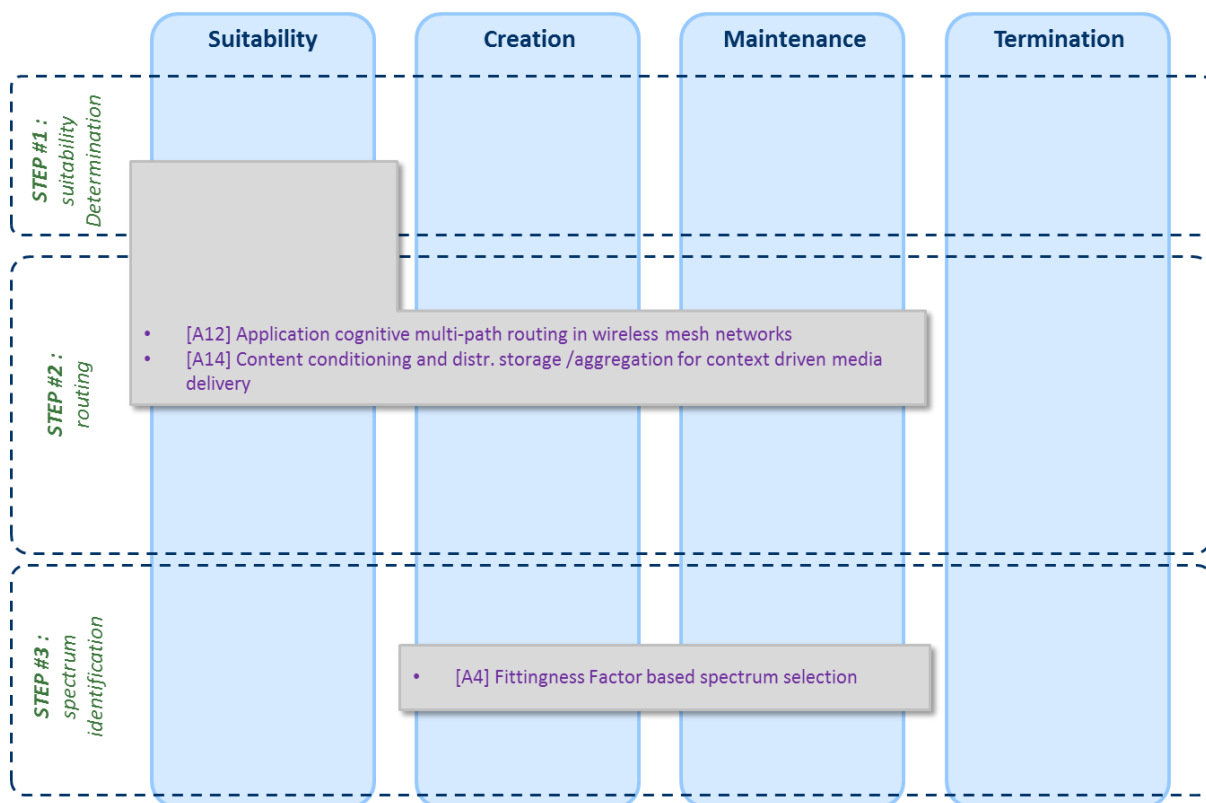


Figure 189 – Mapping of the OneFIT algorithms onto scenario 5 challenges

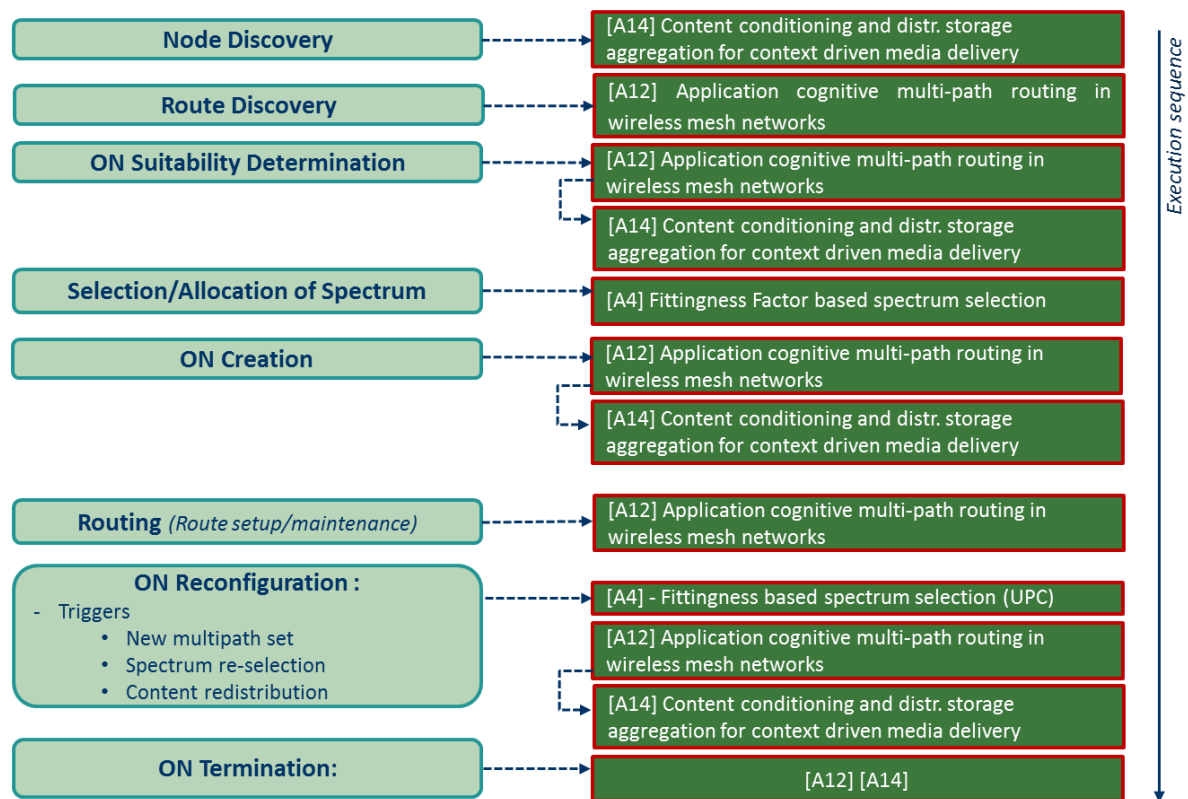


Figure 190 – Example algorithm execution sequence for the OneFIT scenario 5

4.4.1 ON Suitability

After triggers are recognized (process described in the section 2.12 of this document), the ON suitability determination starts. Trigger class is used for selection of proper algorithmic solution or their combination as shown in Figure 191. Two challenges from perspective of opportunistic backhaul resource aggregation are envisioned. The first one is need for opportunistic aggregation of backhaul bandwidth resources. The second one is need for opportunistic backhaul storage resource aggregation for enabling creation of virtual storage pools on level of WMNs. The first challenge is addressed by the application cognitive multipath routing algorithm described in section 2.12, while the second challenge is addressed by the node selection algorithm described in section 2.14. In real applications of multimedia delivery over WMNs both of these challenges will appear at the same time and mentioned algorithms will need to cooperate. Once the trigger class is recognized and a proper algorithm execution sequence is selected, the suitability determination starts, as shown in Figure 192. Order of suitability determination execution will define order of the algorithms for ON creation phase. If in any step of the suitability check ON is found unsuitable for solving the challenge, the default situation handling procedures will be utilized.

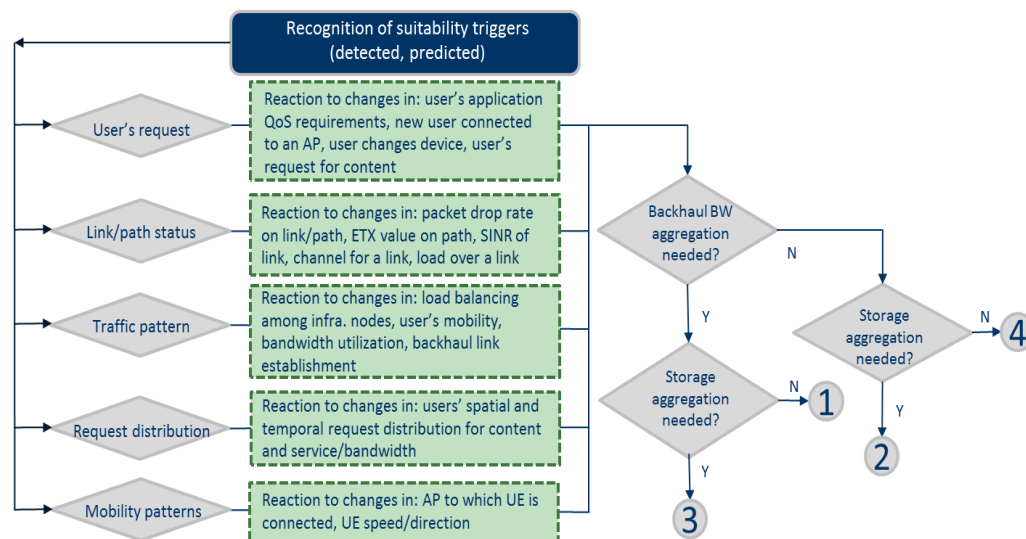


Figure 191 – Trigger detection and initial decision making

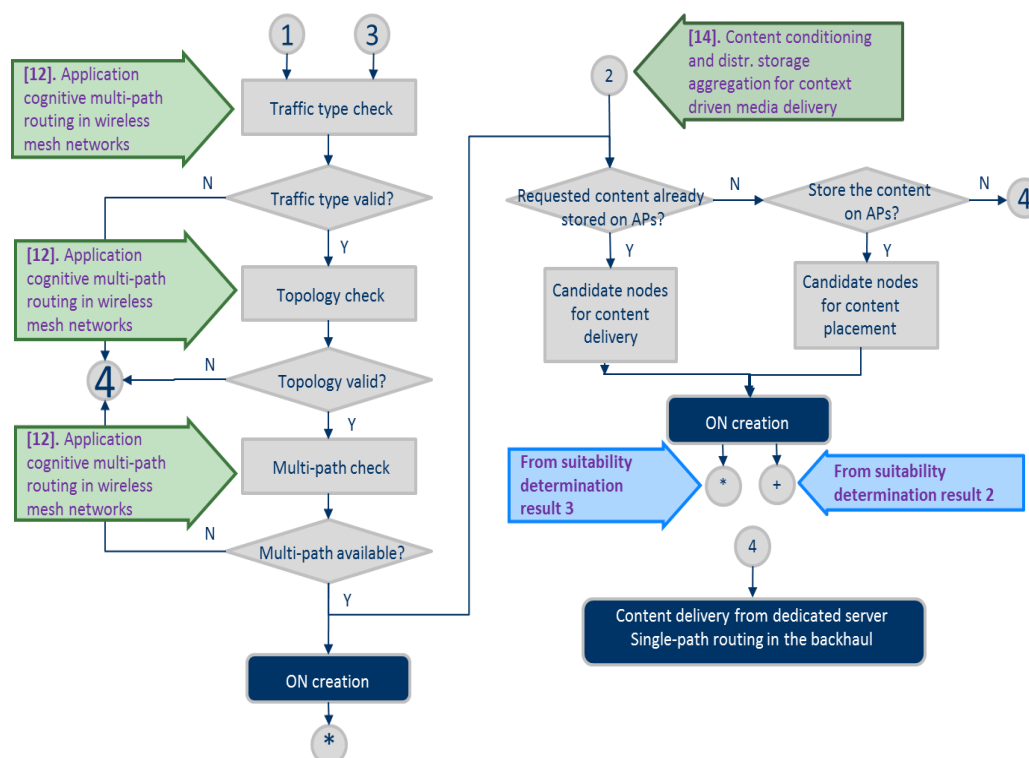


Figure 192 – Selection of nodes and routes in suitability determination phase for OneFIT scenario 5

4.4.2 ON Creation

The ON suitability determination phase provides a list of nodes and routes which can be used for establishment of ONs. During this phase, the algorithm described in 2.4 is used for the selection of proper communication channels (from precisely defined list of available channels) for establishment of backhaul paths/links. The goal of proper channel selection is to minimize interference generated by the ON creation and minimize the impact of other interference sources. The order in which algorithms depicted in Figure 189 are executed during the ON creation phase is presented in Figure 193. The ON creation phase has 3 control points where ON creation can be aborted if found necessary. In the case of ON creation abortion, proper handling of the situation is done to ensure service provision (i.e. using single path routing and/or content delivery from dedicated servers). When the ON is successfully created, the ON maintenance phase starts.

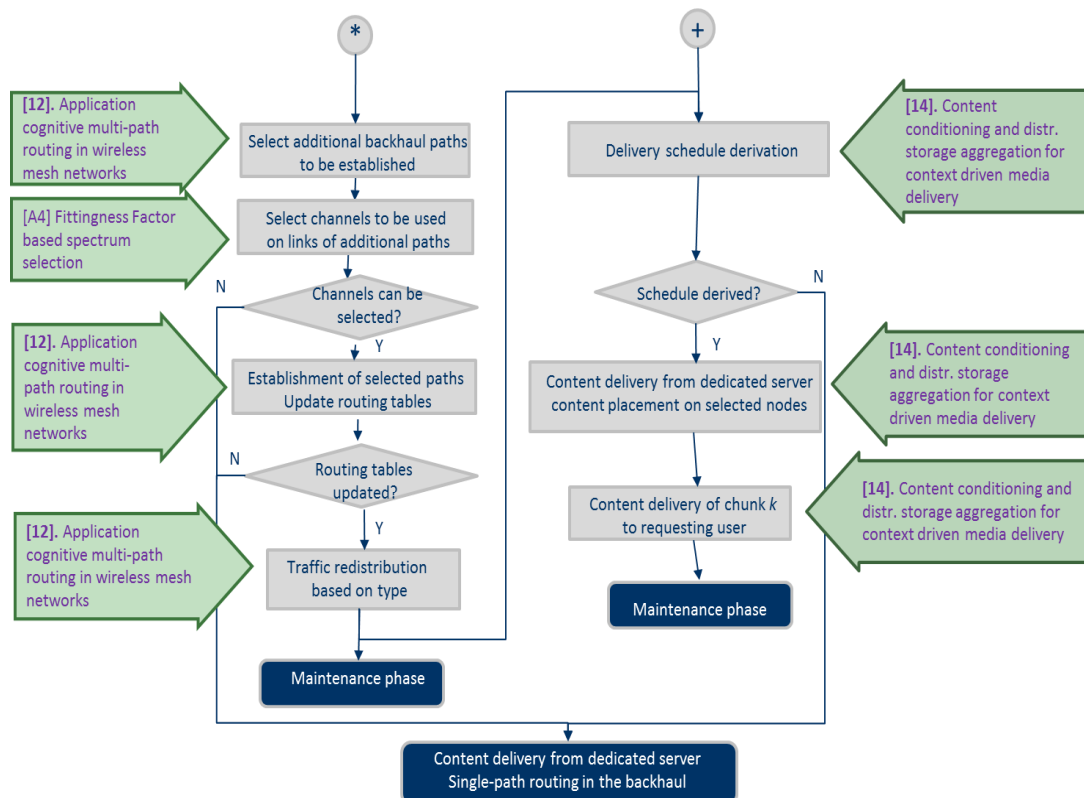


Figure 193 – Selection of nodes, routes and spectrum in ON creation phase of the OneFIT scenario 5

4.4.3 ON Maintenance

Created ONs are maintained through constant system monitoring (underlying WMN and established ONs) and trigger awareness. When reconfiguration triggers are met (see Figure 194), in a way described in the section 2.12 of this document, a new suitability determination phase starts with the goal of finding new ON solutions (modification of existing ON). Execution of the algorithms in the ON maintenance phase is shown in Figure 195.

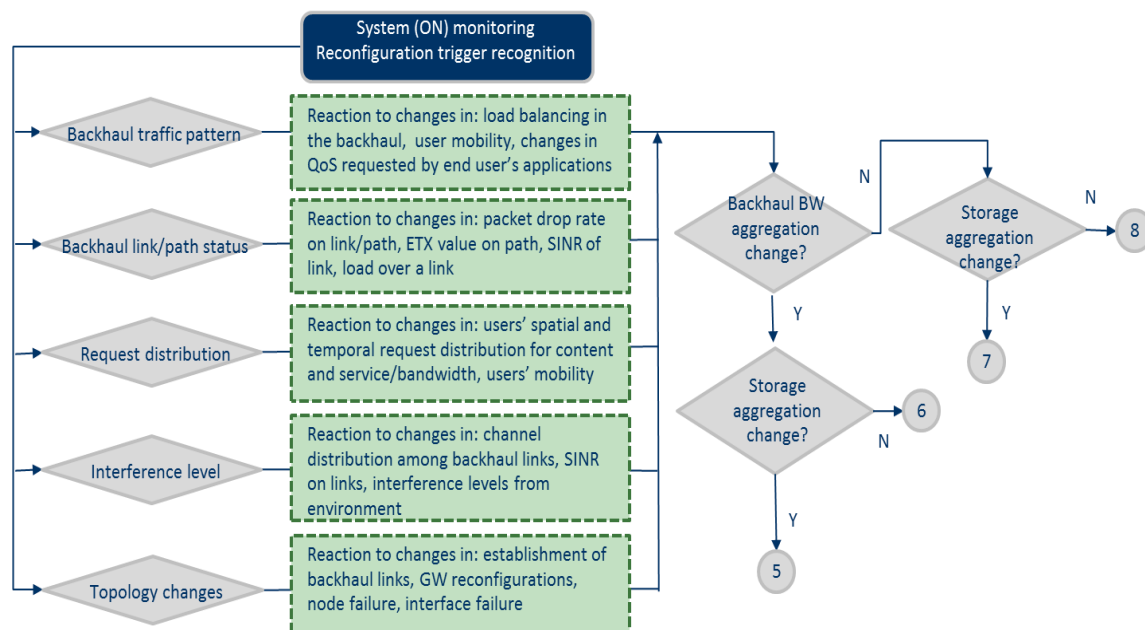


Figure 194 – Maintenance trigger recognition and decision making

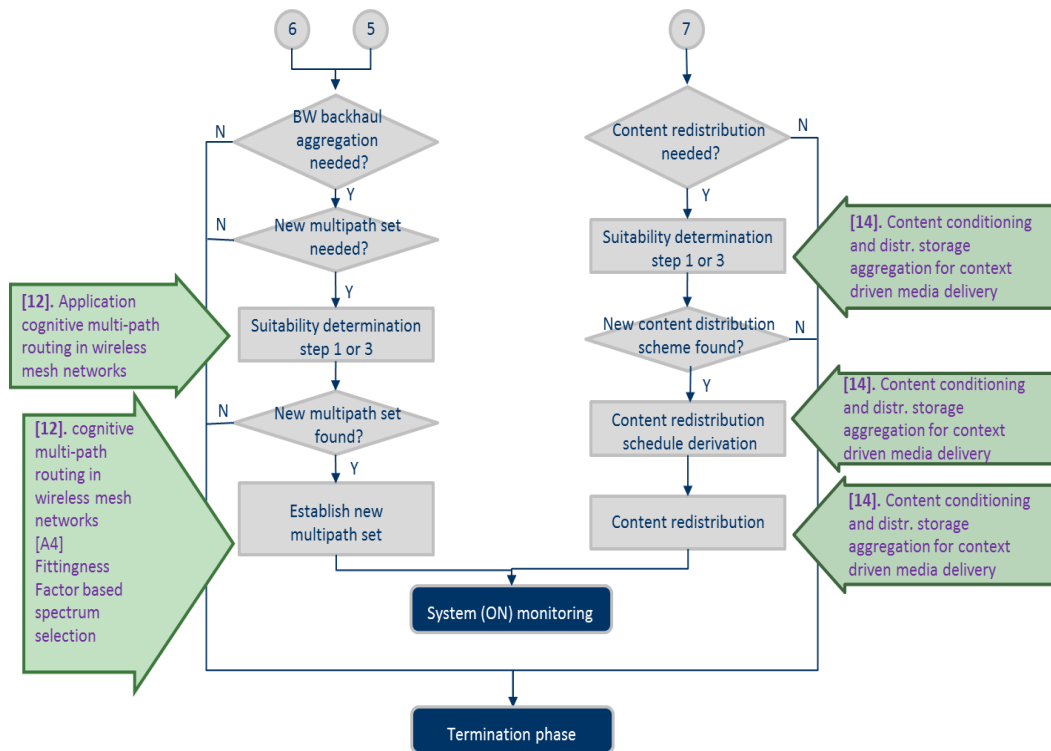


Figure 195 - Selection of nodes, routes and spectrum in ON maintenance phase of the OneFIT scenario 5

4.4.4 ON Termination

If the reconfiguration of the established ON is not possible (not justified) or the ON is no longer needed, it will be terminated and established services/applications will be properly handed-over to the infrastructure configured in line with legacy approaches (i.e. end users will receive requested content from the dedicated server; single path routing in the WMN backhaul). Recognition of ON termination triggers and sequence of termination execution are shown in Figure 196.

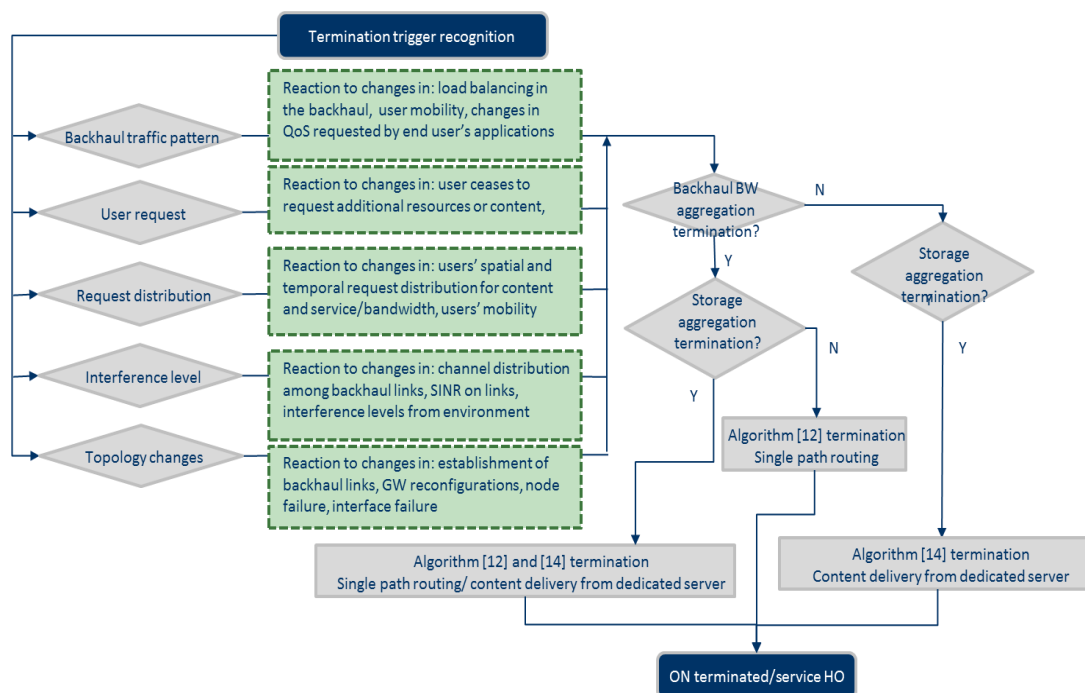


Figure 196 – Recognition of ON termination triggers and termination execution

4.4.5 Performance evaluation

Algorithms [12] and [14] are defined, developed and implemented with specific focus on node and route selection in the Scenario 5 related challenges. All performance evaluation and validation results related to these algorithms, which are presented in sections 2.12 and 2.14 of this document as well as in D4.2 [7] and D5.2 [12], can be considered as justification for algorithms' roles in the comprehensive algorithmic solution for the OneFIT scenario 5.

4.4.5.1 Performance evaluation of algorithm A4

In the following some performance evaluation results are presented to illustrate the performance of the spectrum selection step performed during the ON creation and the ON maintenance stages in this scenario.

Specifically, the mesh type topology considered for the evaluation is plotted in Figure 197, based on the platform presented in section 3.13 of deliverable D4.2 [7]. The scenario includes a number of backhaul links that are already established for communication between different APs and GWs. The capacity of these links is 30 Mb/s. The simulations considered here assume that during the ON suitability determination and ON creation phases the algorithm for application of cognitive multipath routing in wireless mesh networks makes the decision to establish the additional links #1 and #2 marked in red in Figure 197, so that the fittingness factor based spectrum selection is executed to decide on the most suitable spectrum pools, as reflected in the flow diagram of Figure 193. In addition, the algorithm is also executed during the ON maintenance in Figure 195 in case the spectrum assignment needs to be modified.

All the APs deliver their traffic to the L3 switch. The considered configuration assumes the total traffic for each of the APs as indicated in

Table 17. Note that for AP4 and AP5 it is assumed that the load is variable depending on traffic activity, considering two different load levels. Given that the capacities of the already established backhaul links is 30 Mb/s, whenever the traffic load in AP4 is the level 1 of 15 Mb/s, this can be supported with the established backhaul links following the route AP4->AP1->GW1->switch, so the additional link #2 does not need to be established. On the contrary, when the traffic load in AP4 reaches the level 2 of 20 Mb/s, the backhaul link AP1->GW1 becomes the bottleneck because it should support a total of 35 Mb/s (15 Mb/s from AP1 and 20 Mb/s from AP4) while its capacity is only 30 Mb/s. When this occurs, the multipath routing algorithm decides to establish the additional link #2 AP1->AP2 to be able to transmit the remaining 5 Mb/s following the route AP4->AP2->GW2->switch. Correspondingly, the spectrum selection algorithm will need to allocate a spectrum pool able to support 5 Mb/s for link #2.

Similarly, when the traffic load in AP5 is in the level 1 of 5 Mb/s, its traffic can be delivered through the route AP5->AP3->GW2->switch, so additional link #1 is not required. On the contrary, when the traffic load in AP5 is in the level 2 of 15 Mb/s, the link AP3->GW2 becomes the bottleneck since it should support a total of 40 Mb/s (15 Mb/s from AP5 and 25 Mb/s from AP3), while its capacity is limited to 30 Mb/s. In this case, the multipath routing algorithm decides to establish the additional link #1 between AP5 and AP2 so that the remaining 10 Mb/s will follow the route AP5->AP2->GW2->switch. Note that even when both links #1 and #2 have been established, the link AP2->GW2 can support the total traffic coming from these two additional links, so it does not become a bottleneck.

Based on all the above, Table 18 summarizes the capacity requirements for both additional links #1 and #2 depending on the traffic activity level in AP4 and AP5.

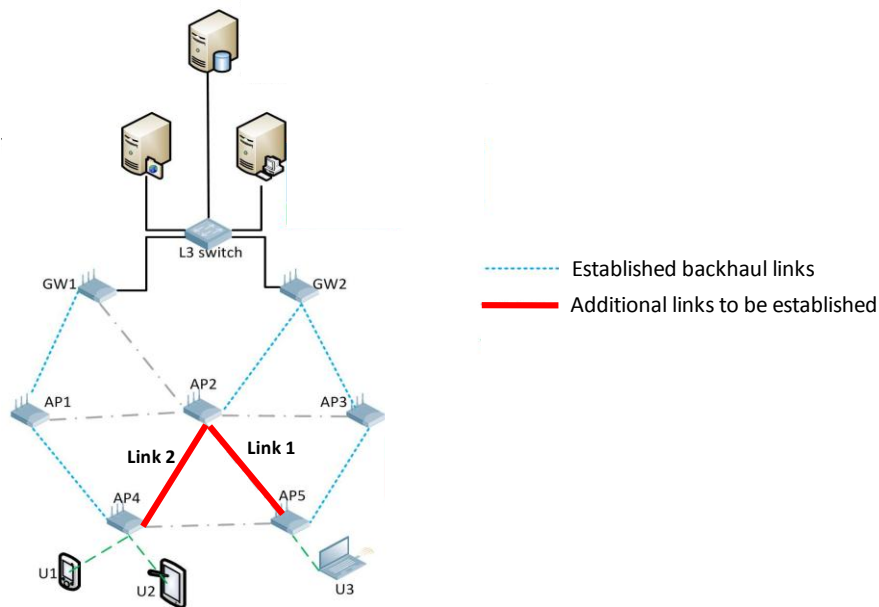


Figure 197: Backhaul links considered in the evaluation

Table 17: Input traffic (Mbit/s) for the different APs

AP1	AP2	AP3	AP4	AP5
15 Mb/s	15 Mb/s	25 Mb/s	Load 1: 15 Mb/s Load 2: 20 Mb/s	Load 1: 5 Mb/s Load 2: 15 Mb/s

Table 18: Required configurations for additional links #1 and #2

	AP4 – Load 1 (15 Mb/s)	AP4 – Load 2 (20 Mb/s)
AP5 – Load 1 (5 Mb/s)	Link 1: --- Link 2: ---	Link 1: --- Link 2: 5 Mb/s
AP5 – Load 2 (15 Mb/s)	Link 1: 10 Mb/s Link 2: ---	Link 1: 10 Mb/s Link 2: 5 Mb/s

The simulation assumes that for AP5 the length of traffic load 2 periods (i.e. the duration of link 1) is $T_{\text{req},1}=10$ min, and the traffic load 1 periods are exponentially distributed with average $1/\lambda_1$. Similarly, for AP4 the length of traffic load 1 periods (i.e. the duration of link 2) is $T_{\text{req},2}=30$ min, and the traffic load 2 periods are exponentially distributed with average $1/\lambda_2$.

The possible candidate spectrum pieces for establishing the backhaul links are the following $P=3$ candidate spectrum pools:

- Pool #1: $BW_1=20$ MHz bandwidth in the 5 GHz band;
- Pool #2: $BW_2=16$ MHz bandwidth in the 600 MHz TV White Space (TVWS) band that can be operated opportunistically;

- Pool #3: $BW_3=20$ MHz bandwidth in a 2.6 GHz licensed band of the Mobile Network Operator (MNO)

Both pools #1 and #2 experience an interference that alternates between two states, either low or high. Behavior is independent for both pools. The average duration of the low interference state is 7.5 h for pool #1 and 2.5 h for pool #2. Similarly, the average duration of the high interference state is 30 min for pool #1 and 10 min for pool #2. Pool #3 relies on interference control mechanisms existing in the operator licensed band and is always assumed to be in the low interference state.

Pools #1 and #2 are considered to have a higher preference factor than pool #3, that is $(\lambda_{i,1} = \lambda_{i,2} = 0.9; \lambda_{i,3} = 0.1), i \in \{1, 2\}$.

Simulation parameters assume the propagation model described in equation (52) of section 3.1.2.2 with a link length of $d=80$ m and traversing $N_w=2$ walls, for both links. In all cases the background noise power spectral density is $N_0=-164$ dBm/Hz. In pool #1 transmit power is 20 dBm, the low interference state is characterized by null external interference while the high interference state corresponds to an interference of $I_0=-156$ dBm/Hz. For pool #2 the transmit power is 10 dBm, the low interference state is characterized by an external interference of $I_0=-149.6$ dBm/Hz and the high interference state by $I_0=-143.7$ dBm/Hz. Finally, in pool #3 the transmit power is 14 dBm and there is no external interference. Based on these considerations, the average achievable bit rates in the different pools are given by Table 19.

Table 19: Average achievable bit rates (Mbit/s) in the different pools

	Pool #1		Pool #2		Pool #3
	State I_0	State I_1	State I_0	State I_1	State I_0
Link #1	40	10	15	5	40
Link #2	40	10	15	5	40

The proposed spectrum selection algorithm that includes both the knowledge management and the spectrum mobility support (SS+KM+SM) is compared against the following baseline strategies:

- Spectrum selection supported by Knowledge Manager (SS+KM): This is the proposed fitness factor-based spectrum selection supported by the KM block but without spectrum handover support.
- *Rand*: This implements only the spectrum selection at ON creation and performs a random selection among available pools.
- *Optim*: This scheme is an upper bound theoretical reference. In each simulation step, the procedure assigns the combinations of pools - active links that maximize the total instantaneous throughput at every time instant subject to keeping preference constraints (see section 4.3.5.1)

Figure 198 plots the global efficiency metric defined in section 3.1.2.2 that includes both the dissatisfaction probability and the regret. The horizontal axis presents the total offered traffic $\lambda_1 T_{req,1} R_{req,1} + \lambda_2 T_{req,2} R_{req,2}$. In turn, the dissatisfaction probability for link #1 is plot in Figure 199. The dissatisfaction probability for link #2 is not plotted because it is all the time satisfied regardless of the allocated pool. In addition, Figure 200 presents the relative regret of link #1 when using the less preferred pool #3. Finally, Figure 201 plots the spectrum handover rate per session with the proposed strategy SS+KM+SM, split between the two causes, namely link release or change in the fitness factor values.

Looking at Figure 198, it can be observed that the proposed strategy (SS+KM+SM) achieves a performance almost identical to the optimum approach. Also the strategy SS+KM without the support of the spectrum mobility achieves a very high efficiency although a bit smaller than SS+KM+SM. In all the cases, there is a significant improvement with respect to the random selection. The reason for this behaviour can be better analysed by looking at the dissatisfaction probability of link #1 shown in Figure 199. SS+KM exhibits a worse behaviour than SS+KM+SM because it may happen at some instants that, at link #1 establishment the pool #1 is already allocated to link #2, so pool #2 is allocated that is not able to provide the desired bit rate. On the contrary, with the support of SM, spectrum handovers will be executed to allocate link #1 to pool #2, thus contributing to reducing the dissatisfaction probability achieving a performance very close to the optimum case.

Focusing on the relative regret in Figure 200, the prioritization of pools #1 and #2 (SS+KM and SS+KM+SM) tends to avoid using pool #3 as much as possible which significantly improves the relative regret level of using it compared to Rand. In this case, SS+KM and SS+KM+SM are performing almost equally because $Usage(1,3)$ is very small for both (pools #1 and #2 are often enough to handle the low traffic load).

From Figure 201 it can be observed how the SpHO rate is quite small under the conditions considered in this scenario. Moreover, most of SpHOs are triggered by changes in $F_{i,p}$ values. The reason is that the average durations of the low- and high-interference states are comparable to the link session duration so it is very likely that a link experiences interference degradation in the allocated pool while being established. In particular, most of SpHOs are triggered to move a dissatisfied link #1 after being assigned pool #2, because (link #1 is always satisfied using the other pools. As traffic increases, it can be noticed how the probability of assigning pool #2 to link #1 decreases since it becomes most of the time used by link #2 at the time of establishing link#1 sessions, so link #1 usually has pool #1 allocated and correspondingly it is satisfied without requiring SpHO.

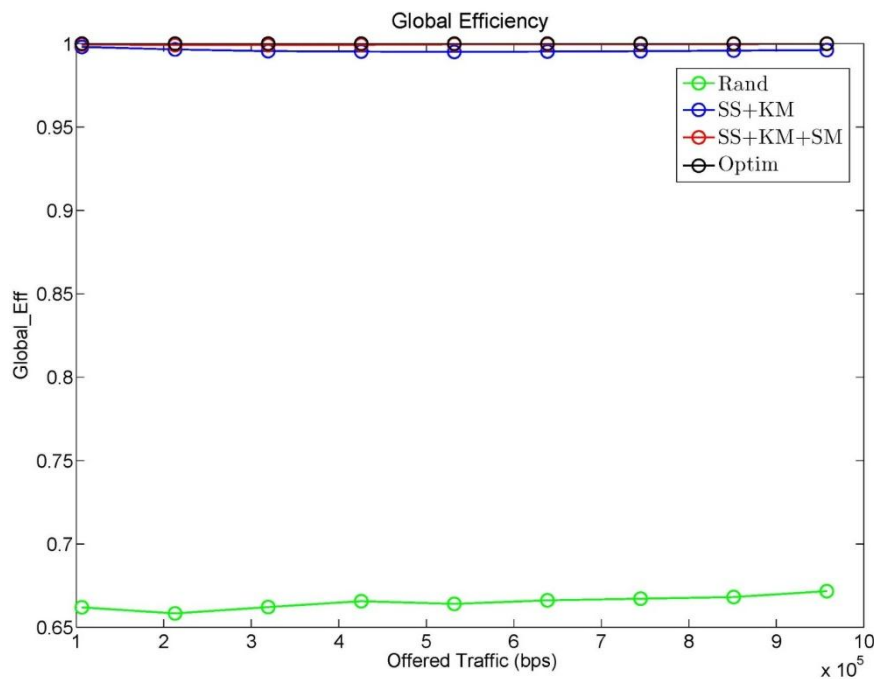


Figure 198: Global efficiency

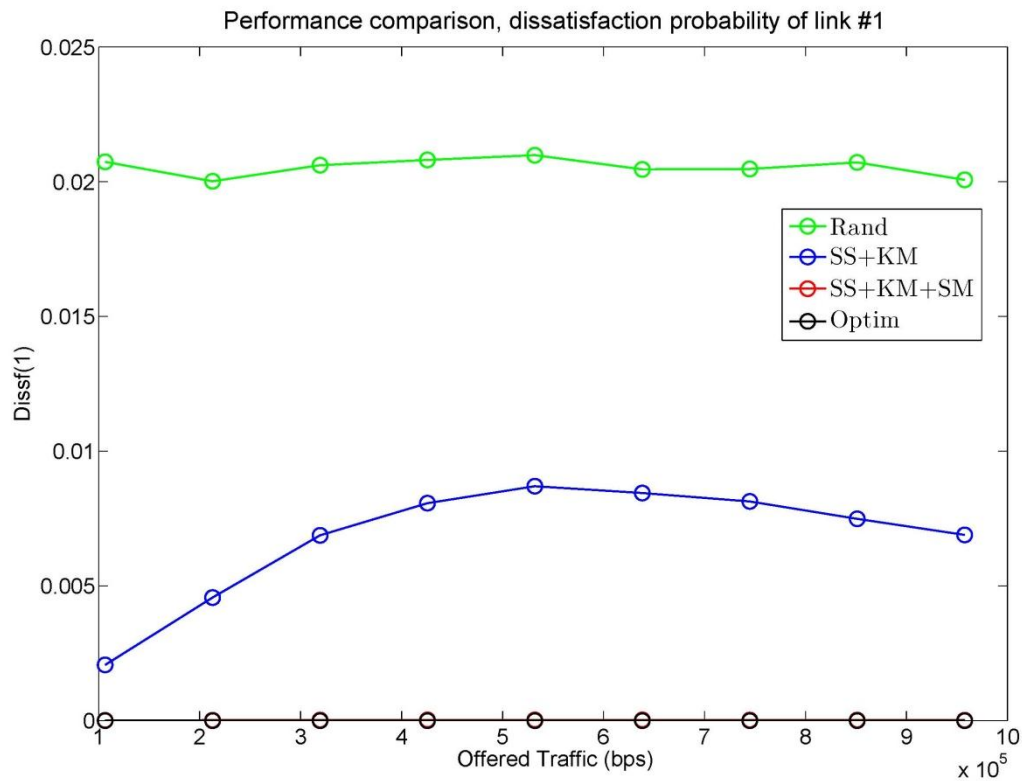


Figure 199: Dissatisfaction probability for link #1

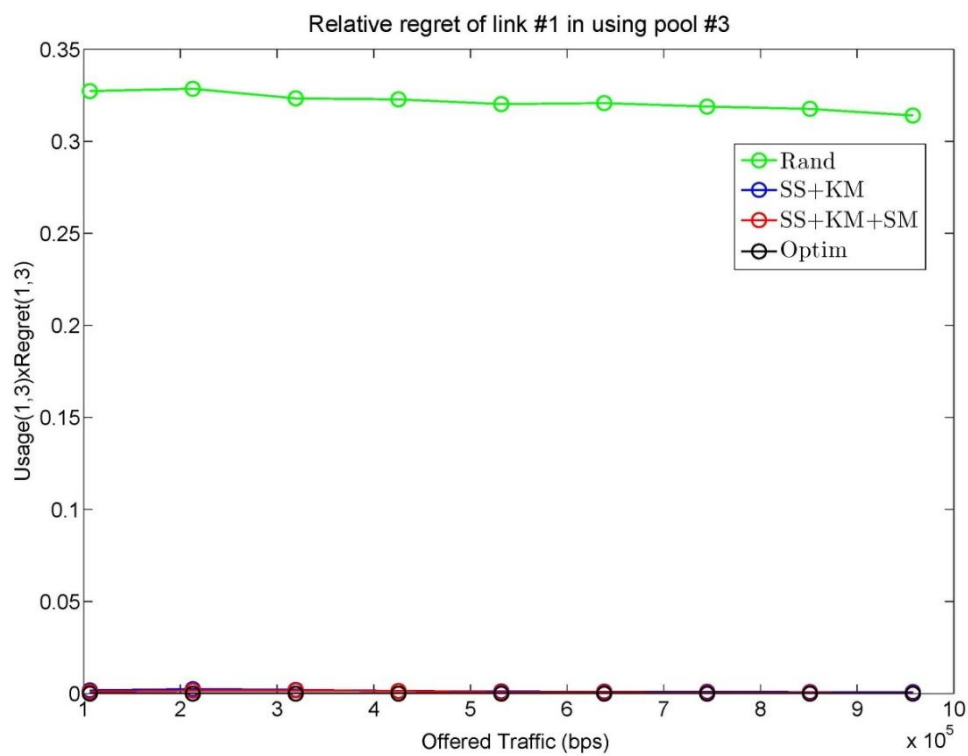


Figure 200: Relative regret in using pool #3 by link#1 (defined as $\text{Usage}(1,3) \times \text{Regret}(1,3)$)

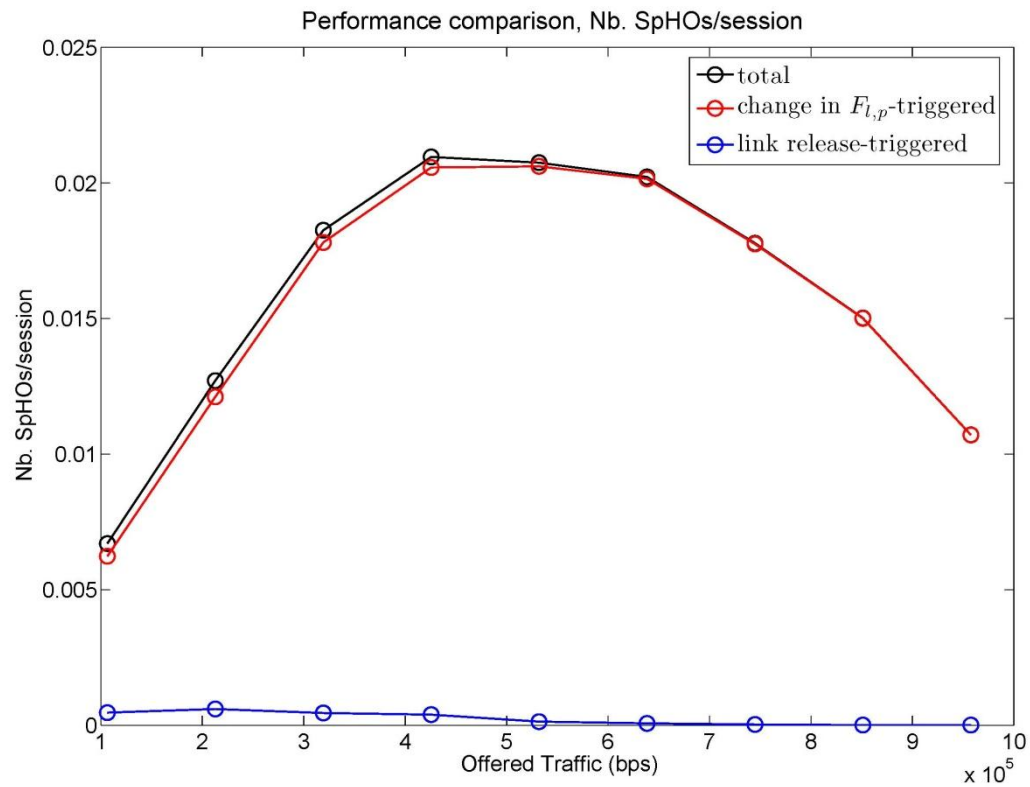


Figure 201: Rate of spectrum handovers per session.

5. Conclusions

This deliverable has presented the final set of the algorithmic solutions and results for enabling the management of opportunistic networks developed in the OneFIT project. Each algorithm has been specified, pointing out its integration within the OneFIT architecture and with the C4MS signalling, and it has been evaluated in representative scenarios. The proposed algorithms include the ON suitability determination in different scenarios, spectrum opportunity identification and selection, nodes and routes selection, and the capacity extension through femtocells. Also included, are aspects related to the practicality of the solutions proposed. Moreover, the last part of the deliverable has focused on the dynamic operation of the ONs, outlining the relationship between the developed algorithmic solutions in each of the considered scenarios. For that purpose, the algorithm execution sequence is provided for each of the scenarios.

In the following, a list of the major conclusions that have been obtained from the evaluations is provided.

- Suitability determination: This document has addressed first the ON suitability determination algorithm, focusing on the “infrastructure supported device-to-device communication” and “coverage extension” scenarios. Two algorithms have been proposed and evaluated. The following conclusions have been obtained:
 - In the case of the direct D2D communication, as it was aforementioned in D4.2 [7], it has been obtained that the probability for direct connectivity largely depends on the probability of the users being in the same “Area of Interest” (AOI) and if the range of the wireless interface is around the size of the AOI or larger. Further on, an ON can be maintained for a long time if the users are not moving, but the possible lifetime of an ON decreases with the velocity of the users.
 - In the opportunistic coverage extension scenario where a user can not directly access the infrastructure, the solution is to establish an opportunistic network with a so called “supporting device” in the neighbourhood, which then relays the traffic to the infrastructure. In D4.2 [7] the probability of finding such a supporting device in dependency of the range of the direct interface and in dependency of the supporting device density was presented.
- Spectrum opportunity identification and selection: One of the relevant technical challenges in the ON formation is the decision on which spectrum has to be assigned to the different links existing in an ON. In this context, this deliverable has considered the following:
 - Modular decision flow approach for selecting frequency, bandwidth and radio access technique for ONs: The proposed multi-objective approach considers joint selection of band, RAT, frequency and bandwidth for one or several ON links and ensures fair operation for whole ON and adequate quality of service for each user. The algorithm is employed when a new link needs to be established in the creation or maintenance phase or when a new spectrum band and RAT needs to be selected in the maintenance phase. The results show that for less data rate demanding applications such as web browsing 2.4 GHz band and 802.11a RAT are selected more often. Results also demonstrate that as the transmission range increases the selected band is IMT and RAT LTE or LTE-A for guaranteeing proper QoS for ON users. The algorithm makes decisions also taking into account users’ capabilities, velocity and application requirements. In the final design, a new learning-aided channel selection feature is introduced where selected band is divided into channels and reciprocity based reinforcement learning is used in channel selection.
 - Fittingness factor-based spectrum selection: The fittingness factor concept has been proposed as a novel metric that captures the time-varying suitability of available

spectrum resources to different applications supported in each ON link. This metric is used in the spectrum selection and spectrum mobility algorithms executed at a decision-making entity at ON creation and ON maintenance stages, respectively. They are supported by a Knowledge Management mechanism that involves a set of advanced statistics capturing the influence of the dynamic radio environment on the fittingness factor stored in a Knowledge Database. The performance evaluation results have shown that the proposed strategy efficiently exploits the knowledge management support and the spectrum mobility functionality to introduce significant gains (ranging from 85% to 100%) with respect to a random selection and to perform very closely to the upper-bound optimal scheme. The strategy has been evaluated from the perspective of its practicality in terms of the rate of required spectrum handovers, revealing a low rate of handovers particularly for low and medium loads, and from the perspective of the measurements exchange to support the context acquisition. In the final design, the algorithm is extended to cope with non-stationary environments, through inclusion of a novel Reliability Tester (RT) component, and, to support different operator preferences in the spectrum pools that can be assigned to each radio link, reflecting e.g. different operator policies or different regulatory constraints related to the band that each specific pool belongs to. The new set of results that rely on the spatial and temporal dimensions are captured, by analysing the robustness to non-stationary environments, indicate that proposed strategy does lead to substantial increases of the global efficiency under varying traffic loads.

- Machine Learning based Knowledge Acquisition on Spectrum Usage: This work has proposed a novel distributed available channel identification algorithm based on machine learning, targeted at modelling the spectrum occupancy and identification of spectrum pools for opportunistic access by secondary network. The secondary nodes do not need any information from the primary system or any neighbouring nodes. Simulation results reveal that the algorithm learns by means of successful interactions with the environment (rather than relying on spectrum usage history) which resources to be used in order to avoid/not cause harmful interference as well as providing a suitable quality of service (QoS). The extended model is demonstrated to solve the interference coordination problem in the downlink of cognitive femtos in heterogeneous networks, under restrictions of minimizing interference to both network tiers. It is shown that the proposed strategies can identify unused resource blocks with high degree of accuracy; co/cross-layer interference can be reduced significantly, thus resulting in average cell throughput increases of 30% and satisfaction probabilities increase of 47%, in the considered scenarios. Such intelligent schemes can provide practical solutions to the main challenge of interference in multi-layer networks.
- Techniques for Aggregation of Available Spectrum Bands/Fragments: This work has considered the spectrum selection with spectrum aggregation capabilities. When it is difficult to support enough bandwidth with contiguous available spectrum, multiple small spectrum fragments (sub-channels) can be exploited to yield a (virtual) single channel by spectrum aggregation. In the proposed algorithm, total throughput, the number of channel switching, and the number of bands for aggregate channels have been considered as the performance metrics with respect to the nature of secondary user's spectrum use and aggregation capability. The utility function for three performances is developed as a weighted sum utility function. A Q-learning approach is used to determine the optimal weights. The proposed approach is shown to be adaptable to changes amongst different objectives of interest. The proposed framework included a learning module allowing

for to management of complex interactions and trade-offs between different metrics without manual intervention.

- As a result of the integration and identification of synergies between the previous approaches, a comprehensive OneFIT solution for dynamic operation of ONs, for spectrum opportunity identification and selection has been presented. It is based on the interaction between a decision making entity and a knowledge management functional block which use the information captured from the radio environment, categorized in terms of context awareness, operator policies and user/application profiles. An evaluation of the different elements of the framework has been presented to consolidate the importance of the different specificities brought by the different solutions. In particular, the importance of knowledge management has been evaluated in a DH scenario under non-stationary conditions, revealing that the joint operation of knowledge management, reliability tester, spectrum selection and spectrum mobility mechanisms is able on the one hand to achieve a very important gain with respect to a random reference and with respect to strategies not including the KM or the SM components. On the other hand, the reliability tester is able to detect relevant changes in the environment and to properly regenerate the statistics stored in the knowledge database. The benefits of including operator preferences with respect to the desired spectrum bands have also been analyzed, revealing that the proper inclusion of these constraints in the spectrum selection is able to avoid the use of the less preferred spectrum pools unless this becomes strictly necessary. The benefits brought by learning-based models and machine-learning techniques in the spectrum opportunity identification and in the spectrum selection have also been pointed out. Finally, the relevance of the joint consideration of the RAT and spectrum selection has been analysed.
- Node and route selection: Another important technical challenge in the ON formation is the capability to select the adequate nodes to form the ON and the most suitable routes between them. This technical challenge has been addressed in this deliverable through different approaches:
 - Algorithm on knowledge-based suitability determination and selection of nodes and routes: This work has considered two main approaches for the ON creation. In the capacity extension scenario the proposed solution for route selection, based on the Ford-Fulkerson maximum flow algorithm, is triggered whenever a congestion situation occurs in the infrastructure and makes decisions on the establishment of routes which will redirect traffic from the congested service area into non-congested ones, assuming that each terminal in the congested service area creates one application flow which can be redirected through neighbouring terminals. In turn, the node selection process during the ON creation is based on a fitness function which is a weighted, linear formula which takes into consideration different parameters of the candidate nodes. The obtained results indicate that terminals which act as intermediate nodes may experience an increase in their transmission power of 19% in average, compared to the situation when the solution is not applied. Extended evaluation results based on new use-cases show that the transmission power of congested base stations is reduced due to ON formation , resulting in less interference whilst that of non-congested base stations tends to remain unchanged.
 - Route pattern selection in ad-hoc networks: This work has proposed an algorithm for routing packets taking into account the constraints associated to the different user application classes and to control and signalling information including C4MS

messages. Qualitative evaluations of the algorithm have revealed the ability of the algorithm to satisfy the end user QoE.

- Multi-flow routes co-determination: This work has presented a first proposal of network coding for multi-flow route co-determination. Optimization on this initial algorithm follows to extend the set of topologies on which the multi-flow route co-determination is applicable, with the introduction of delegated nodes. Results on an implementation of these optimizations have been presented. They show good improvements on the throughputs in terms of number of data packets to be exchanged between pair nodes, thanks to the integration of the delegated nodes.
- Techniques for Network Reconfiguration – topology Design: This work has addressed the creation and maintenance of network topology through enabling coordination in decision making among the nodes taking into account different parameters to establish links/topology with a set of desired global properties and constraints (K connectivity, interference minimization and energy efficiency). Based on adoption of PSO heuristic and column-generation techniques, the enhanced methodology provides near optimal solutions (better approximation) for otherwise infeasible problems through exact methods. The results indicate that the proposed models can provide practical yet optimal power levels to minimize the power utilization while maintaining K connectivity.
- Application cognitive multi-path routing in wireless mesh networks: An algorithm has been proposed to select and establish the appropriate set of multiple paths in the wireless backhaul side of a wireless mesh network. By using multiple paths, we can achieve better utilization of the available backhaul resources and increase the overall capacity of a given WMN. Aggregated backhaul bandwidth, made available to a particular AP, can increase access capacity of that AP several times when compared to the capacity achieved through the best available single path. Compared to legacy multipath approaches, the proposed algorithm provides the same level of backhaul bandwidth aggregation with drastically increased QoS levels. Results of experiments conducted in the open platform WMN test-bed show that jitter levels tend to be lower and more stable in comparison with single path routing. Furthermore, performance evaluations also show that the algorithm provides better load balancing and BW resource utilization than the QoS aware single path routing solutions while maintaining or increasing QoS level offered to the end users. In addition the autonomic trigger recognition feature of the algorithm is presented together with proof of concept results.
- UE-to-UE trusted direct path: In this work the problem addressed is the establishment of a WLAN network between different devices using the connections allowed by another technology. More specifically, it deals with the selection of the candidate AP and the operating channel to fulfil the QoS and power minimization requirements for all the WLAN clients. The extended formulations propose a QoS-constrained, novel joint AP and channel allocation technique for WLAN configuration.
- Content conditioning and distributed storage virtualization/aggregation for context driven media delivery: This work has addressed specifically the scenario 5 “Opportunistic resource aggregation in the backhaul network” and in particular the aggregation of backhaul storage resources in wireless mesh networks. The algorithm provides context aware, knowledge based and opportunistic aggregation of backhaul storage resources in WMNs. Based on contextual data gathered from the environment and end users, it performs the node selection for multimedia content

placement and distribution. Criteria for node selection is based on request distribution, popularity of multimedia content, status of caching storage of WMN nodes and user's behaviour. The evaluation of the proposed algorithm has revealed the advantage of cache-p2p organization of WMN CDN, the robustness of the proposed MILP model with respect to different user mobility models, and that the content provider can save costs with the same QoS. Using smart caching/content placement algorithms, observed features can be further improved (for an example, if local groups of nearby APs contain complementary content, whose union is almost (or to some extent) the copy of source server's content). Shown results are collected from "regular" scenarios, but using the presented MILP model one can examine different experimental setups, modelling and detecting "corner cases" and irregular situations in (parts of) WMN CDN. In addition, the GW selection variant of the algorithm, showing algorithm's ability to make proper decision making regarding selection of WMN nodes, is implemented in the open platform WMN test-bed and practically validated. These two algorithms provide complete solution for node and route identification and selection from perspective of the scenario 5.

- Algorithm on knowledge-based suitability determination and selection of nodes and routes: This work has considered two main approaches for the ON creation. In the capacity extension scenario the proposed solution for route selection, based on the Ford-Fulkerson maximum flow algorithm, is triggered whenever a congestion situation occurs in the infrastructure and makes decisions on the establishment of routes which will redirect traffic from the congested service area into non-congested ones, assuming that each terminal in the congested service area creates one application flow, which can be redirected through neighbouring terminals. In turn, the node selection process during the ON creation is based on a fitness function which is a weighted, linear formula that takes into consideration different parameters of the candidate nodes. The obtained results indicate that terminals, which act as intermediate nodes, may experience an increase in their transmission power of 19% in average, compared to the situation when the solution is not applied. Extended evaluation results based on new use-cases show that the transmission power of congested base stations is reduced within a range from 15% to 25% as more terminals switch to ON. Also, the quality of communication is benefited, as delay of successfully delivered messages drops approximately 15-35%, compared to the congested situation. Furthermore, an average decrease of 15-40% has been achieved in the load of the congested base stations.
- Capacity Extension through Femto-cells: This work has considered the capacity extension of the congested infrastructure by exploiting nearby femtocells. In particular, it has addressed the complex optimization problem of allocation of network resources to reroute macro-terminals to the femtocells and also to allocate the minimum possible power level to these femtocells. The problem is mathematically formulated and solved by means of a novel greedy algorithm denoted as DRA (Dynamic Resource Allocation). Results from testing the proposed algorithm reveal that the proposed algorithm outperforms both Simulated Annealing (SA) and Tabu Search (TS) reference algorithms in terms of solution quality (DRA computed a solution that was 4.7% better than the one of the SA and 2.27% better than the one of the TS) and runtime (DRA algorithm proved to be 36.7% faster than the SA and significantly faster than the TS). In addition, the benefits that derived from the use of femtocells at the energy consumption of a macro BS and the battery lifetime of terminals through 3 indicative test cases were studied. Simulation results depicted that the BS energy consumption is decreasing while the number of femtocells increases. Moreover, it was illustrated that the BS energy consumption increases while the number of terminals increases in the network. Furthermore, the battery lifetime of a macro-terminal, a femto-

terminal and a moving terminal was studied, obtaining that the battery lifetime of a femto-terminal is significantly higher than the one of the macro-terminal. Further work was carried out with the aim of way extending the capacity of the congested networks in an energy-efficient manner. The problem was formulated and solved by means of a novel greedy algorithm (Energy-Efficient Resource Allocation - ERA). The evaluations show that the utilization of femtocells in a congested area can reduce the average delay in packet delivery by around 15%, as well as increase the delivery probability.

- In the support activity to validate ON algorithms on an offloading-oriented real-deployment testbed: This work has described the implementation of the evaluation platform that will be used to assess the performance of a macro-to-femto offloading mechanism in urban areas. First tests on the platform show indicate that the offloading mechanism behaves properly, leading to a scenario with a 70% reduction in the disconnection rate, a 35% increase in the total capacity and a 10% enhancement in the load balance indicator. Further results based on a number of testcases, indicate that after the ON is created, the number of UEs that cannot connect to the macro layer is reduced by 68%, and the average per user DL throughput increases by 33%. The total traffic in the macro layer grows to 29% even when the number of users is 19% less. At the femto layer, the average throughput reduces slightly, as the limited resources have to be shared among existing and incoming UEs, but the total traffic supported at the femto layer is 34% higher.

6. References

- [1] ICT-2009-257385 OneFIT Project, <http://www.ict-onefit.eu/>
- [2] OneFIT Deliverable D2.1 "Business scenarios, technical challenges and system requirements", October 2010
- [3] OneFIT Deliverable D2.2 "Functional and System Architecture", February 2011
- [4] OneFIT Deliverable D2.2.2 "Functional and System Architecture Version 2", December 2012
- [5] OneFIT Deliverable D3.2 "Information definition and signalling flows", September 2011
- [6] OneFIT Deliverable D3.3 Part B "Detailed C4MS Protocol Specification", June 2012
- [7] OneFIT Deliverable D4.2 "Performance assessment & synergic operation of algorithmic solutions enabling opportunistic networks", June 2012
- [8] M. Cesana, F. Cuomo, E. Ekici, "Routing in cognitive radio networks: challenges and solutions," Elsevier Ad Hoc Networks Journal, 9 (3) (2011), pp. 228–248.
- [9] Won-Yeol Lee and Ian F. Akyildiz, "A spectrum Decision Framework for Cognitive Radio Networks," IEEE Transactions on Mobile computing, vol 10, no 2, Feb 2011.
- [10] H. Sarvanko, M. Mustonen, M. Matinmikko, M. Höyhty, J. Del Ser, "Spectrum Band and RAT Selection for Infrastructure Governed Opportunistic Networks," in Proc of the 17th International Workshop on Computer Aided Modeling and Design of Communication Links and Networks, CAMAD, Sept. 2012.
- [11] OneFIT Deliverable D4.1 "Formulation, implementation considerations, and first performance evaluation of algorithmic solutions", May 2011
- [12] OneFIT Deliverable D5.2 "Validation platform implementation description", June 2012
- [13] OneFIT Deliverable D5.3 "Results analysis and validation", December 2012
- [14] M. Logothetis, K. Tsagkaris, P. Demestichas, "Application and Mobility Aware Integration of Opportunistic Networks with Wireless Infrastructures", accepted for publication in Computers and Electrical Engineering, Elsevier, 2012
- [15] A. Georgakopoulos, K. Tsagkaris, V. Stavroulaki, P. Demestichas, "Efficient opportunistic network creation in the context of Future Internet", The Future Internet: Achievements and Technological Promises, vol.6656, pp.293-306, 2011, Springer/Verlag
- [16] A. Georgakopoulos, K. Tsagkaris, V. Stavroulaki, P. Demestichas, "Specification and assessment of a fitness function for the creation of opportunistic networks", in Proc. 20th Future Network and Mobile Summit (FNMS) 2011, Warsaw, Poland, 15-17.06.2011
- [17] A. Georgakopoulos, D. Karvounas, V. Stavroulaki, K. Tsagkaris, M. Tasic, D. Boskovic, P. Demestichas, "Scheme for Expanding the Capacity of Wireless Access Infrastructures through the Exploitation of Opportunistic Networks", Mobile Networks and Applications (MONET), vol.17, no.4, pp.463-478, 2012, Springer
- [18] A. Georgakopoulos, D. Karvounas, V. Stavroulaki, M. Tasic, D. Boskovic, J. Gebert, W. Koenig, P. Demestichas, "Cognitive cloud-oriented wireless networks for the Future Internet", in Proc. IEEE Wireless Communications and Networking Conference (WCNC) 2012, Paris, France, 01-04.04.201
- [19] M. Tasic, M. Cirilovic, O. Ikovic, D. Kesler, S. Dautovic, D. Boskovic, "Impact of different content placement and delivery approaches on streaming capacity of the wireless mesh networks", AdHoc Now 2012, Belgrade.
- [20] Z. Jako, G. Jeney, "Downlink Femtocell Interference in WCDMA Networks", Energy-Aware Communications, 17th International Workshop, EUNICE 2011, September 2011
- [21] 3GPP TR 22.803 "Feasibility Study for Proximity Services (ProSe)"

- [22] A. Keränen, J. Ott, K. Teemu, "The ONE Simulator for DTN Protocol Evaluation" in Proc. SIMUTools'09: 2nd International Conference on Simulation Tools and Techniques, 2009
- [23] Anandkumar, A.; Michael, N.; Tang, A. K. ; Swami, A., "Distributed algorithms for learning and cognitive medium access with logarithmic regret," IEEE Journal on Selected Areas in Communications, 2011.
- [24] E.L. Lehmann, "Testing statistical hypotheses: The story of a book". Statistical Science 12: 48, 1997
- [25] N. Schenker, J.F. Gentleman, "On judging the significance of differences by examining the overlap between confidence intervals". Am Stat. 2001;55; 182-186
- [26] J. Nasreddine, O. Sallent, J. Pérez-Romero, R. Agustí "Positioning-based Framework for Secondary Spectrum Usage", Physical Communication Journal (PhyCom), Elsevier, June 2008, Vol. 1, pp. 121-133
- [27] R. Fraile, J. Gozálvéz, O. Lázaro, J. Monserrat, N. Cardona "Effect of a Two Dimensional Shadowing Model on System Level Performance Evaluation", WPMC, 2004.
- [28] D. Boscovic, F. Vakil, S. Dautovic and M. Tosic, "Greening of Video Streaming to Mobile Devices by Pervasive Wireless CDN", Journal of Green Engineering Vol. 2, pp. 1-27, 2011, River Publisher.
- [29] D.J. Leith, P. Clifford, V. Badarla and D. Malone, "WLAN Channel Selection Without Communication", Technical report, Hamilton Institute, National University of Ireland, 2006.
- [30] 3GPP TS 23.203 v11.6.0, June 2012.
- [31] C. Heegard, J.T. Coffey, S. Gummadi, P. A. Murphy, R. Provencio, E. J. Rossin, S. Schrum, M. B. Shoemake, "High Performance Wireless Ethernet", IEEE Communications Magazine, vol. 39, Issue 11, pp. 64-73, Nov. 2001.
- [32] J. Mikulka, S. Hanus, "Complementary Code keying Implementation in the Wireless Networking", 6th EURASIP Conference on Speech and Image Processing, Multimedia Communications and Services, pp. 315-318, June 2007
- [33] P. Mahasukhon, M. Hempel, H. Sharif, T. Zhou, S. Ci, H. H. Chen, "BER Analysis of 802.11b Networks under Mobility", IEEE ICC, pp. 4722-4727, June 2007.
- [34] G. Alnwaimi, T. Zahir, S. Vahid, and K. Moessner, "Machine learning based knowledge acquisition on spectrum usage for lte femtocells," in *Conference on Communications (ICC), 2013 The IEEE International*, 'under review' 2013.
- [35] G. Alnwaimi, K. Arshad, and K. Moessner, "Dynamic Spectrum Allocation Algorithm with Interference Management in Co-Existing Networks," *Communications Letters, IEEE*, vol. 15, no. 9, pp. 932-934, 2011.
- [36] Michel Bourdellès, Nicolas Ménégale "Routing optimization for network coding" IEEE/ IFIP Wireless Days 2012, Dublin