

A Scheduling Algorithm Based Self-Learning Technique for Smart Grid Communications over 4G Networks

Ayman Hajjawi and Mahamod Ismail

Department of Electronic, Electrical and System Engineering, Faculty of Engineering and Built Environment, The National University of Malaysia (UKM), Bangi, Selangor, Malaysia
Email: ayman_2747@yahoo.com, mahanod@eng.ukm.my

Abstract—The latest 4G wireless technology, Long Term Evolution (LTE), is expected to be a promising wireless network for smart grid communications due to its high data rate and low latency. However, Scheduling mechanism has a crucial role on the smart grid performance since it is responsible for resources allocation process among users taking into account delay, fairness, throughput and channel status. LTE has not been intended for smart grid communications which results in unfair resource allocation between LTE users and smart grid applications. Moreover, the smart grid utilizes real time communications which have to be prioritized and allocated more resources in some cases. Following the requirements for a scheduling algorithm that takes into account smart grid communications and offers a compromise among different classes, this paper proposes a two level scheduling algorithm for smart grid communications in order to meet the requirements of these classes. In first level, the smart grid applications are classified into three classes where the resources are distributed based on the Packet Drop Rate (PDR). In second level, a novel queuing algorithm is proposed to prioritize the RT users with shortest delay. Furthermore, the proposed algorithm's performance is evaluated in term of throughput, delay and fairness index and compared with Frame Level Scheduler (FLS) and Proportional Fairness (PF) algorithms where the proposed approach illustrates higher overall system performance in terms of throughput, delay and fairness.

Index Terms—Smart grid communications, LTE, smart applications, resource allocation, scheduling, 4G

I. INTRODUCTION

Smart grid technology, which widely utilizes wireless communications networks [1], is expected to be the promising electrical technology. It has several advantages over the traditional grid such as auto electricity maintains, damage detection, monitor and control resource generations, substations and user end, electricity illegal use as in Table I [2]. The design of the smart grid communications network is composed of three key networks namely Wide Area Network (WAN), Neighbor Area Network (NAN) and Home Area Network (HAN) [3].

Wireless communications technology has an essential impact on the smart grid since it handles control and monitor command among source power generation, substations and end user [4]. However, several wireless

networks can be employed in the smart grid, such as Worldwide Interoperability networks can be employed in the smart grid, such as Worldwide Interoperability for Microwave Access (WiMAX) and Long Term Evolution (LTE) as shown in Fig. 1. In concept, there is no single wireless network that can serve the entire smart grid applications [5]. LTE network is expected to be an appropriate wireless technology for next generation of the electrical grid, because it demonstrates high performance, security, data rates and low latency [6].

TABLE I: THE MAIN DIFFERENCES BETWEEN TRADITIONAL GRID AND THE SMART GRID

Current grid	Smart grid
Electromechanical	Digital
Centralized Generation	Centralized Or Distributed
Hierarchical	Network
Few Sensors	Sensors throughout
Blind	Self-Monitoring
Manual Restoration	Self-Healing
Manual Check/Test	Remote Check/Test
Limited Control	Pervasive Control
Few customer choices	Many customer choices

Scheduling plays a crucial role in the smart grid communications, because it is in charge of the resources allocation process. Owing to the fact that LTE is not specifically designed for smart grid, novel scheduling algorithms should be proposed to meet the smart grid requirements [7]. The smallest allocated resource unit in LTE is called Resource Block (RB) spanning 180 KHZ in frequency domain and it is divided into two slots in time domain as shown in Fig. 2, each slot length is 0.5 ms. Scheduling decision is updated at each Transmission Time Interval (TTI) which is set to just 1 ms [8].

This paper proposes a novel scheduling algorithm which allocates the resources based on Packet Drop Rate (PDR) which results in reducing the starvation for all classes since the proposed algorithm takes into account all classes including Real-Time applications (RT), Non-Real-Time applications (NRT) and Best Effort (BE). We use self-learning algorithm based on Habian concept which proves short computation time and less complexity, these issues are considered critical in scheduling design. The Packet Loss Ratio (PLR) is reduced due to trend indicators implementation which used to indicate whether the PDR value is increasing or decreasing according the threshold allowed for each class. Moreover, a new queuing algorithm is proposed to prioritize RT users and improve the overall system throughput.

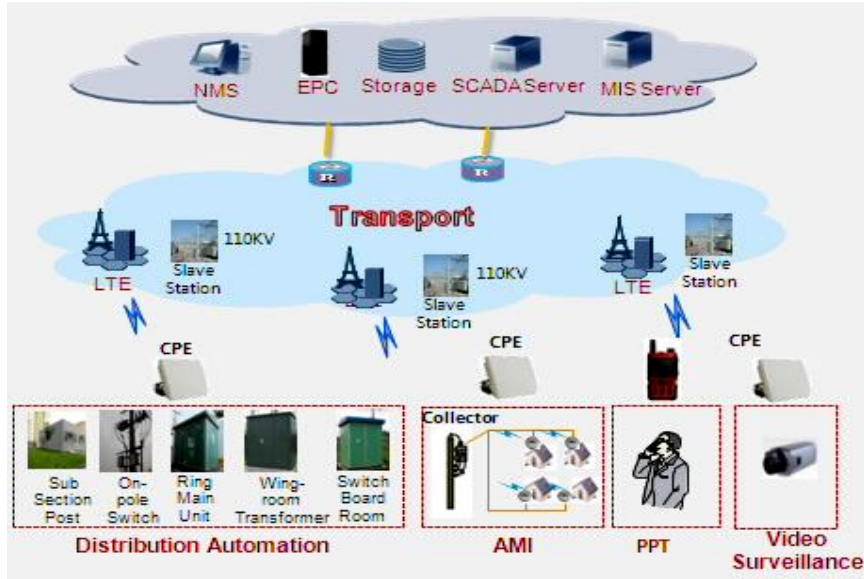


Fig. 1. Smart grid communications over LTE network.

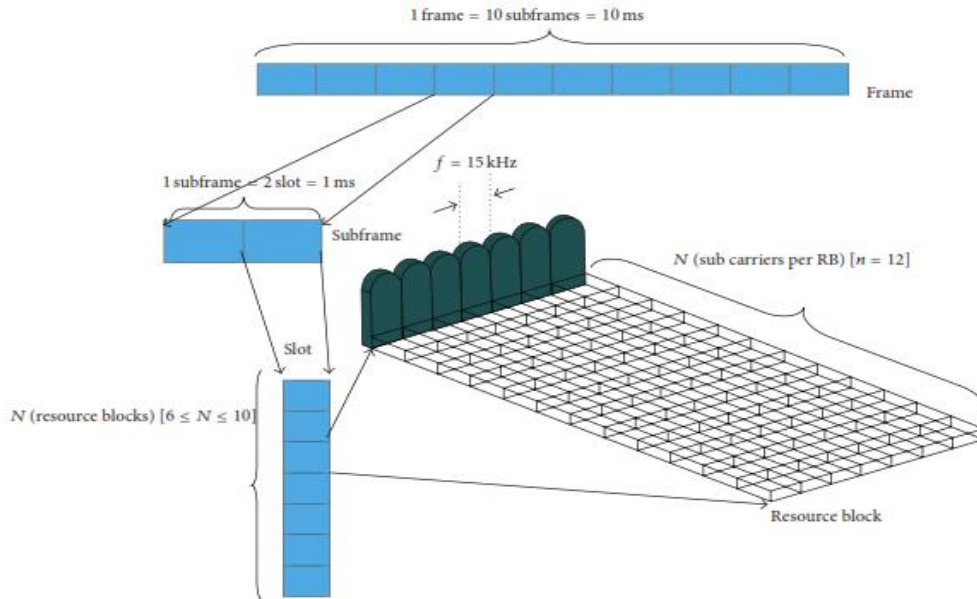


Fig. 2. LTE resource block.

The rest of paper is organized as follows. The related works are explained in details in Section II. Section III is dedicated to present the system model. Section IV describes the simulation scenario. In Section V, the results are deeply discussed and Section VI concludes the paper.

II. RELATED WORKS

Several algorithms have been proposed in order to improve the scheduling for both Real-Time applications (RT) and Non-Real time applications (NRT). Such as Proportional Fairness (PF) which takes into account the experienced channel quality and the past average throughput of the users, which means it aims at improving fairness in term of throughput regardless other QoS requirements [9], [10]. Moreover, it doesn't concern about delay constraint therefore it is suitable only for

NRT applications as in (1). Thus it is not a preferable choice for smart grid which utilizes real-time communications [11].

$$M_{k,i} = \frac{r_{k,i}}{R_i} \quad (1)$$

where $R_i(t)$ is the average throughput for user k . $r_{k,i}(t)$ is the expected throughput for user k .

Another approach which is proposed in [12], namely Frame Level Scheduler (FLS), it is a two level scheduling algorithm which utilizes discrete time control loop (DLCL) in the first level which results in reducing the complexity of computation as shown in (2).

$$u_i(k) = q_i(k) + \sum_{n=2}^{M_i} [q_i(k-n+1) - q_i(k-n+2) - u_i(k-n+1)]c_i(n) \quad (2)$$

where $u_i(k)$ is the corresponds to the amount of data that is transmitted during the k -th frame; $q_i(k)$ is the i -th queue length at time $c_{k,i}$; K is the frame number; c_i is the coefficients which used to determine how the queued data are spread over the consecutive sampling interval; M is the consecutive sampling interval; n is the active traffic flows. And the second level is based on PF. [13] proposed Maximum Throughput approach (MT) has also improved the overall system throughput since it allocates the available resources to users with best channel conditions which results in a huge starvation to the users with bad channel conditions as shown in (3). One of the recent and closest works to our model is proposed in [14]. It is also based on PDR and prioritizes users with tightest delay requirements. But the main weakness in this work is that it prioritizes RT users and makes a huge starvation to NRT users since almost all the time the allocated bandwidth is limited. Here in this paper, we improve the delay sensitivity and improve the overall system throughput.

$$m_{i,k}^{MT} = d_k^i(t) \quad (3)$$

where $d_k^i(t)$ is the maximum achievable throughput.

III. SYSTEM MODEL

A two-level scheduling algorithm is proposed based on self-learning and Real Time (RT) service queuing algorithms. The first level is based on self-learning concept, here is a simple example to clarify the self-learning concept, a system can be trained to detect whether the email message is spam or not so that it can easily detect upcoming messages and classify them into classes. However, Complexity and delay are critical issues which must be considered during the scheduling algorithm design so that the proposed algorithm is intended to have less complexity and shorter computation time compared to other existing algorithms.

A. First Level

The basic concept of Self-Learning Algorithm (SLA) mechanism is to show how much the weight of the connected factors should be increased or decreased. Furthermore, the proposed algorithm is based on Hebbian learning algorithm which is recently utilized in wireless networks such as cognitive radio systems [15]. The proposed algorithm absorbs the required information from the environment and keeps it in synaptic weights as in (4).

$$D_i(m + 1) = D_i(m) + \Delta D_i(m) \quad (4)$$

where $\Delta D_i(m)$ is the m th step, may have positive or negative values. $D_i(m)$ is a synaptic weight. The main idea of the proposed algorithm is to allocate the available resources to Real-Time (RT), Non-Real Time (NRT) and Best Effort (BE) applications based on Packet Drop Rate (PDR) where PDR is calculated as in (5).

$$PDR = \frac{1}{k} \sum_{k=1}^k \frac{n_k^{dropped}}{n_k^{total}} \quad (5)$$

where $n_k^{dropped}$ is the total number of packets dropped for user k . n_k^{total} is the total number of packets arrived for user k and k is the total number of active users. a , b and c represent the weights of RT, NRT and BE respectively as in (6).

$$a + b + c = 1 \quad (6)$$

The values of a , b and c are initially calculated as a ratio of the number of active users in each service to the active users in whole system as shown in (7).

$$a = \frac{A}{A+B+C}, \quad b = \frac{B}{A+B+C}, \quad c = \frac{C}{A+B+C} \quad (7)$$

where A , B and C are the active number of users in RT, NRT and BE applications respectively. The number of allocated resources to the aforementioned applications are represented by α , β and γ respectively and calculated as (8).

$$\begin{aligned} \alpha &= (\text{round off } a)M \\ \beta &= (\text{round off } b)M \\ \gamma &= (\text{round off } c)M \end{aligned} \quad (8)$$

where M is the total number of RBs. The PDR value is calculated for RT and NRT applications at each Transmission Time Interval (TTI) and kept in vectors R_T and N_RT . The calculated PDR values for RT and NRT applications of the current and previous TTIs are compared with the PDR threshold p_{th} (the minimum allowed PDR value).

The resource allocation strategy is changed due to the PDR values changes. But in case when the PDR value change is so small, there is no need to change the resource allocation strategy. To achieve such a goal, the proposed algorithm utilizes trend indicators (I_{RT} and I_{NRT}) which indicate the increases or decreases of the PDR values and change the resource allocation strategy after a specific number of the PDR value changes.

To guarantee higher levels of QoS for each application, the values produced for RT by the comparisons between current TTI, previous TTI and p_{th} are temporarily saved in α_{RT} , β_{RT} and γ_{RT} . Similarly, the PDR value of NRT application also saved in temporary parameters α_{NRT} , β_{NRT} and γ_{NRT} . These temporary values are then used to make final decision on the values of α , β and γ in the next TTI. However, in the first level the proposed algorithm works as follows:

1. PDR value is compared with p_{th} value, if the PDR value in the current TTI is larger than p_{th} , the change in resource allocation is triggered. The allocated resources to either NRT or BE service are decreased by one resource block. NRT application has more priority compared to BE one so that the resources which allocated to BE are decreased. In case BE is not allocated any resources ($\gamma_{NRT} < 0$), the value of β_{NRT} is decreased by one resource block to increase α_{RT} value. One resource block is considered as enough to improve the QoS for RT service since it is 1ms and 180 KH.

2. If $PDR < p_{th}$, all values are not changed since their QoS requirements are satisfied.
3. Compare the PDR values in current and previous TTIs, if PDR in the current TTI is lower than that of the previous TTI then I_{RT} is checked. If $I_{RT} \leq 0$, it is decremented by one stating that the average PDR of RT service is decreasing with time. If $I_{RT} > 0$, its value reset to zero.
4. If I_{RT} value of the current TTI is higher than that of previous TTI but lower than p_{th} , I_{RT} is checked. If $I_{RT} < 0$, it is reset to zero otherwise it is increased by one stating that PDR value is increasing with time.
5. If the PDR values in current and previous TTIs are similar, there are no changes.

B. Second Level

In the second level, we propose a queuing algorithm which focuses on Real-Time applications (RT) and tries to serve the users within their delay constraints. Moreover, the proposed model aims at improving the overall system throughput as in Fig. 3.

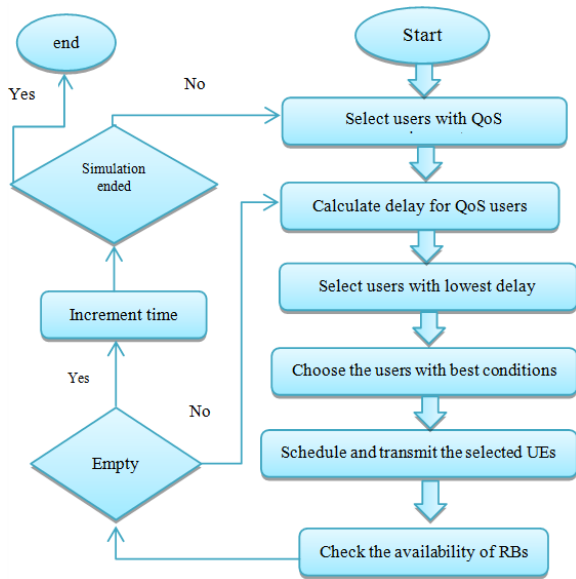


Fig. 3. Illustrates the second level

The most important factors considered during queuing the users based on their delay are $T_{waiting}$ which refers to the waiting time of the packet in the queue and delay budget of the RT service as in (9).

$$Hol_i(t) = T_p(t) - T_{waiting} \quad (9)$$

where $T_{waiting}$ is the waiting time of the packet in the queue and T_p is the time when the packet served. $T_{comparison}(t)$ is defined as the difference in time between budget delay ($T_{b,j}(t)$) and $Hol_i(t)$ delay as in (10).

$$T_{comparison}(t) = T_{b,j}(t) - Hol_i(t) \quad (10)$$

The user with the tightest delay is served first as in (11).

$$m = argmin delay(t) \quad (11)$$

where m is the delay metric which aims at prioritizing user with the shortest delay [16].

IV. SIMULATION SCENARIO

The bandwidth is assumed to be 10 MHz and the number of resources block is 50 with 12 subcarriers each. All the applications users are randomly distributed through the cell coverage. The overall system performance is evaluated in terms of delay, throughput and fairness index. Furthermore, the simulator used in this scenario is LTE-Sim which is an open access. Table II shows the most important parameters used in this scenario:

TABLE II: SIMULATION PARAMETERS

Parameter	Value
System bandwidth	10MHz
RBs No	Centralized Or Distributed
Subcarriers	12
RB bandwidth	180 KHz
TTI	1 ms
Transmission power	125 mw

V. RESULTES AND DISCUSSION

The proposed model is evaluated in terms of throughput, delay and fairness index. Furthermore, we chose such approaches in order to see their performance for overall system taking into account all scheduling classes. Proportional Fairness (PF) approach results in the lowest comparative throughput for RT applications since it takes into account the experienced channel quality and the past average throughput of the users as shown in Fig. 4. Moreover, it aims at improving the fairness level in term of throughput among users regardless the delay requirements. FLS algorithm shows an average throughput performance for RT services where it is designed for real time applications and the goal is to serve RT users within their delay constraints. The proposed algorithm achieves the best performance even for BE services since it guarantees the minimum requirements of data rate as in Fig. 5.

The aforementioned approaches in the literature concern about one or two scheduling criteria, for example some algorithms serve the users with tightest delay first and others concern about throughput, and evaluate their works based on these criteria not in overall network. If such approaches are practically implemented in the network, only users with tightest delay or users with best channel conditions will be served which will results in significant starvation for other services such as bad channel conditions users or non-real time applications. To overcome such a problem, we evaluate the delay for overall system including RT, NRT and BE applications. Also PF shows the lowest performance in term of delay since it doesn't take into account any form of delay measures, whereas FLS performs better up to 40 users to reach approximately 0.2s for 50 users, where this is considered very high delay especially for RT applications

as illustrated in Fig. 6. The proposed algorithm suggests the lowest delay for overload situation which means it serves all users from different classes within their delay requirements. This is due to the fact that the second level of the proposed algorithm is designed to allocate the resources to the users with tightest delay as in (8).

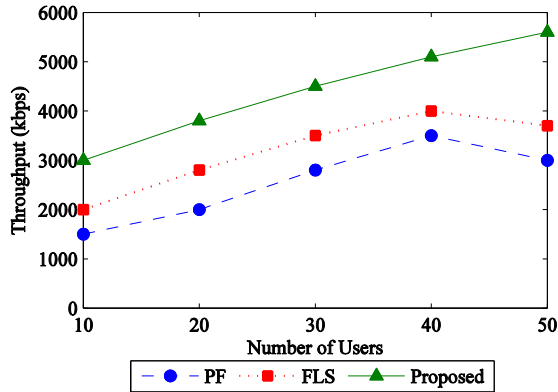


Fig. 4. Throughput of RT applications

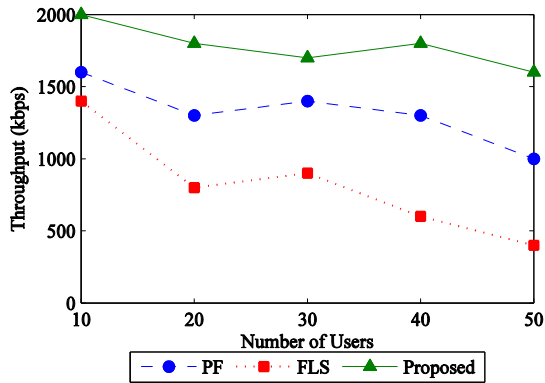


Fig. 5. Throughput of BE applications

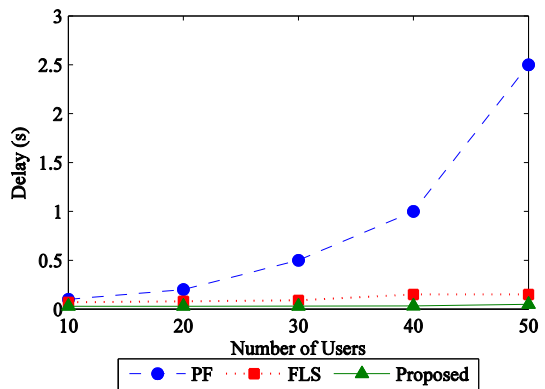


Fig. 6. Overall system delay

Similarly, the fairness index is evaluated for overall system in order to show the real performance of the comparative algorithms for different classes with different QoS requirements. The overall system fairness of the proposed algorithm outperforms similar to other two algorithms for low loaded users (20 users) whereas for overloaded systems the proposed algorithm shows the best performance. This is due to the fact that our proposed algorithm takes into account all classes

including BE and allocates resources to them based on the Packet Drop Rate (PDR) hence overall system fairness considerably increases as in Fig. 7.

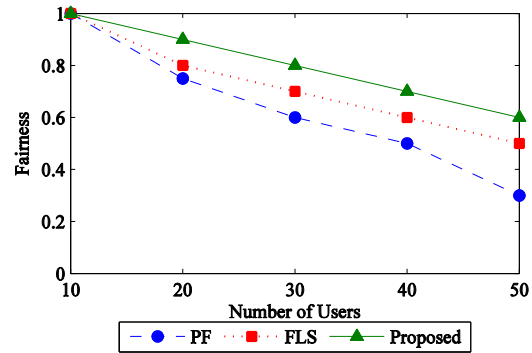


Fig. 7. Overall system fairness

VI. CONCLUSION

The proposed scheduler based self-learning algorithm has illustrated a significant improvement for all classes including RT, NRT and BE. Moreover, the NRT and BE applications have a higher opportunity to be served even in high loaded situations. Whereas the other two algorithms result in low performance when evaluated for overall system performance because they are designed for specific criteria such as PF which concerns about throughput regardless other QoS parameters. Also FLS approach considers only the users with tightest delay nevertheless others criteria such NRT users. This is not practically realistic since every wireless network has all types of classes (RT, NRT and BE) and requires scheduler which offers a compromise among classes. The proposed approach has outperformed other comparative algorithms for overall system performance in term of throughput, delay and fairness index. The proposed algorithm has improved the system throughput up to 6000 kbps for over loaded situation (50 users) and maintained the minimum data rate for NRT and BE applications which results in higher overall system throughput. Similarly, the proposed algorithm has demonstrated the best performance in terms of delay and fairness index, whereas PF has proved the worst performance in all comparative parameters. FLS has illustrated an average performance in all studied parameters.

REFERENCES

- [1] R. Zhu, X. Tan, J. Yang, S. Chen, H. Wang, K. Yu, *et al.*, "An adaptive uplink resource allocation scheme for LTE in smart grid," *Journal of Communications*, vol. 8, pp. 505-511, 2013.
- [2] A. Hajjawi, M. N. Hindia, M. Ismail, K. A. Noordin, and M. Dernaika, "Teleoperation scheduling algorithm for smart grid communications in LTE network," *Applied Mechanics and Materials*, 2014, pp. 340-345.
- [3] J. Huang, H. Wang, Y. Qian, and C. Wang, "Priority-based traffic scheduling and utility optimization for cognitive radio communication infrastructure-based smart grid," *IEEE Transactions on Smart Grid*, vol. 4, pp. 78-86, 2013.
- [4] P. Yi, X. Dong, A. Iwayemi, C. Zhou, and S. Li, "Real-time opportunistic scheduling for residential demand response," *IEEE Transactions on Smart Grid*, vol. 4, pp. 227-234, 2013.

- [5] O. Al-Khatib, W. Hardjawana, and B. Vucetic, "Queuing analysis for smart grid communications in wireless access networks," in *Proc. IEEE International Conference on Smart Grid Communications*, 2014, pp. 374-379.
- [6] A. Hajjawi, M. Ismail, M. N. Hindia, M. S. Jawad, and S. Musleh, "Integration facilitation of scheduling algorithms in the smart grid communications over 4G networks," in *Proc. Third Intl. Conf. on Advances in Computing, Electronics and Electrical Technology*, Malaysia, 2015, pp. 39-41.
- [7] M. N. Hindia, A. W. Reza, and K. A. Noordin, "A novel scheduling algorithm based on game theory and multicriteria decision making in LTE network," *International Journal of Distributed Sensor Networks*, 2014.
- [8] G. Piro, L. A. Grieco, G. Boggia, F. Capozzi, and P. Camarda, "Simulating LTE cellular systems: an open-source framework," *IEEE Transactions on Vehicular Technology*, vol. 60, pp. 498-513, 2011.
- [9] F. Capozzi, G. Piro, L. A. Grieco, G. Boggia, and P. Camarda, "Downlink packet scheduling in LTE cellular networks: Key design issues and a survey," *Communications Surveys & Tutorials, IEEE*, vol. 15, pp. 678-700, 2013.
- [10] N. Saed, K. Wee, C. Siew-Chin, L. T. Hui, and W. Y. Yin, "Low complexity in exaggerated earliest deadline first approach for channel and QoS-aware scheduler," *Journal of Communications*, vol. 9, 2014.
- [11] M. B. Shahab, A. Hussain, and M. Shoaib, "Smart grid traffic modeling and scheduling using 3GPP LTE for efficient communication with reduced RAN delays," in *Proc. 36th International Conference on Telecommunications and Signal Processing*, 2013, pp. 263-267.
- [12] G. Piro, L. A. Grieco, G. Boggia, R. Fortuna, and P. Camarda, "Two-level downlink scheduling for real-time multimedia services in LTE networks," *IEEE Transactions on Multimedia*, vol. 13, pp. 1052-1065, 2011.
- [13] S. Mahato, B. Allen, E. Liu, and J. Zhang, "Performance evaluation of maximum throughput based scheduling of OFDMA LTE networks," in *Proc. Wireless Advanced (WiAd)*, 2012, pp. 1-5.
- [14] R. Kausar, Y. Chen, and K. K. Chai, "Adaptive time domain scheduling algorithm for OFDMA based LTE-advanced networks," in *Proc. 7th International Conference on Wireless and Mobile Computing, Networking and Communications*, 2011, pp. 476-482.
- [15] S. Zhao, H. Guo, J. Xing, and S. Ji, "Performance analysis of cognitive cooperative communication system based on optimal relay selection scheme," *Journal of Communications*, vol. 9, 2014.
- [16] A. Hajjawi and M. Ismail, "QoS-Based scheduling algorithm for LTE-A," *Research Journal of Applied Sciences, Engineering and Technology*, in Press.



Ayman Hajjawi was born in Menbij, Syria, in 1987. He received the B.S. degree in Communications Engineering from Ittihad Private University, Syria, in 2010 and the M.S. degree in Telecommunications Systems from Technical University of Malaysia Malacca (UTEM), Malaysia, in 2013. His research interests include QoS in 4G & 5G networks. He is currently a Ph.D. fellowship student at the Department of Electronic, Electrical and System Engineering, The National University of Malaysia (UKM).



Mahamod Ismail was born in Selangor, Malaysia, in 1959. He received the B.S. degree in Electronics and Electrical Engineering from the University of Strathclyde, United Kingdom, in 1985 and the M.S. degree in Communications Engineering and Digital Electronics from the University of UMIST, Manchester, United Kingdom, in 1987. He is currently a professor at the Department of Electrical, Electronic and System Engineering. His research interests include mobile communications and wireless networking with particular interest on radio resource management for 4G and beyond.