

# JOURNAL OF TELECOMMUNICATIONS AND INFORMATION TECHNOLOGY

2/2009

## Percolation Driven Flooding for Energy Efficient Routing in Dense Sensor Networks

*G. Vakulya and G. Simon*

*Paper*

5

## Multiobjective Design of Wireless Ad Hoc Networks: Security, Real-Time and Lifetime

*Z. Karakehayov*

*Paper*

13

## A Method of Mobile Base Station Placement for High Altitude Platform Based Network with Geographical Clustering of Mobile Ground Nodes

*H. Y. Song*

*Paper*

22

## A Framework for Detection of Selfishness in Multihop Mobile Ad Hoc Networks

*J. Konorski and R. Orlikowski*

*Paper*

34

## Modulo N Backoff Scheme for Effective QoS Differentiation and Increased Bandwidth Utilization in IEEE 802.11 Networks

*T. Janeczak, J. Konorski, J. Woźniak, and K. Pawlikowski*

*Paper*

41

## Empirical Season's Fadings in Radio Communication at 6 GHz Band

*J. Bogucki and E. Wielowieyska*

*Paper*

48

## Laser Beam Attenuation Determined by the Method of Available Optical Power in Turbulent Atmosphere

*L. Dordová and O. Wilfert*

*Paper*

53

## The Design of $4 \times 4$ Multimode Interference Coupler Based Microring Resonators on an SOI Platform

*T.-T. Le and L. W. Cahill*

*Paper*

58

(Contents Continued on Back Cover)

## ***Editorial Board***

Editor-in Chief: ..... ***Paweł Szczepański***

Associate Editors: ..... ***Krzysztof Borzycki***  
***Marek Jaworski***

Managing Editor: ..... ***Maria Łopuszniak***

Technical Editor: ..... ***Ewa Kapuściarek***

## ***Editorial Advisory Board***

Chairman: ..... ***Andrzej Jajszczyk***  
***Marek Amanowicz***  
***Daniel Bem***  
***Wojciech Burakowski***  
***Andrzej Dąbrowski***  
***Andrzej Hildebrandt***  
***Witold Hołubowicz***  
***Andrzej Jakubowski***  
***Alina Karwowska-Lamparska***  
***Marian Kowalewski***  
***Andrzej Kowalski***  
***Józef Lubacz***  
***Tadeusz Łuba***  
***Krzysztof Malinowski***  
***Marian Marciniak***  
***Józef Modelski***  
***Ewa Orłowska***  
***Andrzej Pach***  
***Zdzisław Papir***  
***Michał Pióro***  
***Janusz Stokłosa***  
***Wiesław Traczyk***  
***Andrzej P. Wierzbicki***  
***Tadeusz Więckowski***  
***Józef Woźniak***  
***Tadeusz A. Wysocki***  
***Jan Zabrodzki***  
***Andrzej Zieliński***

ISSN 1509-4553      on-line: ISSN 1899-8852

© Copyright by National Institute of Telecommunications  
Warsaw 2009

Circulation: 300 copies

Sowa - Druk na życzenie, [www.sowadruk.pl](http://www.sowadruk.pl), tel. 022 431-81-40

# JOURNAL OF TELECOMMUNICATIONS AND INFORMATION TECHNOLOGY

## *Preface*

Despite the apparent idealization and separation from many mundane aspects of life, basic research is not completely free of the influence of fashion and politics. Telecommunication, with its aspiration to high practical relevance, may be more prone to such influences than some other areas. Most of us remember the times when it was difficult to work on anything not directly related to ATM networks, which were expected to supersede all other networking concepts. Those days are gone. And a good riddance, too. The square blocks would not fit the round holes no matter how hard we tried. Ultimately, the reality has sealed the fate of that, one has to admit, highly destructive and costly trend.

For some time, it would seem that wireless ad hoc networking was a similarly misguided fashion. The volume of research in this area driven by the popular arguments about the alleged multitude of critical applications, remained for too long in a blatant contrast with its poor materialization in the real world. If all those problems had been addressed and solved so many times over, then why couldn't we see those devices in action? Where was the fulfillment of the promise of ubiquity. Where were the low cost, the ease of deployment, the reliability, the interoperability, the security, the long battery life, and so on?

Fortunately, we seem now to be reaching a turning point: the situation begins to change. Vendors are starting to offer wireless sensor networks for environmental monitoring and security applications. Several academic and industrial projects demonstrate working solutions in the areas of medical monitoring, location tracking, advanced sensing – all of them based on some flavors of ad hoc networking. While many unsettled issues are still being researched and many good solutions still have to be found, we can confidently say that ad hoc networking is coming of age and gaining ground as a practicable concept.

Why has it taken so long? In our opinion, it was necessary to realize that a wireless network is not quite the same thing as a wired system with the wires removed and replaced by wireless channels. Most of the early wireless protocols basically focused on emulating wires over wireless, the common assumption being that one of the prerequisites for successful networking was a way to turn the inherently broadcast and messy wireless medium into a collection of disjoint and reliable point-to-point links. Another destructive legacy concept was the abominable layering of protocols and application software. That was much harder to get rid of than one could imagine. Even having finally realized that all those layers get in the way of small, cheap, simple, reliable, and durable (e.g., in terms of power consumption and battery life), many people continued to see in the wireless networks some “deep structures” and “planes” turning their “solutions” into monstrosities not completely unlike the abominable ATM switches.

Not surprisingly, the workable solutions that we have witnessed recently do not look like monsters at all. They employ simple protocols fully embracing the unkempt nature of the wireless medium, and their software is organized in a holistic manner, whereby the functional interconnection of the different modules does not follow a fixed pattern of layers or planes. They do not insist on providing exact solutions to complex problems within the framework of a small node with a limited set of resources. Instead, the distributed processing carried out by the network is organized in a way that factors in the unavoidable limitations of its cheap and inherently unreliable members. Such a wireless network is not an emulated replica of a wired system. It follows different rules and behaves in a different way. It is what we would call a true ad hoc network.

One of the papers included in the present issue, *Percolation Driven Flooding for Energy Efficient Routing in Dense Sensor Networks*, by Gergely Vakulya and Gyula Simon, deals with flooding lying at the heart of the network's routing scheme. This is a good exemplification of the paradigm shift that we have in mind. Normally, flooding would carry negative connotations. Aren't we supposed to avoid it at all cost? Doesn't it waste bandwidth? Can not we do better by lying those illusory semi-reliable point-to-point links in our wireless network? Well, the answer is that a form of flooding is bound to happen in a wireless setup no matter how much we try to send packets along those imagined paths. If so, then why not turn it into a feature contributing to the effectiveness and reliability of the complete scheme.

In the same spirit, another paper, *Multiobjective Design of Wireless Ad Hoc Networks: Security, Real-Time and Lifetime*, by Zdravko Karakehayov, discusses the holistic character of design in wireless ad hoc network, focusing on the tradeoffs that must be all addressed globally in any practicable solution. The conclusion is quite illuminating: while hierarchical solutions need not be banned (their usefulness as complexity reducers and facilitators of understanding is unquestionable), the ways of arriving at them need not isolate the heavily interacting components of the whole design.

Yet another contribution, *A Method of Mobile Base Station Placement for High Altitude Platform Based Network with Geographical Clustering of Mobile Ground Nodes*, by Ha Yoon Song, addresses the inherent complexity of designing high-bandwidth wireless networks with nontrivial mobility. While not strictly related to ad hoc networking, the flexible clustering scheme for aerial networks presented by Ha Yoon Song demonstrates how to cope with the fuzziness of neighborhoods (and the illusory links) in a high reliability public wireless system. It is useful as an illustration of how the concept of neighborhood has been abused and trivialized in many abstract studies of wireless networks, where many nonchalant assumptions about unrealistically good isolation of neighborhoods have often resulted in practically void performance claims.

The next paper, *A Framework for Detection of Selfishness in Multihop Mobile Ad Hoc Networks*, by Jerzy Konorski and Rafał Orlikowski, describes the need for a fully-distributed fairness enforcement mechanism for wireless ad hoc networks where nodes may exhibit a non-cooperative forwarding behavior. The paper describes a new framework, dubbed DST-SDF, building on Dempster-Shafer theory, with the mathematical background and simulation analysis.

The issue of ensuring service quality (QoS) in radio networks remains important, and, to this end, the contribution *Modulo N Backoff Scheme for Effective QoS Differentiation and Increased Bandwidth Utilization in IEEE 802.11 Networks*, by Tomasz Janczak, Jerzy Konorski, Józef Woźniak and Krzysztof Pawlikowski, deals with improvements in the operational strategies of WiFi networks. The paper presents a new modulo N channel access scheme for wireless local area networks. The solution derives from the distributed coordination function (DCF) of the IEEE 802.11 standard, which was further extended into the enhanced distribution channel access (EDCA) scheme by the 802.11e draft specification. The main innovation concerns improvement of the binary exponential backoff scheme used for collision avoidance in 802.11 networks. The most appealing feature of the new modulo N backoff scheme is that it outperforms the original 802.11 solution in terms of the channel utilization ratio under any traffic conditions. Furthermore, the modulo N proposal can be naturally augmented with QoS differentiation mechanisms, like the 802.11e extensions. The prioritized modulo N scheme achieves better throughput-delay characteristics for multimedia traffic when compared with the original 802.11e proposal. At the same time, the new solution retains backward compatibility and includes all the features that have made IEEE 802.11 networks extremely popular nowadays.

This paper devoted to wireless networking: *Empirical Season's Fadings in Radio Communication at 6 GHz Band*, written by Jan Bogucki and Ewa Wielowieyska, covers the more "traditional" matter of radio propagation, specifically the unavailability of line-of-sight radio

links due to multipath propagation. Multipath fading in the atmosphere is not a permanent phenomenon, being strongly dependent on weather and its seasonal variations. The paper presents investigation results of received radio signal fading in radio links and their seasonal distributions, collected over a period of 5 years.

Atmospheric propagation issues are even more important in open-space optical communication systems. This is the subject of the next paper, *Laser Beam Attenuation Determined by the Method of Available Optical Power in Turbulent Atmosphere*, by Lucie Dordová and Otakar Wilfert, which focusses on the atmospheric turbulence effects and presents a new method for determining optical signal attenuation caused by turbulence. This new method of calculating the power budget of an optical link is based on optical intensity distribution in a laser beam after the beam has passed through a turbulent atmosphere. Results obtained with this method are compared to those produced by the Rytov approximation, which is nowadays the most frequently used method for determining turbulence-related attenuation. The paper presents results for the typical communication wavelengths of 850 nm and 1550 nm, as well as 633 nm.

Development of improved communication systems – optical or otherwise – requires the use of advanced components. An example is presented in *The Design of  $4 \times 4$  Multimode Interference Coupler Based Microring Resonators on an SOI Platform*, by Trung-Thanh Le and Laurence W. Cahill, where the authors propose a novel microring resonator based on  $4 \times 4$  multimode interference (MMI) couplers. The device acts as two separate microring resonators within a single structure. The transfer matrix method and the three-dimensional beam propagation method (3D-BPM) are used to verify the working principle of the device. The device is then designed in silicon using insulator (SOI) technology. This device may be a very promising building block for optical switches, filters, add-drop multiplexers, delay lines and modulators.

Regardless of the transmission technology used, efficient operation of a digital network requires more and more intensive data processing, with extensive computing functionality becoming critical for network performance. This matter is dealt with in the paper titled *Linux Scheduler Improvement for Time Demanding Network Applications, Running on Communication Platform Systems*, by Marcin Hasse and Krzysztof Nowicki. Communication platform systems such as advanced telecommunication computing architecture (ATCA) standard blades located in a standardized chassis, provide high-level communication services between system peripherals. Each ATCA blade brings dedicated functionality to the system, but can as well exist as a stand-alone host responsible for servicing a set of tasks. According to the platform's philosophy these parts of the system can be quite independent from other solutions provided by competitors. Every case of system design poses its specific problems. One of the most difficult ones is system integration with a set of components running on different operating system levels. This paper discusses a possible way of augmenting the Linux scheduler as to make it possible for a user-space application to become a critical part of a hard real-time system, e.g., handling a high-bandwidth network service.

Finally, we must not forget that valuable and sensitive information transmitted over a network or stored and processed in a computer system must be adequately protected. A common and effective tool used for this purpose is cryptography. Steady advances made in interception and decryption techniques need to be matched and countered by ongoing improvements in cryptosystems. This important matter is covered by the last two papers published in this issue. The first of them, *Simple Dynamic Threshold Decryption Based on CRT and RSA*, by Bartosz Nakielski and Jacek Pomykała, describes a simple threshold decryption system based on the RSA cryptosystem. This model avoids the application of the Shamir secret sharing protocol and is based only on the Chinese remainder theorem. The flexibility in the threshold level is attained due to the suitable preparation of the input data. The second part of the article describes a modification of the basic model, which admits the sender's impact on the choice of the real receiver's group.

The last contribution, *Generating Pseudorandom S-Boxes – a Method of Improving the Security of Cryptosystems Based on Block Ciphers*, by Piotr Mroczkowski, presents a general framework for improving the security of the cryptosystem based on the symmetric block cipher. The main idea explores the possibility of exchanging the substitution boxes (called S-boxes) in the encryption/decryption algorithm.

Paweł Gburzyński  
Józef Woźniak  
Guest Editors



# Percolation Driven Flooding for Energy Efficient Routing in Dense Sensor Networks

Gergely Vakulya and Gyula Simon

**Abstract**—Simple flooding algorithms are widely used in ad hoc sensor networks either for information dissemination or as building blocks of more sophisticated routing protocols. In this paper a percolation driven probabilistic flooding algorithm is proposed, which provides large message delivery ratio with small number of sent messages, compared to traditional flooding. To control the number of sent messages the proposed algorithm uses locally available information only, thus induces negligible overhead on network traffic. The performance of the algorithm is analyzed and the theoretical results are verified through simulation examples.

**Keywords**— *flood routing, percolation, sensor network.*

## 1. Introduction

Energy efficiency is a key question in wireless sensor networks made of large number of inexpensive nodes with limited energy reserves and restrained processing and communication capabilities. Such networks have been used in a wide range of applications with various goals, using different platforms and technology [1]–[4]. Independently of the nature of the application, however, certain middleware services are always present, one of them being routing. Limited resources on the sensor nodes require that routing protocols be simple, energy-conserving and robust.

Depending on the actual deployment scenario and available hardware, several ad hoc routing algorithms have been proposed. Location aware routing schemes, e.g., greedy perimeter stateless routing [5], location aided routing [6], distance routing effect algorithm [7] use position information of the nodes obtained from global positioning system (GPS) or other localization services. Other data centric schemes do not use location information, e.g., directed diffusion [8], ad hoc on demand distance vector routing (AODV) [9], dynamic source routing [10], temporarily ordered routing [11], or zone routing [12], just to name a few. To allow sensor networks better cope with scaling, hierarchical protocols were designed, e.g., low energy adaptive clustering hierarchy (LEACH) [13] and threshold sensitive energy efficient sensor network protocol (TEEN) [14].

The simplest routing protocol is flooding [15], which is a useful simple (but inefficient) protocol in itself, and also is a building block for several more sophisticated algorithms, e.g., directed flood routing [16], controlled flood routing [17], tiny ad hoc routing [18], or AODV [9], which is used in Zigbee's routing protocol [19].

The basic flooding protocol is the following:

1. *Start*. When a source node intends to send a message to all other nodes in the network, it broadcasts the message to its neighbors.
2. *Relay*. When a node receives a message for the first time, it rebroadcasts the same message to its neighbors. All other copies of the same message, received later, will be discarded by the recipient node.

Flooding naturally can be used in its basic form for network-wide dissemination of information. In this context the important performance criteria of the algorithm are delivery ratio, latency, and the energy efficiency. Flooding is also used for route discovery [9]. Route discovery protocols find a route from a source node to a destination node in two phases. First the network is flooded by a discovery request message originated from the source, and then a reply message is sent back from the destination to the source node. In this way the network learns the path from the source to the destination and this route will be used in the subsequent communication sessions. In this context the quality – primarily the length – of the discovered route is an important performance criterion as well.

Apart from its simplicity the main advantage of the flood routing protocol lies in its robustness: the implicit redundancy built in the algorithm provides resistance against high degree of message loss and node failures as well. The drawback of the algorithm is the large number of packets transmitted in the network, referred to as broadcast storm: whether it is necessary or not, each node will retransmit each message. In a dense network the unnecessarily high number of messages can cause frequent collisions (and thus performance degradation) and it wastes network energy as well [20].

To reduce the number of routing messages, probabilistic variations of the flood routing were proposed. The main idea behind these algorithms is that a node randomly decides whether to forward a message or not. Several variants of random flood algorithms were proposed. According to the simplest protocol, a node will rebroadcast a message after its first reception with probability  $p$  and discards it with probability  $1 - p$  (clearly,  $p = 1$  results in basic flooding).

Test results verify what intuition suggests: higher  $p$  values provide higher network coverage than small  $p$  values. Moreover, in sufficiently large and sufficiently dense random networks there is an interesting bimodal behavior of

the algorithm: if  $p > p_c$ , where the critical probability  $p_c$  depends on the network topology, the message reaches practically all nodes in the network, otherwise only a small portion of the network receives the message. Thus the optimal choice clearly would be  $p = p_c$ . Modified algorithms try to further increase the performance by various modifications; e.g., premature message death can be avoided by varying  $p$  as a function of hop-distance from the source: nodes close to the source rebroadcast messages with higher probability, while distant nodes use smaller  $p$ . A comprehensive study on random flooding algorithms can be found in [15]. The common problem with these algorithms is that the design parameters (probabilities) depend on actual network layout, no automatic or adaptive solution is known. Especially in a network with varying node density the probabilistic algorithms tuned to work reliably will be suboptimal.

More sophisticated adaptive schemes can successfully handle networks with varying properties as well. Location based algorithms use node positions to optimize retransmissions [21], graph-based algorithms utilize connectivity information up to 2-hop neighbors to construct a dominating set in a distributed way [22]. Other heuristic rules, e.g., message counters and neighbor coverage schemes were also proposed [20]. These schemes can adapt automatically to changing topology at the price of higher complexity.

Sensor networks are often modeled by the Poisson-Boolean model. In this model nodes are scattered on the plane by a Poisson point process with constant density  $\lambda$ , and each node has a fixed communication radius  $r$  (disc model). Two nodes can communicate with each other if their distance is less than  $r$ . The Poisson point process applies reasonably well to controlled, but still random deployments. The simplistic disc model, however, is far from reality, where communication range is not circularly symmetric, but rather has an anisotropic shape. Speaking of connectivity, however, the disc model has an important property: it can be considered the worst case model, since other transmission models provide easier percolation under similar circumstances [23].

Percolation theory [24] has important impact on wireless communication networks. Results in continuum percolation theory show that if the density of transmitting stations (or alternatively, their communication power) in a Poisson-Boolean network is greater than a critical value  $\lambda_c$  then an unbounded connected component is formed with probability one (the network percolates) [23], [25]. On the other hand, if  $\lambda < \lambda_c$  then all connected components are finite with probability one. Thus percolation is an important property of a network if long distance multi-hop communication is required.

In this paper a new percolation inspired solution will be proposed based on probabilistic flood routing algorithms. The new algorithm adaptively sets the rebroadcast probability based on the network topology, using locally available neighborhood information only. The proposed algorithm provides high performance with a low number of

messages and can adapt its behavior dynamically to the network properties. The proposed algorithm is only marginally more complex than the basic flood routing thus it can be successfully used in applications with limited resources as well.

The rest of the paper is organized as follows. In Section 2 related works will be summarized. In Section 3 the proposed algorithm will be introduced. An important property of the algorithm will be proven based on percolation theory: every message broadcast in the network will reach infinite number of nodes almost surely. In practice this property ensures that the message reaches almost all of the nodes in a finite-sized network. The performance evaluation of the algorithm through simulation can be found in Section 4. Conclusions are drawn in Section 5.

## 2. Related Work

Results of percolation theory have been applied to wireless networks, forming the theory of continuum percolation [26]. More realistic network models, i.e., unreliable communication links and anisotropic radiation patterns were studied in [23]. An important result of [23] shows that the debatable disc model actually provides a conservative estimate of the network connectivity: percolation in networks with disc communication shapes is more difficult than in networks with any other convex communication shape with the same surface ("discs are hardest" conjecture). This property justifies the usage of the disc model when connectivity issues are studied.

Results from percolation theory have been infiltrating various recent sensor networking algorithms. As a practical application of this phenomenon, the distributed minimum-link-degree rule was proposed to counterbalance local spatial inhomogeneities in ad hoc networks [25]. To provide network wide connectivity, the transmission powers of the nodes are tuned so that each node has at least a minimum number of neighbors.

In [27] a distributed energy saving mechanism is used by switching the nodes in the network on and off, while providing percolation in the network. The simple rule for the algorithm is to keep the active time ratio of the nodes larger than  $\lambda/\lambda_c$ . This algorithm does not scale well since each node in the network has to know the global parameter  $\lambda$ . The idea was further elaborated in [28], where a practically useable distributed algorithm was proposed for the control of the active time ratio. In this algorithm only the local density, namely the number of neighbors is used. For a node with degree  $k$  the active ratio  $\eta_k$  is defined as follows:

$$\eta_k = \begin{cases} \frac{\phi}{k} & k > \phi \\ 1 & k \leq \phi \end{cases},$$

where  $\phi$  is a density independent design parameter. A similar idea will be used in the proposed flooding algorithm.

### 3. Percolation Driven Flooding

The idea behind the percolation driven flood routing is the following: use probabilistic flooding and set the retransmission probability adaptively at every node so that with high probability the message reaches almost all of the nodes in the network.

The algorithm is the following:

1. The source node broadcasts the message with probability one.
2. After the first reception of a message node  $n$  rebroadcasts it with probability  $p_n$ , and discards it with probability  $1 - p_n$ . Copies of the same message received multiple times are discarded with probability one. Probability  $p_n$  is defined as

$$p_n = \begin{cases} \frac{K_{\min}}{K_n} & K_n > K_{\min} \\ 1 & \text{otherwise} \end{cases}, \quad (1)$$

where:  $K_n$  is the degree of node  $n$  (i.e., the number of neighbors of node  $n$ ), and the  $K_{\min}$  design parameter is the required minimum number of neighbors. Two nodes are neighbors if they can hear each other (i.e., a symmetric link exists between them).

The algorithm requires neighbor discovery to determine the number of neighbors. Each node can gather this local information by transmitting and receiving *hello messages*. In practical situations parameter  $K_{\min}$  is chosen around 7, as will be justified by the simulation examples.

**Theorem 1:** If in the infinite random network, generated by a Poisson-Boolean process with density  $\lambda$  and communication radius  $r$ , there exists design parameter  $K_{\min} < \infty$  independent of  $\lambda$ , such that when node  $n$  with degree  $K_n$  transmits a message with probability according to Eq. (1), then each message reaches infinite number of nodes with probability one.

*Proof:* Let us denote the graph representing the network by  $G(\lambda, r, 1)$  where  $\lambda$  is the density,  $r$  is the communication radius, and the third parameter is a probabilistic value used in the generator process (used and discussed later) and is simply set to one at the moment. Using scaling, we define another process with density  $\lambda' = \lambda r^2$  and communication radius 1. This process yields the same graph  $G(\lambda', 1, 1)$ . Now let us use the following Theorem 2 [28].

**Theorem 2:** Given  $G(\lambda', 1, 1)$  with  $\lambda' > \lambda'_c$ , there exists  $2 < \phi < \infty$ , independent of  $\lambda$ , such that when each node with degree  $k$  is active with probability:

$$p(k) = \begin{cases} \frac{\phi}{k} & k > \phi \\ 1 & k \leq \phi \end{cases},$$

then  $G(\lambda', 1, p(\cdot))$  is percolated. The proof of Theorem 2 can be found in [28].

If in the network we decide *before* a particular message is broadcasted which nodes will retransmit the message (when

they receive it) by activating/deactivating nodes with probability (1) we get  $G(\lambda', 1, p(\cdot))$ . In this way we construct a virtual network for each message and we can apply Theorem 2 to it. Thus there exist  $K_{\min} = \phi < \infty$  so that  $G$  percolates. From this it follows that the message will reach infinite number of nodes [24]. ■

In real-world finite-size networks Theorem 1 means that with an appropriate choice of  $K_{\min}$  each message reaches practically all nodes in the network. Further properties of the algorithm will be analyzed through simulation examples.

### 4. Performance Evaluation

In this section the performance of the algorithm will be evaluated through simulation results. First the performance metrics and the network/communication models will be introduced, followed by the simulation environment and the simulation scenarios. Finally, the results of the simulations will be presented to illustrate the behavior of the proposed algorithm.

#### 4.1. Performance Metrics

The main performance criteria of broadcast algorithms are the message delivery ratio and the number of sent messages. If the number of nodes in the network is  $N$  and the number of nodes receiving the message is  $N_{rec}$  then the coverage ratio is defined as  $C_r = \frac{N_{rec}}{N}$ .

Clearly,  $C_r$  close to 1 is required for a good quality of service.

Another quality metric is the total number of sent packets  $M$  while a message is propagated in the network. Trivially, for basic flood routing  $M = N$ , if  $C_r = 1$ . We expect lower number of messages and still high delivery ratio for a good performance algorithm. For easier comparison, we will use the normalized number of messages defined as

$$M_{norm} = \frac{M}{N}.$$

If flooding is used for route discovery, an important performance metric is the length  $L$  of the discovered route.

Other possible metrics are message delay/latency, length of flood period, and number of collisions. The actual low level details, e.g., implementation, hardware and media access control (MAC) layer properties are not investigated in detail in this paper, but their effect is examined through high level parameters they have impact on, i.e.,  $C_r$ ,  $M_{norm}$ , and  $L$  will be studied.

#### 4.2. Network Model

To model random deployment of nodes or possibly mobile networks, in the simulations nodes are placed at random, according to a uniform distribution on a two-dimensional area. Finite communication distances are represented by the disc model, where a communication link is assumed to exist between two nodes if they are less than the communication radius apart. According to the results published

in [24], this – otherwise rather simple and idealistic – model can be considered a worst case model.

Our communication channel model does not deal with specific details of the physical layer or the MAC layer; rather we use a probabilistic model: a message can reach a neighbor within the communication radius with probability  $p_{rec} < 1$ . This high-level model represents message losses due to collision, fading, or other disturbances as well. The model does not distinguish between individual phenomena but rather incorporates the different sources in one parameter, thus some aspects of reality (e.g., inter-message dependencies) are neglected, but the model is faithful enough to provide useful and easy means for testing.

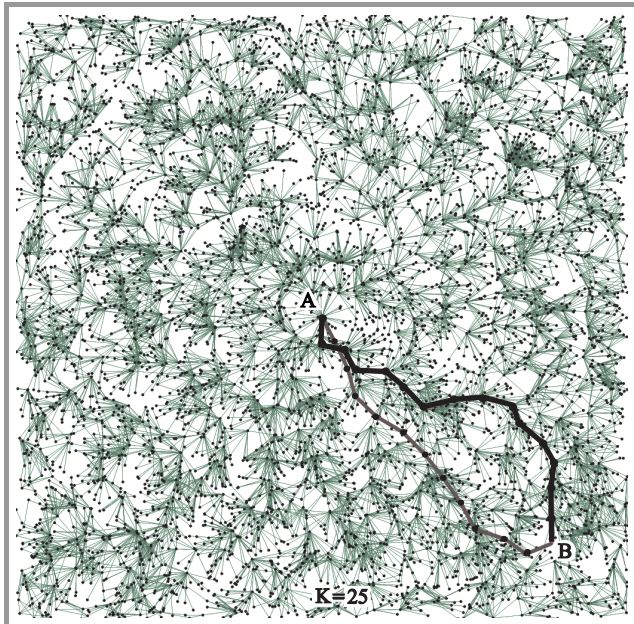
#### 4.3. Simulation Environment

To perform high-speed simulation, a simulator was written in C to validate the efficiency of the proposed algorithm. The program places nodes randomly according to the different test scenarios considered. The number of nodes  $N$ , the communication radius  $r$ , and  $K_{min}$  are input parameters.

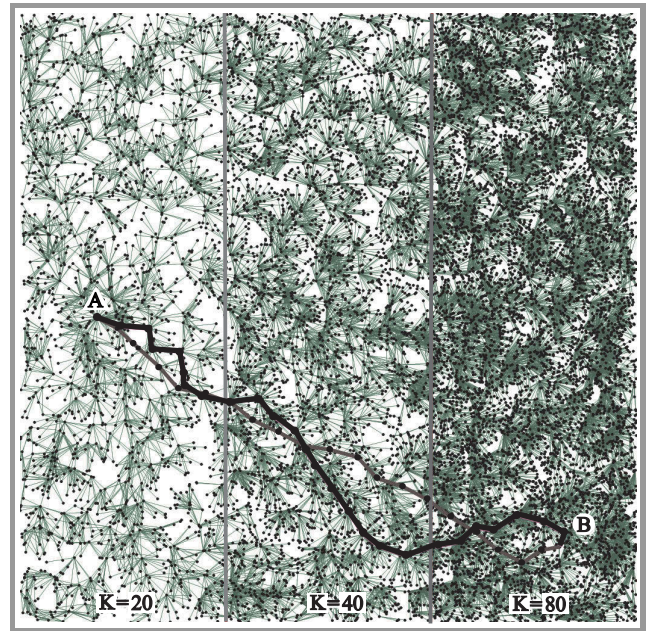
In each simulation, a new placement is generated, and a source node starts transmission. The simulator returns the size of connected component containing the source node, the number of nodes receiving the packet, and the total number of messages. Optionally, the software can visualize the topology of the network, the active data paths, and the nodes' reception status.

#### 4.4. Test Scenarios

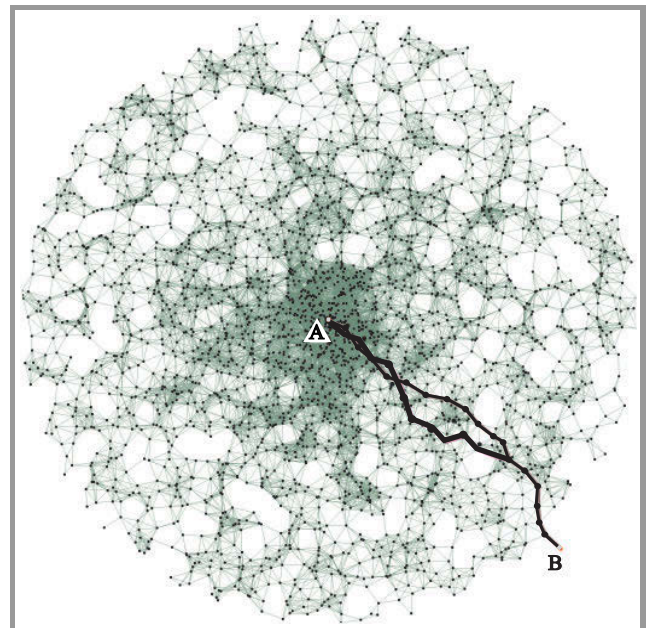
In the test we used 3 different scenarios to model sensor network setups. The first scenario is a random uniform distribution (with constant density), while in the second



**Fig. 1.** Uniform random distribution scenario with 3000 nodes. The average number of neighbors is 25. Thin and thick lines show examples for the message paths between nodes A and B, discovered by flood and percolation driven flood, respectively.



**Fig. 2.** The 3-density scenario. The average numbers of neighbors in the regions from left to right are 20, 40, and 80. Thin and thick lines show examples for the message paths between nodes A and B, discovered by flood and percolation driven flood, respectively.



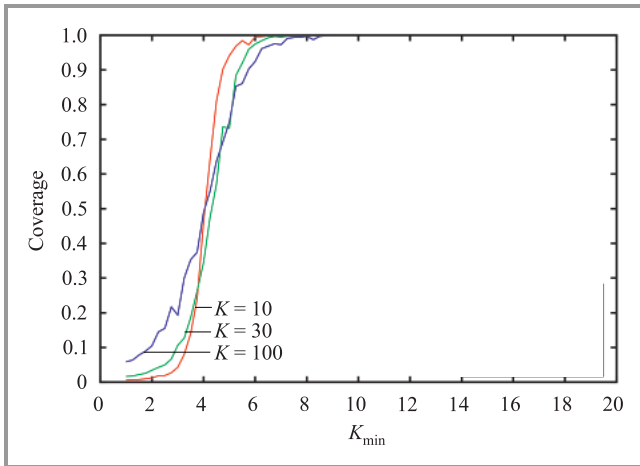
**Fig. 3.** The circular scenario with 3000 nodes. The average number of neighbors is 100 in the centre and 2 at the perimeter. Thin and thick lines show examples for the message paths between nodes A and B, discovered by flood and percolation driven flood, respectively.

scenario we used three regions; in each region nodes were placed with uniform random distribution, but with different densities. In the third scenario the density is increasing towards the centre, thus there is continuous change and extreme deviation in local density. Typical examples for the test networks can be seen in Figs. 1, 2 and 3.

#### 4.5. Test Results

To test coverage ratio  $C_r$  and the number of messages  $M$  we placed 3000 nodes in a square-shaped area, as shown in Fig. 1.

We varied the communication radius to provide different network densities: the average number of neighbors  $K$  was set to 10, 30, and 100. The  $K_{\min}$  value varied from 1 to 20 in 0.25 steps. We run the simulation 100 times for each communication radius,  $K_{\min}$ , and  $p_{rec}$  values, thus each point in the subsequent graphs is the average of 100 experiments.

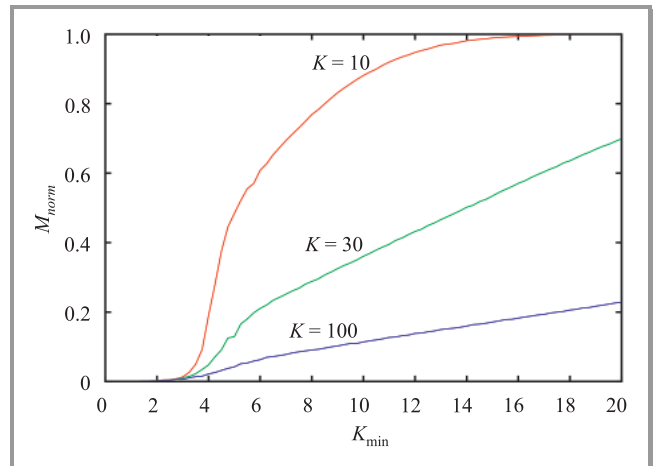


**Fig. 4.** Coverage versus  $K_{\min}$  for different network densities in the uniform distribution scenario (constant network density),  $p_{rec} = 1$ .

Figure 4 presents coverage versus  $K_{\min}$ , while  $p_{rec}$  was set to constant 1 (error-free communication). The graphs for different network densities are similar: when  $K_{\min}$  is low ( $K_{\min} < 3$ ), the message delivery ratio is very low. If  $K_{\min}$  is increased, the coverage is increasing very quickly, the critical value being around  $K_{\min} \approx 4.5$ . Coverage reaches 90% at  $K_{\min} \approx 5$ . If  $K_{\min} > 7$ , practically all nodes in the network receive the message. Figure 4 illustrates that percolation driven flood can indeed be used in networks with different densities:  $K_{\min}$  is independent of the actual network density.

In Fig. 5 the normalized number of sent messages  $M_{norm}$  is shown versus  $K_{\min}$ , for the previous experiments. In case of the conventional flood routing  $M_{norm} = 1$ , because all nodes relay the received message. As Fig. 4 illustrates,  $K_{\min} > 7$  gives almost perfect delivery ratio. Around this value  $K_{\min} \approx 7$  the total number of messages is greatly decreased, as presented in Fig. 5. Depending on the node density, in the tests the total number of messages were reduced by 30..95%, higher reduction rates belonging to higher densities.

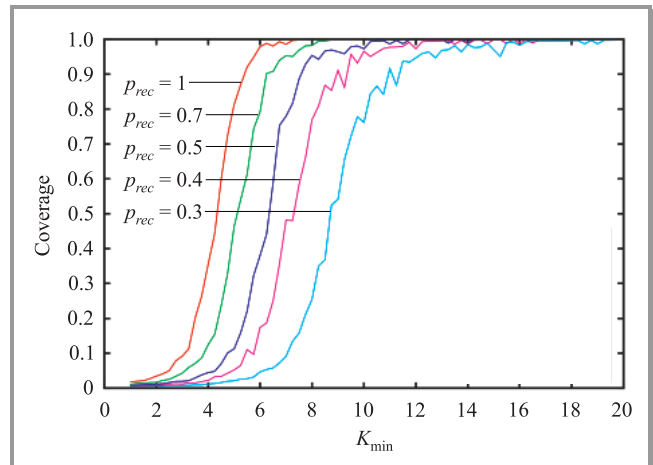
Clearly, in networks, where the average node degree  $K$  is only slightly higher than  $K_{\min}$  the number of messages can be decreased only moderately, while in dense networks a much lower number of messages is enough to provide good delivery ratio, as shown in Fig. 5. According



**Fig. 5.** Total number of messages versus  $K_{\min}$  for different network densities in the uniform distribution scenario (constant network density),  $p_{rec} = 1$ .

to Figs. 4 and 5, in dense networks the percolation driven flood algorithm can effectively reduce the number of message while maintaining good coverage.

Figure 6 illustrates the effect of unreliable communication links. In the experiment constant density ( $K = 30$ ) was used, while the  $p_{rec}$  reception probability was varied between 0.3 and 1. The figure shows that the algorithm



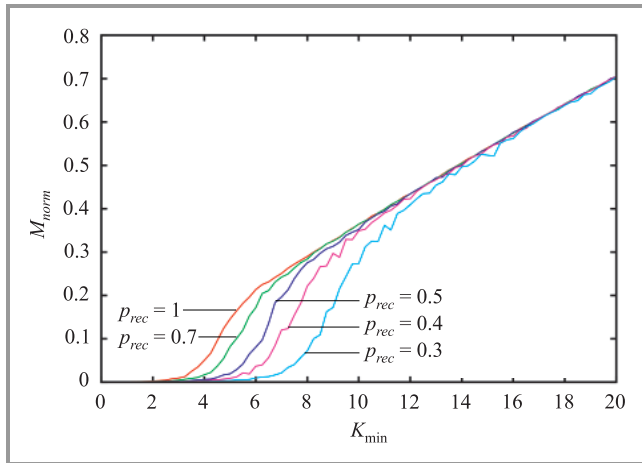
**Fig. 6.** Coverage versus  $K_{\min}$  for different  $p_{rec}$  in the uniform distribution scenario (constant network density),  $K = 30$ .

can provide high coverage even in the presence of bad quality communication links, but as intuitively expected, higher  $K_{\min}$  is necessary as the communication channel degrades.

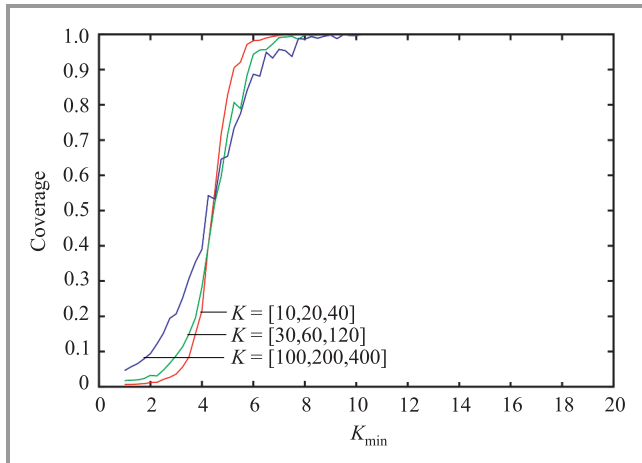
Figure 7 presents the associated number of messages for the unreliable communication experiment.

Percolation driven flood routing algorithm is capable of handling varying network densities. To test this property we used the scenario illustrated in Fig. 2, where three areas with different densities are present. Clearly, a constant retransmission probability would either be suboptimal (set to provide sufficient coverage in the less dense area

as well) or would not provide good coverage in the sparse regions of the network.



**Fig. 7.** Number of total messages versus  $K_{\min}$  for different  $p_{\text{rec}}$  in the uniform distribution scenario (constant network density),  $K = 30$ .

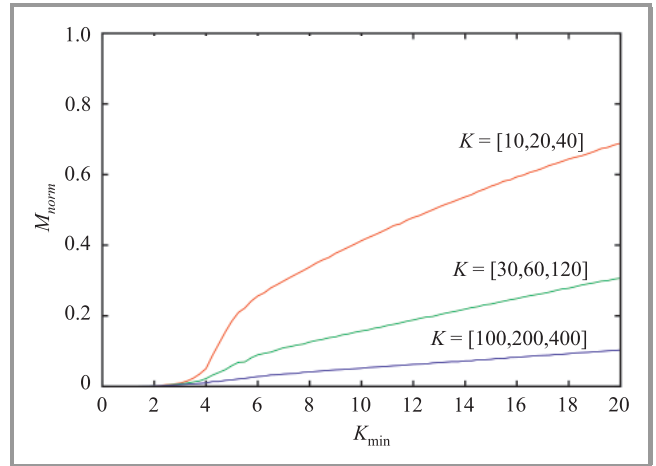


**Fig. 8.** Coverage versus  $K_{\min}$  for different network densities in the varying network density scenario,  $p_{\text{rec}} = 1$ .

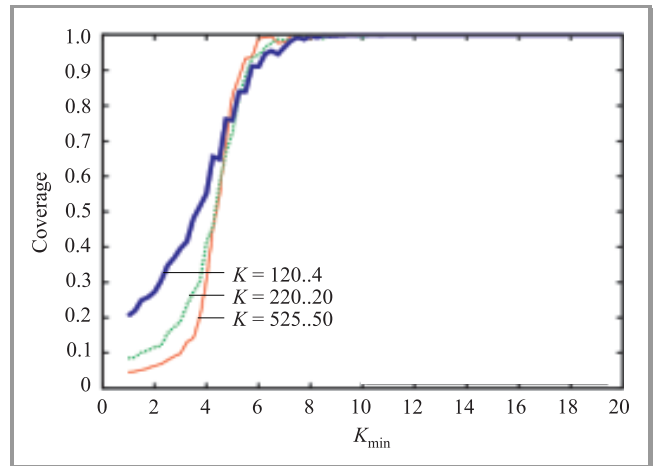
Figure 8 illustrates the effectiveness of the percolation driven flood algorithm, showing the coverage versus  $K_{\min}$  when the densities of regions from left to right were set to  $[10, 20, 40]$ ,  $[30, 60, 120]$ , and  $[100, 200, 400]$ ; and  $p_{\text{rec}} = 1$ . The behavior is quite similar to that of the first scenario (constant density): percolation happens around the same  $K_{\min} \approx 5$ , and practically full coverage can be provided if  $K_{\min} > 7$ .

The associated total numbers of messages are shown in Fig. 9, where the gain with respect to the basic flooding algorithm is apparent.

In Fig. 10 the number of coverage versus  $K_{\min}$  is illustrated in the circular scenario. In the experiments the number of neighbors varied between  $[120..4]$ ,  $[220..20]$  and  $[525..50]$ , respectively. The figure clearly presents the excellent performance of the algorithm even if there are extreme differences in the number of neighbors of each node. The ideal value for  $K_{\min}$  is around 7, in this case the algorithm gives



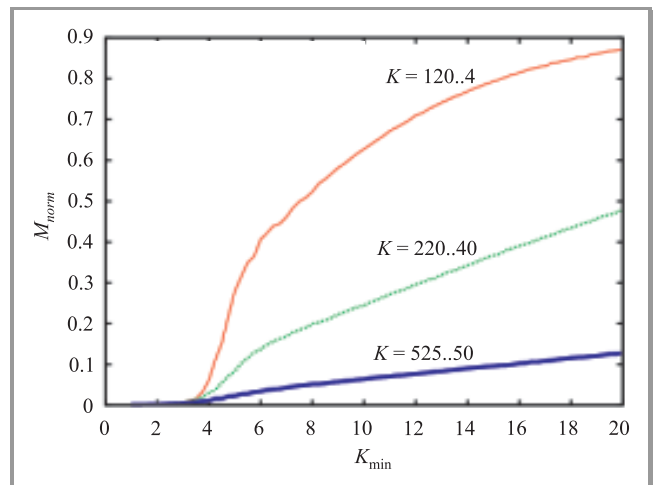
**Fig. 9.** Number of total messages versus  $K_{\min}$  for different network densities in the varying network density scenario,  $p_{\text{rec}} = 1$ .



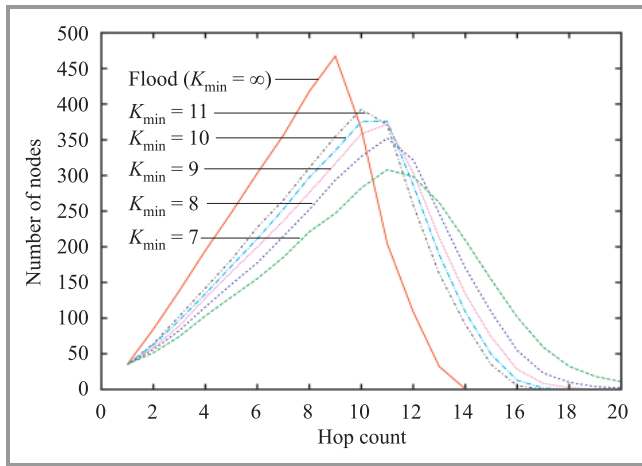
**Fig. 10.** Coverage versus  $K_{\min}$  for different network densities in the circular scenario,  $p_{\text{rec}} = 1$ .

at least 95% delivery ratio. In case of  $K_{\min} = 8$  practically full coverage is provided.

In Fig. 11 the normalized number of messages is shown for the three different density intervals. In case of  $K_{\min} = 8$ ,



**Fig. 11.** Number of total messages versus  $K_{\min}$  for different network densities in the circular scenario,  $p_{\text{rec}} = 1$ .



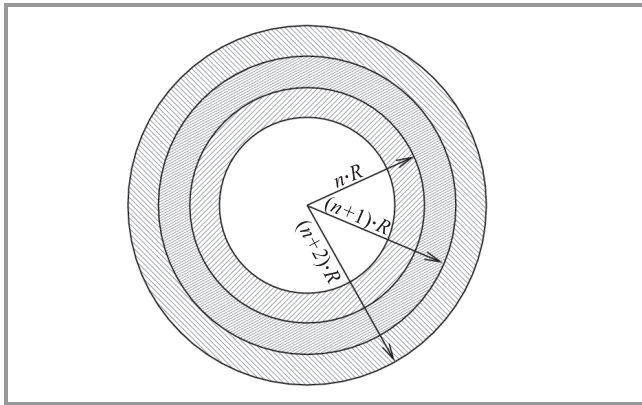
**Fig. 12.** Distribution of the hop counts in the uniform distribution network for different  $K_{\min}$  values,  $p_{rec} = 1$ .

the algorithm gives a gain of 55–95%, with respect to the original flooding.

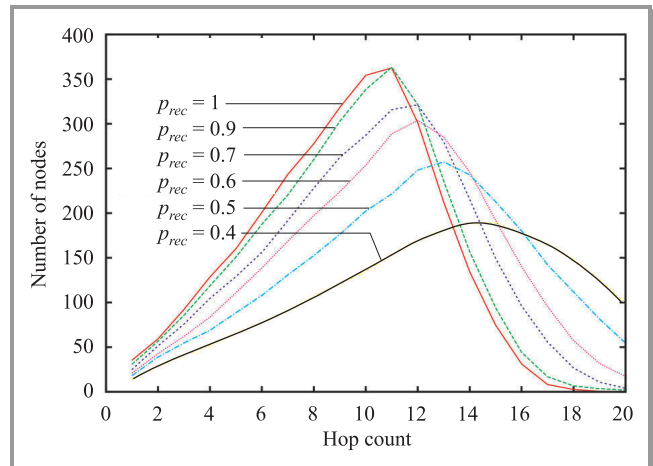
The length  $L$  of the discovered route is important when flooding is used for route discovery. To test the properties of percolation driven flooding in this respect we measured hop distances discovered by flooding and percolation driven flooding, for different  $K_{\min}$  values. In the tests the uniform distribution (constant density) scenario was used with  $K = 35$  and  $p_{rec} = 1$ . The source node was placed to the center of the field and hop distances to all other nodes were measured. The histogram created from averaging 100 independent experiments is illustrated in Fig. 12.

The distribution is linear for small hop distances. The explanation is shown in Fig. 13 for the ideal case: messages are propagated in belt-shaped increments. The areas of the subsequent belts increase linearly, thus the number of nodes in each belt is linearly increasing. If the node density is smaller, the number of nodes in each belt is smaller, thus the slope of the histogram will be smaller. According to Fig. 12, smaller  $K_{\min}$  also causes decrease in the slope: e.g., for  $K_{\min} = 9$ , the slope is 40% smaller than in the basic flooding case. This means that the detected route is 40% longer than that of the basic flooding.

The declining part at higher hop distances is due to the finite size of the network: the messages reach the network



**Fig. 13.** Ideal propagation of a message in a dense network.



**Fig. 14.** Distribution of the hop counts in the uniform distribution network for different  $p_{rec}$  values,  $K_{\min} = 9$ .

perimeter and eventually die. The tail of the distribution presents the maximum hop distance from the source, which is in the square setup ideally  $\sqrt{2}$  times the hop distance measured at the maximum of the histogram.

The effect of unreliable links is shown in Fig. 14. The histogram shows the average of 100 experiments conducted in the uniform test scenario with  $K = 35$ , and  $p_{rec} = 0.4..1$ . As intuitively expected, unreliable links cause smaller slope, i.e., higher hop distances. For example, link quality  $p_{rec} = 0.7$  results in the increase of hop distance by 18%, while for the rather unreliable link quality  $p_{rec} = 0.4$  the increase was as high as 63%.

As simulation examples show, percolation driven flooding results in considerably higher hop distances than basic flood routing when small  $K_{\min}$  parameters are used. Naturally, low link quality also causes larger hop distances. This effect may be an undesired side effect in route discovery protocols, thus in these applications basic flooding may be more advantageous.

## 5. Conclusions

A percolation driven flood routing algorithm was proposed, which uses only locally available neighborhood information to reduce broadcast storm. The algorithm is able to massively reduce the number of messages in the network and maintaining high delivery ratio at the same time. Theoretical results prove the usefulness of the algorithm: it is able to provide high coverage, if the network density is high enough.

Simulation tests were performed to validate the performance of the algorithm. The proposed algorithm reduced the total number of messages in the network by 30–95%, depending on the network density, while the coverage was almost 100%, with appropriate choice of  $K_{\min}$  ( $K_{\min} > 7$ ). The percolation driven flood algorithm is adaptive to changing node density thus provides high coverage with low number of messages in all scenarios. The algorithm is robust in the presence of unreliable links as well.

According to test result, the percolation driven flood routing algorithm results considerably (even 50%) longer hop distances when used for route discovery purposes.

The overhead of the algorithm is very low since only the number of neighbors must be known by each node. This information can be gained by simple neighborhood discovery protocols, e.g., sending and receiving *hello messages*.

For route discovery purposes the application of the traditional flooding is more advantageous, if the side effect of the longer routes is undesirable. On the other hand, the proposed algorithm is a superior and robust alternative to traditional flooding algorithm in dense networks for message dissemination purposes.

## References

- [1] A. Mainwaring, D. Culler, J. Polastre, R. Szewczyk, and J. Anderson, "Wireless sensor networks for habitat monitoring", in *Proc. WSN Conf.*, Atlanta, USA, 2002, pp. 88–97.
- [2] A. Arora *et al.*, "A line in the sand: a wireless sensor network for target detection, classification, and tracking", *Comput. Netw.*, vol. 46, no. 5, pp. 605–634, 2004.
- [3] A. Ledeczi *et al.*, "Countersniper system for urban warfare", *ACM Trans. Sen. Netw.*, vol. 1, no. 2, pp. 153–177, 2005.
- [4] T. He *et al.*, "VigilNet: an integrated sensor network system for energy-efficient surveillance", *ACM Trans. Sen. Netw.*, vol. 2, no. 1, pp. 1–38, 2006.
- [5] B. Karp and H. T. Kung, "Greedy perimeter stateless routing (GPSR) for wireless networks", in *Proc. ACM MobiCom Conf.*, Boston, USA, 2000, pp. 243–254.
- [6] Y. B. Ko and N. H. Vaidya, "Location-aided routing (LAR) in mobile ad hoc networks", in *Proc. ACM MobiCom Conf.*, Dallas, USA, 1998, pp. 66–75.
- [7] S. Basagni, I. Chlamtac, V. R. Syrotiuk, and B. A. Woodward, "A distance routing effect algorithm for mobility (DREAM)", in *Proc. ACM MobiCom Conf.*, Dallas, USA, 1998, pp. 76–84.
- [8] C. Intanagonwiwat, R. Govindan, D. Estrin, J. Heidemann, and F. Silva, "Directed diffusion for wireless sensor networking", *IEEE/ACM Trans. Netw.*, vol. 11, no. 1, pp. 2–16, 2003.
- [9] C. E. Perkins and E. M. Royer, "Ad-hoc on-demand distance vector routing", in *Proc. 2nd IEEE Worksh. Mob. Comput. Syst. Appl.*, New Orleans, USA, 1999, pp. 90–100.
- [10] D. B. Johnson and D. A. Maltz, *Dynamic Source Routing in Ad Hoc Wireless Networks*. Boston: Kluwer, 1996.
- [11] V. Park and M. S. Corson, "A highly adaptive distributed routing algorithm for mobile wireless networks", in *Proc. IEEE Infocom Conf.*, Kobe, Japan, 1997, pp. 1405–1413.
- [12] Z. Haas and M. Pearlman, "The performance of query control schemes for the zone routing protocol", in *Proc. ACM SIGCOMM Conf.*, Vancouver, Canada, 1998, pp. 167–177.
- [13] W. Heinzelman, A. Chandrakasan, and H. Balakrishnan, "Energy-efficient communication protocol for wireless sensor networks", in *Proc. Hawaii Int. Conf. Syst. Sci.*, Hawaii, USA, 2000.
- [14] A. Manjeshwar and D. P. Agrawal, "TEEN: a protocol for enhanced efficiency in wireless sensor networks", in *Proc. 1st Int. Worksh. Paral. Distrib. Comput. Iss. Wirel. Netw. Mob. Comput.*, San Francisco, USA, 2001.
- [15] Z. J. Haas, J. Y. Halpern, and L. Li, "Gossip-based ad hoc routing", *IEEE/ACM Trans. Netw.*, vol. 14, no. 3, pp. 479–491, 2006.
- [16] M. Maróti, "Directed flood-routing framework for wireless sensor networks", in *Proc. 5th ACM/IFIP/USENIX Int. Conf. Middlew.*, Toronto, Canada, 2004, pp. 99–114.
- [17] L. H. Costa, M. D. De Amorim, and S. Fdida, "Reducing latency and overhead of route repair with controlled flooding", *Wirel. Netw.*, vol. 10, no. 4, pp. 347–358, 2004.
- [18] P. Gburzynski and W. Olesinski, "On a practical approach to low-cost ad hoc wireless networking", *J. Telecommun. Inform. Technol.*, no. 1, pp. 29–42, 2008.
- [19] Zigbee Alliance [Online]. Available: <http://www.zigbee.org>
- [20] Y. C. Tseng, S. Y. Ni, and E. Y. Shih, "Adaptive approaches to relieving broadcast storms in a wireless multihop mobile ad hoc network", in *Proc. 21st Int. Conf. Distrib. Comput. Syst.*, Phoenix, USA, 2001, pp. 481–488.
- [21] J. Yang, B. Kim, M.-T. Sun, and T.-H. Lai, "Location-aided broadcast in wireless ad hoc networks", *J. Inform. Sci. Eng.*, vol. 23, no. 3, pp. 869–884, 2007.
- [22] J. Wu, F. Dai, M. Gao, and I. Stojmenovic, "On calculating power aware connected dominating set for efficient routing in ad hoc wireless networks", *J. Commun. Netw.*, vol. 5, no. 2, pp. 169–178, 2002.
- [23] M. Franceschetti, L. Booth, M. Cook, R. Meester, and J. Bruck, "Percolation in multi-hop wireless networks", in *Caltech Paradise, ETR055, 2003* [Online]. Available: <http://www.paradise.caltech.edu/papers/etr055.pdf>
- [24] G. Grimmet, *Percolation*. New York: Springer, 1989.
- [25] I. Glauche, W. Krause, R. Sollacher, and M. Greiner, "Continuum percolation of wireless ad hoc communication networks", *Physica A*, vol. 325, no. 3–4, pp. 577–600, 2003.
- [26] J. Jonasson, "Optimization of shape in continuum percolation", *Ann. Probab.*, vol. 29, no. 2, pp. 624–635, 2001.
- [27] O. Dousse, P. Mannersalo, and P. Thiran, "Latency of wireless sensor networks with uncoordinated power saving mechanisms", in *Proc. 5th ACM Int. Symp. MobiHoc'04*, Tokyo, Japan, 2004, pp. 109–120.
- [28] Z. Kong and E. M. Yeh, "Distributed energy management algorithm for large-scale wireless sensor networks", in *Proc. 8th ACM Int. Symp. MobiHoc'07*, Montreal, Canada, 2007, pp. 209–218.



**Gergely Vakulya** received his M.Sc. in information technology from the University of Pannonia, Hungary, in 2008. Currently he is pursuing his Ph.D. in information science and technology at the University of Pannonia. His research interest is sensor networking.

e-mail: [vakulya@dcs.uni-pannon.hu](mailto:vakulya@dcs.uni-pannon.hu)  
University of Pannonia  
Egyetem u. 10  
8200 Veszprém, Hungary



**Gyula Simon** received his M.Sc. and Ph.D. in electrical engineering from the Budapest University of Technology, Hungary, in 1991 and 1998, respectively. Currently he is an Associate Professor at the University of Pannonia. His main research interest includes adaptive signal processing and sensor networks.

e-mail: [simon@dcs.uni-pannon.hu](mailto:simon@dcs.uni-pannon.hu)  
University of Pannonia  
Egyetem u. 10  
8200 Veszprém, Hungary

# Multiobjective Design of Wireless Ad Hoc Networks: Security, Real-Time and Lifetime

Zdravko Karakehayov

**Abstract**—This paper deals with the tradeoffs between security, real-time and lifetime performance. Due to the multihop nature of communication wireless ad hoc networks are very vulnerable to attacks. Malicious nodes included in a routing path may misbehave and organize attacks such as black holes. Scaling the number of hops for a packet delivery we trade off energy efficiency against security and real-time communication. To study the multihop communication we propose a hierarchical communication model. The REWARD (receive, watch, redirect) algorithm for secure routing is employed as a main example for corrective actions. Symmetrical routing is a distinguish feature of protocols such as REWARD and we outline the threshold of conflict between power-efficient partitioning of communication links and symmetrical routing.

**Keywords**—*ad hoc networks, low-power routing, multihop communication, secure routing.*

## 1. Introduction

Ad hoc networks have a wide spectrum of military and commercial applications. Ad hoc networks are employed in situations where installing an infrastructure is too expensive, too vulnerable or the network is transient. The interaction between the nodes is based on wireless communication. Packets are forwarded in a multihop manner. Nodes have a limited radio footprint and when a node receives a packet it applies a routing algorithm to select a neighbor for forwarding.

There is a class ad hoc networks, sensor networks, where the requirements for lifetime and size of the nodes are driven to extremes. A wireless sensor network consists of a large number of nodes that may be randomly and densely deployed. Sensor nodes are capable of sensing many types of information such as temperature, light, humidity and radiation. Sensor networks must collect data in an area of interest for months or years. Since the energy is a scarce and usually non-renewable resource, the network's functionality must be viewed from a low-power perspective. Sensor network nodes execute three major tasks: sensing, computation and communication.

Communication energy dominates the overall energy budget. The greater than linear relationship between transmit energy and distance promises to reduce the energy cost when the radio link is partitioned. Nodes calculate the distance and tune their transmit power accordingly. Consequently, it would be beneficial to use several hops to reach a node within the transmission radius instead of

a direct link. Along with available locations of the nodes, a multihop optimization requires an appropriate power model. For some applications it is not necessary nodes to have real coordinates. Instead, nodes may have virtual coordinates: hop-distances to other nodes.

Moreover, some applications require the network to influence the environment via actuators. Synchronization between input and output demands real-time traffic. Real-time forwarding of packets under multihop communication scheme is a serious challenge. When we factor in security, the outlook becomes even more grim. Packets travel over several nodes and malicious attacks are easy to organize. To detect malicious influence and wage corrective actions the nodes must spend extra energy. Consequently, the multihop nature of ad hoc networks, while beneficial for energy reduction, brings the packets delivery time up. The dynamic nature of the network and the power-efficient partitioning of communication links in particular, often result in unpredictable traffic timing parameters. Enemy nodes included in a routing path may misbehave and any attempt to make the network less vulnerable requires extra energy and affects the lifetime, thus closing the loop.

## 2. Related Work

Different medium access control (MAC) protocols are discussed in [1]–[6]. Energy efficiency is the primary goal of the research. While a power saving technique, termed Span [1], dynamically splits the nodes into sleeping nodes and forwarding nodes, S-MAC, a MAC protocol [2], establishes a low duty cycle operation in all nodes. Extremely opportunistic routing (ExOR) is a routing method developed to reduce the total number of transmissions taking into account the actual packet propagation [3]. Data transmission algebra (DTA) has been developed to generate complex transmission schedules based on collision-free concurrent data transmissions [5]. In related research we proposed ALS-MAC, a medium access control protocol where contention-based advertising slots are mapped to scheduled-based transmission slots [6]. The energy model employed in this paper has been adopted from [7], [8]. Despite there being a plethora of sensing and MAC papers, comparatively little has been published on the companion task of actuation and real-time requirements. Sensor-actuator networks are discussed in [9], [10]. The problem of obtaining virtual coordinates is addressed in [11].

Different aspects of node architectures and capabilities can be found in [12]–[17]. The power reduction methods discussed in [15]–[17] are not confined to computation energy of network nodes. They can be applied, also, in other cases where voltage-scalable or speed-scalable central processing units (CPUs) follow the current requirements and save energy. Another approach to reduce the power consumption is to remove hardware used for localization, such as global positioning system (GPS), and utilize receive signal strength (RSS). The resulting accuracy and impact factors are investigated in [14].

Methods for energy efficient multihop communication are discussed in [18]–[22]. A detailed investigation for simple settings is available in [19]. In related research we studied multihop optimization for non-regular topologies [6], [10]. An Aloha type access control mechanism for large, multihop, wireless networks is defined in [21]. The protocol optimizes the product of the number of simultaneously successful transmissions per unit of space, spatial reuse, by the average range of each transmission.

A review of routing protocols for wireless ad hoc networks is available in [23]. The problem of radio irregularity is discussed in [24]. Later in Section 5, we compare distances with the communication range. Due to radio irregularity some neighbors located within the transmission disk may be inaccessible while some remote nodes, outside the disk, will be capable to communicate. Since quite a few processor architectures vie for attention in the realm of sensor networks, target-aware modeling of routing algorithms helps to evaluate important timing properties [25]. Security of wireless sensor networks is in focus in [26]–[31]. Two papers, [22] and [30], emphasize the fact that multiobjective design is needed. Listening to neighbor transmissions to detect black hole attacks is discussed in [32]–[36].

### 3. Communication Model

The communication model describes a packet forwarding from a source to a destination. The destination is within the communication range of the source. The communication model  $C$ , has three components: a set of the locations of nodes  $L$ , a medium access control model  $M$ , and an energy model  $E$ :

$$C = \{L, M, E\}. \quad (1)$$

#### 3.1. Medium Access Control Model

Medium access control mechanism has a significant impact on the energy efficiency [2], [4], [6]. Currently available MAC protocols for wireless sensor networks can be broken down into two major types: contention-based and scheduled-based. While under contention-based protocols nodes compete among each other for channel access, scheduled-based schemes rely on prearranged collision-free links between nodes. There are different methods to assign collision-free links to each node. Links may be assigned as time slots, frequency bands, or spread spectrum codes. However, size and cost constraints may not permit allocat-

ing complex radio subsystems for the node architecture. Logically, time-division multiple access (TDMA) scheduling is the most common scheme for the domain of wireless sensor networks. The limited communication range of network nodes provides an extra opportunity for collision-free interaction, space division access [5], [6], [21].

##### 3.1.1. Assume Scheduled Links

In order to save energy nodes should stay in a sleeping mode as long as possible. Ideally, nodes should have prearranged collision-free links and wake up only to exchange packets. This MAC approach can be termed assume scheduled links (ASL). The ASL model has two parameters: a packet length in bits  $p$  and a bit rate  $B$ :

$$M = \{ASL, p, B\}. \quad (2)$$

While ASL is a theoretical concept, it helps to outline the floor of the energy required for communication.

##### 3.1.2. Beacon Advertise Transmit

Beacon advertise transmit (BAT) model is a widespread MAC mechanism [4]. Beacons are employed to synchronize internode communications. A beacon period  $T_B$  includes two major sections. The period begins with a traffic indication window  $T_A$ . During  $T_A$  all nodes are listening and pending packets are advertised. The nodes addressed till the end of  $T_A$  send acknowledgements and receive data packets. Data transmissions are followed by acknowledgement frames to confirm successful reception. Figure 1 illustrates a beacon period.

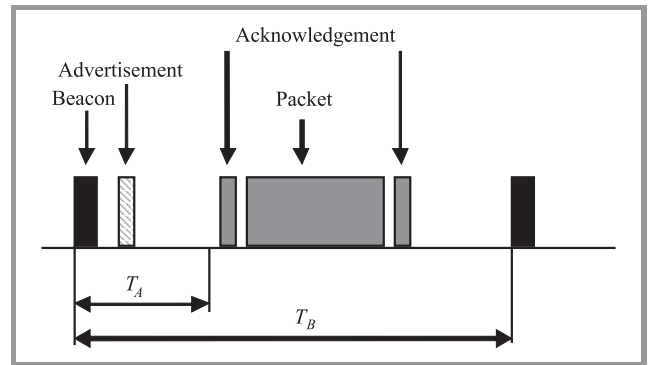


Fig. 1. Beacon period.

The BAT model has five parameters:  $T_A$ ,  $T_B$ , a data packet length in bits  $p$ , a control packet length in bits  $q$ , and a bit rate  $B$ :

$$M = \{BAT, T_A, T_B, p, q, B\}. \quad (3)$$

#### 3.2. Energy Model

The energy used to send a bit over a distance  $d$  via radio communication may be written as

$$E = ad^n + b, \quad (4)$$

where  $a$  is a proportionality constant [7], [8]. The radio parameter  $n$  is a path loss exponent that describes the rate

at which the transmitted power decays with increasing distance. Typically,  $n$  is between 2 and 4. The  $b$  constant is associated with specific receivers, CPUs and computational algorithms. Thus the model emerges as

$$\mathbf{E} = \{a, n, b, P_R\}, \quad (5)$$

where  $P_R$  is the power consumption of a turned on receiver.

## 4. Real-Time Behavior

Using the BAT model and counting the beacon periods nodes are in position to calculate the packets delivery time. While this completely applies for destination nodes, intermediate nodes can use the actual packet propagation time and virtual coordinates to foresee the overall delivery time.

In the large, energy versus real-time tradeoffs can be resolved via different values assigned for the beacon period. In the small, at each hop nodes decide whether to include an extra intermediate node for power efficiency or to forward the packet as fast as possible. The local decision is based on the actual propagation of the packet measured in number of beacon periods and the remaining number of hops.

## 5. Lifetime

An ad hoc network lifetime can be measured by the time when the first node runs out of energy, or a network can be declared dead when a certain fraction of nodes die. Alternatively, the system lifetime can be measured by application-specific parameters, such as the time until the system can no longer provide acceptable quality of service. Clearly, the higher the energy efficiency is, the longer the network will survive. The energy efficiency can be optimized at three levels.

### 5.1. Node Architecture

A typical node is built around a low-power microcontroller [12], [13], [15]. Wireless transceivers create physical links between nodes. Hardware provides the following low-power mechanisms. The receiver and transmitter can be individually enabled and disabled. The transmit power can be adjusted gradually. For many applications nodes are capable of determining their coordinates. Voltage-scalable systems may apply dynamic voltage or clock frequency scaling to reduce the power consumption.

### 5.2. Multihop Routing Service

Once the routing protocol has provided the next relay another neighbor can be considered to partition the link. The number of hops is increased to save energy. As an additional benefit, the reduced transmit power allows better spatial reuse.

Figure 2 shows how an intermediate node can be used to break down the link between a source  $S$  and a destination  $D$  into two hops.

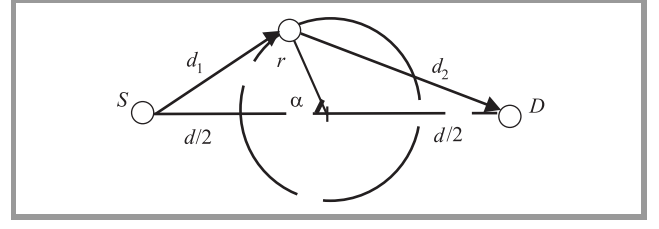


Fig. 2. Routing via an intermediate node.

**Theorem 1:** Let  $\mathbf{C} = \{\mathbf{L}\{\text{ASL}, p, B\}, \{a, 4, b, P_R\}\}$  be the communication model of a wireless ad hoc network. If the distance between the source  $S$  and the destination  $D$  is  $d \geq ((8b + (p/B)P_R)/7a)^{1/4}$  and the distance between an intermediate node and the halfway point between  $S$  and  $D$  is  $r \leq (-0.75d^2 + 0.25(9d^4 - a^{-1}(8b - 7ad^4 + (p/B)P_R)))^{1/2}$ , the two-hop communication requires less energy than the direct link.

*Proof:* We must prove when the following inequality holds

$$ad_1^4 + b + ad_2^4 + b + 2(p/B)P_R \leq ad^4 + b + (p/B)P_R. \quad (6)$$

Taking into account that

$$d_1 = (d^2/4 - dr\cos\alpha + r^2)^{1/2}, \quad (7)$$

$$d_2 = (d^2/4 + dr\cos\alpha + r^2)^{1/2}. \quad (8)$$

We get

$$16ar^4 + 8ad^2(1 + 2\cos^2\alpha)r^2 + 8b - 7ad^4 + (p/B)P_R \leq 0. \quad (9)$$

The inequality has solutions if and only if  $d \geq ((8b + (p/B)P_R)/7a)^{1/4}$ . Since the threshold value for the distance  $r$  will vary with  $\alpha$ , we take the worst case,  $\cos\alpha = 1$ .

Using the quadratic formula

$$r \leq (-0.75d^2 + 0.25(9d^4 - a^{-1}(8b - 7ad^4 + (p/B)P_R)))^{1/2}. \quad (10)$$

□

Figure 3 shows plots for the radius  $r$  compared with half of the distance. This example assumes two bit rates, 1 Mbit/s and 0.5 Mbit/s,  $a = 0.2$  fJ/m<sup>4</sup>,  $b = 1$  nJ,  $P_R = 10$  mW and  $p = 128$  bit.

**Theorem 2:** Let  $\mathbf{C} = \{\mathbf{L}\{\text{BAT}, T_A, T_B, p, q, B\}, \{a, 4, b, P_R\}\}$  be the communication model of a wireless ad hoc network. Let the average number of neighbors listening to a beacon transmission be  $D$ . If the distance between the source  $S$  and the destination  $D$ :

$$d \geq ((b(3q + p) + P_R B^{-1}(q + p) + P_R D T_A) a^{-1} (3.5625q + 0.875p)^{-1})^{1/4} \quad (11)$$

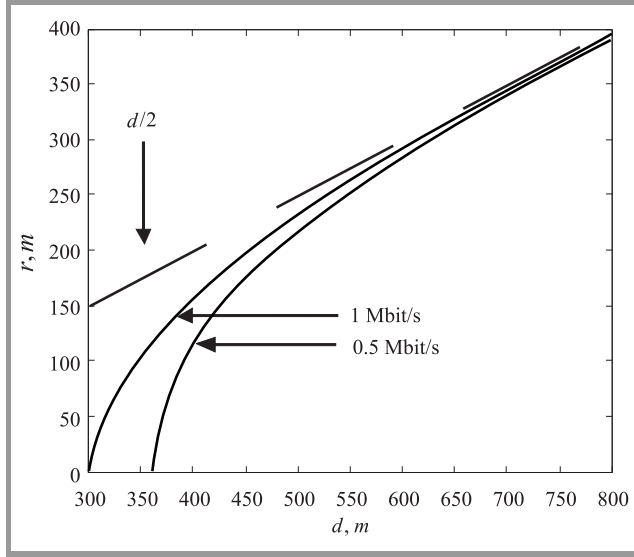


Fig. 3. Radius  $r$  scales with the distance for two bit rates.

and the distance between the intermediate node and the halfway point between  $S$  and  $D$

$$r \leq (-0.25d^2(10.5q + 3p + 0.5qd)(3q + p + qd)^{-1} + 0.5a^{-1}(3q + p + qd)^{-1}(0.25a^2d^2(10.5q + 3p + 0.5qd)^2 - 2a(3q + p + qd)(-ad^4(3.5625q + 0.875p) + b(3q + p) + P_R B^{-1}(q + p) + P_R D T_A))^{1/2})^{1/2} \quad (12)$$

the two-hop communication requires less energy than the direct link.  $\square$

The radius  $r$  for a given distance  $d$  indicates application-specific opportunities for power-efficient partitioning of

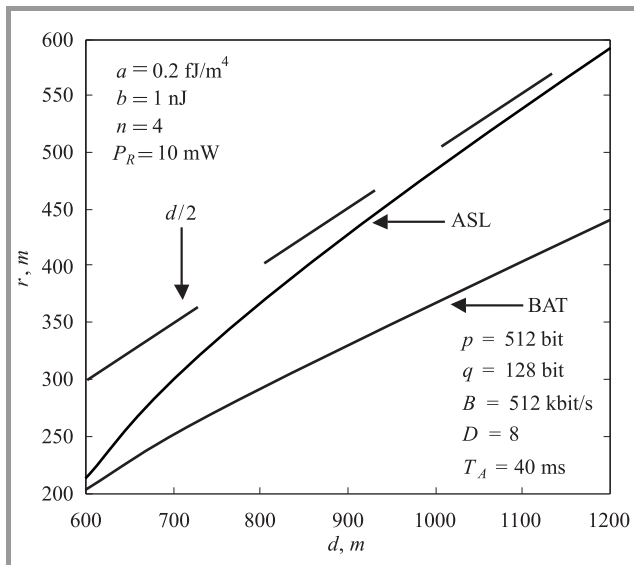


Fig. 4. Radius  $r$  scales with the distance for two MAC models.

communication links. Figure 4 compares ASL and BAT MAC models for a bit rate of 512 kbit/s.

### 5.3. Routing Algorithms

Routing algorithms can be based on two major approaches: topology-based and position-based routing [23]. The topology-based algorithms can be further split into table-driven and demand-driven. The main idea behind the table-driven routing protocols is to create a clear picture of all available routes from each node to every other node in the network. In contrast to the table-driven protocols, the demand-driven algorithms create routes via route discovery procedures only when a necessity arises.

Position-based routing algorithms utilize the physical positions of the participating nodes [19], [21], [23]. Position-based or geographic routing does not require each node to have the locations of all other nodes. Each node keeps track of the coordinates of its neighbors and their neighbors. A greedy routing algorithm based on geographic distance selects the closest to the destination neighbor for the next hop [19].

Assume that the nodes of a wireless ad hoc network are members of the following set  $\mathbf{N} = \{N_1, N_2, N_3, \dots, N_{n(N)}\}$ . The nodes are placed in a rectangular region of  $X$  by  $Y$ . The distance between node  $i$  and node  $j$  is  $d(i, j)$ . The distance between node  $k$  and the halfway point between node  $i$  and node  $j$  is  $d(k, m_{i,j})$ .

Routing algorithms are employed to determine the next hop of  $N_i$ ,  $N_i^{+1}$ . The distance between  $N_i$  and its next hop  $N_i^{+1}$  is  $d(i, +1)$ . Likewise, the distance between  $N_k$  and the halfway point between  $N_i$  and  $N_i^{+1}$  is  $d(k, m_{i,+1})$ . A statement **power** ( $d(i, j)$ ) in the pseudocode listing adjusts the transmit power according to the distance  $d(i, j)$ . A statement **send** ( $N_i \rightarrow N_j$ ) indicates a packet forwarding from node  $i$  toward node  $j$ .

Algorithm 1 describes the procedure to determine the set  $N_i^R$ , which includes the one-hop neighbours of  $N_i$ . The  $R$  denotes the communication range.

---

#### Algorithm 1: $N_i^R \leftarrow \text{OneHop}(N_i)$

---

```

1  $N_i^R = \emptyset$ 
2 for  $1 \leq j \leq n(N)$ ,  $j \neq i$  do
3   if  $d(i, j) \leq R$ 
4      $N_i^R = N_i^R \cup N_j$ 
5   end if
6 end for

```

---

Algorithm 2 applies the greedy routing algorithm to find the next relay of  $N_i$ .

**Algorithm 2:**  $N_i^{+1} \leftarrow \text{NextHop}(N_i, N_D, N_i^R)$ 


---

```

1 if  $N_D \in N_i^R$ 
2   return  $N_D$ 
3 end if
4  $s = (X^2 + Y^2)^{\frac{1}{2}}$ 
5 for  $1 \leq j \leq n(N)$ ,  $j \neq i$  do
6   if  $N_j \in N_i^R$  and  $d(j, D) < s$ 
7      $N_i^{+1} = N_j$ ,  $s = d(j, D)$ 
8   end if
9 end for

```

---

The multihop service can be integrated into the routing algorithm.

Algorithm 3 applies Theorem 1 or 2 to partition the communication link until suitable intermediate nodes are found. The procedure results in one forwarding.

**Algorithm 3:**  $\text{MultiHop}(N_i, N_i^{+1})$ 


---

```

1 do
2   MULTI = 0
3    $d = d(i, +1)$ 
4    $s = (X^2 + Y^2)^{\frac{1}{2}}$ 
5   for  $1 \leq j \leq n(N)$ ,  $j \neq i$  do
6     if  $d(j, m_{i,+1}) \leq \text{MIN}(r, s)$ 
7        $s = d(j, m_{i,+1})$ ,  $N_i^{+1} = N_j$ , MULTI = 1
8     end if
9   end for
10 while MULTI
11 power ( $d(i, +1)$ )
12 send ( $N_i \rightarrow N_i^{+1}$ )

```

---

Algorithm 4 describes the successive approximation routing. The interaction between the routing procedure and the low-power forwarding is implemented via successive approximations. As soon as the routing algorithm determines the next hop, multihop optimization is applied to select

**Algorithm 4:**  $\text{Send}(N_S, N_D)$ 


---

```

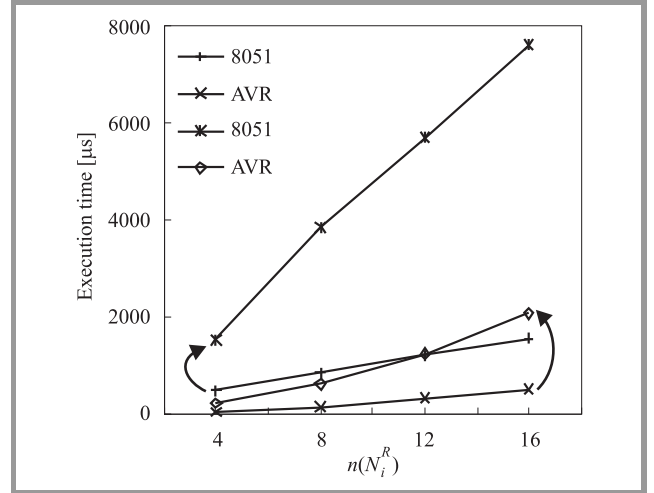
1  $N_i = N_S$ 
2 do
3    $\text{NextHop}(N_i, N_D, N_i^R)$ 
4    $\text{MultiHop}(N_i, N_i^{+1})$ 
5 while  $N_i \neq N_D$ 

```

---

an intermediate node. As soon as the packet is sent to the intermediate node, the routing algorithm is executed again. The multihop service algorithm itself is a successive approximation procedure as well.

In a two-hop distance approach, each node maintains a table of all immediate neighbors as well as each neighbor's neighbors. The number of hops taken into account determines the vulnerability of the routing in case of topology holes. However, considering more hops will require longer execution times. Figure 5 shows how the transition from



**Fig. 5.** Execution time to select the next relay.

a single hop to two hops brings the execution time up. The code has been written in C and compiled for two CPUs: 8051 and Atmel AVR [25].

## 6. Security

The network functional partitioning into sensing, computation and communication can be used to deal with possible avenues of attacks. First, a misbehaving node may provide false sensor readings. In general, this kind of attack is not effective. Collected data is aggregated and a small number of malicious nodes can not change the profile of the physical event. However, a false alarm, an input has reached a threshold, will wake up several nodes and attack the batteries. Another attack related to the environment is a wrong location. Sensing is useful only in the context of where the data has been measured.

In contrast to sensing, a well placed enemy may successfully attack via wrong calculations. Aggregation is important for power efficiency and nodes that aggregate data packets are in a good position to attack.

Communication is what makes ad hoc networks most vulnerable and the multihop forwarding of packets unrolls ample possibilities for attackers. Once a malicious node has been included on the routing path, it will be in position to change the content of the packets. Along with data, packets may convey code. Mobile agent-based sensor networks distribute the computation into the participating leaf

nodes [28], [29]. Since agents may visit a long path of nodes, a single modified packet can force several nodes to execute enemy code. Another axis along which packets can be affected relates to timing. A scheduling attack would change the number of past beacon periods a packet carries. Another form of a scheduling attack is delayed packets. An extreme type of this attack, termed black hole, is observed when a malicious node consumes packets. In a special case of black hole, an attacker could create a gray hole, in which it selectively drops some packets but not others. For example, the malicious node may forward control packets but not data packets.

## 7. REWARD Algorithm

The REWARD (receive, watch, redirect) is a routing method that provides a scalable security service for geographic ad hoc routing [33]–[35].

### 7.1. Black Holes Data Base

The algorithm creates a distributed data base for detected black hole and scheduling attacks. The data base keeps records for suspicious nodes and areas. The REWARD security service provides alternative paths for the geographic routing in an attempt to avoid misbehaving nodes and regions of detected attacks. The algorithm utilizes two types of broadcast messages, MISS (material for intersection of suspicious sets) and SAMBA (suspicious area, mark a black-hole attack), to recruit security servers. Security servers are nodes that keep records of the distributed data base and modify the geographic forwarding of packets to bypass insecure nodes and regions.

Assume that a demand-driven protocol performs a route discovery procedure. When the destination receives the query, it sends its location back and waits for a packet. If the packet does not arrive within a specified period of time, the destination node broadcasts a MISS message. The destination copies the list of all involved nodes from the query to the MISS message. Since the reason for not receiving the packet is most likely a black hole attack, all nodes listed in the MISS message are under suspicion. Nodes collect MISS messages and intersect them to detect misbehaving participants in the routes. The detected malicious nodes are excluded from the routing if other paths are available.

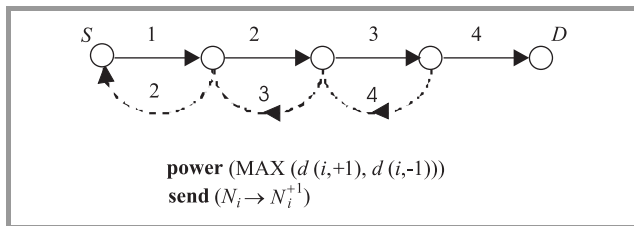


Fig. 6. Transmissions must be received by two nodes.

Radio is inherently a broadcast medium and nodes can detect black hole attacks if they listen to neighbor transmissions [32]. Figure 6 shows an example. Each node tunes the transmit power to reach both immediate neighbors,  $N_i^{+1}$  and  $N_i^{-1}$ . We call this type of forwarding symmetrical. The nodes transmit packets and watch if the packets are forwarded properly. If a malicious node does not act as expected, the previous node in the path will broadcast a SAMBA message.

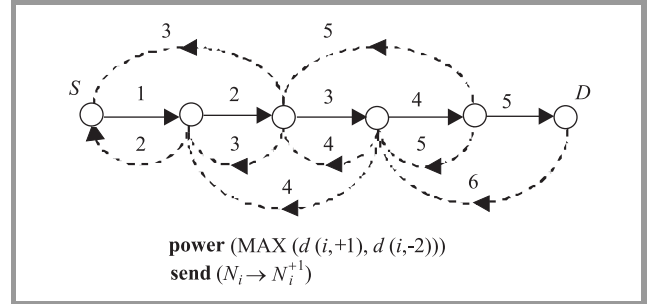


Fig. 7. REWARD against two black holes.

Figure 7 presents an example routing with the assumption that two malicious nodes would attempt a black hole attack. In this case the algorithm requires the nodes to listen for two retransmissions. Figure 8 indicates the exact positions of two black holes in the path. The first malicious node forwards the packet using the required transmit power to deceive two nodes backward. The second malicious node drops the packet, however the attack is detected by the last node before the black holes. The missing transmission is shown by a dot line in Fig. 8. An extra black hole in the path would mask the attack.

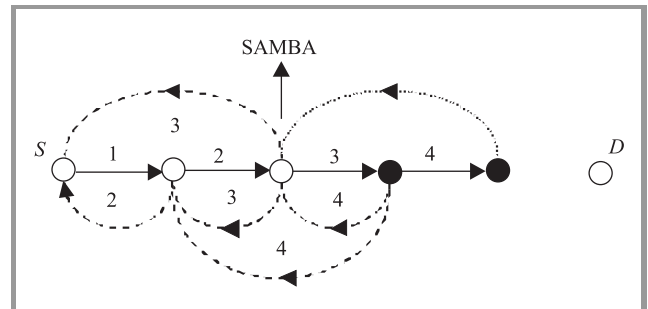


Fig. 8. REWARD detects the second black hole.

In order to determine the effectiveness of REWARD we used ANTS (ad hoc networks traffic simulator) [34], [35]. We assume that all nodes are stationary throughout the simulation. Figure 9 illustrates simulation results of the throughput, 100 packets routing for eight example deployments. Each deployment has a density of 100 nodes randomly located in a square kilometer. The maximum communication range of the nodes is 100 m. Also, the simulation results are obtained at 10% misbehaving nodes. MISS servers are recruited in a rectangular region.

The source and destination locations define the diagonal of the rectangle.

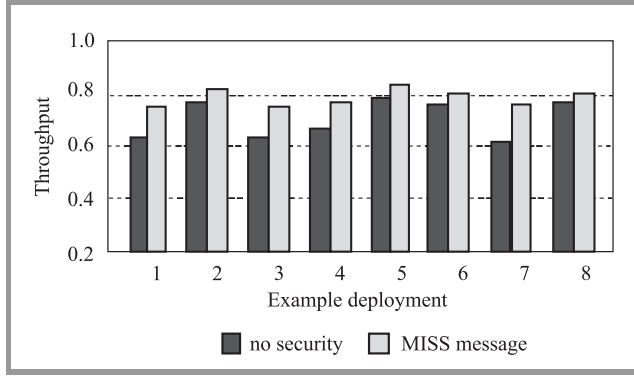


Fig. 9. The fraction of packets received for eight examples.

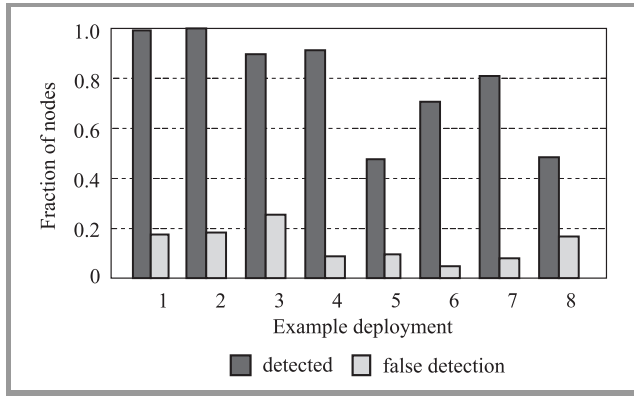


Fig. 10. Detected malicious nodes against false detection.

Figure 10 shows the fraction of malicious nodes detected against false detection. False detection is associated with nodes excluded from the network as malicious when in fact they are not. For the current simulation, nodes that are listed in two or more MISS messages are marked as malicious.

## 7.2. Energy Overhead

We distinguish between two types of security energy overhead. Static overhead is the additional energy required to watch for attacks. Dynamic overhead is the extra amount of energy spent to detect compromised nodes and mitigate routing misbehavior. While the dynamic overhead will vary from application to application, the static overhead is a constant and an inevitable item in the energy budget.

Since secure routing protocols such as REWARD require symmetrical forwarding, the power efficiency is declined. Figure 11 presents symmetrical routing for an example deployment. Three cases must be considered according to the distances:

$$d(i, -1) \leq (d(i, +1))/2 - r. \quad (13)$$

There is no security overhead in this case:

$$(d(i, +1))/2 - r < d(i, -1) \leq (d(i, +1))/2 + r. \quad (14)$$

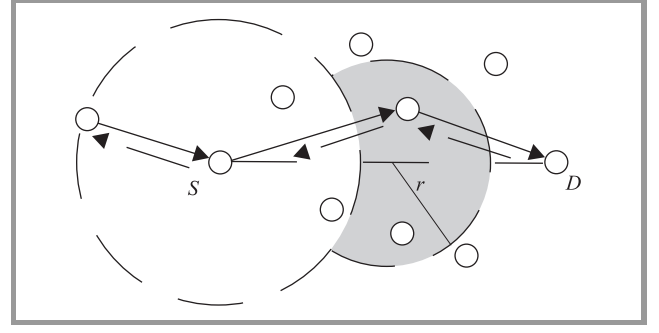


Fig. 11. Symmetrical routing.

Again, there is no single-hop security overhead. Opportunities for partitioning of the link remain if neighbors are located within the shaded area (Fig. 11):

$$d(i, -1) > (d(i, +1))/2 + r. \quad (15)$$

Symmetrical routing may not increase the energy, however, partitioning of the link is not power efficient in this case.

Algorithm 5 provides multihop optimization for symmetrical routing.

---

### Algorithm 5: MultiHopSym ( $N_i, N_i^{+1}$ )

---

```

1  $s = (X^2 + Y^2)^{\frac{1}{2}}$ 
2 if  $d(i, -1) > (d(i, +1))/2 + r$ 
3   power ( $\text{MAX}(d(i, +1), d(i, -1))$ )
4   send ( $N_i \rightarrow N_i^{+1}$ )
5   return
6 end if
7 if  $d(i, -1) \leq (d(i, +1))/2 - r$ 
8   for  $1 \leq j \leq n(N), j \neq i$  do
9     if  $d(j, m_{i+1}) \leq \text{MIN}(r, s)$ 
10       $s = d(j, m_{i+1}), N_i^{+1} = N_j$ 
11    end if
12  end for
13  power ( $d(i, +1)$ )
14  send ( $N_i \rightarrow N_i^{+1}$ )
15  return
16 end if
17 for  $1 \leq j \leq n(N), j \neq i$  do
18   if  $d(j, m_{i+1}) \leq \text{MIN}(r, s)$  and  $d(i, j) \geq d(i, -1)$ 
19     $s = d(j, m_{i+1}), N_i^{+1} = N_j$ 
20  end if
21 end for
22 power ( $d(i, +1)$ )
23 send ( $N_i \rightarrow N_i^{+1}$ )

```

---

Theorems 3 and 4 are companion proofs of Theorems 1 and 2, respectively, for symmetrical routing.

*Theorem 3:* Let  $\mathbf{C} = \{\mathbf{L}\{\text{ASL}, p, B\}, \{a, 4, b, P_R\}\}$  be the communication model of a wireless ad hoc network which applies symmetrical routing. If the distance

$$d(i, +1) \geq ((8b + (p/B)P_R)/7a)^{\frac{1}{4}}, \quad (16)$$

the distance

$$\begin{aligned} d(i, -1) \leq & (d(i, +1))/2 - (-0.75(d(i, +1))^2 \\ & + 0.25(9(d(i, +1))^4 - a^{-1}(8b - 7a(d(i, +1))^4 \\ & + (p/B)P_R))^{\frac{1}{2}} \end{aligned} \quad (17)$$

and the distance between an intermediate node and the halfway point between  $S$  and  $D$ :

$$\begin{aligned} r \leq & (-0.75(d(i, +1))^2 + 0.25(9(d(i, +1))^4 \\ & - a^{-1}(8b - 7a(d(i, +1))^4 + (p/B)P_R))^{\frac{1}{2}} \end{aligned} \quad (18)$$

the two-hop communication requires less energy than the direct link.

*Proof:* From Theorem 1 the shortest distance between  $S$  and a power efficient intermediate node would be

$$\begin{aligned} & (d(i, +1))/2 - (-0.75(d(i, +1))^2 + 0.25(9(d(i, +1))^4 \\ & - a^{-1}(8b - 7a(d(i, +1))^4 + (p/B)P_R))^{\frac{1}{2}} \end{aligned} \quad (19)$$

Since, this distance is greater or equal to the distance  $d(i, -1)$ , the symmetrical routing does not affect the power efficient partitioning of the link. Any intermediate node closer to the halfway point between  $S$  and  $D$  than

$$\begin{aligned} & (-0.75(d(i, +1))^2 + 0.25(9(d(i, +1))^4 \\ & - a^{-1}(8b - 7a(d(i, +1))^4 + (p/B)P_R))^{\frac{1}{2}} \end{aligned} \quad (20)$$

will decrease the energy.  $\square$

*Theorem 4:* Let  $\mathbf{C} = \{\mathbf{L}\{\text{BAT}, T_A, T_B, p, q, B\}, \{a, 4, b, P_R\}\}$  be the communication model of a wireless ad hoc network which applies symmetrical routing. If the distance

$$\begin{aligned} d(i, +1) \geq & ((b(3q + p) + P_R B^{-1}(q + p) \\ & + P_R D T_A) a^{-1} (3.5625q + 0.875p)^{-1})^{\frac{1}{4}}, \end{aligned} \quad (21)$$

the distance

$$\begin{aligned} d(i, -1) \leq & (d(i, +1))/2 - (-0.25(d(i, +1))^2 (10.5q + 3p \\ & + 0.5qd(i, +1))(3q + p + qd(i, +1))^{-1} + 0.5a^{-1}(3q \\ & + p + qd(i, +1))^{-1} (0.25a^2(d(i, +1))^2 (10.5q \\ & + 3p + 0.5qd(i, +1))^2 - 2a(3q + p \\ & + qd(i, +1))(-a(d(i, +1))^4 (3.5625q + 0.875p) \\ & + b(3q + p) + P_R B^{-1}(q + p) + P_R D T_A))^{\frac{1}{2}} \end{aligned} \quad (22)$$

and the distance between an intermediate node and the halfway point between  $S$  and  $D$ :

$$\begin{aligned} r \leq & (-0.25(d(i, +1))^2 (10.5q + 3p + 0.5qd(i, +1))(3q \\ & + p + qd(i, +1))^{-1} + 0.5a^{-1}(3q + p \\ & + qd(i, +1))^{-1} (0.25a^2(d(i, +1))^2 (10.5q + 3p \\ & + 0.5qd(i, +1))^2 - 2a(3q + p \\ & + qd(i, +1))(-a(d(i, +1))^4 (3.5625q + 0.875p) \\ & + b(3q + p) + P_R B^{-1}(q + p) + P_R D T_A))^{\frac{1}{2}} \end{aligned} \quad (23)$$

the two-hop communication requires less energy than the direct link.  $\square$

## 8. Conclusion

This paper manifests wireless ad hoc networks need multiobjective design. The multihop communication approach brings tradeoffs between security, real-time and lifetime. We proposed a hierarchical communication model and employed it to compare how two MAC models are capable of partitioning the communication link for non-regular topologies. The proofs can be used to organize look-up tables in the nodes memory and streamline the selection of the best next relay. We evaluated the static energy overhead associated with algorithms for secure routing, such as REWARD, which will help to reassess the lifetime of the network.

## References

- [1] B. Chen, K. Jamieson, H. Balakrishnan, and R. Morris, "Span: an energy-efficient coordination algorithm for topology maintenance in ad hoc wireless networks", *ACM Wirel. Netw. J.*, vol. 8, no. 5, pp. 481–494, 2002.
- [2] W. Ye, J. Heidemann, and D. Estrin, "Medium access control with coordinated adaptive sleeping for wireless sensor networks", *IEEE/ACM Trans. Netw.*, vol. 12, no. 3, pp. 493–506, 2004.
- [3] S. Biswas and R. Morris, "Opportunistic routing in multi-hop wireless networks", *ACM SIGCOMM Comput. Commun. Rev.*, vol. 34, iss. 1, pp. 69–74, 2004.
- [4] D. Dewasurendra and A. Mishra, "Design challenges in energy-efficient medium access control for wireless sensor networks", in *Handbook of Sensor Networks: Compact Wireless and Wired Sensing Systems*, M. Ilyas and I. Mahgoub, Eds. Boca Raton: CRC Press LLC, 2005, pp. 28–1–28–25.
- [5] V. Zadorozhny, D. Sharma, P. Krishnamurthy, and A. Labrinidis, "Tuning query performance in mobile sensor databases", in *Proc. 6th Int. Conf. Mob. Data Manage.*, Ayia Napa, Cyprus, 2005, pp. 247–251.
- [6] Z. Karakehayov and N. Andersen, "Energy-efficient medium access for data intensive wireless sensor networks", in *Proc. Int. Worksh. Data Intens. Sens. Netw. 8th Int. Conf. Mob. Data Manage.*, Mannheim, Germany, 2007, pp. 116–120.
- [7] J. L. Gao, "Energy efficient routing for wireless sensor networks", Ph.D. thesis, University of California, Los Angeles, 2000.
- [8] J. M. Rabaey, M. J. Ammer, J. L. Silva, D. Patel, and S. Roundy, "PicoRadio supports ad hoc ultra-low power wireless networking", *IEEE Computer*, vol. 33, pp. 42–48, July 2000.
- [9] I. F. Akyildiz and I. H. Kasimoglu, "Wireless sensor and actor networks: research challenges", *Ad Hoc Netw.*, no. 2, pp. 351–367, 2004.

- [10] Z. Karakehayov, "Low-power communication for wireless sensor-actuator networks", in *Proc. Fifth IASTED Int. Conf. Commun. Syst. Netw.*, Palma de Mallorca, Spain, 2006, pp. 1–6.
- [11] T. Moscibroda, R. O'Dell, M. Wattenhofer, and R. Wattenhofer, "Virtual coordinates for ad hoc and sensor networks", in *Proc. ACM Joint Worksh. Found. Mob. Comp.*, Philadelphia, USA, 2004.
- [12] D. Puccinelli and M. Haenggi, "Wireless sensor networks: applications and challenges of ubiquitous sensing", *IEEE Circ. Syst. Mag.*, pp. 19–29, 3rd quart. 2005.
- [13] Z. Karakehayov, K. S. Christensen, and O. Winther, *Embedded Systems Design with 8051 Microcontrollers*. New York: Dekker, 1999.
- [14] T. Stoyanova, F. Kerasiotis, A. Prayati, and G. Papadopoulos, "Evaluation of impact factors on RSS accuracy for localization and tracking applications", in *Proc. 5th ACM Int. Worksh. Mob. Manage. Wirel. Acc.*, Chania, Greece, 2007, pp. 9–16.
- [15] V. Swaminathan, Y. Zou, and K. Chakrabarty, "Techniques to reduce communication and computation energy in wireless sensor networks", in *Handbook of Sensor Networks: Compact Wireless and Wired Sensing Systems*, M. Ilyas and I. Mahgoub, Eds. Boca Raton: CRC Press LLC, 2005, pp. 29–1–29–34.
- [16] M. T. Schmitz, B. M. Al-Hashimi, and P. Eles, *System-Level Design Techniques for Energy-Efficient Embedded Systems*. Boston: Kluwer, 2004.
- [17] Z. Karakehayov, "Dynamic clock scaling for energy-aware embedded systems", in *Proc. IEEE Fourth Int. Worksh. Intell. Data Acquis. Adv. Comp. Syst.*, Dortmund, Germany, 2007, pp. 96–99.
- [18] H. Takagi and L. Kleinrock, "Optimal transmission ranges for randomly distributed packet radio terminals", *IEEE Trans. Commun.*, vol. 32, no. 3, pp. 246–257, 1984.
- [19] I. Stojmenovic and X. Lin, "Power aware localized routing in wireless networks", *IEEE Trans. Parall. Distr. Syst.*, vol. 12, no. 11, pp. 1122–1133, 2001.
- [20] Z. Karakehayov, "Low-power design for Smart Dust networks", in *Handbook of Sensor Networks: Compact Wireless and Wired Sensing Systems*, M. Ilyas and I. Mahgoub, Eds. Boca Raton: CRC Press LLC, 2005, pp. 37–1–37–12.
- [21] F. Baccelli, B. Blaszczyzyn, and P. Muhlethaler, "An Aloha protocol for multihop mobile wireless networks", in *Proc. ITC Spec. Sem. Perform. Eval. Wirel. Mob. Syst.*, Antwerp, Belgium, 2004.
- [22] Z. Karakehayov, "Security – lifetime tradeoffs for wireless sensor networks", in *Proc. 12th IEEE Int. Conf. Emerg. Technol. Fact. Autom.*, Patras, Greece, 2007, pp. 646–650.
- [23] E. M. Royer and C. Toh, "A review of current routing protocols for ad hoc mobile wireless networks", *IEEE Pers. Commun.*, vol. 6, no. 2, pp. 46–55, 1999.
- [24] G. Zhou, T. He, S. Krishnamurthy, and J. Stankovic, "Models and solutions for radio irregularity in wireless sensor networks", *ACM Trans. Sens. Netw.*, vol. 2, no. 2, pp. 221–262, 2006.
- [25] Z. Karakehayov and Z. Monov, "Target-aware timing modelling for wireless ad-hoc networks", in *Proc. Int. Sci. Conf. Comput. Sci. 2006*, Istanbul, Turkey, 2006, pp. 54–59.
- [26] N. Sastry, U. Shankar, and D. Wagner, "Secure verification of location claims", in *Proc. 2003 ACM Works. Wirel. Secur.*, San Diego, USA, 2003.
- [27] C. Karlof and D. Wagner, "Secure routing in wireless sensor networks: attacks and countermeasures", *Ad Hoc Netw.*, no. 1, pp. 293–315, 2003.
- [28] H. Qi, S. S. Iyengar, and K. Chakrabarty, "Multiresolution data integration using mobile agents in distributed sensor networks", *IEEE Trans. Syst. Man Cyber. Part C: Appl. Rev.*, vol. 31, no. 3, pp. 383–391, 2001.
- [29] Q. Wu, N. S. V. Rao, R. R. Brooks, S. S. Iyengar, and M. Zhu, "Computational and networking problems in distributed sensor networks", in *Handbook of Sensor Networks: Compact Wireless and Wired Sensing Systems*, M. Ilyas and I. Mahgoub, Eds. Boca Raton: CRC Press LLC, 2005, pp. 25–1–25–17.
- [30] D. D. Hwang, B. C. Lai, and I. Verbauwhede, "Energy-memory-security tradeoffs in distributed sensor networks", in *Ad-Hoc, Mobile, and Wireless Networks*, I. Nikolaidis, M. Barbeau, and E. Kranakis, Eds., Lecture Notes in Computer Science, vol. 3158. Berlin-Heidelberg: Springer, 2004, pp. 70–81.
- [31] Z. Karakehayov, "Design of distributed sensor networks for security and defense", in *Proc. of the NATO Advanced Research Workshop on Cyberspace Security and Defense: Research Issues* (Gdańsk, September 6–9, 2004), J. S. Kowalik, J. Gorski and A. Sachenko, Eds., NATO Science Series II, vol. 196. Dordrecht: Springer, 2005, pp. 177–192.
- [32] S. Marti, T. J. Giuli, K. Lai, and M. Baker, "Mitigating routing misbehavior in mobile ad hoc networks", in *Proc. 6th Int. Conf. Mob. Comput. Netw. MOBICOM-00*, New York, USA, 2000, pp. 255–265.
- [33] Z. Karakehayov, "Using REWARD to detect team black-hole attacks in wireless sensor networks", in *Proc. Worksh. Real-World Wirel. Sens. Netw. REALWSN'5*, Stockholm, Sweden, 2005.
- [34] Z. Karakehayov and I. Radev, "REWARD: a routing method for ad-hoc networks with adjustable security capability", in *Proc. NATO Adv. Res. Worksh. Secur. Embed. Syst.*, Patras, Greece, 2005, pp. 180–187.
- [35] Z. Karakehayov and I. Radev, "A scalable security service for geographic ad-hoc routing", *Int. Sci. J. Comp.*, vol. 4, iss. 2, pp. 124–132, 2005.
- [36] Z. Karakehayov, "Wireless ad hoc networks: where security, real-time and lifetime meet", in *Proc. Int. Multiconf. Comput. Sci. Inform. Technol. Worksh. Wirel. Unstruct. Netw.*, Wisla, Poland, 2008, pp. 861–868.



**Zdravko Karakehayov** received the Ph.D. degree from the Technical University of Sofia, Bulgaria. He is an Associate Professor in the Department of Computer Systems at the Technical University of Sofia. Formerly he was with the Technical University of Denmark, Lyngby and the University of Southern Denmark,

Sønderborg. He co-authored five books in the field of embedded systems and holds eight patents. His research field includes low-power design for embedded systems, low-power and secure routing for wireless sensor networks. He served as a reviewer for the "Journal Transactions on Embedded Computing Systems" and several international conferences. He is a senior member of the IEEE Computer Society and a Distinguished Visitors Program speaker. Dr. Karakehayov currently chairs the Computer Chapter, IEEE Bulgaria.

email: zkg@tu-sofia.bg

Department of Computer Systems

Technical University of Sofia

Kliment Ohridski st 8

Sofia 1000, Bulgaria

# A Method of Mobile Base Station Placement for High Altitude Platform Based Network with Geographical Clustering of Mobile Ground Nodes

Ha Yoon Song

**Abstract**—High altitude platforms (HAPs) such as unmanned aerial vehicles (UAVs) which can be deployed as stratospheric infrastructures enable a sort of new configurations of wireless networks. Ground nodes must be clustered in multiple sets and one dedicated UAV is assigned to each set and act as an mobile base station (MBS). For the intra-set nodes, UAVs must communicate each other in order to establish network links among intra-set nodes. Here we find a geographical clustering problem of networking nodes and a placement problem of MBSs. The clustering technique of mobile ground nodes can identify the geographical location of MBSs as well as the coverage of MBSs. In this paper we proposed a clustering mechanism to build such a configuration and the effectiveness of this solution is demonstrated by simulation. For a selected region with a relatively big island, we modeled mobile ground nodes and showed the result of dynamic placement of MBSs by our clustering algorithm. The final results will be shown graphically with the mobility of ground nodes as well as the placement of MBSs.

**Keywords**—clustering geographical, clustering of mobile ground nodes, HAP based network, high altitude platforms, MBS placement, mobile base stations.

## 1. Introduction

The recent development of unmanned aerial vehicles (UAV) leads to interests in high altitude platforms (HAPs) as mobile base stations (MBSs) of wireless wide area networks. UAVs usually carry wireless network equipments and require less management and thus have been regarded as one of the best method to deploy wireless network over wide area without typical ground network equipments. This new sort of network configuration requires a lot of unsolved research topics from the physical layer to transportation layer as well as the ideal configuration of network. Many researches have been concentrated on the communication between stratospheric UAVs and mobile ground nodes. These topics include establishments of communication links among UAVs as well.

The goal of HAP based network is to cover as wide area as possible with deployment of multiple UAVs, i.e., ultimate goal of HAP network is to deploy as many MBSs to cover dedicated area in order to construct a network structure. In this configuration, UAVs can act as mobile base stations

for the network. This sort of network configuration raises a new configuration problem of UAVs and mobile ground nodes.

Here we suggest an idea. Mobile ground nodes consist a number of clusters in order to be served by MBSs and each MBS covers a dedicated cluster of mobile ground nodes. MBSs cooperate each other in order to support communication of mobile ground nodes in the whole area. The HAP based wireless network usually regarded as a viable solution for the networks with minimal ground infrastructures or countries under development and so on.

In this paper, we will show a method to deploy MBSs for a dedicated area with possible number of MBSs and to cover dedicated area efficiently. By adopting a clustering mechanism for mobile ground nodes, we can solve a placement problem of MBSs as well as the coverage of each MBS. For an island area we set mobility model and simulated the dynamic clustering of mobile ground nodes and find proper locations for MBSs and coverage.

This paper is structured as follows. In Section 2 we will discuss research basis for this paper. In Section 3 we will see the basic configuration concept of HAP based network assumed in this paper. In the following Section 4 we will explain our clustering algorithm for mobile ground nodes and MBS placement. Then we will show simulation experiment and results in Section 5. The final Section 6 will conclude this paper and will discuss the future research direction.

## 2. Related Works

The basis of this research is usually categorized into two basic parts. The first one is HAP based networks. The second one is clustering algorithms. We will discuss about this two topics simply.

### 2.1. HAP Based Network

The HAP based network is one of the most recent research topics while there have been a long idea about stratospheric platforms as a media of networking. With the development of network technologies as well as aerial industries, HAP can be actual one and the idea of HAP based network spawns into real world. Nowadays, most of countries

research on this topic. Some of them regard this sort of network as national project while others act as an individual company bases. The Republic of Korea researches on this topic as a national project as shown in [1] with the fruitful research at Electronics and Telecommunication Research Institute (ETRI).

For EU nations or EU companies, there are projects such as HeliNet and CAPANINA. England was one of the early starter in this topic and most of EU nations are participated in those projects [2]. Japan has the similar situation. Japanese Aerospace Exploration Agency [3] has an ongoing effort with HAP based network, named as SkyNet project. Also in the US, projects such as Sky Station, High Altitude Long Operation, SkyTower, Staratellite, Weather Balloon HAPs are undergoing ones [4].

The characteristics of this project are that they are still assuming a single HAP platform and are mostly concentrated on communication links establishment.

For example, World Radiocommunication Conference (WRC) of ITU started the first specification of frequency allocation for HAP network since the year 1997 and several bandwidth frequencies are allotted for specific nations [5]. The most recent decision made by WRC 2007 can be found in [6].

There have been several researches on the communication link establishment between mobile ground nodes and HAP base station [7], [8]. Also a protocol standard such as 802.16 can be a possible candidate for HAP networks [9]. For the inter-HAP links, there also have been researches such as [10].

Apart from these researches, this paper assumes an environment of multiple heterogeneous HAPs. We regard each HAP as a MBS of mobile ground nodes and tried to cluster mobile ground nodes in order to solve the placement problem of MBS. As far as we know, no such topic has been dealt until the first submission of this paper. The most similar one is in [11] but it is on a positioning control for HAP station.

## 2.2. Clustering Algorithms

In a field of data mining, there have been researches regarding clustering of data in a large database [12], and multiprocessor versions can be found in [13]. They are usually concentrated on the patterns of related data while we need a geographical clustering based on the ground node coordinates. Thus we screened out several candidates and tried combinations of those clustering algorithms.

The most prominent candidates of clustering algorithms are K-mean, BIRCH, EM, PROCLUS, and SUBCLUS drawn from [12]. The K-mean is one of the most popular algorithm in the clustering world. Since we are dealing with mobile ground nodes, we cannot identify the exact coordinates but we only assume the probability of node coordination. In this aspect, we concentrated on expectation maximization (EM) with a probability of node coordinates.

Considering the limitations of number of nodes in a cluster, that is actually a bandwidth limitation of wireless routers equipped on HAP MBS, we need to split a cluster or merge clusters in order to balance the number of nodes in a cluster. The BIRCH and their successors can cope with such a requirement, thus we choose it as a possible candidate.

Even though EM and BIRCH are nice candidates for our application, these two cannot deal ideal initial clustering and the results are usually not a geographical one. Thus we need a multiphase clustering algorithm. Usually initial clusters can be drawn by K-mean algorithm and then the core cluster algorithm such as BIRCH or EM can work. For geographical consistency, a postprocessing of merging or splitting must be done. Two algorithms of SUBCLUS and PROCLUS are typical examples of multiphase clustering algorithms that attracts our interest.

However, apart from this existing algorithms, we will study the most appropriate clustering algorithms for HAP MBS placement. In this paper, we will show a basic result based on K-mean algorithm [14].

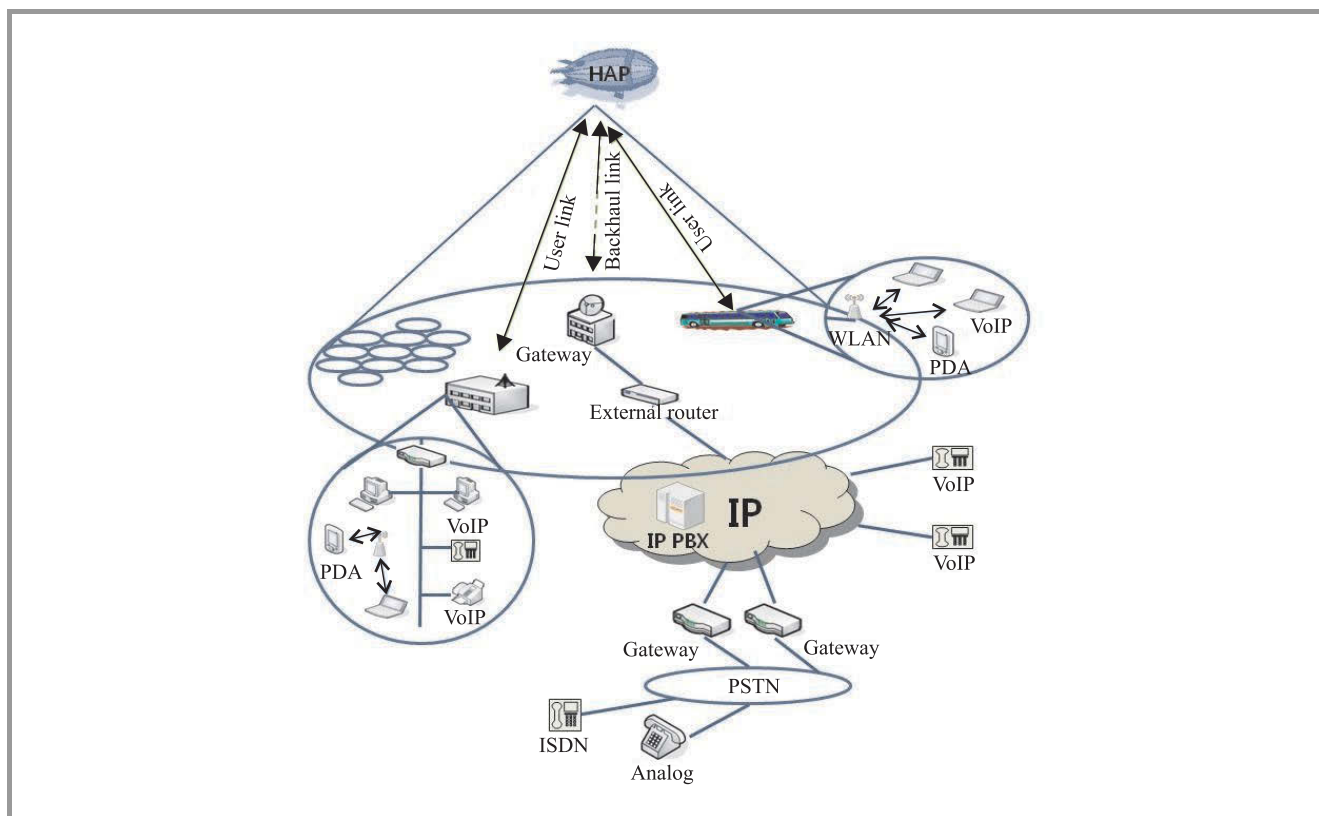
## 3. Basic Concept of HAP Based Network Configuration

The concept of HAP networking is similar to that of satellite communication. Characteristics such as broadness, simultaneousness, flexible configuration, and broadband-widthness of HAP based network are similar to that of satellite network, however, on-time supplements, ease of management, short communication distance, low power mini handset, low round trip time are typical characteristics of HAP network that cannot be found in satellite network. The key point here is low communication delay and flexible network configuration. In this paper, we spotted on flexible network configuration and will start the discussion from the basic configuration of HAP network. Figure 1 shows a basic configuration of HAP based network. HAP usually resides on stratospheric area, about 21 km altitude with very low speed of winds as reported by NASA, USA.

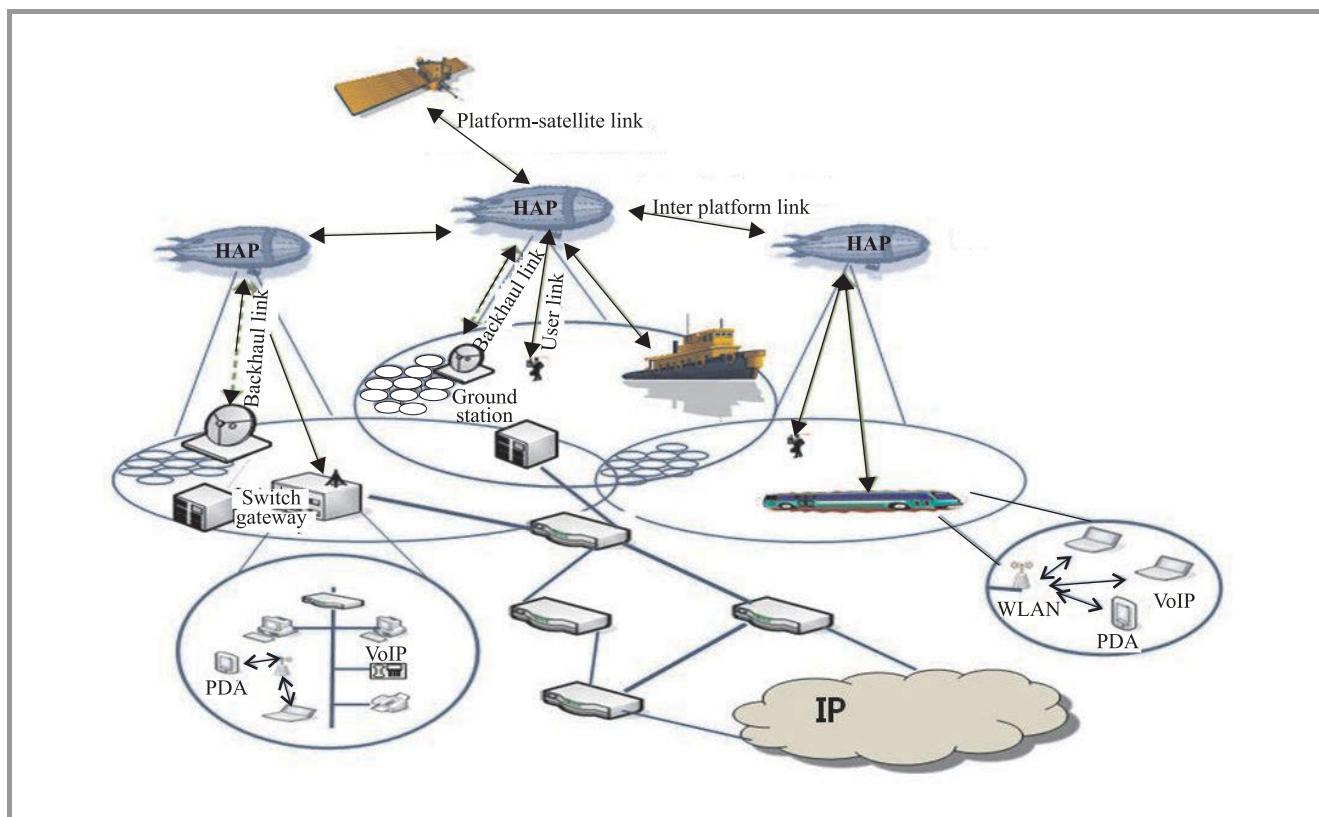
Even though Fig. 1 presents only one HAP device, however there could be a lot of plans for multiple HAPs. Figure 2 shows the configuration of HAP based network with multiple HAPs deployed on more wide area.

This configuration with multiple HAPs will cause a interplatform link study and reapplication of traditional frequency reuse problem. It has been assumed that a radio or a optical communication for interplatform link and it is another research area. Several research regarding optical links, including the CAPANINA projects, issued results regarding optical communication links for inter-HAP links as well as downlinks [15], [16], [17]. One of the interesting approach includes downlinks of satellite as well [18].

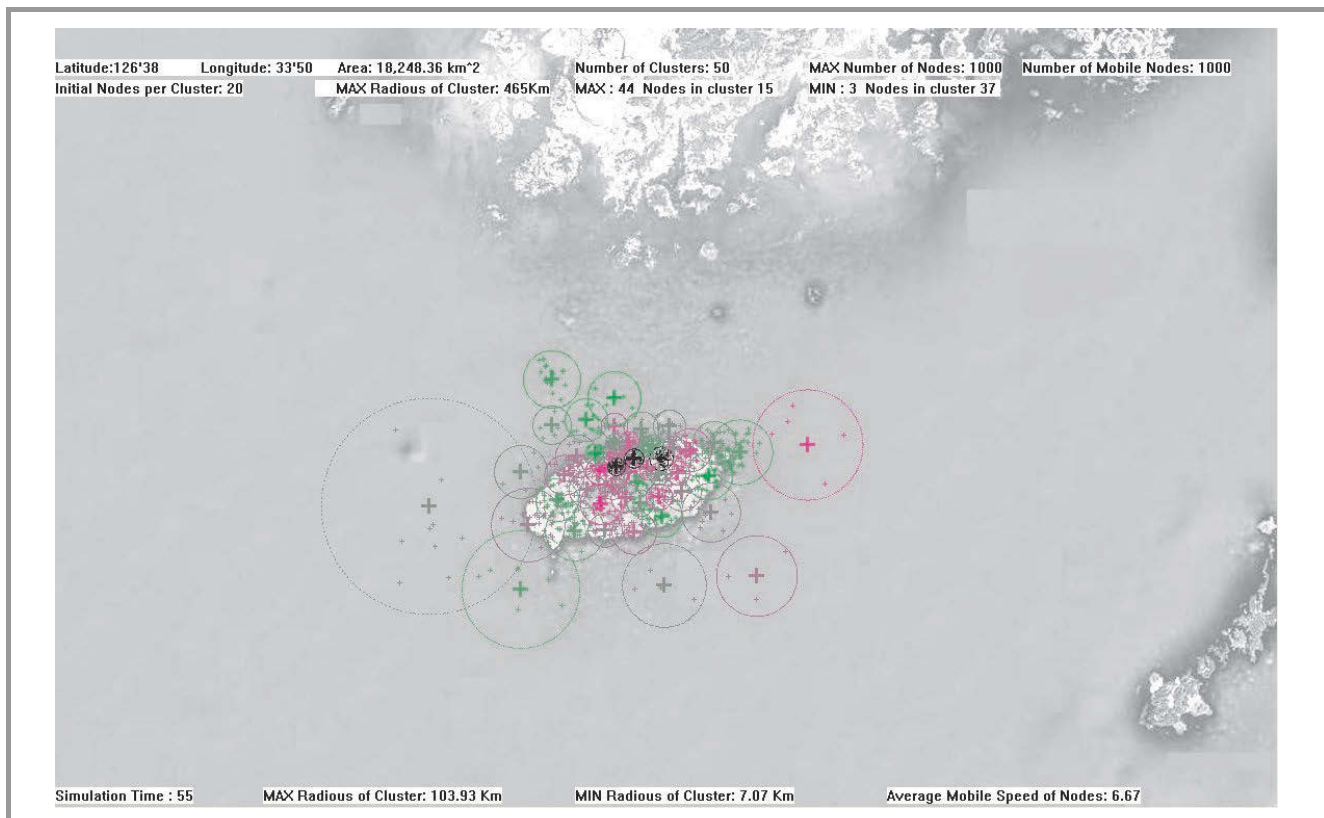
Also a routing problem can be arisen while there are a lot of prepared routing protocol such as OLSR, TBRPF, DSDV, CGSR, WR, AODV, DSR, LMR, TORA, ABR, SSR, and ZRP.



**Fig. 1.** Basic configuration of HAP based network: singular HAP case. Explanations: IP PBX – Internet protocol private branch exchange, ISDN – integrated services digital network, PDA – personal digital assistant, PSTN – public switched telephone network, VoIP – voice over Internet protocol, WLAN – wireless local area network.



**Fig. 2.** Basic configuration of HAP based network: multiple HAP case.



**Fig. 3.** Geographical simulation results mapped on Cheju Island and neighboring seas.

Therefore we will concentrate on the multiple HAPs network. The first problem of this multiple HAPs network is a clustering of mobile ground nodes and finds an optimal location for multiple HAPs as a clusterhead or mobile MBS. We expect the final result looks like in Fig. 3 where we can see 50 MBSs work as clusterhead and the coverage of each MBS are defined by the population density.

For the environment of this multiple HAPs network, we have similar problems in wireless network listed below:

- uncertainty of mobile ground node coordinates;
- link error rate, lower than satellite network;
- network topology variance, frequent;
- network security.

We can consider parameters below for network configuration:

- total number of mobile ground node, network size  $N$ ;
- network connectivity;
- topological rate of change;
- link capacity;
- fraction of unidirectional links;
- mobility of ground nodes;
- fraction and frequency of sleeping nodes.

#### 4. Clustering of Ground Nodes for MBS Placement

For multiple HAP network, mobile ground nodes are clustered into several clusters, and each HAP acts as MBS of a cluster or is regarded as clusterhead. Since the number of nodes in a cluster is a major parameter, another parameter of cluster coverage is a dependent variable. Each HAP has capability to adjust its coverage individually by adjusting its angle of elevation. These parameters can vary dynamically according to the mobility of ground nodes. We can assume the following scenarios of dynamic clustering.

- **Initial clustering.** The very first stage of network configuration. Almost a static clustering.
- **Service oriented reclustering.** New influx of mobile ground nodes to a specific cluster cause a overloaded router on a HAP and thus requires a reclustering. Maybe a new UAV can be added.
- **Geographical reclustering.** The mobility of ground nodes can cause more area to be covered by HAP based network. This situation causes reclustering.
- **Airship backup reclustering.** Failure of one or more UAV can cause a reclustering situation.
- **Shrinking clustering.** Decrement of the number of mobile ground nodes will cause a non-mandatory reclustering.

From the scenarios above, we assume a geographical reclustering situation in this paper. However, all of the above scenarios require reclustering algorithm according to the geographical network coverage or mobile ground node distribution. In order to provide such an algorithm, we can assume the following requirements for network parameters.

1. An UAV as a HAP can equip MBS facilities.
2. An UAV can identify number of ground nodes served.
3. An UAV can identify the location of mobile ground nodes.
4. An UAV can identify the location of itself.
5. An UAV can vary the geographical network coverage by adjusting the angle of elevation.
6. An UAV can communicate with ground base stations.

These assumptions can be realized with the help of aero engineering, electro engineering, antenna technology or other related modern technology. Under these assumptions, the following lists requirements for clustering algorithms of mobile ground nodes.

1. The sum of bandwidth requirement from mobile ground nodes in a cluster is equal to or smaller than the total bandwidth of an MBS for the cluster.
2. Each UAV can move in a limited speed. Thus the reclustering must consider the speed of UAVs.
3. Cluster algorithm requires realtimeness for continuous network service.

Here we can show a simple clustering algorithm based on K-mean clustering algorithm.

---

**Algorithm 1: Clustering and MBS placement**

---

**Require:** B : Router bandwidth capacity on a HAP  
 S : Total sum of required bandwidth by nodes in a cluster  
 C : A set of coordinates of mobile ground nodes  
 A : Total area size to be covered by network  
 M : Maximum area can be covered by a HAP

**Ensure:** L : List of clusters  
 P : List of cluster centroids

```

2:   SpB = S/B //bandwidth requirement
4:   ApM = A/M //geographical requirement
6:   Calculate the number of cluster K = min(SpB,ApM)
8:   L = result of K-mean clustering algorithm with input C
10:  P = calculated centroid of each cluster.
```

---

The exact location of each MBS for a cluster is a centroid of the each cluster. The centroid of a cluster is median value of geographical coordinates of mobile ground nodes in the cluster. Instead of mean value, we use median value in order to guarantee the efficiency of MBS coverage. Thus we can use this algorithm for MBS placement.

Even though we use K-mean algorithm, this algorithm can be replaced any of proper clustering algorithm. We experienced several clustering algorithms produce inclusive clusters. To resolve such problem, we need postprocessing of clusters that merges inclusive clusters to including cluster. We regard this problem as a future topic and will discuss basically in the conclusions and future research section.

The time complexity of Algorithm 1 is  $O(TN^2)$ , where  $N$  is total number of mobile ground nodes and  $T$  is number of reclustering. Simply speaking, K-mean clustering takes time of  $O(N^2)$  for  $N$  data entries and the reclustering happens, and the total count of reclustering is  $T$ . Whenever the constraints of HAP network clustering are violated, a reclustering is done by K-mean algorithm.

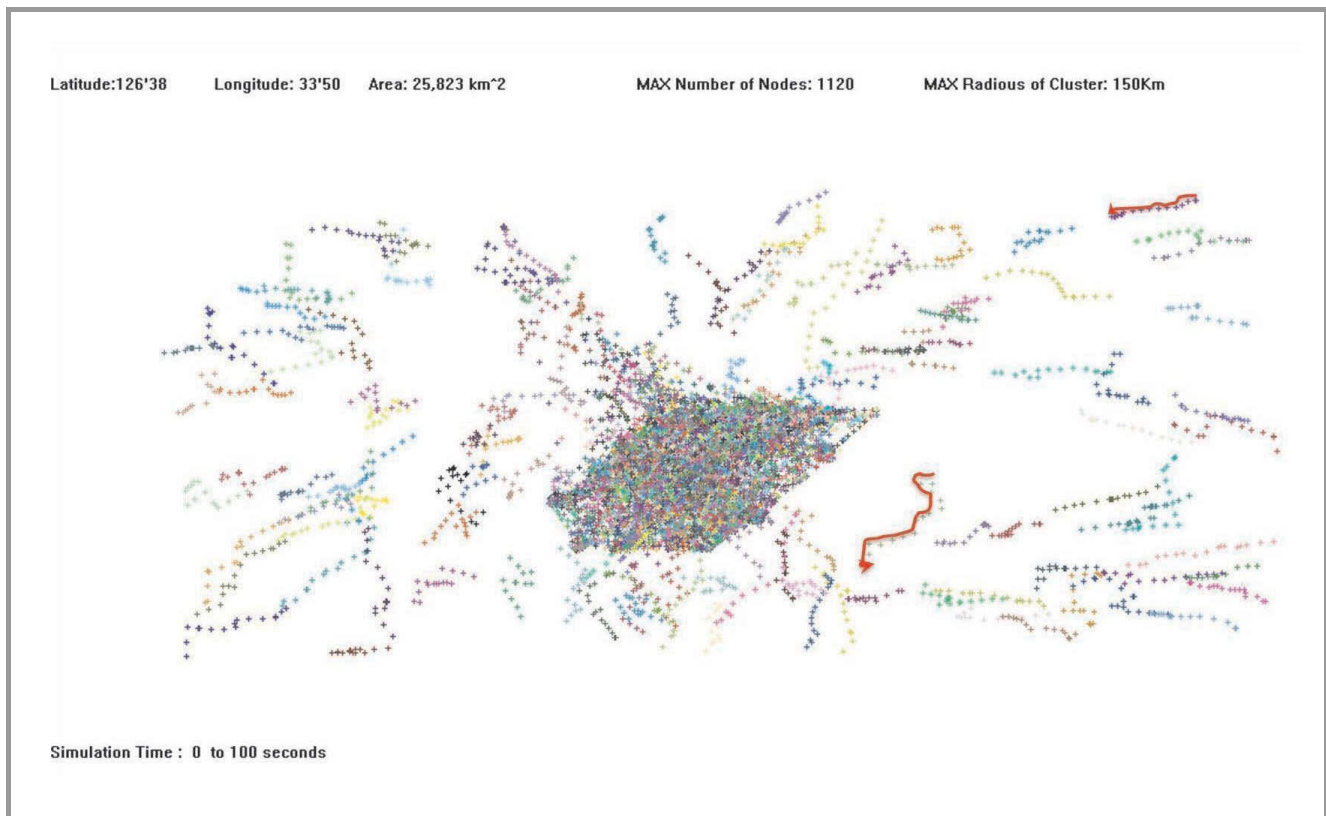
## 5. Simulation of HAP MBS Placement

We do simulation experiment based on assumptions and algorithms for HAP based network. Our aim is to show simulation results geographically in order to identify clustering result and HAP placement result graphically. The simulation is done first by NS-2 (network simulator 2) with mobility models for nodes. With the NS-2 trail output, we parse them in order to identify node mobility in time and the parsed results are used by clustering algorithms. Finally a clustering results are visualized.

### 5.1. Simulation Environments and Parameters

For the geographical environment for network simulation, we choose Cheju Island of the Republic of Korea and its neighboring seas. 1000 nodes are distributed according to population distribution on Cheju Island and several marine vehicles were also assumed. The reason why we choose this area is as follows.

- Two cities have denser population than other area.
- The other area has lower population.
- There is high mobility to or from two cities to other area.
- Almost all sort of mobility can be included. There are various type of mobility including human walk, ground vehicle, horse rider, and marine vehicle.
- There are frequent mobility between Cheju Island and its two neighboring small islands which have lower population.
- Lots of travelers show very high mobility.



**Fig. 4.** Node mobility snapshot.

In addition we wanted to look at the phenomena that non-island areas are connected from island area by HAP MBS. The simulation parameters are as follows:

- 28,246 km<sup>2</sup> coverage area;
- total number of mobile ground node is 1000;
- number of initial clusters is 50;
- number of average nodes in a cluster is 20;
- random waypoint (RW) mobility model;
- speed of mobility is 4~7 km/s;
- node population is proportional to actual population density;
- simulation time up to 60 min;
- maximum 465 km for inter-hap links;
- maximum 150 km for cluster radius as specified by ITU [19].

Figure 4 shows a cumulated snapshot of mobile nodes trajectory for these simulation experiments.

## 5.2. Simulation Results and Analysis

With this simulation parameters, simulation results are shown geographically in selected Figs. 5–13.

Each figure stands for the result of simulation from 15 to 55 min after the start of simulation. In each figure, the radius of a cluster is magnified by the factor of 1.4 for visibility reason. The dots are mobile ground nodes while crosses are centroids of clusters and thus locations of HAP MBSs. Two cities with high population show a lot of clusters with small coverage. This is due to the restriction of network bandwidth provided by on HAP MBS. In order to guarantee minimum bandwidth for individual nodes, the maximum number of nodes per cluster is restricted since HAP MBS has limited bandwidth as a wireless router. The neighboring seas are covered by relatively larger clusters and smaller number of HAP MBS. The maximum number of nodes in a cluster is up to 47 while the minimum number is 2 for a cluster. The largest cluster has radius of 154 km in marine area while the smallest cluster has 6.5 km in city area. Note that 154 km of cluster radius slightly exceeds a maximum radius of HAP MBS coverage specified by ITU.

Even though some nodes look like multiple participant in a cluster, they are actually a member of specific single clusters. This sort of overlapping can be overcome by frequency allocation and reuse [20]. Some clusters with dedicated MBSs look like an orphan. However, the presumably orphan MBSs are within 930 km of other connected HAP MBS, i.e., orphans are within the range of inter-hap links.

Some postprocessing algorithms required in order to merge or to split clusters with abnormal number of nodes. Clusters with larger number of nodes must be split and clusters

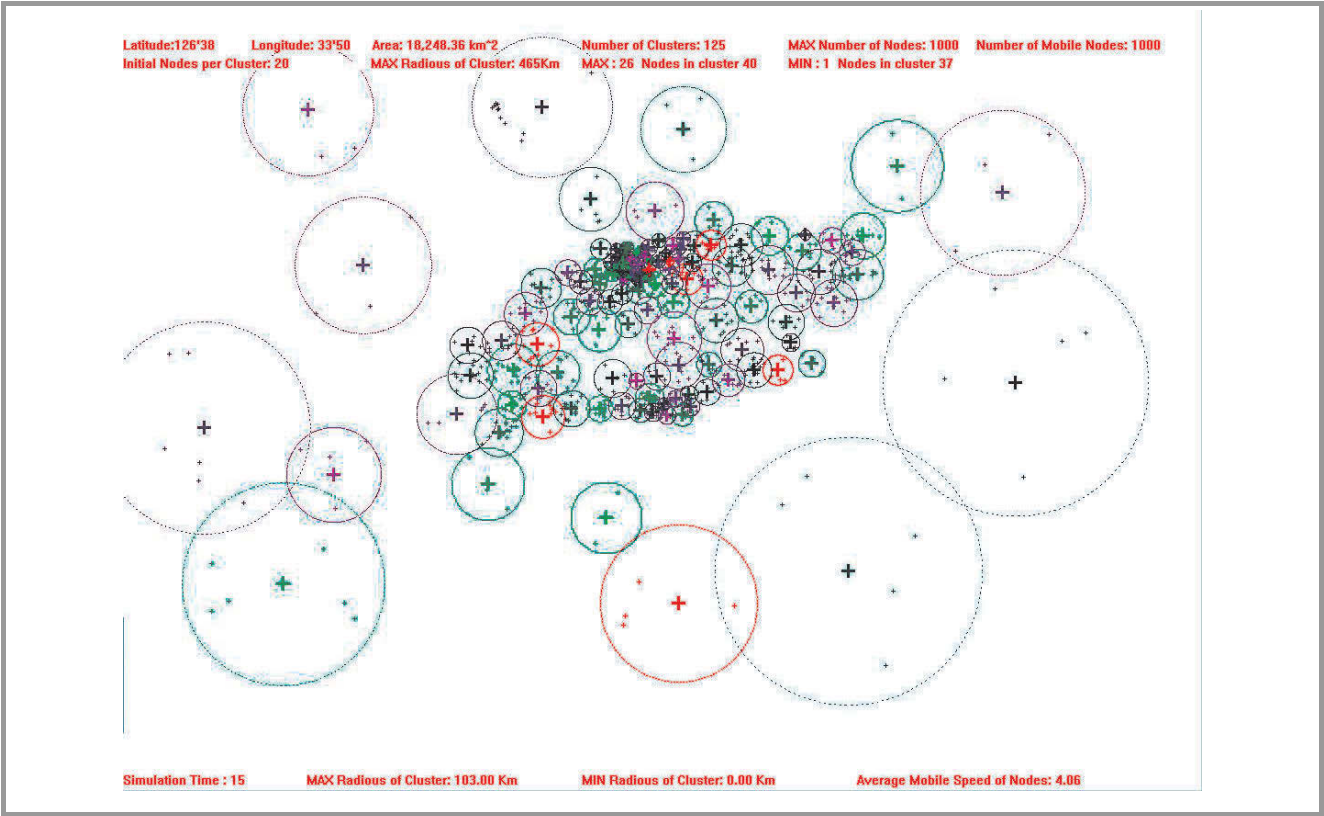


Fig. 5. Placement and coverage of MBSs after 15 min.

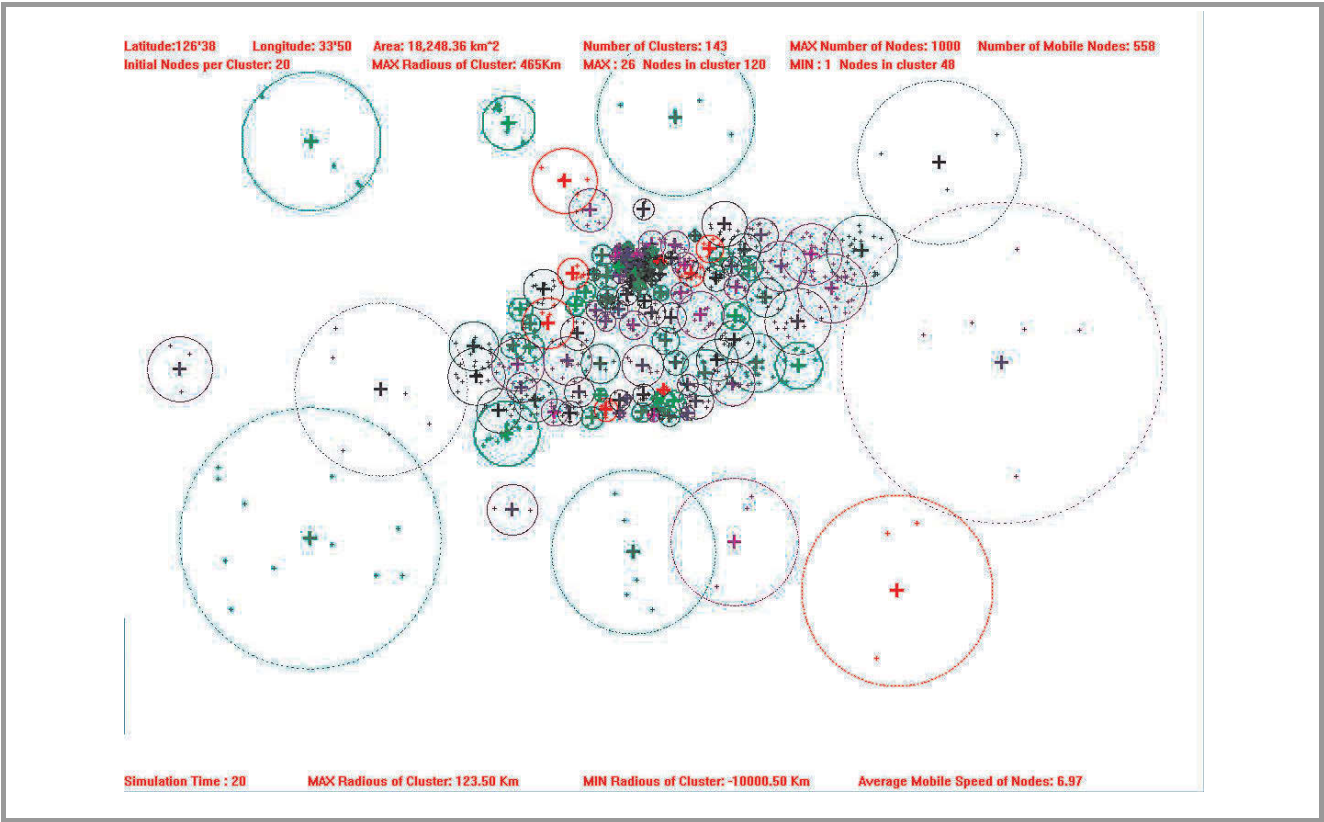


Fig. 6. Placement and coverage of MBSs after 20 min.

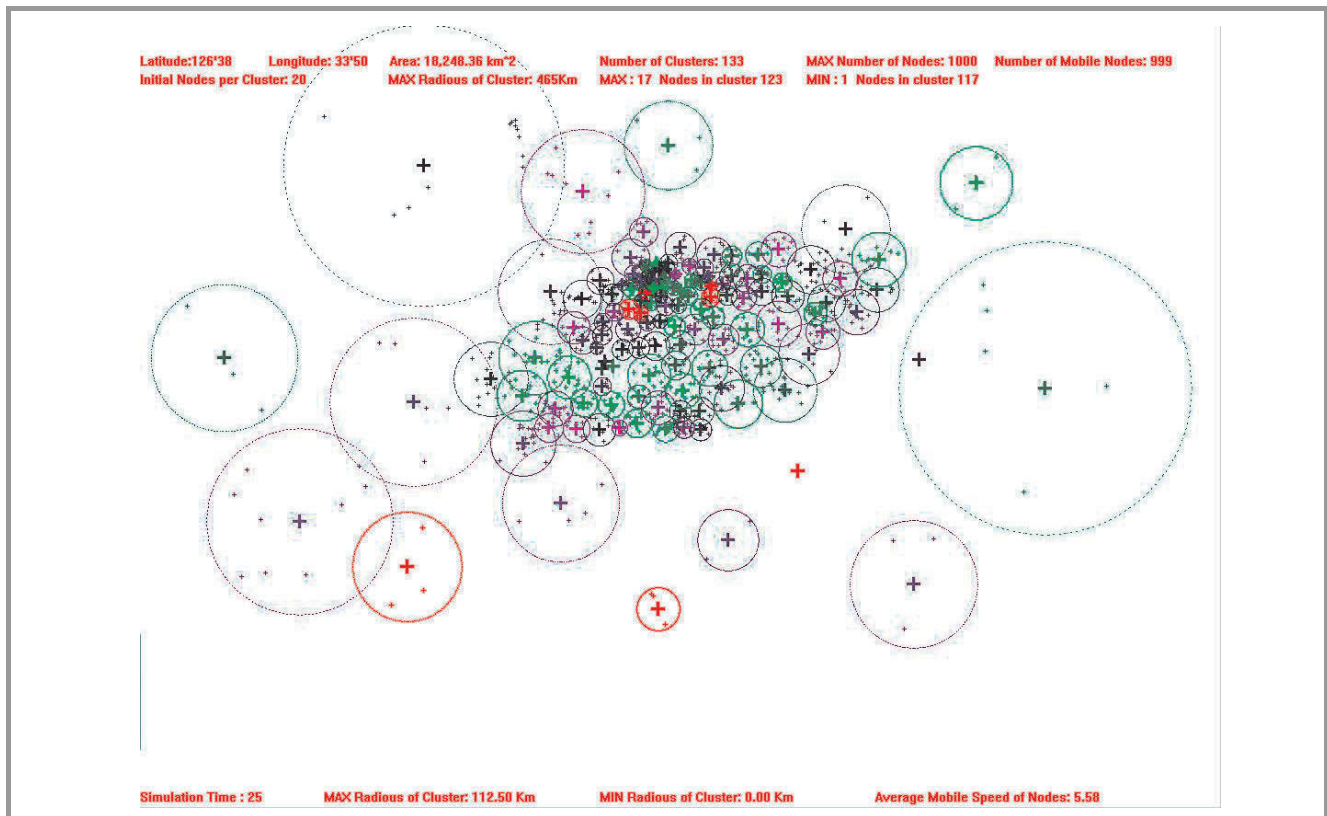


Fig. 7. Placement and coverage of MBSs after 25 min.

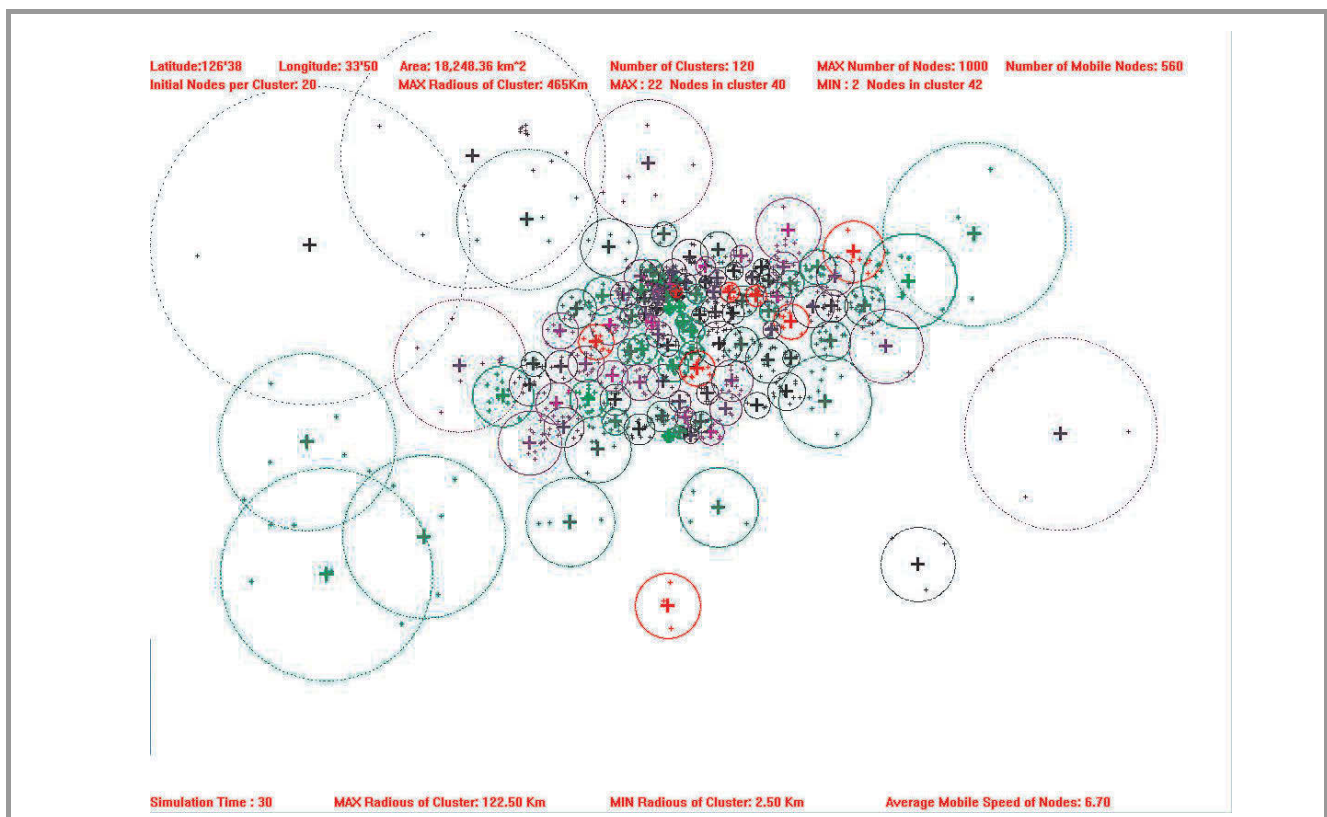


Fig. 8. Placement and coverage of MBSs after 30 min.

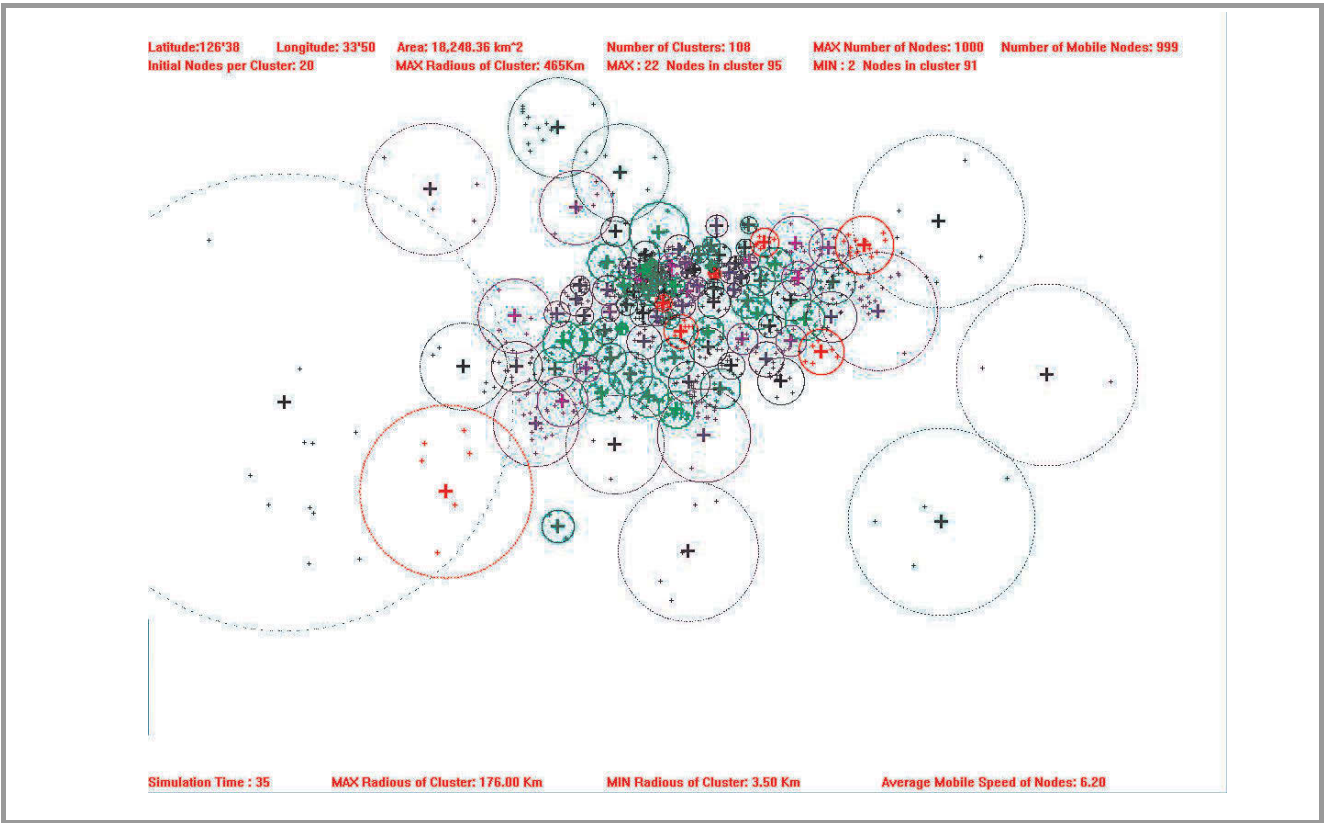


Fig. 9. Placement and coverage of MBSs after 35 min.

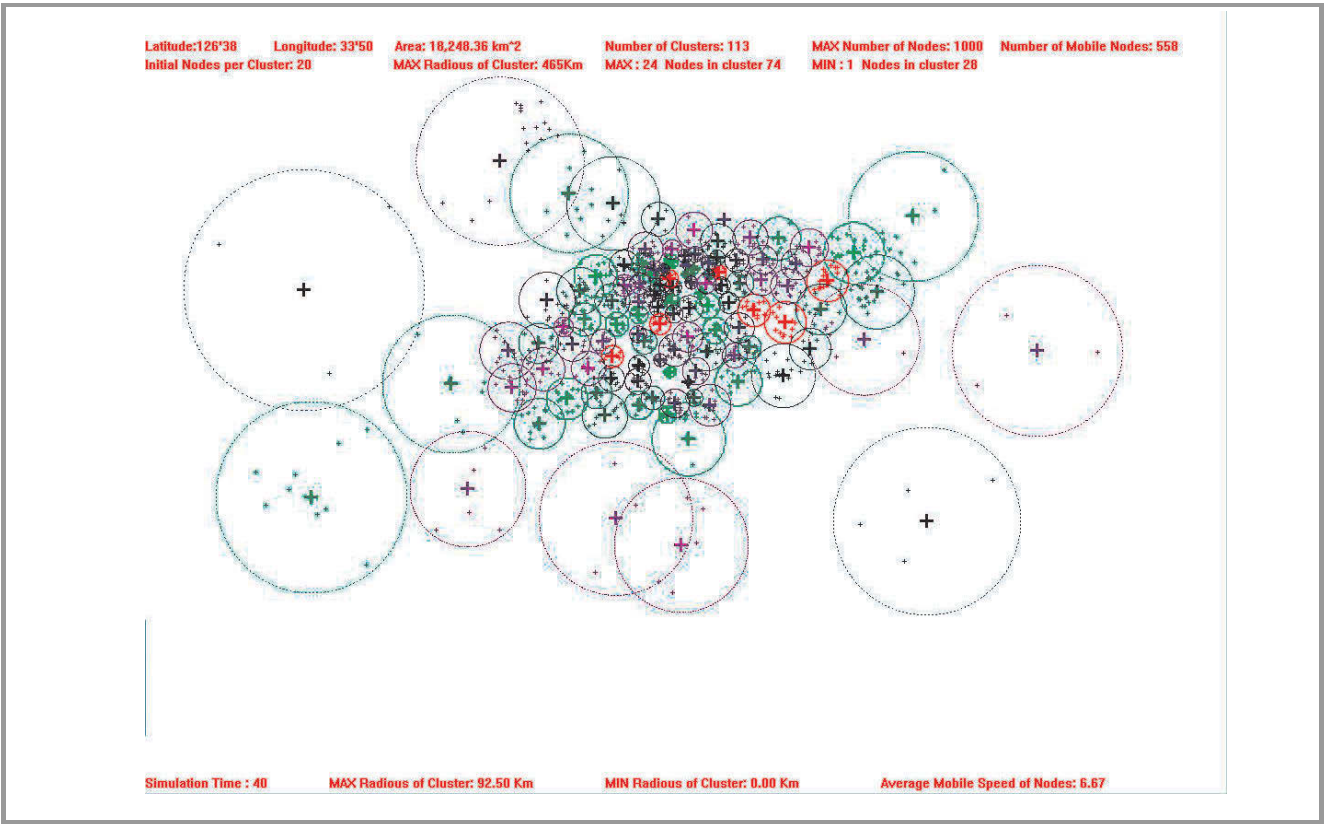


Fig. 10. Placement and coverage of MBSs after 40 min.

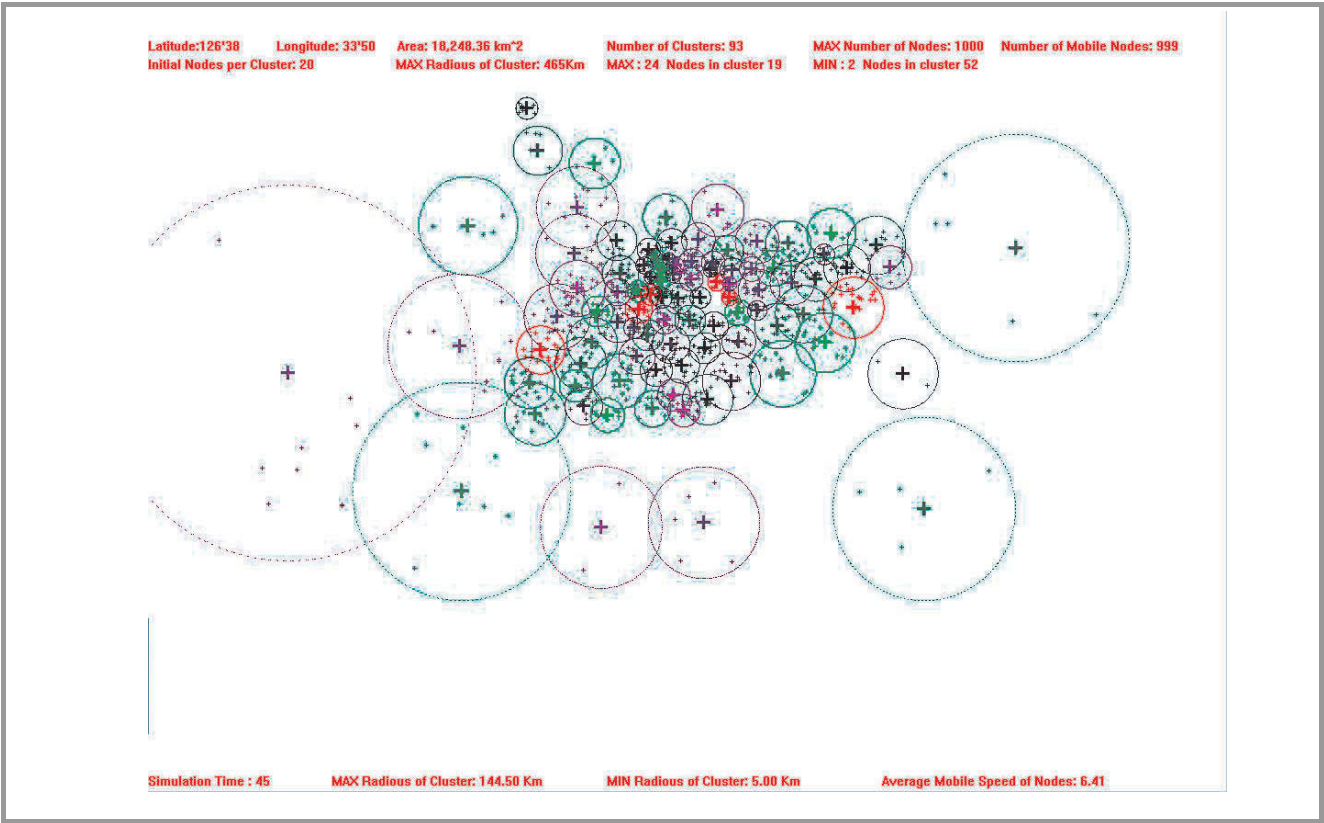


Fig. 11. Placement and coverage of MBSs after 45 min.

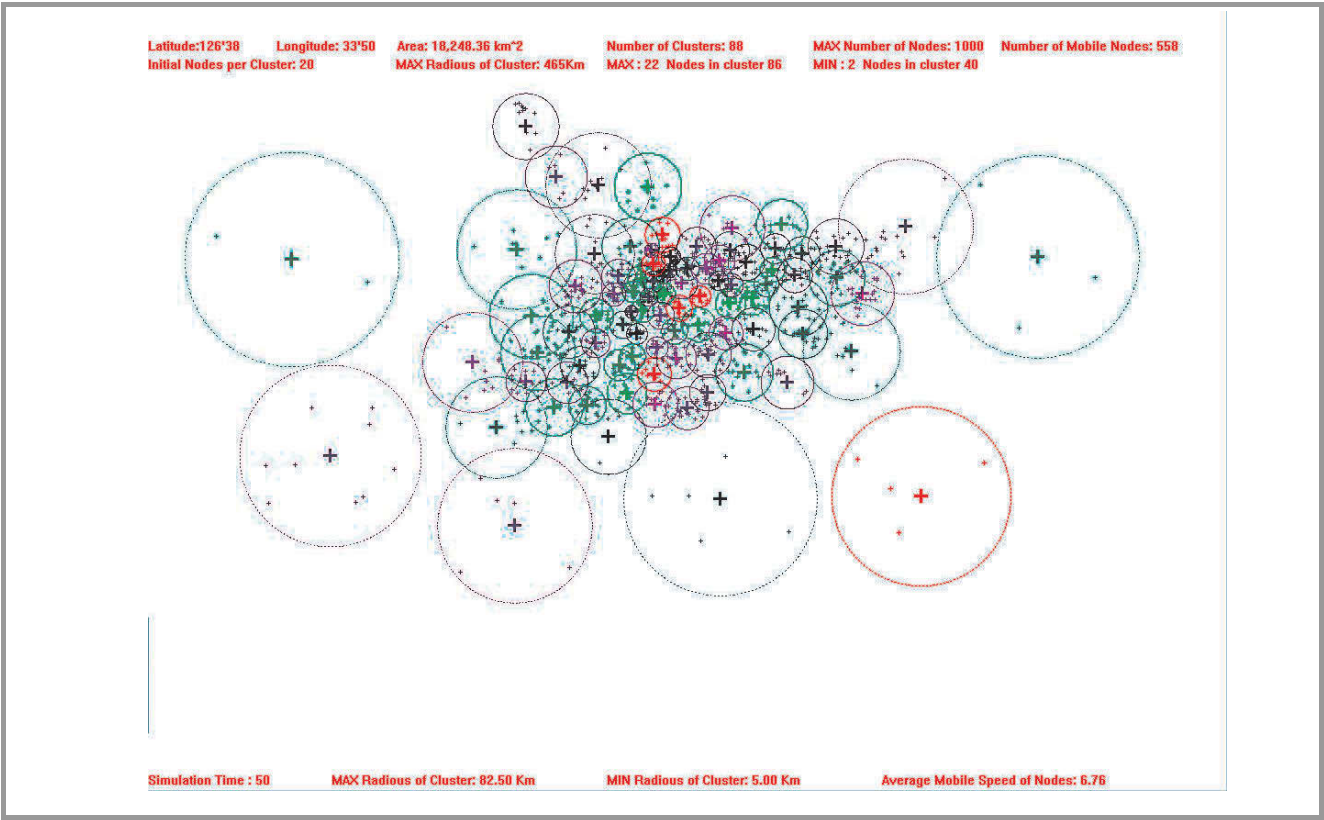


Fig. 12. Placement and coverage of MBSs after 50 min.

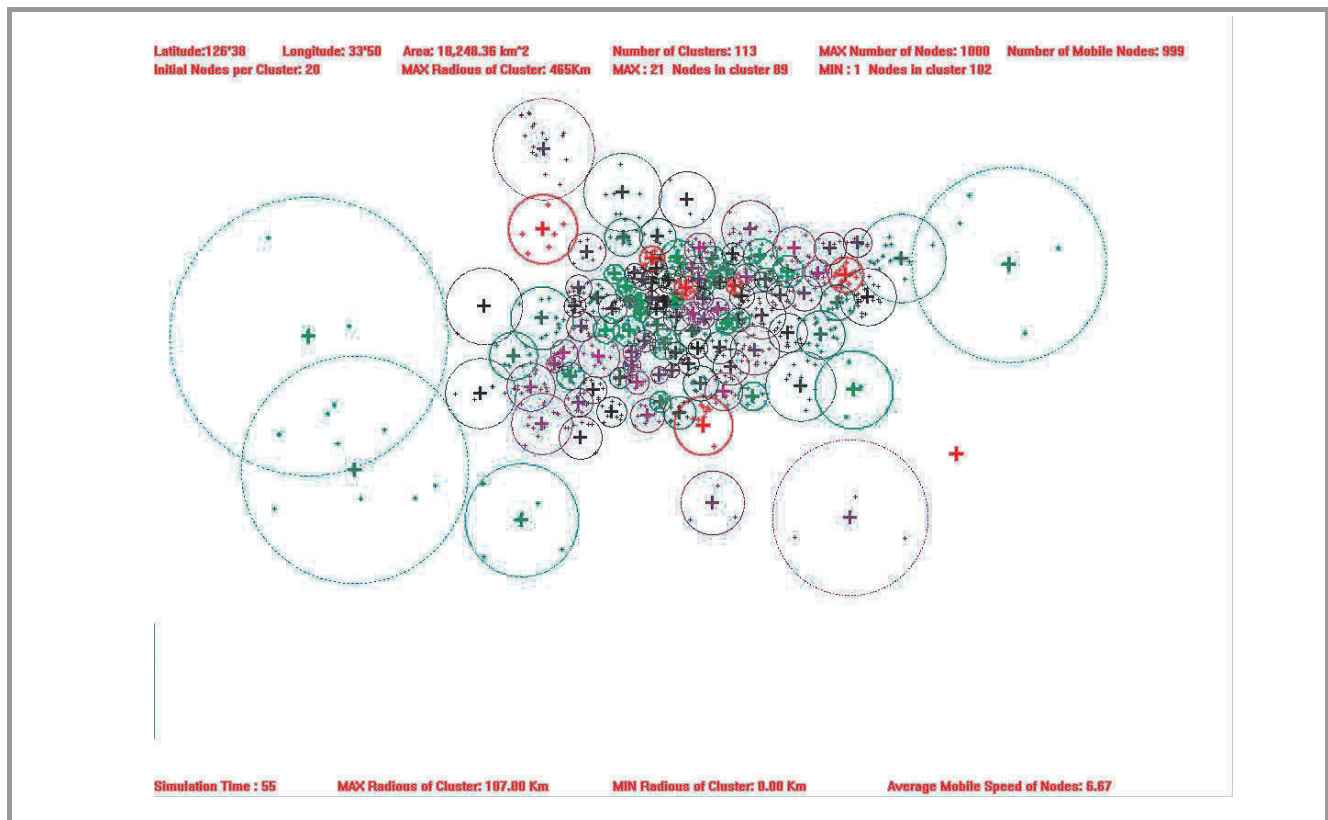


Fig. 13. Placement and coverage of MBSs after 55 min.

with smaller number of nodes must be merged. The ultimate goal of postprocessing is to assign adequate number of nodes equally likely to clusters. We experienced the overloaded clusterhead (MBS) is usually a problem for network configuration. We will provide a simple postprocessing algorithm in Section 6.

We experienced less than 1 s for each clustering with usual Pentium based personal computers. We believe this guarantees realtimeness of our clustering algorithm.

## 6. Conclusion and Future Research

In this paper, we present a HAP MBS placement solution over multiple HAP based wireless network. Considering the number of ground nodes and area coverage, mobile ground nodes are clustered in order to prepare placement location. As a result, each HAP based MBS can be placed at geographically suitable location in order to serve as a BS of a cluster. We experiment these scenarios by simulation and the results are visualized in a specific area of Cheju Island.

However, we can find several problems regarding the clustering results.

First, we must suppress of mobility of MBSs for stable network service. Even if we assumed mobility of HAP MBSs, high mobility of existing MBSs would cause instable network service.

Second, other than K-mean algorithm, more sophisticated algorithm must be introduced. We are now focusing on EM and BIRCH [12], [13] as a core algorithm for HAP MBS placement. For uncertain coordinates of highly mobile nodes we can adjust probabilities for such nodes to be members of a specific cluster with the speed of mobile node or so. Then this probability set will be directly applied to EM based clustering.

Third, the number of nodes per cluster must be distributed evenly for stable network bandwidth allocation. Since there are no clustering algorithms that calibrate the number of elements per cluster, a new clustering algorithm must be introduced. And we sometimes observe inclusive clusters and semi-inclusive (highly overlapped) clusters in simulation results. BIRCH has the abilities of splitting or merging clusters by use of CF tree. These capabilities can be a useful method to assign evenly distributed number of nodes to clusters.

These combination of clustering algorithms ultimately leads to a multiphase clustering algorithm. We are now considering the following process.

1. Preprocessing with static clustering algorithms such as K-mean.
2. Main clustering with dynamic clustering algorithms such as EM.
3. Postprocessing for split and merge of clusters with algorithms such as BIRCH.

We believe the outmost clustering algorithm would be BIRCH since it requires another clustering algorithm for core clustering features.

For postprocessing as mentioned in Section 4, we have the following idea. We can merge a cluster into outer cluster when the cluster is inclusive under the condition:

$$|C_1 - C_3| + R_3 < R_1 \quad \text{and} \quad N_1 + N_3 < N$$

where  $C_i$  is a coordinate for cluster  $i$ ,  $R_i$  is a radius for cluster  $i$ , and  $N_i$  is number of nodes in a cluster  $i$ .

The first term stands for cluster  $C_3$  is totally inclusive in cluster  $C_1$  and the second term restricts merge if the total number of nodes in two clusters exceeds a desired number of nodes  $N$  in a cluster.

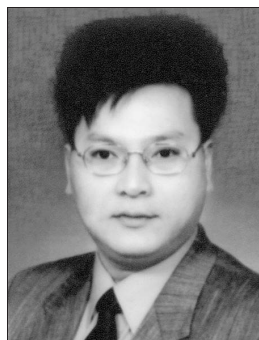
Similarly we can merge semi-inclusive clusters where most part of a cluster, e.g., 75%, is covered by another cluster.

## Acknowledgements

The author is deeply indebted to all helpers of the papers quoted, which testify many fruitful long-lasting cooperations. This work was supported by the Korea Research Foundation Grant (KRF-2009-013-D00106).

## References

- [1] "Korea radio promotion association" [Online]. Available: <http://www.rapa.or.kr>
- [2] J. Thornton, D. Grace, C. Spillard, T. Konefal, and T. C. Tozer, "Broadband communications from a high-altitude platform", *Electron. Commun. Eng. J.*, vol. 13, no. 3, pp. 138–144, 2001.
- [3] "Japan Aerospace Exploration Agency" [Online]. Available: <http://www.jaxa.jp>
- [4] A. K. Widiawan and R. Tafazolli, "High altitude platform station (HAPS): a review of new infrastructure development for future wireless communications", *Wirel. Pers. Commun.*, vol. 42, no. 3, pp. 387–404, 2007.
- [5] "Dudley Lab's list of frequency allocations" [Online]. Available: <http://www.dudleylab.com/freqaloc.html>
- [6] *Final Acts – WRC-07*. Geneva: ITU, 2007.
- [7] D. Grace, C. Spillard, J. Thornton, and T. C. Toze, "Channel assignment strategies for a high altitude platform spot-beam architecture", in *13th IEEE Int. Symp. Pers. Indoor Mob. Radio Commun.*, Lisbon, Portugal, 2002, vol. 4, pp. 1586–1590.
- [8] D. S. Dasha, A. Durrezi, and R. Jain, "Routing of VoIP traffic in multi-layered satellite networks", in *Proc. SPIE Perform. Contr. Next Generat. Commun. Netw.*, Orlando, USA, 2003, vol. 5244, pp. 65–75.
- [9] F. De Rango, A. Malfitano, and S. Marano, "PER evaluation for IEEE 802.16-SC and 802.16e protocol in HAP architecture with user mobility under different modulation schemes", in *IEEE Globecom Symp. Wirel. Commun. Netw.*, San Francisco, USA, 2006, pp. 1–6.
- [10] R. A. J. Purvinskis, "Interplatform links". University South Australia, Adelaide, 2003.
- [11] T. Tsujii, J. Wang, L. Dai, and C. Rizos, "A technique for precise positioning of high altitude platforms system (HAPS) using a GPS ground reference network", in *14th Int. Tech. Meet. Satel. Div. U.S. Inst. Navig.*, Salt Lake City, USA, 2001, pp. 1017–1026.
- [12] J. MacQueen, "Some methods for classification and analysis of multivariate observations", in *Proc. Fifth Berkeley Symp. Math. Statist. Prob.*, Berkeley, USA, 1967, vol. 1, pp. 281–297.
- [13] I. S. Dhillon and D. S. Modha, "A data-clustering algorithm on distributed memory multiprocessors", in *Large-Scale Parallel Data Mining*, Lecture Notes in Computer Science, vol. 1759. Berlin: Springer, 1999, pp. 245–260.
- [14] V. Faber, "Clustering and the continuous K-means algorithm", *Los Alamos Sci.*, no. 22, pp. 138–144, 1994.
- [15] J. Horwath, D. Grace, D. Giggenbach, M. Knappek, and N. Perlot, "Optical communication from HAPs – overview of the stratospheric optical payload experiment (STROPEX)", in *22nd AIAA Int. Commun. Satel. Syst. Conf. Exhib. ICSSC*, Monterey, USA, 2004.
- [16] D. Giggenbach and J. Horwath "Optical free-space communications downlinks from stratospheric platforms – overview on STROPEX, the optical communications experiment of CAPANINA", in *Proc. 2005 IST Mob. Sum. Conf.*, Dresden, Germany, 2005.
- [17] E. Leitgeb, K. Zettl, S. S. Muhammad, N. Schmitt, and W. Rehm, "Investigation in free space optical communication links between unmanned aerial vehicles (UAVs)", in *9th Int. Conf. Transp. Opt. Netw. ICTON*, Rome, Italy, 2007, vol. 3, pp. 152–155.
- [18] D. Giggenbach, B. Eppe, J. Horwath, and F. Moll, "Optical satellite downlinks to optical ground stations and high-altitude platforms", in *Lecture Notes in Electrical Engineering*, vol. 16. Berlin-Heidelberg: Springer, 2008, pp. 331–349.
- [19] S. Karapantazis and F. N. Pavlidou, "The role of high altitude platforms in beyond 3G networks", *IEEE Wirel. Commun.*, vol. 12, no. 6, pp. 33–41, 2005.
- [20] M. Asvial, "Traffic model and performance analysis for HAPs networks", in *Proc. 2007 IEEE Int. Conf. Telecommun. Malaysia Int. Conf. Commun.*, Penang, Malaysia, 2007, pp. 278–282.



**Ha Yoon Song** earned degrees of B.S. and M.Sc. from the Seoul National University, Korea, 1991 and 1993, respectively, and earned Ph.D. degree from the Computer Science Department, University of California, Los Angeles, USA, in the year 2001. Currently he is an Associate Professor of Department of Computer Engineering,

School of Information and Computer Engineering at the Hongik University, Seoul, Korea. His research interests include network simulation, especially for satellite and high altitude platforms based wireless networks plus mobility model, and mobile sensor networks as well, especially for the localization techniques of mobile sensor nodes. One of his current interest, in cooperation with many companies is data broadcasting technology with a set-top-box. At present he is at the Institute of Computer Technology, Vienna University of Technology, Austria, for the year 2009 as a visiting scholar, and is working for international research cooperation between European Nations and the Republic of Korea.

Department of Computer Engineering  
Hongik University  
72-1 Sangsu, Mapo, Seoul, Republic of Korea  
e-mail: [song@ict.tuwien.ac.at](mailto:song@ict.tuwien.ac.at)  
Institute of Computer Technology  
Vienna University of Technology  
Gusshausstrasse 27-29/E384  
A-1040 Vienna, Austria

# A Framework for Detection of Selfishness in Multihop Mobile Ad Hoc Networks

Jerzy Konorski and Rafał Orlikowski

**Abstract**—The paper discusses the need for a fully-distributed selfishness detection mechanism dedicated for multihop wireless ad hoc networks which nodes may exhibit selfish forwarding behavior. The main contribution of this paper is an introduction to a novel approach for detecting and coping with the selfish nodes. Paper describes a new framework based on Dempster-Shafer theory-based selfishness detection framework (DST-SDF) with some mathematical background and simulation analysis.

**Keywords**—*reputation system, selfish behavior, wireless ad hoc network.*

## 1. Introduction

Mobile ad hoc networks (MANETs) and ad hoc wireless sensor networks (WSNs) are collections of mobile nodes that exchange packets over a wireless transmission medium. There may be pairs of nodes out of each other's reception range, for which the only way of exchanging data is via in-range nodes acting as packet forwarders, i.e., agreeing to relay packets on behalf of other nodes. However, packet forwarding costs extra energy and bandwidth, each being a scarce resource in wireless ad hoc devices. Rational nodes try to save energy and bandwidth as much as possible, and the most obvious way of doing it is by refusing to relay packets. Such non-cooperative behavior is usually called *selfish*. Without a mechanism preventing it, MANETs and/or ad hoc WSNs become unreliable. Selfishness is to be distinguished from malicious behavior, a type of non-cooperative behavior that brings no tangible benefit to the perpetrators.

Prevention, detection and/or mitigation of selfishness, as well as enforcement of cooperative behavior among MANET or WSN nodes have recently received considerable attention. Currently there are a large number of solutions addressing these goals. A promising class of solutions are reputation-based systems, where the cooperation goals are achieved by way of determination and sharing reputation values among all the network nodes or within groups thereof.

In this work we propose a new approach for detection of non-cooperative (selfish) behavior in the wireless mobile ad hoc networks. The solution is a framework which can be used by the reputation-based systems to detect selfishness. It can replace standard, very often faulty selfishness detection mechanisms (e.g., based on the well-known

*watchdog* mechanism). Because our framework is based on Dempster-Shafer theory (DST) [1]–[4] we call it Dempster-Shafer theory-based selfishness detection framework (DST-SDF).

The rest of the paper is organized as follows: Section 2 discusses related work and outlines some of the well-known methods of selfishness evaluation. Section 3 describes the general concept of our approach, Section 4 contains a brief introduction to Dempster-Shafer theory and the methods of evidence combinations with uncertain information. Section 5 describes DST-SDF in more detail. Sample performance evaluation results are reported in Section 6. Finally, Section 7 states conclusions and outlines future work.

## 2. Related Work

Enforcement of cooperative behavior in MANETs has been the subject of a number of works. Basically, two types of solutions dealing with non-cooperative (malicious as well as selfish) nodes are being proposed. The first type are schemes based on virtual currency, e.g., Nuglets [5] or Sprite [6], that use a form of micropayments to build incentives for cooperation. These are usually quite complex and hard to implement in real networks, typically require tamper-proof hardware in each node or a trusted third party to ensure transaction security.

More promising type of solutions are reputation-based schemes. The most popular ones include cooperation of nodes fairness in dynamic ad hoc networks (CONFIDANT) [7], collaborative reputation mechanism (CORE) [8], secure and objective reputation-based incentive scheme (SORI) [9], observation-based cooperation enforcement in ad hoc networks (OCEAN) [10] and reputation-based mechanism for isolating selfish nodes in ad hoc networks [11], locally aware reputation system (LARS) [12].

The concepts of all of the above reputation-based systems are very similar. The key functional aspects they all share are as follows. Each network node:

- gathers information about the other nodes' behavior;
- calculates reputation values associated with each other node based on direct behavioral information and possibly additional indirect information (in the form of recommendations) received from third-party nodes;

- shares evaluated reputation values or direct behavioral information with all the other nodes (in the case of global reputation systems) or within the immediate neighborhood (in the case local reputation systems);
- tries to enforce cooperative behavior of the other ones by introducing different kinds of punishment (e.g., isolation of non-cooperative nodes from the network);
- excludes nodes it considers non-cooperative from paths used by the packets it forwards taking advantage of standard route selection processes.

Currently existing reputation systems have a number of drawbacks, which our solution aims at overcoming, and which can be summarized as follows:

- **Lack of reliable non-cooperative behavior detection mechanisms.** Gathering information about the other nodes' behavior involves additional external mechanisms. All of the above mentioned reputation-based solutions (besides the one described in [11]) use the watchdog mechanism for this purpose. Therefore each network node is obliged to promiscuously overhear transmissions by its neighbors to determine their cooperative or non-cooperative behavior. It is commonly recognized that watchdog is a faulty tool by nature. Obviously there are other approaches of non-cooperative behavior detection like in [11], but there are persistent problems with distinguishing real from apparent non-cooperative behavior.
- **Lack of robustness against false indirect behavioral information.** Current reputation-based systems cannot effectively cope with indirect behavioral information (recommendations) dictated by ill will, such as denial of service (DoS) attacks or collusion.
- **Ineffective distribution of indirect behavioral information.** Known reputation-based systems introduce significant communication overhead related to the distribution of recommendation messages.

### 3. Solution Overview

The DST-SDF is dedicated for MANETs based on standard routing like dynamic source routing (DSR) [13]. The main concept relies on end-to-end packet acknowledgments in the following way: every time a source node sends a packet to a destination node, it waits for a certain predefined time for an acknowledgement of the packet. If one arrives within the predefined time, the source node has reason to claim that all nodes on the path are cooperative (none is selfish). Otherwise if there are no other indications of faultiness on the path (e.g., RERR messages), the source node knows that there are selfish nodes on the path. Whenever an acknowledgment does or does not arrive in time, a special *recommendation message* is sent out to inform the other

nodes about the detected situation (selfish or cooperative behavior on the path, respectively). Every node in the network is equipped with a dedicated component executing a DST-based algorithm that uses received recommendation messages to evaluate the selfishness of each node. The resulting values can be used as routing metrics while selecting packets' routes in the near future. A more detailed description of the proposed solution is presented further in Section 5.

The DST-SDF differs from the existing ones in the following main respects:

- There is no need to overhear immediate neighbor nodes' transmissions to detect their cooperative or non-cooperative behavior – no additional tools (e.g., watchdogs) to cover this functionality are needed.
- Communication overhead is significantly reduced through an economy of scale – no recommendation message pertains to a single node; rather, each one pertains to a set of nodes, namely a path.
- Determination of nodes' selfishness is based on consistent evidence received both directly (as derived from the successive packet acknowledgments or lack thereof) and indirectly via recommendation messages.
- DST is used to determine selfishness.

Further we describe our approach in more detail, but before we do, we give some introduction to DST and the methods it uses to combine pieces of uncertain information into new information, and give some arguments for employing the theory as the basis of DST-SDF.

### 4. Overview of Dempster-Shafer Theory

The Dempster-Shafer theory, developed by A. P. Dempster and G. Shafer in the 1960s and 1970s [1]–[4], offers an alternative to classical probability as a formal representation of uncertainty. It is in fact a mathematical theory of evidence based on the so-called belief functions and plausible reasoning, and may be used to combine separate and independent pieces of evidence to quantify the belief in a given statement, further reflected as an *evidence value*. DST is a potentially valuable tool for the evaluation of risk and reliability in engineering applications when it is not possible to obtain precise measurements from experiments, or when knowledge is independently elicited from a number of experts. Instead of giving a thorough exposition of the mathematical basics of the theory, we only focus on those of its aspects used in our DST-SDF approach.

Statements in DST are related to some *universal set*  $\Theta$  and take the form of claims that a particular element  $x$  of  $\Theta$  belongs to a set  $X \subseteq \Theta$ . Belief in a statement derives from a DST primitive called *basic probability assignment*. It is a function mapping the powerset of  $\Theta$  onto the inter-

val  $[0, 1] : m : 2^\Theta \rightarrow [0, 1]$ , with the normalization constraint satisfied over the entire powerset. That is, with each  $X \subseteq \Theta$  (i.e.,  $X \in 2^\Theta$ ) is associated a real number  $m(X)$  between 0 and 1 that measures the amount of trust we put in the claim that  $x \in X$ , and there is no reason to believe that  $x \in X'$  for any  $X' \subset X$  (i.e., no evidence supports a stronger statement), with  $m(\emptyset) = 0$ , and

$$\sum_{X \in 2^\Theta} m(X) = 1. \quad (1)$$

Belief, or evidence value, associated with  $X$  is then defined as

$$ev(X) = \sum_{X' \in 2^\Theta | X' \subseteq X} m(X'), \quad (2)$$

i.e., is the arithmetic sum of basic probability assignments to statements at least as strong as the one in question.

As an example, consider a network node that can be designated as *SELFISH* or *NONSELFISH*. Thus we have the universal set  $\Theta = \{SELFISH, NONSELFISH\}$ . Assuming that there is enough information to claim that the node is *SELFISH* with probability 0.1 and *NONSELFISH* with probability 0.9, we can write down the following basic probability assignment:

$$m(X) = \begin{cases} 0.1, & X = \{SELFISH\}, \\ 0.9, & X = \{NONSELFISH\}. \end{cases} \quad (3)$$

This resembles classical probability distribution over  $\Theta$  and results in the distribution of evidence values identical with Eq. (3). However, one might just as well assign a basic probability of 0.9 to not knowing at all whether the node is *SELFISH* or *NONSELFISH*. In that case we get

$$m(X) = \begin{cases} 0.1, & X = \{SELFISH\}, \\ 0.9, & X = \{SELFISH, NONSELFISH\}, \end{cases} \quad (4)$$

and the resulting distribution of evidence values becomes  $ev(\{SELFISH\}) = 0.1$  and  $ev(\{NONSELFISH\}) = 0$  (note that they need not sum up to 1).

A useful feature of DST is the formalism to express the basic probability assignment associated with a subset of  $\Theta$  through other basic probability assignments associated with subsets of  $\Theta$ ; this enables, e.g., combination of (possibly conflicting) pieces of evidence obtained from multiple sources into a new piece of evidence in the course of knowledge updating. Although several evidence combination rules exist, dealing in different ways with conflicting evidence, hereafter we stick to a simple one known as Dempster's combination rule. Given two pieces of evidence in the form of basic probability assignments  $m_1$  and  $m_2$  over  $2^\Theta$ , the resulting basic probability assignment for a set  $X \subseteq \Theta$  is defined as

$$m(X) = (m_1 \oplus m_2)(X) = \frac{\sum_{Y, Z \in 2^\Theta | Y \cap Z = X} m_1(Y)m_2(Z)}{1 - C}, \quad (5)$$

where the factor  $C$  represents the total basic probability mass associated with conflicting evidence and is given by

$$C = \sum_{Y, Z \in 2^\Theta | Y \cap Z = \emptyset} m_1(Y)m_2(Z). \quad (6)$$

Coming back to our example, let  $m_1$  be as in Eq. (4) and

$$m_2(X) = \begin{cases} 0, & X = \{SELFISH\}, \\ 0.5, & X = \{NONSELFISH\}, \\ 0.5, & X = \{SELFISH, NONSELFISH\}, \end{cases} \quad (7)$$

then

$$\begin{aligned} & (m_1 \oplus m_2)(\{SELFISH\}) \\ &= \frac{m_1(\{SELFISH\})m_2(\{SELFISH, NONSELFISH\})}{1 - m_1(\{SELFISH\})m_2(\{NONSELFISH\})} \\ &= \frac{0.1 \cdot 0.5}{1 - 0.1 \cdot 0.5} \approx 0.053. \end{aligned} \quad (8)$$

The main reasons to advocate DST in our framework are as follows:

- It is able to cope with two kinds of uncertainty that can be expected in a mobile ad hoc environment: *aleatory* uncertainty, resulting from the fact that network nodes can behave in a random way (e.g., perform selective or random packet dropping) and *epistemic* uncertainty, resulting from the lack of knowledge about the behavior of other nodes (recall that there is no direct transmission overhearing mechanism to control nodes' behavior, such as a watchdog, hence, when detecting possible non-cooperative behavior one has to rely on incomplete information based on evidence originating from different sources).
- There are many sources of information on which to base the evaluation of selfishness; as a consequence, there inevitably arise ambiguities and conflicting information (possibly, but not necessarily due to false recommendations).

## 5. The DST-SDF Details

### 5.1. Assumptions and Implementation

Each time a source node  $S$  wishes to send a packet to a destination node  $D$ , a path selection process according to DSR is performed to determine an appropriate path  $p_{S,D}$  from  $S$  to  $D$  for the packets. Let us assume that the selected path  $p_{S,D}$  consists of the set  $N_{S,D}$  of intermediate nodes, whose cardinality (i.e., the length of  $p_{S,D}$ ) is  $L_{S,D}$ . As regards routing, the only restriction we place on our solution is that a source node should know beforehand the identities of all the intermediate nodes on the path being selected for any packet (note that on-demand distance vector routing (AODV)-like protocols are therefore unsuitable as they do not reveal intermediate nodes to a source node). Although DST-SDF can cope both with single-path and multipath routing protocols, to simplify the description we further assume that MANET nodes only employ a single-path routing protocol like DSR.

Every network node implements a dedicated component (Fig. 1) responsible for maintaining information about

the other nodes' behavior. We call it the evidence manager component (EMC). Its only task is to detect selfish nodes based on provided input information of two types:

- direct, i.e., nodes' own observations (arrival/lack of arrival of packets' acknowledgements);
- indirect, i.e., information spread all over the network in the form of recommendation messages.

The output data of EMC can then be fed into the routing protocol's path selection mechanism in a standard way typical of traditional reputation-based systems.

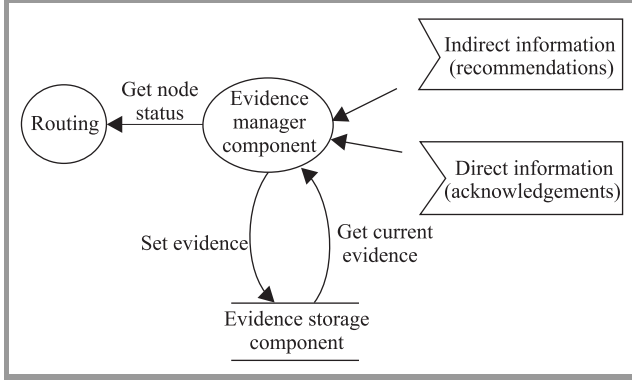


Fig. 1. Node's internal dataflow diagram.

Inside the EMC, behavioral data for each node are converted to and maintained as evidence values. Current evidence values evaluated by EMC for all the nodes are stored in an evidence storage component (ESC). When a node becomes operational (i.e., joins the network) and before it receives input information (direct or indirect) for the first time, an arbitrary initial basic probability assignment is created. Throughout the node's operational lifetime within the network, it is updated according to subsequent input events (i.e., reception of direct or indirect behavioral information regarding other nodes).

### 5.2. Direct Information

At the outset, every network node maintains an initial basic probability assignment regarding all the other nodes:

$$init\_m_{ij}(X) = \begin{cases} 0.5, & X = \{SELFISH\}, \\ 0.5, & X = \{NONSELFISH\}, \end{cases} \quad (9)$$

where  $init\_m_{ij}$  denotes the initial basic probability assignment at node  $i$  regarding node  $j$ 's status (*SELFISH*, *NONSELFISH*, or not known to be *SELFISH* or *NONSELFISH*). These initial assignments simply tell node  $i$  to consider node  $j$  *SELFISH* and *NONSELFISH* with the same uncertainty, by setting the probabilities of these two designations to 0.5 (node  $i$  has no information about node  $j$ ). As described earlier, every time a source node  $S$  sends a packet to a destination node  $D$ , it waits for an acknowledgment of the packet. If it arrives in time, node  $S$  is certain that all nodes along the selected path  $p_{S,D}$  have

behaved cooperatively (there is no selfish node in  $N_{S,D}$ ). When the source node  $S$  receives in time an acknowledgment for a packet sent over  $p_{S,D}$ , it creates the following new basic probability assignments regarding each node  $j \in N_{S,D}$ :

$$init\_m_{Sj}(X) = \begin{cases} 0, & X = \{SELFISH\}, \\ 1, & X = \{NONSELFISH\}, \end{cases} \quad (10)$$

and updates according to Eq. (5) its basic probability assignment:

$$curr\_m_{Sj} := init\_m_{Sj} \oplus new\_m_{Sj}, \quad (11)$$

where  $curr\_m_{Sj}$  is the current basic probability assignment at node  $S$  regarding node  $j$ , and  $\oplus$  denotes Dempster's evidence combination operator as in Eq. (5).

If no acknowledgment for the packet arrives within the predefined time, the source station  $S$  can only claim that there are selfish nodes in  $N_{S,D}$ . Node  $S$  does not know exactly which one of the nodes in  $N_{S,D}$  is *SELFISH*, it does not even know how many *SELFISH* nodes there are, it is just certain that there is at least one such node. While one can imagine making any kind of assumptions as to the conjectured number of *SELFISH* nodes in  $N_{S,D}$ , our approach relies on the following simplest assumption: *if no acknowledgment for a packet sent over  $p_{S,D}$  has arrived in time, only one SELFISH node is conjectured to be in  $N_{S,D}$* . It is probably appropriate to stress, in view of this somewhat arbitrary and simplifying assumption, that our approach is expected to provide efficient detection of selfishness in the first place, generality and conceptual elegance being secondary considerations.

The next simplification of ours is taking the classical Bayesian approach whereby some probabilities can be assigned to a concrete node being *SELFISH*, and finally restricting our attention to uniform probabilities. That is, given there is exactly one *SELFISH* node in  $N_{S,D}$ , and because the source node  $S$  has no knowledge as to exactly which node it is, it assumes that all nodes in  $N_{S,D}$  are *SELFISH* with the same probability  $P = 1/L_{S,D}$  (recall that  $L$  is the length of  $p_{S,D}$ ). The following new basic probability assignments are then created at node  $S$  regarding each node  $j \in N_{S,D}$ :

$$new\_m_{Sj}(X) = \begin{cases} P, & X = \{SELFISH\}, \\ 1-P, & X = \{SELFISH, NONSELFISH\}. \end{cases} \quad (12)$$

Node  $S$  next updates its initial or (if it already exists) current basic probability assignments regarding each node  $j \in N_{S,D}$ , i.e., according to Eq. (11) or to

$$curr\_m_{Sj} := curr\_m_{Sj} \oplus new\_m_{Sj}. \quad (13)$$

### 5.3. Indirect Information

Whenever a packet's source node receives an acknowledgment for a packet sent over  $p_{S,D}$  or observes the predefined time for acknowledgment arrival expired, it spreads

a recommendation message all over the network. The message lists the set  $N_{S,D}$  and contains an indication of the respective path's behavior status that can assume one of two values: *SELFISH* (if the acknowledgment has arrived) or *NONSELFISH* (otherwise). An important point to note is that unlike in traditional reputation-based systems, only packets' source nodes ever spread out recommendation messages. When a given node  $i$  receives from another node a recommendation message, it builds basic probability assignments regarding all the nodes listed therein, i.e.,  $j \in N_{S,D}$ , based on the path behavior indication. If the path behavior indication is *NONSELFISH* then

$$new\_m_{ij}(X) = \begin{cases} u, & X = \{NONSELFISH\}, \\ 1-u, & X = \{SELFISH, NONSELFISH\}, \end{cases} \quad (14)$$

whereas if the path behavior indication is *SELFISH* then

$$\begin{aligned} new\_m_{ij}(X) &= \begin{cases} uP, & X = \{NONSELFISH\}, \\ 1-(1-u)P, & X = \{SELFISH, NONSELFISH\}. \end{cases} \end{aligned} \quad (15)$$

The factor  $u \in [0, 1]$  present in Eqs. (14) and (15) accounts for the possibility that the recommendation messages can be faked or modified by malicious intermediate nodes; it is needed in order to represent uncertainty created by recommendation messages and weigh their influence upon current basic probability assignments. In other words,  $u$  is the value reflecting how much trust a recipient of the recommendation message puts in it. The value of  $u$  can be different for each recommendation message (e.g., depending on its source node). Node  $i$  next updates its initial *init\_m* or current *curr\_m* (if it already exists) basic probability assignments regarding all nodes in  $N_{S,D}$  analogously with Eqs. (11) or (13).

Not exactly according to Eq. (2), but in the spirit of DST, we assume that node  $j$  is considered by node  $i$  as:

- selfish, if  $curr\_m_{ij}(\{SELFISH\}) \geq T$ ,
- nonselfish, if  $curr\_m_{ij}(\{NONSELFISH\}) \geq T$ ,
- undefined, if  $curr\_m_{ij}(\{SELFISH\}) < T$   
and  $curr\_m_{ij}(\{NONSELFISH\}) < T$ ,

where  $T \in (0.5, 1]$  is a selfishness threshold. It is very important to come up with an appropriate  $T$  value. Too low a value contributes to false accusations, whereas too high one lengthens the time needed to detect selfish nodes and in the worst case can prevent DST-SDF from determining nodes' selfishness at all.

## 6. Simulation

In this section we investigate via simulation the robustness and efficiency of the proposed DST-SDF for detection of node selfishness in a mobile ad hoc network. We try to address the questions how long it takes to detect all selfish nodes and what is the communication overhead

introduced by DST-SDF. The proposed mechanism is implemented and evaluated using the J-Sim tool [14] in a simulation environment composed of IEEE 802.11-based ad hoc networks. The simulated scenario features 100 nodes arranged on a grid with each node pair's reception range confined to one hop. To demonstrate the robustness of our reputation system, we let 10% of the network nodes behave selfishly, i.e., refuse to forward packets.  $T$  is set to 0.8 and  $u$  to 0.9. The DST-SDF efficiency is presented in Fig. 2. Four test scenarios are analyzed with packets' paths of uniform lengths  $L$ .

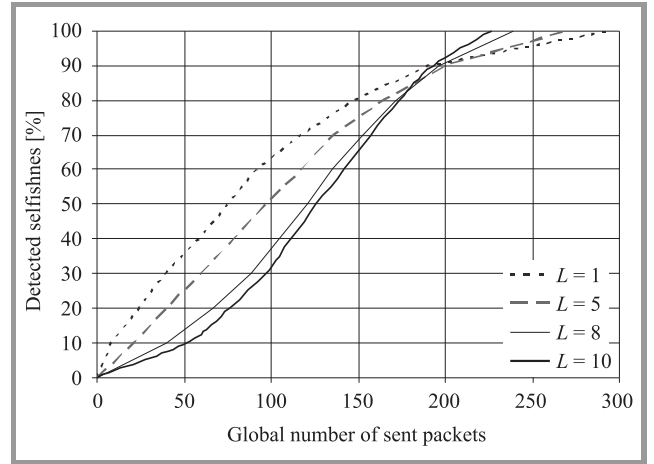


Fig. 2. Efficiency of selfishness detection.

The simulations show that in order to detect all selfish nodes only around 300 packets in total are needed to be sent by all the network nodes. The selfishness detection process can be divided into two phases. The first one covers the time up to about 90% of detected selfish nodes and the second one the remaining percentage. Clearly, the shorter the paths, the higher is the probability  $P$  that a given node along the path whose behavior indication is *SELFISH* has behaved selfishly. In Eqs. (12) and (15), the evidence built is stronger than in the case of longer paths where the probability of selfish behavior is spread among more nodes. DST-SDF needs less strong evidence (less certainty) to take a decision in the case of shorter paths. Conversely, the longer paths, the more uncertain information (weaker evidence) DST-SDF is getting and in order to evaluate selfishness it needs more time than it does in the case of stronger evidence. Nevertheless, the time to detect 90% selfish nodes in our simulation environment turns out to be largely independent of the path length. This apparent anomaly is due to the particular path selection process implemented. Our simulation environment only features end-to-end connections between node pairs at a constant distance  $L$  from each other ( $L = 1, 5, 8$ , or  $10$ ). Hence, the shorter the path, the lower the probability that it passes through a selfish node, and the more paths exist that only pass through cooperative nodes; consequently, more time is required to detect all the selfish nodes. At the level of 90% detected selfish nodes, this effect upon the selfish-

ness detection time happens to almost precisely compensate for the differences in  $P$ .

One of the most outstanding issues in all existing indirect reputation-based systems is the communication overhead they induce. It also affects DST-SDF as a result of the dissemination of recommendation messages. Since the envisaged future DST-SDF implementation may use acknowledgement mechanisms inherent in higher layers of the open system interconnection (OSI) reference model, e.g., TCP, one can argue that ultimately, packets' acknowledgements should not be regarded as extra communication overhead. The total communication overhead induced by DST-SDF in comparison with a generic theoretical reputation-based solution (TRBS) that uses indirect behavioral information is presented in Fig. 3 as a function of the average path length  $L$ . The overhead is expressed as the percentage of the total number of data packets needed to be sent in order to discover all the selfish nodes.

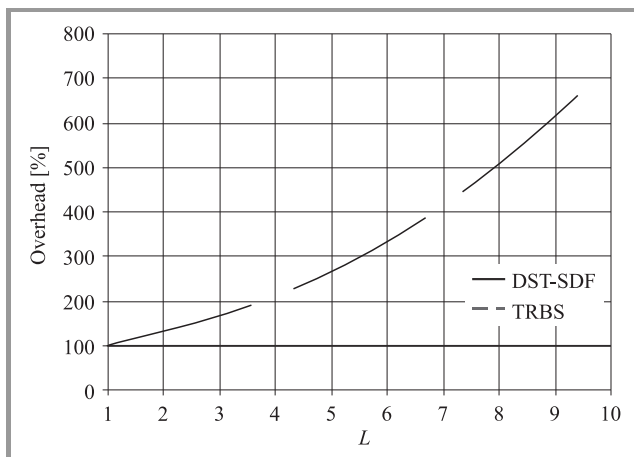


Fig. 3. Recommendation messages overhead.

It is easy to notice that longer paths result in a very distinct advantage of ours over existing reputation systems with respect to the communication overhead. DST-SDF communication overhead stays steady at the 100% level for different  $L$  values, meaning that the number of recommendation messages is equal to the total packets sent. The difference between DST-SDF and TRBS stems from the way recommendation messages are generated. In DST-SDF, they can only originate from packets' source nodes, while in TRBS every node (including intermediate nodes on packets' paths) can originate recommendation messages. In a watchdog-based TRBS, each time the watchdog detects a particular (cooperative or selfish) behavior of an immediate neighbor, a recommendation message is originated. Moreover, a recommendation message, whether containing direct or indirect reputation information, typically pertains to just one node. In DST-SDF, a recommendation message pertains to the whole path, typically containing more than one node, and is sent only by the source node according to whether an acknowledgement for a packet has been received within the predefined time (positive recommendation) or not (negative recommendation).

## 7. Conclusion and Future Work

This paper investigates and presents some aspects of detecting and evaluating selfish node behavior in multihop mobile ad hoc networks. A novel approach to selfishness detection called DST-SDF has been proposed. Preliminary simulations show that DST-SDF does allow to detect fairly quickly all selfish nodes in the network at the cost of definitely lower communication overhead compared to traditional reputation-based systems based on the watchdog mechanism.

Nevertheless, there are still a number of impediments to be overcome. In particular, more work needs to be done on:

- robustness against malicious or colluding nodes (i.e., coping with false accusations or fake positive recommendations);
- reliability and security of recommendation message distribution (e.g., assigning proper weights to recommendations);
- proper configuration of DST-SDF (e.g., of the  $T$  parameter) to ensure higher efficiency;
- the possibility of combining DST-SDF with protocols like the anonymous packet forwarding and congestion control mechanism proposed in a previous paper [15].

## Acknowledgment

Effort sponsored by the Air Force Office of Scientific Research, Air Force Material Command, USAF, under grant FA8655-08-1-3018, also supported in part by the MNiSW grant PBZ-MNiSW-02/II/2007.

## References

- [1] G. Shafer, *A Mathematical Theory of Evidence*. Princeton: Princeton University Press, 1976.
- [2] L. A. Zadeh, "A simple view of the Dempster-Shafer theory of evidence and its implication for the rule of combination", *AI Mag.*, no. 7, pp. 85–90, 1986.
- [3] L. Zhang, "Representation, independence, and combination of evidence in the Dempster-Shafer theory", in *Advances in the Dempster-Shafer Theory of Evidence*, R. R. Yager, J. Kacprzyk, and M. Fedrizzi, Eds. New York: Wiley, 1994.
- [4] K. Sentz and S. Ferson, "Combination of evidence in Dempster-Shafer theory", Tech. Rep. SAND2002-0835, New Mexico, Sandia National Laboratories, 2002.
- [5] L. Buttyan and J.-P. Hubaux, "Nuglets: a virtual currency to stimulate cooperation in self organized mobile ad hoc networks", Tech. Rep. DSC/2001, EPFL, Lausanne, 2001.
- [6] S. Zhong, J. Chen, and R. Yang, "Sprite: a simple, cheatproof credit-based system for mobile ad hoc networks", in *Proc. IEEE Infocom'03 Conf.*, San Francisco, USA, 2003, pp. 1987–1997.
- [7] S. Buchegger and J.-Y. Le Boudec, "Performance analysis of the confidant protocol: cooperation of nodes – fairness in distributed ad-hoc networks", in *Proc. 3rd ACM Int. Symp. Mob. Ad Hoc Netw. Comp. MobiHoc 2002*, Lausanne, Switzerland, 2002, pp. 226–236.
- [8] P. Michiardi and R. Molva, "CORE: a collaborative reputation mechanism to enforce node cooperation in mobile ad hoc networks", in *Proc. IFIP Commun. Multimed. Secur. Conf.*, Portoroz, Slovenia, 2002, pp. 107–121.

- [9] Q. He, D. Wu, and P. Khosla, "SORI: a secure and objective reputation-based incentive scheme for ad hoc networks", in *Proc. IEEE Wirel. Commun. Netw. Conf.*, Pittsburgh, USA, 2004, pp. 825–830.
- [10] S. Bansal and M. Baker, "Observation-based cooperation enforcement in ad hoc networks", Tech. Rep. Arxiv preprint cs.NI/0307012, 2003.
- [11] T. M. Refaei, V. Srivastava, L. Dasilva, and M. Eltoweissy, "A reputation-based mechanism for isolating selfish nodes in ad hoc networks", in *Proc. Second Ann. Int. Conf. Mob. Ubiqu. Syst. Netw. Serv. MobiQuitous'05*, San Diego, USA, 2005, pp. 3–11.
- [12] J. Hu and M. Burmester, "LARS: a locally aware reputation system for mobile ad hoc networks", in *Proc. 44th Ann. South. Reg. Conf.*, Melbourne, USA, 2006, pp. 119–123.
- [13] D. B. Johnson, D. A. Maltz, and Y.-C. Hu, "The dynamic source routing protocol for mobile ad hoc networks (DSR), 2004 [Online]. Available: <http://www.ietf.org/internet-drafts/draft-ietf-manet-dsr-10.txt>
- [14] J-Sim [Online]. Available: [www.j-sim.org](http://www.j-sim.org)
- [15] J. Konorski and R. Orlikowski, "Distributed reputation system for multihop mobile ad hoc networks", in *Proc. 5th Polish-German Telegraf. Symp. PGTS'08*, Berlin, Germany, 2008, pp. 161–167.



**Jerzy Konorski** received his M.Sc. degree in electrical engineering from the Technical University of Gdańsk, Poland, in 1976 and the Ph.D. degree in computer science from the Institute of Computer Science, Polish Academy of Sciences, Warsaw, in 1984. He is currently with the Department of Teleinformatics, Gdańsk Uni-

versity of Technology. He teaches probability theory, operational research, and computer networking, as well as conducts research in wireless networks, performance evaluation, and distributed information systems. He has worked on a number of Ministry-, EU-, and U.S.-sponsored projects, authored over 100 papers published in international journals or conference records, and co-authored another 20. His current work focuses on the application of game theory to medium access control and packet forwarding in wireless networks.

e-mail: [jekon@eti.pg.gda.pl](mailto:jekon@eti.pg.gda.pl)  
 Gdańsk University of Technology  
 G. Narutowicza st 11/12  
 80-952 Gdańsk, Poland



**Rafał Orlikowski** received his M.Sc. degree in telecommunications in 2003 from the Gdańsk University of Technology, Poland. Since 2004 he has been working for R&D Marine Technology Centre at Gdynia, Poland, as a senior software engineer. He is currently working on his Ph.D. thesis devoted to reputation systems in wireless

networks. His research interests include security and non-cooperative behavior in mobile ad hoc networks.

e-mail: [Rafal.Orlikowski@ctm.gdynia.pl](mailto:Rafal.Orlikowski@ctm.gdynia.pl)  
 Research & Development Marine Technology Centre  
 Dickmana st 62  
 81-109 Gdynia, Poland

# Modulo N Backoff Scheme for Effective QoS Differentiation and Increased Bandwidth Utilization in IEEE 802.11 Networks

Tomasz Janczak, Jerzy Konorski, Józef Woźniak, and Krzysztof Pawlikowski

**Abstract**—The paper presents a new modulo N channel access scheme for wireless local area networks (WLANs). The novel solution derives from the distributed coordination function (DCF) of the IEEE 802.11 standard, further elaborated as enhanced distribution channel access (EDCA) by the 802.11e draft specification. The main innovation concerns improvement of the binary exponential backoff scheme used for collision avoidance in 802.11 networks. The most appealing feature of the new modulo N backoff scheme is that it outperforms the original 802.11 solution in terms of channel utilization ratio under any traffic conditions. Furthermore, the modulo N proposal can be naturally augmented with QoS differentiation mechanisms like 802.11e extensions. The prioritized modulo N scheme achieves better throughput-delay characteristics for multimedia traffic when compared with the original 802.11e proposal. At the same time, the new solution retains backward compatibility and includes all features which have made IEEE 802.11 networks extremely popular nowadays.

**Keywords**—channel access, MAC, performance analysis, random backoff, WLAN.

## 1. Introduction

Wireless local area networks (WLANs) have rapidly gained market acceptance over the last few years. The reasons are both growing demand for cable-free communications, as well as advances in portable computers and technology. Although the early WLAN solutions were merely intended as cordless replacement for Ethernet networks, it has now become evident that they must offer wider functionality, and in particular support multimedia traffic.

A significant milestone was marked by the development of the 2nd generation public wireless networks. With the emergence of 3rd generation mobile networks, broadband wireless access becomes possible. The 3rd generation systems, such as universal mobile telecommunications system (UMTS), provide enough bandwidth to support both the existing multimedia applications like speech and upcoming ones, like video-conferencing. One also observes an increased role of wireless networks in providing high-speed Internet access. Wireless LANs are often envisioned as a key element of 4th generation solutions for busy spots

such as airports or commerce centers. In addition, the more and more popular vision of wireless homes opens up even more market opportunities for WLAN appliances. Capability of transferring high-volume multimedia streams becomes a primary goal in the design of a new generation of WLANs.

Growing demand for multimedia traffic calls for efficient bandwidth management over the scarce wireless medium. Due to the scarcity of radio resources, WLAN solutions must cope with stringent bandwidth limits, unlike their fixed counterparts. New channel access algorithms are needed to govern radio resource sharing in a way that meets multimedia application requirements while achieving high wireless medium utilization. As the speed of wireless transmission increases, the latter becomes a hot issue. At present, the medium access control (MAC)-layer protocol overhead in IEEE 802.11 networks becomes so huge that it can consume as much as 50% of available bandwidth or more [1].

This paper presents a new wireless channel access scheme, built on the basis of the IEEE 802.11 [2] and IEEE 802.11e solutions [3]. The novel proposal, called *modulo N* backoff, aims at increasing the overall utilization of a radio channel, while ensuring firm quality of service (QoS) guarantees. The novel proposal significantly outperforms 802.11e as far as the overall channel utilization ratio is concerned. Depending on traffic conditions, the modulo N backoff scheme increases the overall channel utilization ratio from 5% to 30% as compared with 802.11e enhanced distribution channel access (EDCA). The bandwidth gain depends on the number of active stations in each access cycle and the average packet size. The most appealing feature of modulo N is that under no conditions does it perform worse than its 802.11 predecessor. Furthermore, the prioritized variant of the modulo N scheme enables very effective QoS differentiation, more flexible than the original EDCA proposal from 802.11e.

The paper is organized as follows: Section 2 outlines the original channel access scheme in 802.11 networks. Section 3 introduces the concept of modulo N backoff scheme. In Section 4, optimal parameters of modulo N operation are sought. Section 5 augments the pure modulo N scheme with QoS differentiation mechanisms. Section 6 concludes the paper.

## 2. The IEEE 802.11 Backoff Scheme

The IEEE 802.11 standard covers the two lowest layers of the open system interconnection (OSI) model, namely the physical (PHY) and the data link layer. This paper focuses on the MAC sublayer of the latter, as it governs channel access.

The IEEE 802.11e draft [3] specification defines two operating modes for the 802.11 MAC protocol: EDCA and hybrid coordination function (HCF) controlled channel access (HCCA). EDCA is the basic and mandatory operational mode. It implements a fully distributed channel access algorithm and directly derives from the IEEE 802.11 distributed coordination function (DCF) [2]. Like its DCF predecessor, EDCA employs the carrier sense multiple access (CSMA) scheme that differs from classical Ethernet in that collision avoidance (CA) replaces collision detection.

The DCF/EDCA collision avoidance relies on the truncated binary exponential backoff (BEB) strategy, originally employed in IEEE 802.3/Ethernet networks. When an Ethernet station has a frame to transmit, it first senses the channel carrier. Once a station detects any foreign transmission, it defers until the transmission ends and then, after a fixed-duration interframe space, sends its own DATA frame. A collision occurs if two or more stations simultaneously resume transmission after deferring. Ethernet networks allow easy detection of collisions by observing changes in the signal voltage. When a transmitting station detects a collision, it delays the next transmission attempt by an integer number of slot times. The number of backoff slots is drawn from a uniform distribution from a contention window  $< 0, 2^n - 1 >$ , where  $n$  represents the number of the current retransmission attempt. The contention window is doubled (hence *binary* exponentiation) upon each consecutive collision, up to the predefined maximum window size (hence *truncated* BEB).

The 802.11 wireless stations implement the Ethernet backoff scheme with a modification enforced by the wireless nature of the medium. Since Ethernet-like collision detection is not possible there, IEEE 802.11 stations use a nonzero contention window from the very first transmission attempt. The contention window spans the interval  $< 0, 2^{c_{\min}+n} - 1 >$ , where  $c_{\min} > 0$  accounts for the necessary collision avoidance in the first transmission attempt.

Once an IEEE 802.11 station senses the channel idle during DCF interframe space (DIFS) interval after previous access cycle, it defers its own DATA frame transmission for a random number of  $k$  backoff slots to mini-

mize the probability of collisions with other senders. This scheme features a linear relationship between the random value  $k$  (backoff counter) and the number of backoff slots (Fig. 1).

Though very simple and robust, the linear backoff scheme becomes a source of significant protocol overhead in wireless networks. In the IEEE 802.11a, 802.11b and 802.11g standards, the initial contention window ( $CW_{\min} = 2^{c_{\min}} - 1$ ) has 15 slots. For example, it takes  $104 \mu s$  to transmit a 512-byte packet with all MAC and PHY headers in an 802.11a [4] network operating at a 54 Mbit/s data rate. A comparable period of time is “wasted” for a single DIFS interval ( $34 \mu s$ ) along with 7.5 backoff slots (each  $9 \mu s$  long) corresponding to an average backoff time for a channel access cycle with just a single active station.

A straightforward solution aimed at reducing backoff-related overhead would be to minimize the duration of DIFS and backoff slots. Unfortunately, these time constants cannot be decreased at will, since they are determined by the propagation delay and receiver/transmitter switchover time.

Another option is to minimize the number of backoff slots. This can be achieved by:

- adapting the contention window range and/or size to current traffic conditions, or
- changing the way the backoff counter value is encoded and communicated to other stations.

While the former approach has been extensively studied (see, e.g., [5], [6]), there is little work concerning backoff coding schemes. This paper fills this gap by describing a new backoff coding scheme, called modulo N. It is specifically intended for radio environments such as IEEE 802.11 networks, where it can significantly reduce backoff overhead.

## 3. Modulo N Backoff Scheme

An optimal backoff algorithm should have the following properties:

- low best-case backoff length to take advantage of light-load traffic conditions;
- small average backoff length for typical multi-station channel access scenarios;
- robustness in the sense of keeping a moderate frame collision rate under heavy-load traffic conditions.

The modulo N scheme satisfies all the above requirements. It features the best case close to DCF/enhanced DCF (EDCF), but at the same time significantly improves the worst case. It also achieves a reduced average backoff overhead.

A wireless node supporting the modulo N scheme follows the BEB strategy in order to reduce the risk of DATA frame

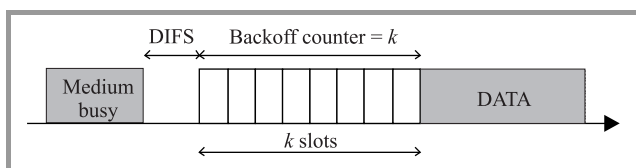


Fig. 1. Linear backoff for collision avoidance in 802.11 networks.

collisions, like it does under DCF. However, it employs a different backoff coding to inform other stations about the backoff counter value it has selected at random. Figure 2 illustrates the principle of modulo N operation. When a station has a DATA frame ready for transmission and senses the medium busy, it selects a random backoff counter value  $k$ . This value is next divided modulo  $N$  into an integer part  $k_{/N}$  and a remainder part  $k_{\%N}$ , so that

$$k = k_{/N} \cdot N + k_{\%N}.$$

After the previous channel access cycle is finished (e.g., after DIFS, like in IEEE 802.11), a station senses the medium for the duration of  $k_{/N}$  slots. If it remains idle, a station broadcasts a one-slot *busy signal*. Next, it waits for  $k_{\%N}$  idle slots before it finally commences a DATA frame transmission. If the station detects any foreign signal during the idle slots, it is inhibited from transmission, i.e., gives up and waits until the next access cycle. This may only happen if another station has won the contention by selecting a shorter backoff. Like in the original DCF, the inhibited station decrements its backoff counter by the number of elapsed slots.

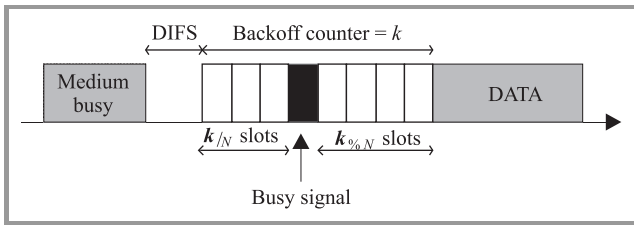


Fig. 2. Modulo N backoff encoding.

In the original IEEE 802.11 standard, the number of backoff slots is always equal to the backoff counter value. In contrast, in the modulo N scheme, the required number of backoff slots for a given backoff counter value  $k$  can be expressed as  $k_{/N} + k_{\%N} + 1$ .

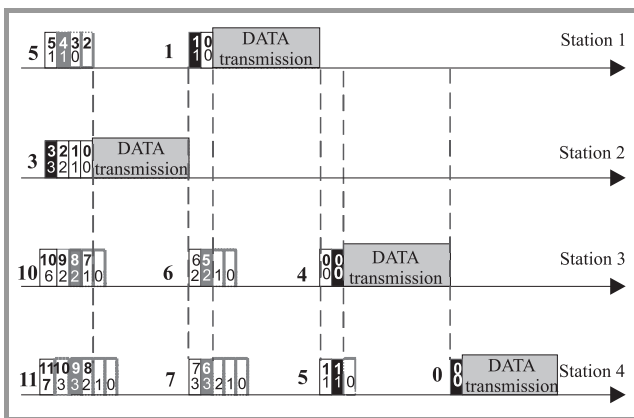


Fig. 3. Example network operation with modulo 4 backoff.

Figure 3 depicts an example network operation with four wireless stations using the modulo 4 backoff scheme (a semiformal specification of the proposed channel access

algorithm can be found later, cf. Fig. 10). There are three types of backoff slots distinguished in Fig. 3:

- white boxes indicate idle slots (a station is listening);
- black boxes represent busy signal slots;
- gray boxes signify slots that would carry busy signals had a station not been inhibited by a foreign transmission.

The numbers shown on the diagram represent the value of a backoff counter:

- the numbers to the left of the slot boxes correspond to backoff counter values at the beginning of a new channel access cycle (i.e., after DIFS);
- the upper numbers inside the boxes represent current of backoff counter values at the end of a slot;
- the lower numbers inside the boxes represent theoretical of backoff counter values had a station not been inhibited by a foreign transmission; note that the upper and lower numbers are equal in a winning station.

Station 2 has the lowest starting backoff counter and it wins the first access cycle. This station starts with a one-slot busy signal, since its initial backoff counter 3 divided modulo 4 gives  $k_{/N} = 0$ . All the other stations start with listening slots, and they are inhibited by the busy signal. In next three slots all stations proceed according to the original DCF algorithm. They decrement their backoff counters by one in every slot before station 2 finally commences transmission. The transmission phase includes also acknowledge (ACK) and interframe spaces. Such a transmission phase is considered a slot in DCF/EDCA, and the remaining stations decrement their backoff counters at the end of the transmission phase.

The second access cycle starts with an immediate busy signal from station 1. Again, the busy signal inhibits stations 3 and 4, which still have their backoff counters higher than 4. Note that no station decrements its backoff counter during the busy signal slot, hence station 1 listens for one more of the  $k_{\%N}$  slots before it zeroes its backoff counter.

Both stations 3 and 4 begin the third access cycle with one of the  $k_{/N}$  listening slots. As they do not detect any transmission during this slot, they both decrease their backoff counters by  $N$  (here,  $N = 4$ ). Next, they both have the backoff counter lower than  $N$  so they announce transition to the  $k_{\%N}$  slots by sending a one-slot busy signal. Station 3 has  $k_{\%N} = 0$  (backoff value 4 modulo 4 gives a remainder of 0), and it starts data transmission immediately after the busy signal slot. Station 4 listens for one of the  $k_{\%N}$  slots and detects signal from station 3. This inhibits station 4 from commencing its own transmission.

Station 4 enters the last access cycle with the backoff counter equal to 0. Even in such a case a station has to send the busy signal. If station 4 did not send the busy signal

first, its transmission could collide with a busy signal from any other station that has a backoff counter less than 4.

By examining the lower numbers inside the boxes, one can compare modulo 4 with the original DCF/EDCA backoff scheme. For example, station 1 would need only 3 slots to announce the backoff counter value of 5, and station 4 would need 6 slots to announce value of 11. On the other hand station 2 needs 4 slots with backoff counter 3, and station 4 transmits a one-slot busy signal even if its backoff counter is equal to zero. The next section provides more detailed analysis of the modulo  $N$  scheme and its comparison with DCF/EDCA schemes.

## 4. Optimal Modulo $N$ Parameters

The modulo  $N$  scheme can be subject to numerous parameterizations. The value  $N$  itself is an apparent parameter to manipulate. As  $N$  increases, the maximum backoff length decreases, but the impact upon the average backoff is not obvious given that a large  $N$  leaves little room for collision avoidance based on the integer parts of the backoff counter values, especially when multiple stations compete in successive access cycles.

Figure 4 compares modulo 4 and DCF backoff schemes assuming default IEEE 802.11a settings: slot time =  $9 \mu s$ , DIFS =  $34 \mu s$ ,  $CW_{min} = 15$ ,  $CW_{max} = 1023$ , and 54 Mbit/s data rate. The simulation results are provided for an ideal radio channel with all the stations within each other's range (no hidden terminals), and DATA frame errors occur only due to collisions, transmission errors being negligible.

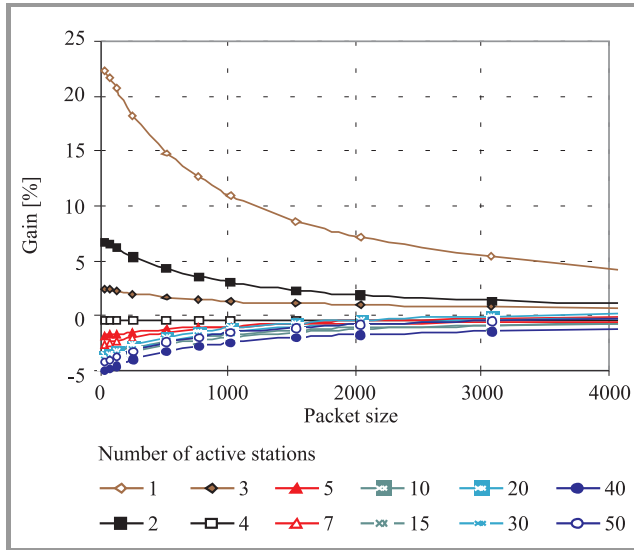


Fig. 4. Modulo 4 throughput gain over legacy EDCA.

The curves represent simulation runs for various numbers of active stations under saturation conditions (i.e., each station always has a DATA frame ready for transmission). Let  $S_x$  denote the saturation throughput for a backoff

scheme  $x$ . The percentage gain shown on the y-axis is defined as

$$\text{gain} = \frac{S_{\text{mod } N} - S_{\text{DCF}}}{S_{\text{DCF}}}.$$

From Fig. 4 it follows that the modulo 4 scheme outperforms DCF if fewer than four stations contend for channel access at one time. The highest gain is achieved in the case where only one station is active in each access cycle. This can be easily explained: with one contending station and  $CW_{min} = 15$ , an average DCF backoff is 7.5 slots whereas an average modulo 4 backoff is only 4 slots.

Under extremely heavy load, with more than 20 active stations in each access cycle, an average DCF backoff becomes less than one slot, which is the lower bound for modulo  $N$ . Therefore, DCF performs better than pure modulo  $N$  under extreme traffic conditions.

An appealing feature of modulo  $N$  is that the maximum backoff window is bounded by  $(CW_{max}/N) + N$ , which is a significant improvement over  $CW_{max}$  in EDCA. Assuming  $N = 4$ , the maximum backoff time is 260 slots in modulo 4 as compared with 1023 slots in EDCA. It makes sense, therefore, to manipulate other protocol parameters. In IEEE 802.11 networks, the contention window ranges between  $CW_{min}$  and  $CW_{max}$ , which are interrelated as follows:

$$CW_{max} = (c_{inc})^{c_{max}} \cdot (CW_{min} + 1) - 1,$$

where the IEEE 802.11a defaults are:  $CW_{min} = 15$ ,  $c_{inc} = 2$ ,  $c_{max} = 6$ , and  $CW_{max} = 1023$ . Thus, in order to get the maximum backoff duration of 1023 slots, one could configure a modulo 4 scheme with  $c_{inc} = c_{max} = 4$ . Hereafter this combination will be denoted *modulo  $N|c_{inc}^{c_{max}}$* .

As illustrated in Figs. 5, 6, and 7, higher  $N/c_{inc}$  values generally lead to better channel utilization. Unfortunately, a serious drawback of modulo 5|5<sup>5</sup> is the maximum contention window  $CW_{max}$ , reaching  $16 \cdot 5^6$ , or 250 000. Considering that the backoff length is 5 times shorter, it is still 50 000 backoff slots in the worst case. Even though

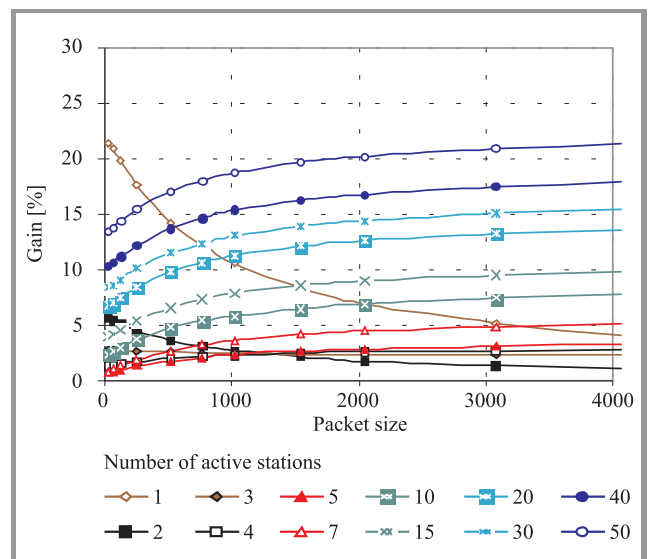


Fig. 5. Modulo 3 throughput gain over legacy EDCA.

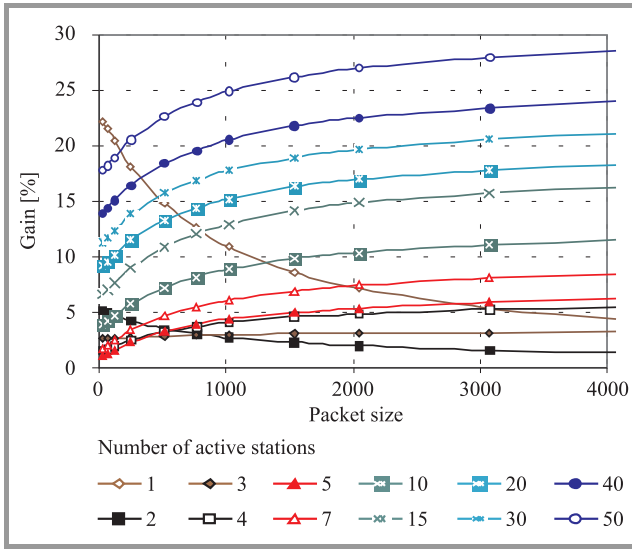


Fig. 6. Modulo 4|4<sup>6</sup> throughput gain over EDCA.

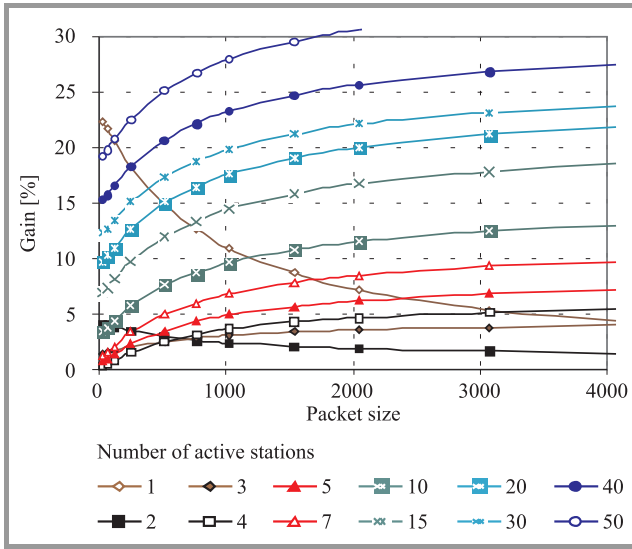


Fig. 7. Modulo 5|5<sup>5</sup> throughput gain over EDCA.

the probability of reaching this limit is very small, such long-tailed distributions should be avoided in real networks. Yet similar performance results can be achieved under modulo 5|5<sup>4</sup>, where the maximum backoff length is limited to 2000 slots, comparable with that under DCF. This limit can be reduced even more under modulo 4|4<sup>4</sup>, which ensures maximum backoff length of 1024 slots, very much like under DCF.

Modulo 4|4<sup>4</sup> seems a reasonable configuration choice, bearing in mind that predictable network operation is more valuable than fine-tuning the protocol parameters under very heavy load (with over 10 stations competing in each access cycle).

From Fig. 8 it can be seen that modulo 4|4<sup>4</sup> configuration outperforms EDCA for all traffic scenarios. The explanation is that modulo 4 allows broadening the contention window beyond the  $CW_{max}$  limit defined for 802.11a, while retaining the maximum backoff duration of 1023 slots.

Clearly, the increased  $CW$  range results in fewer collisions; as a consequence, better channel utilization can be achieved.

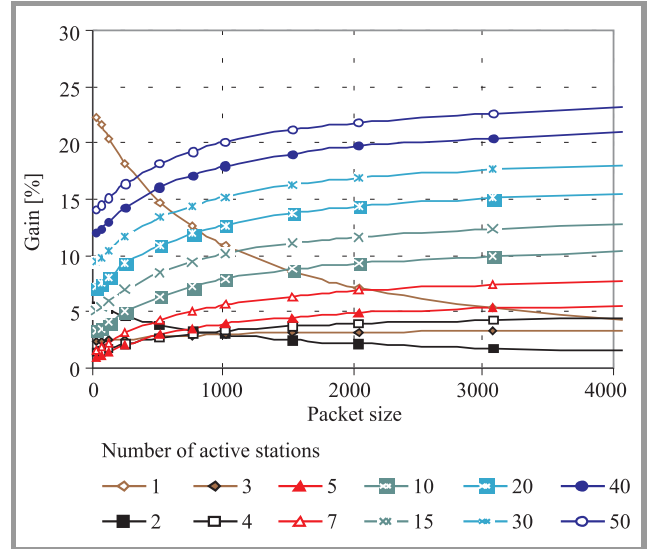


Fig. 8. Modulo 4|4<sup>4</sup> throughput gain over EDCA.

The lowest gain is achieved in a scenario when two stations are active in each access cycle. This nicely tones in with existing research reports, which indicate that default DCF/EDCA parameters are optimal for two-station scenarios [5]. Notably, even in such a case modulo N performs better than the original DCF/EDCA.

## 5. Prioritized Modulo N Backoff

The pure modulo N scheme does not allow for prioritization of traffic streams. Nevertheless, the scheme can be easily augmented with QoS differentiation as illustrated in Fig. 9.

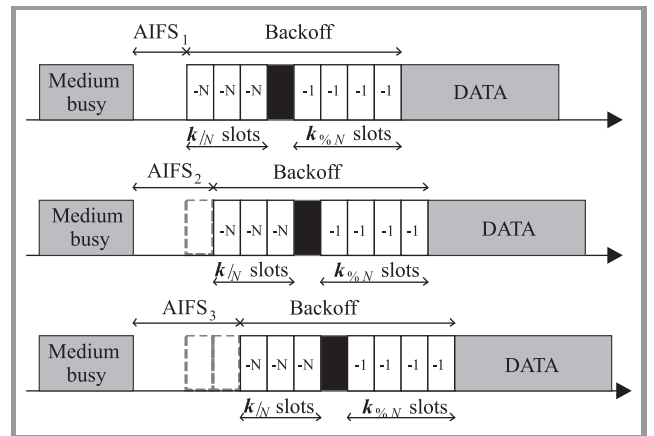


Fig. 9. Prioritized modulo N backoff encoding concept.

Like in IEEE 802.11e EDCA, the basic idea is to replace the DIFS interval with the arbitration interframe space (AIFS) intervals defined on a per-class basis, as well as to use per-class contention window ranges.

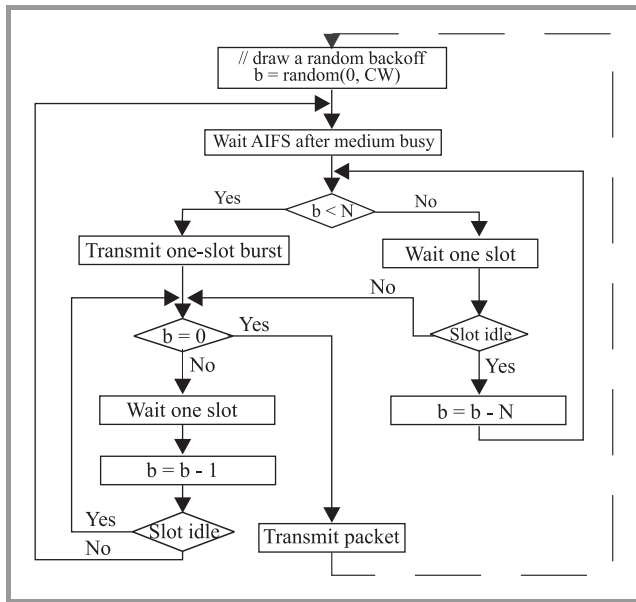


Fig. 10. Prioritized modulo N channel access algorithm.

Figure 10 gives a flowchart description of the prioritized modulo N channel access scheme. Like in EDCA, a station decrements its backoff counter both at the end of an idle slot as well as during a foreign transmission period (i.e., a DATA frame exchange sequence started by another station). Unlike in EDCA, however, decrementing should take place at the beginning of a foreign transmission period (i.e., one slot after transmission starts), and not when AIFS expires as described in the IEEE 802.11e draft. Furthermore, DATA frame transmission should be commenced immediately after the backoff counter reaches 0 (in IEEE 802.11e, a station starts transmission one slot later). In that sense, the prioritized modulo N resembles more IEEE 802.11e draft version 4.2 than the more recent version 8.0. Nonetheless, both approaches are functionally equivalent and one described below facilitates a simple implementation of the modulo N scheme.

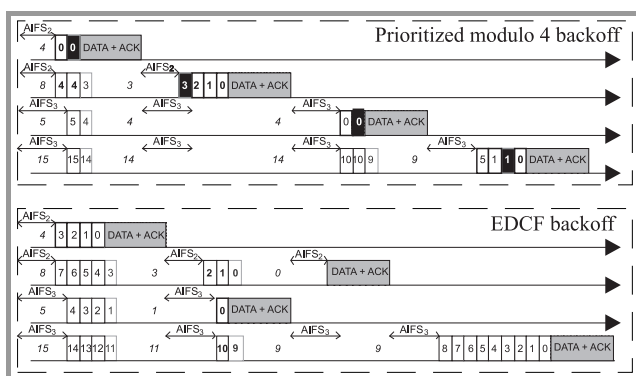


Fig. 11. Example network operation with prioritized modulo N backoff.

Figure 11 shows an example network scenario with four wireless stations using the prioritized modulo 4 and an original EDCA backoff. White and black boxes indicate idle and busy slots, respectively. The numbers to the left

of the slot boxes show backoff counters at the beginning of an access cycle. Those inside the boxes correspond to the current (end of slot for modulo N and start of slot for EDCA) backoff values.

The prioritized modulo N provides more stringent AIFS-based QoS differentiation as compared with original EDCA. The reason is that in modulo N a one-slot difference in AIFS intervals corresponds to an N-slot difference in backoff counters. Consider high-priority stations ( $AIFSN = 2$ ) and low-priority stations ( $AIFSN = 3$ ) depicted in Fig. 11. In the very first DATA frame exchange, a high-priority station 2 decrements its backoff counter by 5, while low-priority stations 3 and 4 decrement their backoff counters just by one. In the original EDCA, the low-priority stations would decrement their backoff counters by 4, which indeed is not much less than 5 in a high-priority station.

In general, a high-priority EDCA station is unaffected by low-priority stations only if its backoff counter is already equal to 0 when AIFS expires. In contrast, transmission from a high-priority modulo N station does not depend on the presence of low-priority stations for all backoff counters less than N. Consider for instance the second DATA frame exchange in Fig. 11. The high-priority station 2 has a backoff counter equal to 3, which is enough to prevent low-priority stations from transmission in this access cycle (even though they had lower backoff counters initially). Similar behavior is not possible in the original EDCA. Consider again the second DATA frame exchange under EDCA. We see that a high-priority station (backoff = 3) loses in competition with a low-priority station (backoff = 1).

## 6. Conclusions

The paper describes the new modulo N backoff scheme. Both a semiformal description of the new channel access scheme and simulation results that compare the new scheme with existing ones like IEEE 803.11 DCF and EDCA, have been presented.

The description of modulo N reveals that its complexity is comparable with legacy schemes. At the same time, the obtained performance results show that the new scheme increases the overall channel utilization between 5% and 30% as compared with IEEE 802.11 DCF.

Furthermore, the paper describes the prioritized variant of the modulo N scheme. This variant enables very effective QoS differentiation, which is also more flexible than the original EDCA scheme of IEEE 802.11e.

## References

- [1] A. Doufexi, S. Armour, M. Butler, A. Nix, D. Bull, J. McGeehan, and P. Karlsson, "A comparison of the HIPERLAN/2 and IEEE 802.11a wireless LAN standards", *IEEE Commun. Mag.*, vol. 40, no. 5, pp. 172–180, 2002.
- [2] "IEEE Standard for Information Technology – Telecommunications and Information Exchange between System – Local and Metropolitan Area Networks – Specific Requirements" – Part 11: "Wireless LAN Medium Access Control (MAC) and Physical Layer (PHY) Specification", IEEE 802.11-1999 [Online]. Available: [www.ieee.org](http://www.ieee.org)

- [3] "Draft Supplement 8.0 to Standard for Telecommunications and Information Exchange between Systems – LAN/MAN Specific Requirements" – Part 11: "Wireless Medium Access Control (MAC) and Physical Layer (PHY) Specifications: Medium Access Control (MAC) Enhancements for Quality of Service (QoS)", IEEE Draft 802.11e-2005.
- [4] "Supplement to IEEE Standard for Information Technology Telecommunications and Information Exchange between System – LAN/MAN Specific Requirements" – Part 11: "Wireless LAN Medium Access Control (MAC) and Physical Layer (PHY) Specifications: High-Speed Physical Layer in the 5 GHz Band", IEEE 802.11a-1999 [Online]. Available: [www.ieee.org](http://www.ieee.org)
- [5] F. Cali, M. Conti, and E. Gregori, "Dynamic tuning of the IEEE 802.11 protocol to achieve a theoretical throughput limit", *IEEE/ACM Trans. Netw.*, vol. 8, no. 6, pp. 785–799, 2002.
- [6] M. Natkaniec and A. Pach, "An analysis of modified backoff mechanism in IEEE 802.11 networks", in *Proc. Polish-German Telegraf. Symp. PGTS'2000*, Dresden, Germany, 2000



**Tomasz Janczak** received the M.Sc. degree in computer science in Gdańsk University of Technology (GUT), Poland, in 1999. In 2001, he also received the M.Sc. degree in management and economics from GUT. Since 1999, he has been working in GUT towards the Ph.D. degree which he received in 2004. His research work fo-

cused on wireless local area networks, including topics covered by IEEE 802.11, HIPERLAN and Bluetooth solutions. He has published several technical papers on supporting quality of service and service fairness in wireless networks with dynamically changing topologies. He is a senior software architect at Intel Technology Poland. At present, he works as a system architect at Intel Corporation, at the R&D networking site located in Gdańsk.

e-mail: [tomasz.janczak@intel.com](mailto:tomasz.janczak@intel.com)

Intel Technology Poland

Słowackiego st 173

80-298 Gdańsk, Poland



**Józef Woźniak** received his Ph.D. and D.Sc. degrees in telecommunications from the Gdańsk University of Technology, Poland, in 1976 and 1991, respectively. At present he is a Full Professor in the Faculty of Electronics, Telecommunications and Computer Science at the Gdańsk University of Technology. He is the author or

co-author of more than 200 journal and conference papers. He has also co-authored 4 books on data communications,

computer networks and communication protocols. In the past he participated in research and teaching activities at Politecnico di Milano, Vrije Universiteit Brussel and Aalborg University, Denmark. In 2006 he was Visiting Erskine Fellow at the Canterbury University in Christchurch, New Zealand. He has served in technical committees of numerous national and international conferences, chairing or co-chairing several of them. He is a member of IEEE and IFIP, being the Vice Chair of the WG 6.8 (Wireless Communications Group) IFIP TC6 and the Chair of an IEEE Computer Society Chapter (at the Gdańsk University of Technology). His current research interests include modeling and performance evaluation of communication systems with the special interest in wireless and mobile networks.

e-mail: [jowoz@eti.pg.gda.pl](mailto:jowoz@eti.pg.gda.pl)

Gdańsk University of Technology

G. Narutowicza st 11/12

80-952 Gdańsk, Poland



**Krzysztof Pawlikowski** received a Ph.D. degree in computer engineering from the Gdańsk University of Technology, Poland, and worked at that university until Feb. 1983. At present he is a Professor of computer science at the University of Canterbury, in Christchurch, New Zealand. He is the author of over 150 journal

and conference papers and four books, has given invited lectures at over 80 universities and research institutes in Asia, Australia, Europe and North America. He is the principal developer of Akaroa2, a universal controller of quantitative stochastic simulation, which is used at universities in over 70 countries around the world. He was the Alexander-von-Humboldt Research Fellow (Germany) in 1983-84 and 1999, and a Visiting Professor at universities in Australia, Italy, Germany and the USA. His research interests include multimedia telecommunication networks and their performance modeling, as well as methodology of discrete-event computer simulation and distributed processing. He is a senior member of the IEEE, member of the ACM and the Royal Society of New Zealand.

e-mail: [K.Pawlikowski@ieee.org](mailto:K.Pawlikowski@ieee.org)

Network Research Laboratory: Protocols,

Distributed Processing and Simulation

Department of Computer Science

and Software Engineering

University of Canterbury

Private Bag 4800

Christchurch, New Zealand

**Jerzy Konorski** – for biography, see this issue, p. 40.

# Empirical Season's Fadings in Radio Communication at 6 GHz Band

Jan Bogucki and Ewa Wielowieyska

**Abstract**—This paper covers unavailability of line-of-sight radio links due to multipath propagation. Multipath fading in the atmosphere is not permanent phenomenon. The five year investigation results of the received radio signal fading in the radio links and their season empirical distributions are presented.

**Keywords**—*line-of-sight radio links, multipath, propagation.*

## 1. Introduction

The use of digital microwave radio-link systems is widely recognized as flexible, reliable and economical means of providing point-to-point communication [1], [2], [3]. These radio systems, when used with appropriate multiplex equipment, can carry from a few up to a large number of voice, video and data transmissions. They can also be arranged to carry additional wide-band for high-speed data, Internet, multimedia wireless or high-quality audio and high definition TV channels.

Comparative cost studies usually prove the radio microwave systems to be the most economical means for providing communication transmission where there are no existing cable lines to be expanded. For temporary facilities and other applications where installation time is severely limited the advantages of the radio technique are obvious.

Many fixed broadband wireless radio links are designed to be available essentially all the time. "Available" means that bit error rate (BER) or frame error rate (FER) is at or below given threshold level. Conversely, "outage" is the time when the link is not available; for example, BER/FER value is above the quality threshold level. In the fix-link, service availability of 99.99% for the worst month is usually a target that means an outage of only 53 minutes a year.

Nearly all radio systems are the subject to regulation by the government of the country, where the system is to be located. In general, each country allocates specific sub-bands of frequencies for specific services or users. Within Poland the Office of Electronic Communications (UKE) is the controlling authority for all the radio-communication systems except those operating in the frequency bands where simplified or no frequency coordination procedures are applied [4].

In Poland the 6 GHz band is meant for high-capacity long distance radio links. The radio signal of this frequency range is susceptible to some kinds of fading due to the changes in atmosphere.

One kind of fading essential in this band is multipath propagation fading [5].

This paper presents the problems of unavailability of line-of-sight radio links due to multipath propagation phenomena. In the National Institute of Telecommunications (NIT) the radio links have been used to investigate the propagation fading in the 6 GHz sub-band. The five year investigation results of the received radio signal fading in the radio links mentioned above and their season empirical distributions are presented below [6]–[9].

## 2. The Multipath Fading

The beam of microwave energy is not a single ray, but wavefront extending in considerable space along the center line. Since the refraction index for normal atmospheric conditions is lower at the top of wavefront and higher at the bottom, and since the wave velocity is inversely proportional to refraction index, the upper portion of wavefront will travel slightly faster, with result that the top of the wavefront is tilted. Since the direction of beam travel is always perpendicular to wavefront, the beam itself will be tilted downward. The degree of the tilt is actually very slight on percentage basis, but is sufficient to cause significant variation of the fading phenomena. It is normal propagation situation.

But sometimes at certain atmospheric situations there can be even "greater" than normal negative  $N$  gradient, or other in which  $N$  gradient becomes less negative and even positive. In the latter situation lower part of wavefront will travel faster, and the beam will be bent upward, reducing apparent clearance.

Most of time the vertical profile of these gradients in the lower atmosphere are essentially linear. These linear variations affect clearance and are also important, when the path is reflective, but they do not produce atmospheric multipath situations.

However, when gradients are nonlinear, it is possible for multiple paths, in addition to direct path, to exist within the atmosphere itself, independently of any reflecting surface on ground. These situations in atmosphere occur when stratified layers with different gradients lie on top of one another. Such conditions strongly depend on seasons of the year.

The incidence of multipath fading varies not only as function of path length and frequency, but also as function of climate and terrain conditions.

The treatment of multipath fading is based largely on experience.

### 3. The Index of Refraction for the Troposphere

The index of refraction for the troposphere air is very close to that of vacuum. Due to that, radio refractivity is used instead of index of refraction:

$$N = (n - 1) \cdot 10^{-6}, \quad (1)$$

where:  $n$  – index of refraction,  $N$  – radio refractivity.

The  $N$  term would be zero in free space and value on order of 300 at the earth surface. An empirical formula for  $N$  is:

$$N = \frac{77.6}{T} \left( p + 4810 \frac{e_H}{T} \right), \quad (2)$$

where:  $T$  – temperature [K],  $p$  – total air pressure [hPa],  $e_H$  – water vapour pressure [hPa].

Water vapour pressure corresponds to relative humidity of air:

$$e_H = H \frac{6.1121 \exp\left(\frac{17.502t}{t + 240.97}\right)}{100}, \quad (3)$$

where:  $H$  – relative humidity of air [%],  $t$  – temperature [°C].

Since  $p$ ,  $e$ , and  $T$  all are functions of height, consequently  $N$  is also function of height. For normal atmosphere, standard – well mixed, the variation of  $N(h)$  with height is:

$$\frac{dN(h)}{dh} = -40 \left[ \frac{1}{\text{km}} \right], \quad (4)$$

$$N(h) = 315 \cdot e^{-0.136h}, \quad (5)$$

where:  $h$  – height above earth surface [km].

Multipath propagation occurs when there is more than one ray reaching the receiver. It is the main cause of fading in 6 GHz band. Multipath can only happen when  $\frac{dN}{dh}$  varies with height.

### 4. The Measurement System

There were six radio links at 6 GHz band with the length from 36.6 km to 69.8 km [6]. Sites of four radio links were located near Warsaw and two of the longest paths were situated farther north of Warsaw – see Fig. 1. Self-operating measuring position was set-up to cooperate with radio link receivers. The measurements were carried out during ordinary operation radio links. Received signals were sampled each 0.2 s during high attenuation and 5 min in the other time. The system measured only result of multipath, not the reason.

Layers of the atmosphere with different gradients of refractivity may cause detrimental effects to received signal. The radio wave rays, that normally would have been lost in the troposphere may be refracted towards receiving antenna, where they are added to wanted signal. The phase and amplitude relationships between multipath signals determine resulting input signal at receiver. The example of input level as a function of time during the fading event is shown in Fig. 2.



Fig. 1. The locations of experimental links.

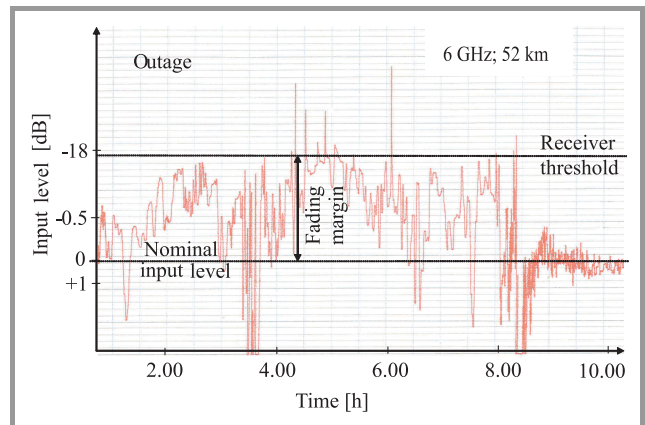


Fig. 2. An example of multipath fading of 6 GHz terrestrial path.

### 5. The Measurement Results

The occurrence of multipath mainly depends on weather conditions such as temperature, wind, humidity, air pressure, and these weather phenomena can be described only

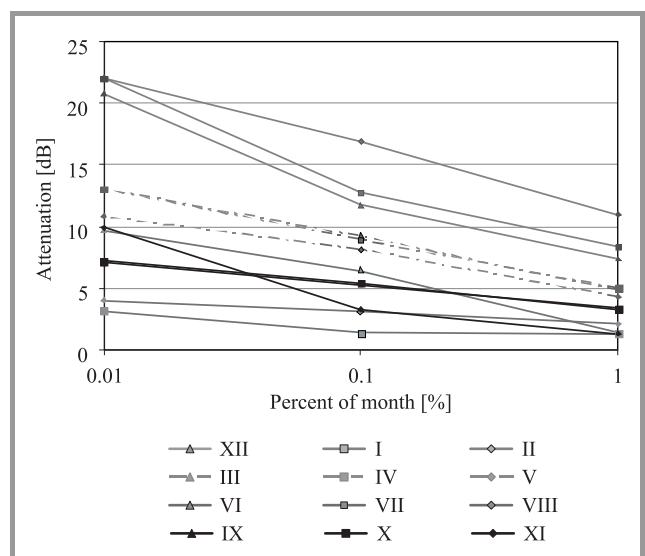
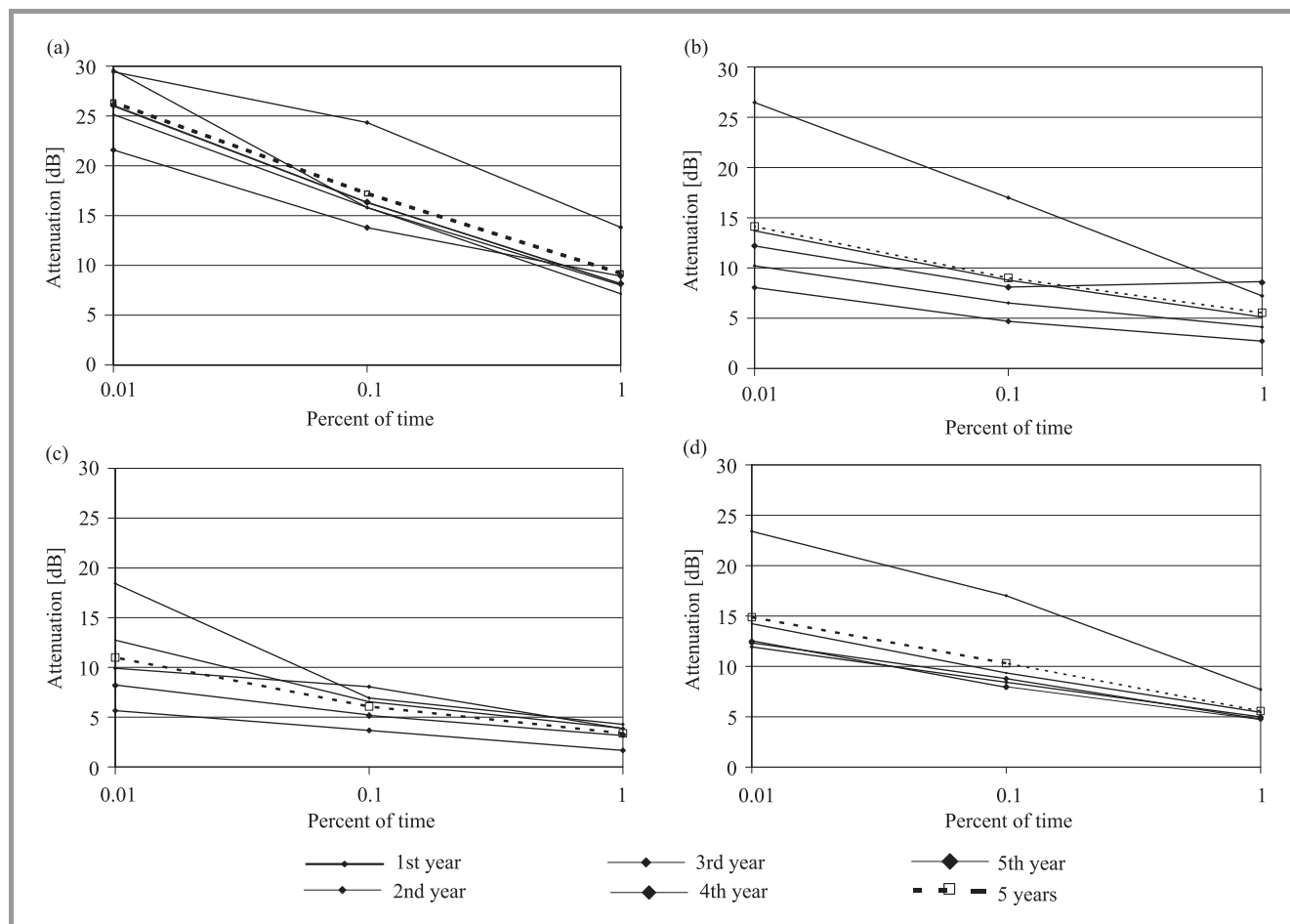


Fig. 3. The measured monthly 4th year distributions of attenuation at 6 GHz.



**Fig. 4.** The measured seasons distribution for: (a) summer (June, July, and August); (b) autumn (September, October, and November); (c) winter (December, January, and February); (d) spring (March, April, and May).

statistically. Therefore the changes of path loss, which are very important in the design of the radio-link system, can also be described only statistically.

The results of the measurement investigations allowed to determine among other things season's behavior of radio channel at 6 GHz band. In point-to-point link immersed in time varying propagation medium the received signal power varies with time even when transmitter power remains constant. Measuring the probability fades of particular magnitude occur with, will lead directly to the probability of outage and hence the link availability probability.

Received signal samples were used for computation of monthly and annual fading distributions as well as distributions for the worst months and then season's fading distributions were obtained. The important problem of "zero" level during multipath effects was solved assuming that monthly 50% level is "zero" level for multipath fading [10]. Monthly distributions of attenuation for each path are the basis for cross-sectional statistical analysis.

Figure 3 shows the distribution of attenuation for 12 months obtained from a one year measurements on the 41.3 km path.

It illustrates how attenuation varies with months. For example, attenuation reaches 22 dB in July and August, and

only a few decibels in January and February, for 0.01% of time.

Changing of season's attenuation distributions will be shown on the basis of measurements results from radio link at 5885.04 GHz and 41.3 km path length.

Figure 4 shows season's distribution of attenuation due to multipath obtained from five years of measurements on the 41.3 km path.

The attenuation changes a lot during a year. Summer season average at 0.01% of time is 26 dB and only 14 dB for season's autumn average.

Our studies indicate that there were many differences between season's average attenuation obtained from 5 years and season's average attenuation in individual year. For example, maximum attenuation difference is 12 dB on autumn at 0.01%, 7 dB on summer at 0.1%, 7.4 dB on winter at 0.01% and 8.5 dB on spring at 0.01%. There are statistics, empirical results.

Figure 5 compares average season, annual and the worst month distributions obtained from 5 year measurements. It shows that average autumn season attenuation is similar to average spring season. Maximum attenuation difference is 0.8 dB at 0.01% and 0.65 dB at 0.1%.

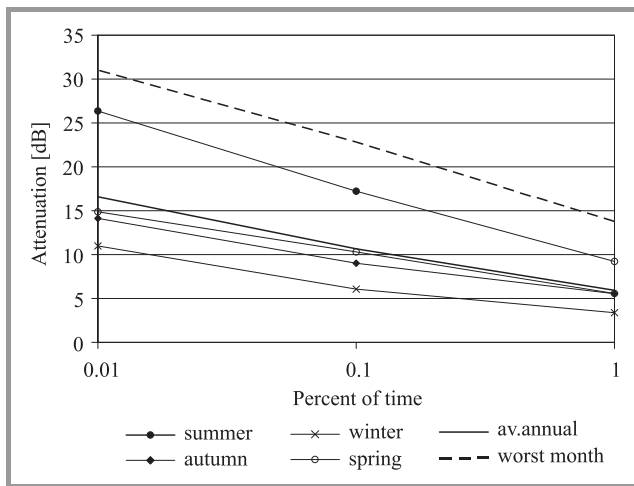


Fig. 5. The average season, annual and worst month distributions for 5 years.

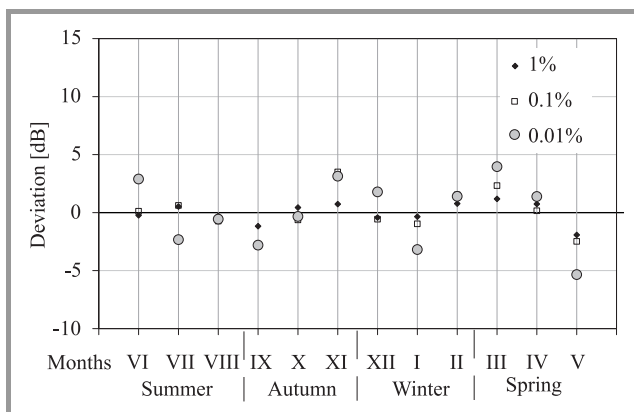


Fig. 6. The deviations of attenuation for average months in the year from appropriate average season for 1, 0.1, and 0.01%.

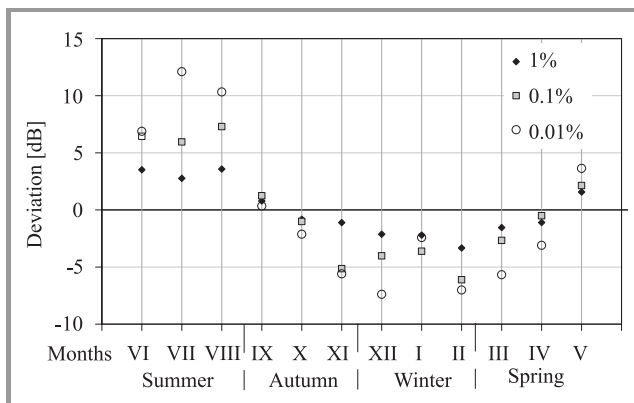


Fig. 7. The deviations of attenuation for average months in the year from average annual attenuation for 1, 0.1, 0.01%.

The order of calculation to obtain the input data for Fig. 6:

- 12 distributions of average month attenuation for the period of 5 years;
- 4 distributions of average season attenuation for the period of 5 years;

- the difference between average month attenuation and appropriate average season attenuation for 1, 0.1, 0.01%.

The result of these calculations seems not to be obvious. It can be explained that in summer average value is high, but there are less unpredicted events. The smallest deviations for 0.01% are for summer (maximum 2.9 dB), and are bigger in autumn (3.1 dB), in winter (–3.2 dB), and are the biggest in spring (–5.5 dB).

The data for Fig. 7 were obtained in an analogous way to these for Fig. 6; the average month distributions were compared with average annual distribution.

The comparison between Figs. 6 and 7 indicate that average season distributions were more accurate than average annual distribution, particularly in summer. The difference between average month and average annual attenuation is 12.1 dB in July while deviation of average month from average season's attenuation never exceeded  $\pm 5.5$  dB.

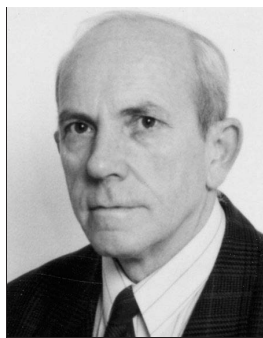
## 6. Conclusion

The knowledge of fading statistics is extremely important for the design of wireless systems. Microwave radio links can be properly and precisely engineered to overcome potentially detrimental propagation effects. One of characteristics that must be taken into consideration is multipath attenuation.

The gathered empirical data at 6 GHz of seasonal statistical distributions broaden our knowledge about the signal changes with weather variations. This knowledge can improve our interpretation of the phenomena that appear in installed modern radio systems.

## References

- [1] R. H. Anderson, *Fixed Broadband Wireless System Design*. Chichester: Wiley, 2003.
- [2] R. K. Crane, *Propagation Handbook for Wireless Communications System Design*. London: CRC Press, 2003.
- [3] C. Salema, *Microwave Radio Links: From Theory to Design*. New Jersey: Wiley, 2003.
- [4] "Office of electronic communications" [Online]. Available: <http://www.en.uke.gov.pl/>
- [5] M. Grabner and V. Kvicera, "Refractive index measurement at TV tower Prague", *Radioengineering*, no. 1, pp. 5–7, 2003.
- [6] J. Bogucki, J. Jarkowski, and E. Wielowieyska, "Propagacyjna zmienność sezonowa systemów radiowych zakresu 6 GHz", in *Proc. KKRRiT Conf.*, Wrocław, Poland, 2008, pp. 245–248 (in Polish).
- [7] J. Bogucki and E. Wielowieyska, "Multipath in line-of-sight links – prediction vs. reality", in *Proc. 16th Int. Czech-Slovak Sci. Conf. Radioelektr. 2006*, Bratislava, Slovakia, 2006.
- [8] J. Bogucki and E. Wielowieyska, "Reliability of line-of-sight radio-relay systems", *J. Telecommun. Inform. Technol.*, no. 1, pp. 87–92, 2006.
- [9] J. Bogucki and E. Wielowieyska, "Wielodrogowość w horyzontowych liniach radiowych – prognoza i rzeczywistość", in *Proc. KKRRiT Conf.*, Kraków, Poland, 2005, pp. 396–399 (in Polish).
- [10] A. Kawecki, "Charakterystyki zaników sygnału, wywołanych propagacją wielodrogową w doświadczalnych liniach mikrofalowych 11,5 i 18,6 GHz", *Prace Instytutu Łączności*, no. 101, pp. 59–83, 1993 (in Polish).



**Jan Bogucki** was born in Warsaw, Poland. He graduated Eng. degree at the Technical University of Warsaw in 1972. Since 1973 he has been employed at the National Institute of Telecommunications, Warsaw, where he has been engaged in digital radio links, digital television, microwave propagation in the troposphere, and

electromagnetic compatibility.

e-mail: J.Bogucki@itl.waw.pl

National Institute of Telecommunications

Szachowa st 1

04-894 Warsaw, Poland



**Ewa Wielowieyska** was born in Warsaw, Poland. She finished the Mathematics Faculty of the Warsaw University. Since 1981 she has been employed at the National Institute of Telecommunications, Warsaw, where she has been engaged in microwave propagation in the troposphere, propagation digital radio signals on short, medium

and long waves.

e-mail: E.Wielowieyska@itl.waw.pl

National Institute of Telecommunications

Szachowa st 1

04-894 Warsaw, Poland

# Laser Beam Attenuation Determined by the Method of Available Optical Power in Turbulent Atmosphere

Lucie Dordová and Otakar Wilfert

**Abstract**—This work is focused on the atmospheric turbulence effect and a new method for determining optical signal attenuation caused by turbulence is presented here. A new method of power budget of optical links comes from optical intensity distribution in a laser beam after the beam passed through turbulent atmosphere. Results given by this method are compared with Rytov approximation which is nowadays the most frequently used method for determining turbulent attenuation. Results for communication wavelength of 850 nm and 1550 nm are presented as well as the results for a wavelength of 633 nm.

**Keywords**—*method of available power, optical wireless link, Rytov approximation, turbulent atmosphere.*

## 1. Introduction

Nowadays the information is a valuable element in industry, science, education, medical sphere, banking system or in common household. Requirements for transmission rate increase every day. It is necessary to ensure not only high bit rate, but high quality signal as well. In the past metallic cables were adequate but with increasing demands on transmitted volume of information new communication systems were developed and existing communication systems were modified. Currently optical links occur as the best solution for high speed communication links with high reliability [1].

Several years ago optical fiber links became very popular and wide spread in the telecommunication industry. This technology is placed in the backbones networks as well as in local area networks (LANs). One of few disadvantages

of optical fibers is cable laying, because it is necessary to build cable infrastructure. When quick installation demanded optical wireless links (OWL) are suitable. Scheme of optical wireless link is depicted in Fig. 1.

As a transmitter laser diode or light emitting diode is used. The second one is placed in optical links with limited short range (up to about 100 m) or in indoor communication systems. Laser diodes are applied in optical links with longer range (more than 100 m, up to 10 km). Laser beam divergence is set by transmitting lens. Typical values of divergence in free space optics (FSO) are 3–8 mrad. Cover windows are installed to avoid debasement of transmitter and receiver optical components. In receiving optical head receiving lens (most frequently Fresnel lens) is placed to focus laser beam on the active area of photodetector. Photodetector PIN is commonly used, but in special application avalanche photodiode is placed. Interferential filter serves to transmit useful wavelength only.

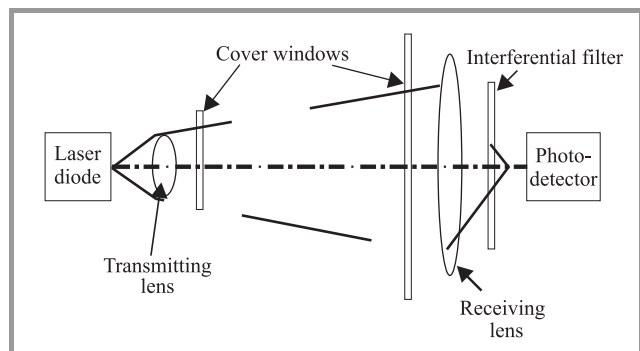
We consider link budget as stationary parameter in optical wireless link as well as transmission rate, bit error rate or link range. Atmospheric transmission media characteristics have statistical nature. The most significant statistical effects in free space optics are atmospheric attenuation caused by molecules and aerosols and of course atmospheric turbulences. Importance of the turbulence phenomena will be specified in this paper later.

## 2. Horizontal and Vertical Optical Wireless Links

Nowadays we differentiate between horizontal and vertical optical wireless links. Model of the horizontal links is shown in Fig. 2.

These links occur in most cases in the urban centers, so typical phenomena influence their operation. Smog, dust, fog, snow or rain are typical events affecting horizontal OWL in towns as well as laser the beam interruption caused by flying birds. Turbulent atmosphere hampers optical beam in the whole path length.

Vertical OWLs are still under development. Communication proceeds between the ground station and the high altitude platform (HAP), which can be placed a few tens of kilometers over the ground level (Fig. 3) [2]. In this case optical beam path is about tens of kilometers long, so it is



**Fig. 1.** Optical wireless link with laser diode, photodetector, cover window, interferential filter, transmitting and receiving lenses.

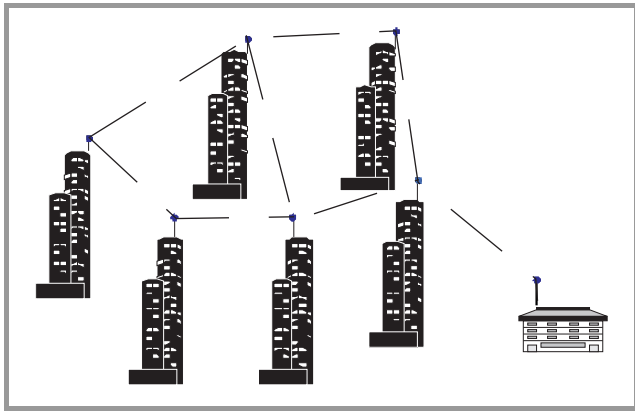


Fig. 2. Model of the horizontal optical wireless links.

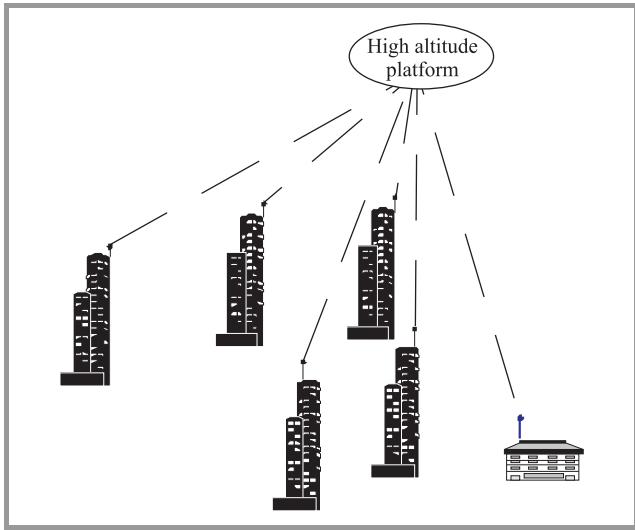


Fig. 3. Model of the vertical optical wireless links.

necessary to design precision tracking system for the signal detection, laser beam divergence should be set in the order of microradians.

### 3. Atmospheric Turbulence Theory

Atmospheric turbulence, generated by a temperature differential between the Earth's surface and the atmosphere, causes effects on optical waves which have been of great interest to scientists for many years. During daytime, the Earth is hotter than the air, causing the air nearest the ground to be hotter than that above. This negative temperature gradient causes light rays parallel to the earth to bend upwards. If the negative temperature gradient is sufficiently strong, it can result in an inverted image known as a mirage. Temperature gradients are positive during nighttime hours, resulting in downward bending of light rays. In addition, atmospheric turbulence disrupts the coherence of laser radiation and optical wave. Wave front distortions in the optical wave induced by atmospheric turbulence result in a broadening of the beam, random variations of

the position of the beam centroid called beam wander, and redistribution of the beam energy within a cross section of the beam leading to irradiance fluctuations [3].

Atmospheric temperature variations and wind speed fluctuations create local unstable air masses, causing them eventually to break up into turbulent eddies or cells of many different scale sizes. These inhomogeneities range in size from a microscale to a macroscale, and hence, in effect form a continuum of decreasing "eddy" size [3].

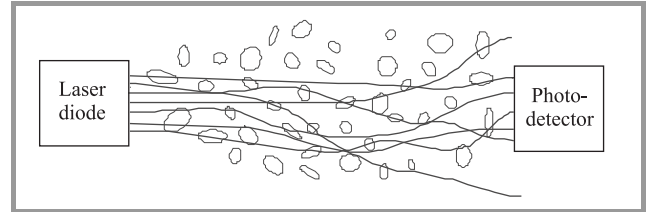


Fig. 4. Turbulent cells through the whole optical path.

Turbulent cells can occur in the whole path length (Fig. 4) or just in the few sections (Fig. 5). Variance in the refraction index is characteristic for atmospheric transmission media with atmospheric turbulences. Refraction in-

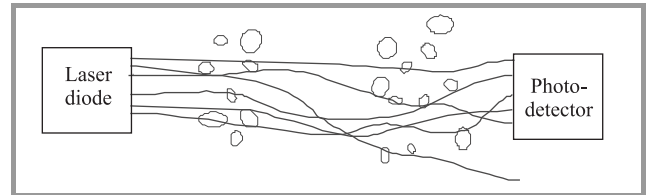


Fig. 5. Turbulent cells in some section of optical path.

index value depends on the local temperature, atmospheric pressure or particle density at specific location. Statistical character of refraction index is measured as refraction index structure function  $D_n$  [4]:

$$D_n = \langle [n(A, t) - n(B, t)]^2 \rangle, \quad (1)$$

where  $n(A, t)$  and  $n(B, t)$  are refractive indexes in the points  $A$ ,  $B$  and in time  $t$ . Refraction index structure parameter is in the relation with the distance  $r$  between points  $A$  and  $B$  according to Kolmogorov model [4]:

$$D_n = \begin{cases} C_n^2 \cdot r^{2/3} & l_0 \ll r \ll L_0 \\ C_n^2 \cdot l_0^{-4/3} r^2 & r \ll l_0 \end{cases}, \quad (2)$$

where  $C_n^2$  is refractive index structure parameter [ $\text{m}^{-2/3}$ ],  $l_0$  presents inner scale of the turbulent cell [m] and  $L_0$  signifies turbulent cell outer scale [m]. There exists relationship between  $C_n^2$  and volume of atmospheric turbulences, which is shown in Table 1.

It is clear that volume of atmospheric turbulences increases with the rising refractive index structure parameter.

Table 1  
Refractive index structure parameter influence  
on the atmosphere

| $C_n^2$ [ $\text{m}^{-2/3}$ ] | Atmospheric turbulences |
|-------------------------------|-------------------------|
| $10^{-16}$                    | Weak                    |
| $10^{-15}$                    | Mean                    |
| $10^{-14}$                    | Strong                  |
| $10^{-13}$                    | Very strong             |

Atmospheric turbulences approves also by the variance of optical intensity in detected signal. Relative variation of optical intensity  $\sigma_{Ir}^2$  can be expressed as [1]

$$\sigma_{Ir}^2 = \frac{\langle I^2 \rangle - \langle I \rangle^2}{\langle I \rangle^2} = \frac{\langle I^2 \rangle}{\langle I \rangle^2} - 1, \quad (3)$$

when  $I$  is optical intensity of the received signal and  $\langle \rangle$  signifies mean value.

When  $\sigma_{Ir}^2 \ll 1$ , then we can use Rytov approximation which ties relative variation of optical intensity and refractive index structure parameter [5]:

$$\sigma_{Ir}^2 = K \cdot C_n^2 \cdot k^{7/6} \cdot L^{11/6}, \quad (4)$$

where  $K$  represents constant for plane wave 1.23 or 0.5 for spherical wave,  $k$  is wave number and  $L$  means distance between transmitter and receiver. This relation is valid only for homogenous distribution of atmospheric turbulences. Optical signal attenuation is caused not only by scattering and absorption on particles, molecules and hydrometeors but also by turbulent atmosphere. Rytov relationship deals with this attenuation due to relation [6]:

$$\alpha_t = 2 \cdot \sqrt{23.17 \cdot k^{7/6} \cdot C_n^2 \cdot L^{11/6}}. \quad (5)$$

Optical signal attenuation and atmospheric turbulences are stronger during sunny days.

#### 4. Available Optical Power

All theories work with the idea that atmospheric turbulences are homogenous and don't take into account laser beam geometric profile. Designed theory of available

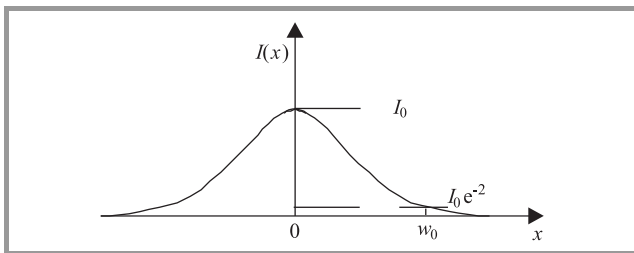


Fig. 6. Gaussian beam distribution.

power looks at optical intensity distribution and there is not necessary to regard homogenous refractive index structure parameter through the laser beam path. We count with mean  $C_n^2$  in the optical path of the laser beam. We also consider Gaussian beam only in this work (Fig. 6). In case of turbulent atmospheric transmission media there occurs optical intensity level fluctuation through the laser beam profile as shown in Fig. 7.

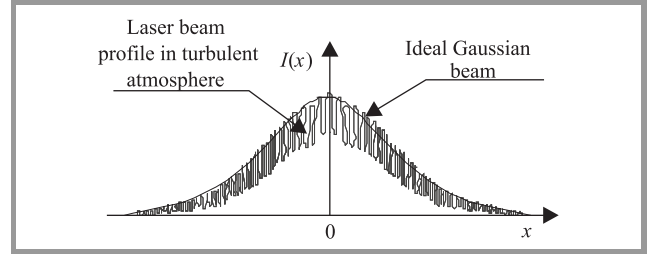


Fig. 7. Affected laser beam by turbulent atmosphere.

It is evident, that turbulent atmosphere debases optical signal properties, and optical intensity level is unstable in the time. This fluctuation has its limits for concrete refractive index structure parameter. We specified turbulent envelope (see Fig. 8) as limits for optical intensity fluctuation.

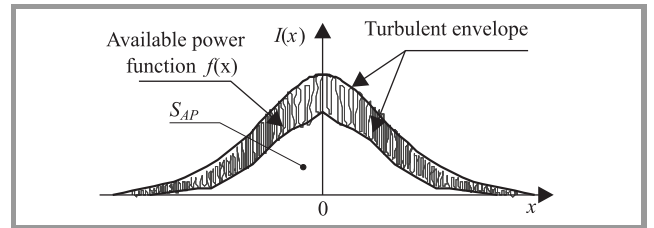


Fig. 8. Turbulent envelope and available power function of Gaussian laser beam.

Available power function  $f(x)$  is defined as lower limit of turbulent envelope. We suppose decreasing available power function with increasing volume of turbulences. According to Fig. 8 we determine area of available power  $S_{AP}$  by following relation:

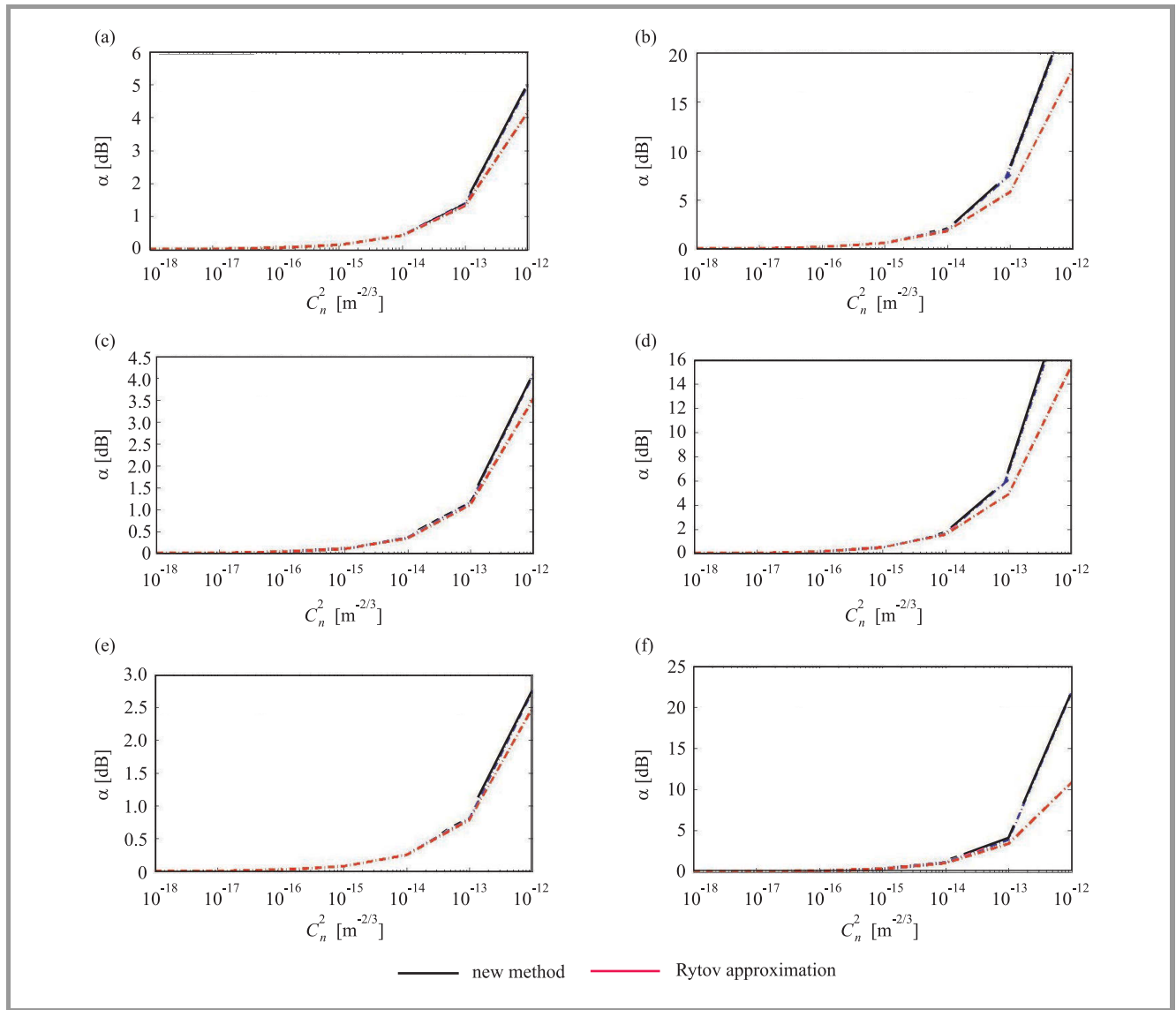
$$S_{AP} = \int_x f(x) dx. \quad (6)$$

Upper limit of  $S_{AP}$  is defined as available power function  $f(x)$  and as lower limit we consider  $x$ -axis (zero). In fact laser beam is 3 dimensional  $(x, y, I(x, y))$  so we set up the volume of available power, for short available power by relation:

$$V_{AP} = \iint_{x, y} f(x, y) dx dy. \quad (7)$$

In non-turbulent area we suppose that available power function is identical to Gaussian optical intensity distribution so we express available power as

$$V_{AP0} = \iint_{x, y} f_0(x, y) dx dy. \quad (8)$$



**Fig. 9.** Comparison results obtained by new method of available power with results given by Rytov approximation for: (a)  $\lambda = 633$  nm,  $L = 50$  m; (b)  $\lambda = 633$  nm,  $L = 250$  m; (c)  $\lambda = 850$  nm,  $L = 50$  m; (d)  $\lambda = 850$  nm,  $L = 250$  m; (e)  $\lambda = 1550$  nm,  $L = 50$  m; (f)  $\lambda = 1550$  nm,  $L = 250$  m.

In case of very high turbulences we consider that available power function is equal to  $x, y$  plane so available power is

$$V_{AP\max} = \iint_{x,y} f_{\max}(x,y) dx dy = 0. \quad (9)$$

To simplify the situation we introduce relative available power  $V_{APr}$  by equation:

$$V_{APr} = \frac{V_{AP}}{V_{AP0}}, \quad (10)$$

where we divide general available power by “non-turbulent” available power. It is evident that relative available power for non-turbulent atmosphere is evaluated by number 1, which goes from the following relation:

$$V_{AP0r} = \frac{V_{AP0}}{V_{AP0}} = 1 \quad (11)$$

and the value of relative available power is equal to 0 for very high volume of turbulences

$$V_{AP\max r} = \frac{V_{AP\max}}{V_{AP0}} = 0. \quad (12)$$

Volume of atmospheric turbulences can be evaluated by the parameter of  $V_{APr}$  from interval  $<0; 1>$ . Finally, we want to express atmospheric turbulence attenuation by method of available area. The next equation is adequate to quantify wanted magnitude:

$$\alpha_{AP} = 10 \cdot \log V_{APr}. \quad (13)$$

We are able to determine turbulence attenuation when the laser beam profile is available at the receiver side.

## 5. Results

We compared results obtained by new method of available power with the results given by Rytov approximation. Available power represents the worst case of turbulent atmospheric attenuation. Characteristics (Fig. 9) show that results for available power and Rytov approximation don't vary much for  $C_n^2$  values smaller than  $10^{-14}$ . The results became different for very strong turbulences. According to the method of available power we can't calculate optical signal attenuation for wavelength 633 nm and  $L = 250$  m as well as for 850 nm and  $L = 250$  m, because relative available powers have value 0, so attenuation is theoretically infinite. This means that there isn't guaranteed detection of the optical signal for this link parameter. Method of available power also says that turbulence attenuation is smaller for higher wavelength.

## 6. Conclusion

A new method of available power is presented in the paper. This method arises from a laser beam optical intensity profile, which is original in atmospheric turbulence phenomena study. On the basis of this method it will be possible to design an appropriate laser beam shape to maximally eliminate negative effect of turbulent atmosphere. For the present, basic theoretical analysis and initiative experiments with Gaussian beam have been provided. According to obtained results, the method of available power is the perspective method for determining optical signal attenuation which is due to turbulence in the atmosphere.

## Acknowledgement

This research has been supported by the research programs of Brno University of Technology MSM21630513 and by the grant Agency of the Czech Republic under contracts No. 102/08/H027.

## References

- [1] L. Dordová and O. Wilfert, "Optimal laser diode operating mode with unstable operating temperature in turbulent atmosphere", in *Proc. Semicond. Laser Dynam. SPIE Eur. Photon. Eur. 2008*, Strasbourg, France, 2008, pp. 1–11.
- [2] O. Tošovský and L. Dordová, "Free space optical channel parameters estimation for high altitude platform system", in *Proc. 18th Int. Conf. Radioelektr. 2008*, Prague, Czech Republic, 2008, pp. 1–5.
- [3] B. E. A. Saleh and M. C. Teich, *Základy Fotoniky*. Prague: Matfyz Press Praha, 1994 (in Czech).

- [4] L. Andrews, R. Phillips, and C. Hopen, *Laser Beam Scintillation with Applications*. Washington: SPIE Press, 2001.
- [5] E. Korevaar, I. Kim, and B. McArthur, "Atmospheric propagation characteristics of highest importance to commercial free space optics", *Proc. SPIE (Atmospheric Propagation)*, vol. 4976, pp. 1–12, 2003.
- [6] A. Naboulsi, M. Sizun, and F. De Fornel, "Propagation of optical and infrared waves in the atmosphere", in *Proc. XXVIIIth URSI Gener. Assem.*, New Delhi, India, 2005, pp. 1–4.



**Lucie Dordová** received the Ing. (M.Sc.) degree from the Brno University of Technology, Czech Republic, in 2006. Since 2006 she has been a Ph.D. student at the Department of Radio Electronics, since 2008 she has been employed as technician. At present she works on dissertation "Method for determination of atmospheric transmission media properties in optical spectrum". Currently she solves a Development Fund of Czech Universities (FRVS) grant and she is a member of IEEE student section. e-mail: xdordo00@stud.feec.vutbr.cz  
Institute of Radio Electronics  
Brno University of Technology  
Purkynova 118  
Brno 612 00, Czech Republic



**Otakar Wilfert** received the Ing. (M.Sc.) degree in electrical engineering in 1971 and CSc. (Ph.D.) degree in applied physics in 1984, both from the Military Academy Brno, Czech Republic. Currently he is a Professor at the Department of Radio Electronics, Brno University of Technology. Areas of his research interest include optical wireless communications and laser radar systems. He is a member of IEEE, SPIE, European Optical Society, Czech and Slovak Photonics Society and an official member of Electronics and Photonics Section of URSI. e-mail: wilfert@feec.vutbr.cz  
Institute of Radio Electronics  
Brno University of Technology  
Purkynova 118  
Brno 612 00, Czech Republic

# The Design of $4 \times 4$ Multimode Interference Coupler Based Microring Resonators on an SOI Platform

Trung-Thanh Le and Laurence W. Cahill

**Abstract**—This paper would like to propose a novel microring resonator based on  $4 \times 4$  multimode interference (MMI) couplers. The device acts as two separate microring resonators just in one structure. The transfer matrix method and the three dimensional beam propagation method (3D-BPM) are used to verify the working principle of the device. The device is then designed on silicon on insulator (SOI) technology. This device may be a very promising building block for optical switches, filters, add-drop multiplexers, delay lines and modulators.

**Keywords**—integrated optics, multimode interference couplers, optical logic gates.

## 1. Introduction

Microring resonators have been used as a basic building block for optical signal processing applications such as optical switches, filters, modulators, and add-drop multiplexers [1]. Almost all of the reported works on microring resonator structures have used directional couplers as the coupling element [2]. In order to meet a variety of requirements for high-speed signal processing, the coupling coefficient needs to be adjusted arbitrarily and the directional couplers can meet this requirement. However, for applications requiring high quality factor  $Q$  of the resonators, i.e., for high speed operation, the separations between two waveguides in the directional coupler must be very small. As a result, high loss due to conversion loss of modes is occurred as shown in [3] recently. Moreover, the directional coupler has a large size and small fabrication tolerance. Therefore, multimode interference (MMI) couplers are used in such structures instead of directional couplers due to their advantages of compactness, ease of fabrication, large fabrication tolerance and ease of cascaded integration [4].

A microring resonator based on a  $2 \times 2$  MMI coupler was demonstrated on silicon on insulator (SOI) channel waveguides for the first time in [2]. However, the coupling ratios of the conventional MMI couplers are very limited and there are only four available coupling ratios 0:100, 50:50, 85:15, and 72:28 if using only one conventional MMI coupler [5]. Therefore, it is highly desired to implement the couplers with variable coupling ratios. To do so, the common approach is to use Mach-Zehnder interferometer (MZI) structures, in which phase shifters are added to the MZI arms to control the phase of the propagation signals as shown in [6], [7].

In this paper, a novel microring resonator based on  $4 \times 4$  MMI couplers is given for the first time. The most interesting characteristic of the proposed device is that the device acts as two separate microring resonators based on  $2 \times 2$  MMI couplers. The coupling ratios can be varied by using two separate phase shifters for the two microring resonators. In order to verify the working principle of the device, the transfer matrix and beam propagation method are used. The device is designed on an SOI platform.

## 2. Microring Resonators Based on $4 \times 4$ MMI Couplers

A microring resonator based on  $4 \times 4$  MMI couplers (MMI-MZI structure) is shown in Fig. 1. The two  $4 \times 4$  MMI couplers have the same width  $W_{MMI}$  and length  $L_{MMI} = \frac{3L_\pi}{2}$ , where  $L_\pi = \frac{\pi}{\beta_0 - \beta_1}$  is the beat length of the MMI coupler. In order to make tunable devices, two phase shifters  $\Delta\phi_1$  and  $\Delta\phi_2$  based on the thermo-optic effect are used in the two arms. Alternatively, passive phase shifters could be used to provide a desired (fixed) power splitting ratio.

In order to analyze the device, the transfer matrix of the identical MMI couplers needs to be derived first and then the total transfer matrix of the device can be determined. The  $4 \times 4$  MMI coupler can be described by a transfer matrix [8]:

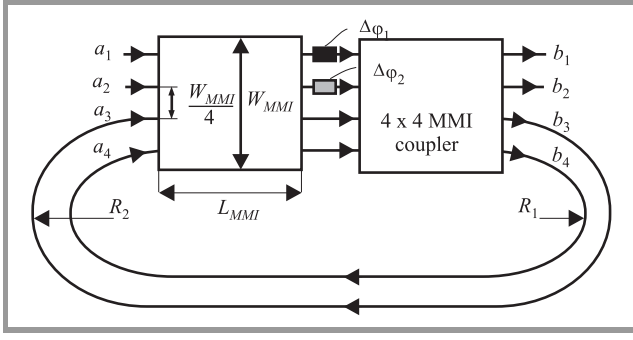
$$M_{4 \times 4} = \begin{bmatrix} m_{11} & m_{12} & m_{13} & m_{14} \\ m_{21} & m_{22} & m_{23} & m_{24} \\ m_{31} & m_{32} & m_{33} & m_{34} \\ m_{41} & m_{42} & m_{43} & m_{44} \end{bmatrix}, \quad (1)$$

where  $m_{ij}$  ( $i, j = 1, \dots, 4$ ) are complex coefficients calculated by using the modal propagation method. At the length  $L = \frac{3L_\pi}{4}$ , the phases of the equal output signals at the output waveguides can be calculated from

$$\phi_{ij} = \phi_0 + \pi + \frac{\pi}{16}(j-i)(8-j+i), \text{ for } i+j \text{ even}$$

$$\text{and } \phi_{ij} = \phi_0 + \frac{\pi}{16}(i+j-1)(8-j-i+1), \text{ for } i+j \text{ odd.}$$

Here, input ports  $i$  ( $i = 1, \dots, 4$ ) are numbered from down to up and the output ports  $j$  ( $j = 1, \dots, 4$ ) are numbered from up to down in the MMI coupler and  $\phi_0 = -\beta_0 L_{MMI} - \frac{\pi}{2}$  is a phase constant factor that is associated with the MMI ge-



**Fig. 1.** The structure of a microring resonator based on  $4 \times 4$  MMI couplers.

ometry and therefore can be neglected in the following analyses.

The  $4 \times 4$  MMI coupler at a length of  $L_1 = \frac{3L_\pi}{4}$  is described by the following transfer matrix:

$$M_{4 \times 4} = \frac{1}{2} \begin{bmatrix} -1 & -e^{j\frac{3\pi}{4}} & e^{j\frac{3\pi}{4}} & -1 \\ -e^{j\frac{3\pi}{4}} & -1 & -1 & e^{j\frac{3\pi}{4}} \\ e^{j\frac{3\pi}{4}} & -1 & -1 & -e^{j\frac{3\pi}{4}} \\ -1 & e^{j\frac{3\pi}{4}} & -e^{j\frac{3\pi}{4}} & -1 \end{bmatrix}. \quad (2)$$

If the length is doubled to  $L_{MMI} = 2L_1 = \frac{3L_\pi}{2}$ , a new  $4 \times 4$  MMI coupler is formed and its transfer matrix is

$$M_{new} = (M_{4 \times 4})^2 = \frac{1}{2} \begin{bmatrix} 1-j & 0 & 0 & 1+j \\ 0 & 1-j & 1+j & 0 \\ 0 & 1+j & 1-j & 0 \\ 1+j & 0 & 0 & 1-j \end{bmatrix}. \quad (3)$$

This matrix can be considered as consisting of two separate submatrices which describe two  $2 \times 2$  the 3 dB MMI couplers, both having the transfer matrix:

$$\mathbf{M}_2 = \frac{1}{2} \begin{bmatrix} 1-j & 1+j \\ 1+j & 1-j \end{bmatrix} = \frac{1}{\sqrt{2}} e^{-j\frac{\pi}{4}} \begin{bmatrix} 1 & j \\ j & 1 \end{bmatrix}. \quad (4)$$

If two  $4 \times 4$  MMI couplers are connected in the MZI structure of Fig. 1, then the relations between the complex amplitudes at the input ports and output ports can be expressed in terms of the transfer matrices of the 3 dB MMI couplers and the phase shifters as follows:

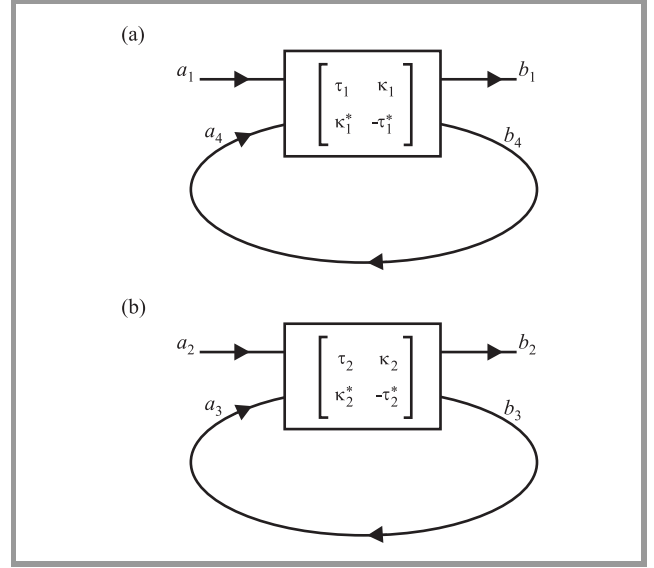
$$\begin{bmatrix} b_1 \\ b_2 \end{bmatrix} = e^{j\frac{\Delta\phi_1}{2}} \begin{bmatrix} \tau_1 & \kappa_1 \\ \kappa_1^* & -\tau_1^* \end{bmatrix} \begin{bmatrix} a_1 \\ a_2 \end{bmatrix}, \quad (5)$$

$$\text{where } \tau_1 = \sin\left(\frac{\Delta\phi_1}{2}\right) \text{ and } \kappa_1 = \cos\left(\frac{\Delta\phi_1}{2}\right), \quad (6)$$

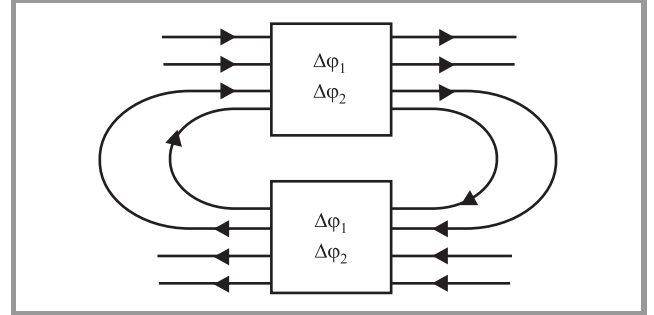
$$\begin{bmatrix} B_1 \\ B_2 \end{bmatrix} = e^{j\frac{\Delta\phi_2}{2}} \begin{bmatrix} \tau_2 & \kappa_2 \\ \kappa_2^* & -\tau_2^* \end{bmatrix} \begin{bmatrix} A_1 \\ A_2 \end{bmatrix}, \quad (7)$$

$$\text{where } \tau_2 = \sin\left(\frac{\Delta\phi_2}{2}\right) \text{ and } \kappa_2 = \cos\left(\frac{\Delta\phi_2}{2}\right). \quad (8)$$

Therefore, the whole device can be viewed as consisting of two independent microresonators having different power coupling ratios as shown in Fig. 2. This means that two independent switches and filters can be made by using this structure.



**Fig. 2.** Two separate microresonators created from the  $4 \times 4$  MMI structure: (a) microresonator using input ports 1 and 4, and (b) microresonator using input ports 2 and 3.



**Fig. 3.** Two separate add-drop multiplexers based on microresonators using  $4 \times 4$  MMI-MZI structures.

Another useful structure for implementing add-drop multiplexing functions is shown in Fig. 3, where the microresonator is coupled to a second MMI-MZI structure. The coupling ratios can be controlled by the phase shifters. Similarly, two independent add-drop filters can be obtained from this structure.

### 3. Simulation Results and Discussions

In this section, in order to verify the working principle of the devices, the three dimensional beam propagation method (3D-BPM) [9], [10] is used to optimize the designs for MMI devices. It is well known that the finite difference time-domain (FDTD) method is a general method

to solve Maxwell's partial differential equations numerically in the time domain. Simulation results for devices on the SOI channel waveguide using the 3D-FDTD method can achieve a very high accuracy. However, due to the limitation of computer resources and memory requirements, it is difficult to apply the 3D-FDTD method to the modeling of large devices on the SOI channel waveguide. Meanwhile, the 3D-BPM was shown to be a quite suitable method [11], [12] that has sufficient accuracy for simulating devices based on SOI channel waveguides [13], [14].

The waveguide structure used in the designs is shown in Fig. 4. Here,  $\text{SiO}_2$  ( $n = 1.46$ ) is used as the upper cladding material. An upper cladding region is needed for devices using the thermo-optic effect in order to reduce loss due to metal electrodes. Also, the upper cladding region is used to avoid the influence of moisture and environmental temperature.

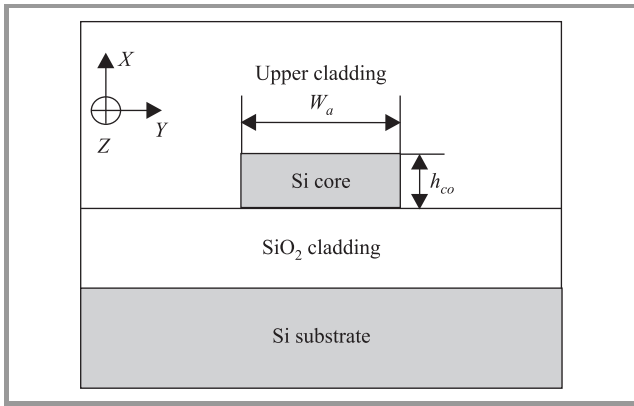


Fig. 4. Waveguide structure used in the simulations.

The parameters used in the designs are as follows: the waveguide has a standard silicon thickness of  $h_{co} = 220$  nm and access waveguide widths are  $W_a = 0.48$   $\mu\text{m}$  for single mode operation. The refractive index of the silicon core is  $n_{\text{Si}} = 3.45$ . It is assumed that the designs are for the TE (transverse electric) polarization at a central optical wavelength  $\lambda = 1550$  nm. The width of the MMI is  $W_{\text{MMI}} = 6$   $\mu\text{m}$ . The access waveguide is tapered to a width of  $W_{tp} = 800$  nm to improve device performance.

The 3D-BPM simulations for a  $4 \times 4$  MMI coupler having a length of  $L_1 = \frac{3L_\pi}{4}$  are shown in Fig. 5a for a signal presented at input port 1 and Fig. 5b for a signal at input port 2. The optimized length of the MMI coupler calculated by using 3D-BPM is  $L_1 = 71.70$   $\mu\text{m}$ . The excess loss calculated is 0.35 dB for both cases.

If two  $4 \times 4$  MMI couplers with the same length of  $L_{\text{MMI}} = \frac{3L_\pi}{4}$  are cascaded together, then a  $4 \times 4$  MMI coupler having a length of  $L_{\text{MMI}} = \frac{3L_\pi}{2}$  is formed. The transfer matrix of this new  $4 \times 4$  MMI coupler is given by Eq. (4). The 3D-BPM will now be used to verify this prediction. The 3D-BPM simulations for this  $4 \times 4$  MMI coupler are shown in Fig. 6. The 3D-BPM simulations show that the power splitting ratios for each MMI coupler used in these microresonators are 0.42/0.43. The optimized length of

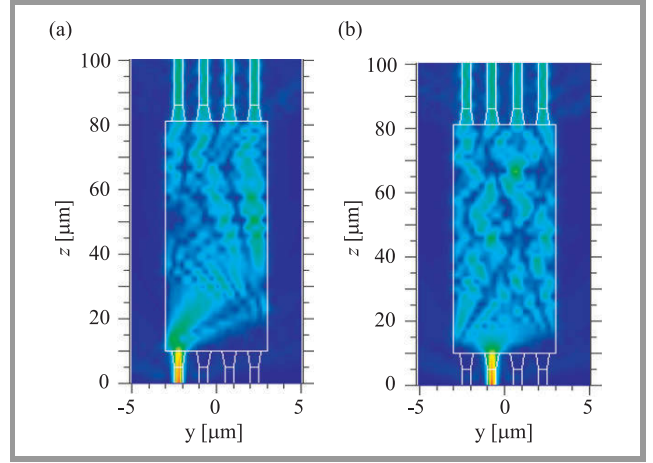


Fig. 5. The 3D-BPM simulations for a  $4 \times 4$  MMI coupler at length  $L_1 = \frac{3L_\pi}{4}$  with input signal at (a) port 1 and (b) port 2.

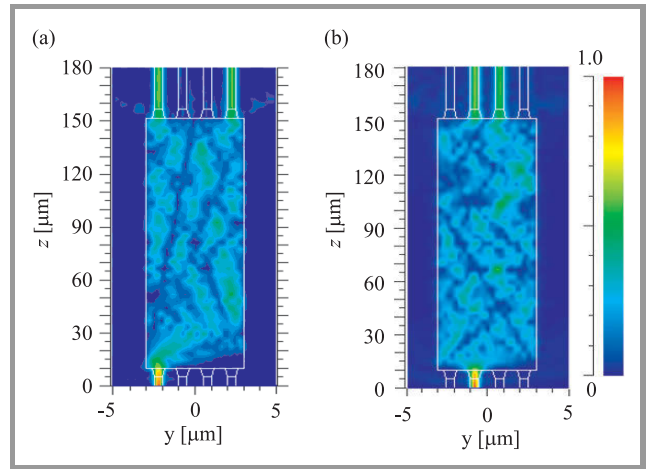
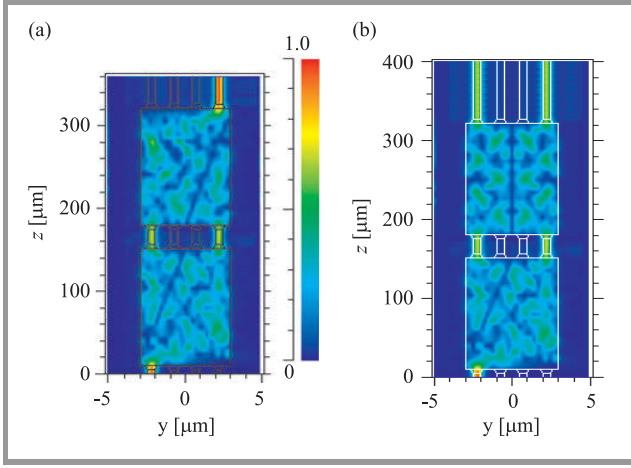


Fig. 6. The 3D-BPM simulations for a  $4 \times 4$  MMI structures with the signal at (a) input port 1 and (b) input port 2.

each MMI coupler is found to be  $L_{\text{MMI}} = 141.7$   $\mu\text{m}$ . For both cases, the excess losses are 0.7 dB and the imbalances are 0.1 dB.

By connecting the above MMI structures (each MMI with length of  $L_{\text{MMI}} = \frac{3L_\pi}{2}$ ) together as shown in Fig. 1, double-microresonators can be achieved. Phase shifters can be introduced in the linking arms of the MZI structure in order to vary the coupling coefficients of the couplers. For example, if phase shifts of  $\Delta\phi_1 = \Delta\phi_2 = 0$  are introduced at the linking arms of the MZI, then the 3D-BPM simulation (Fig. 7a) shows that the normalized output powers (for signal at input port 1) at output ports 1, 2, 3 and 4 are 0, 0, 0 and 0.8, respectively. The computed excess loss is 0.96 dB.

If phase shifts  $\Delta\phi_1 = \frac{\pi}{2}$  and  $\Delta\phi_2 = 0$  are introduced at the linking arms, the 3D-BPM simulation for the signal at input port 1 is shown in Fig. 7b. The normalized output powers at output ports 1, 2, 3 and 4 calculated to be 0.41, 0, 0 and 0.39, respectively. The excess loss is 0.9 dB and the imbalance is 0.2 dB in this case.



**Fig. 7.** The 3D-BPM simulations of cascaded MMI couplers used for the microresonator structure when the signal is at input port 1 with (a) phase shifts  $\Delta\phi_1 = \Delta\phi_2 = 0$  and (b)  $\Delta\phi_1 = \frac{\pi}{2}$  and  $\Delta\phi_2 = 0$ .

The quality factor ( $Q$ ) of a single microresonator based on a  $2 \times 2$  MMI coupler is given by [2]

$$Q \approx \frac{\pi n_g L_R \sqrt{\alpha\tau}}{\lambda(1 - \alpha\tau)}. \quad (9)$$

Here,  $n_g$ ,  $L_R$ , and  $\lambda$  are the group index, length of the racetrack waveguide and wavelength, respectively.  $\tau$  is the transmission coefficient of the coupler. In order to achieve a high  $Q$  microresonator, the coupling coefficient  $|\kappa| = \sqrt{\alpha^2 - \tau^2}$  needs to be small. Note that the bend radius also has a strong effect on the  $Q$ -factor. However, if the bend radius increases, then the transmission loss will increase, while the bend loss does not significantly decrease. Here  $\alpha^2$  is the power loss factor introduced by MMI couplers and racetrack waveguides (including both bend loss and transmission loss). By varying the phase shifts at the linking arms, the power coupling ratios can be tuned.

The complete device presented in this paper is equivalent to two separate  $2 \times 2$  MMI-based microresonators. Each microresonator may have different transmission characteristics such as different quality ( $Q$ ), different free spectral range ( $FSR$ ), finesse ( $F$ ) and different bandwidth ( $BW$ ).

In the following examples, the  $Q$ -factors of the microring resonators will be determined. Consider a  $4 \times 4$  microring resonator with following parameters: bend radii of the first microring resonator and second microring resonator (see Fig. 1) are  $R_1 = 5 \mu\text{m}$  and  $R_2 = 20 \mu\text{m}$ . If the 3 dB MMI couplers with the same length of ( $L_{MMI} = 141.7 \mu\text{m}$  (Fig. 6) are used in the structure, then the calculated power transmission coefficients are  $|\tau_1|^2 = 0.42$  and  $|\tau_2|^2 = 0.42$ . The estimated  $Q$ -factor for the first microring resonator is  $Q_1 \approx 1000$  and for the second microring resonator is  $Q_2 = 1650$ .

If two  $4 \times 4$  MMI couplers having the same length of  $L_{MMI} = 141.7 \mu\text{m}$  are connected together in order to produce  $2 \times 2$  tunable MMI couplers (Fig. 7), then power transmission ratios of the two MMI couplers  $|\tau_1|^2$  and  $|\tau_2|^2$

can be varied by adjusting the phase shifts  $\Delta\phi_1$  and  $\Delta\phi_2$  in the linking arms, respectively. For example, if phase shifts  $\Delta\phi_1 = 0.9\pi$  and  $\Delta\phi_2 = 0.5\pi$ , then the 3D-BPM shows that the normalized output powers at output ports 1, 2, 3 and 4 are 0.7, 0, 0 and 0.1, respectively. The calculated power transmission coefficients are  $|\tau_1|^2 = 0.7$  and  $|\tau_2|^2 = 0.41$ . Therefore, the estimated  $Q$ -factors are  $Q_1 \approx 4300$  for the first microring resonator with bend radius  $R_1 = 5 \mu\text{m}$  and  $Q_2 \approx 2500$  for the first microring resonator with bend radius  $R_2 = 20 \mu\text{m}$ .

It is interesting to note that provided that the phase shifters operate with sufficient speed, functions of switching, modulating and add-drop multiplexing can be achieved in both microresonators.

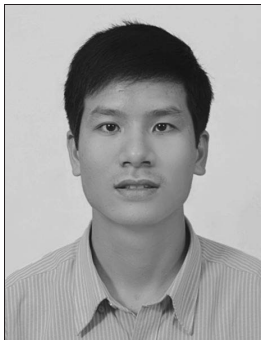
## 4. Conclusions

In this paper, we have presented a novel microring resonator structure based on  $4 \times 4$  MMI couplers. This device acts as two separate microring resonators. The design of the device on the SOI platform has been presented by using the transfer matrix method and the working principle is verified using the 3D BPM method. If the phase shifters operate fast enough, then the device may be a very promising building block for optical switches, optical modulators, and optical add-drop multiplexers.

## References

- [1] D. G. Rabus, *Integrated Ring Resonators – The Compendium*. Berlin: Springer-Verlag, 2007.
- [2] D.-X. Xu, A. Densmore, P. Waldron, J. Lapointe, E. Post, and A. Delage, "High bandwidth SOI photonic wire ring resonators using MMI coupler", *Opt. Expr.*, vol. 15, no. 6, pp. 3149–3155, 2007.
- [3] F. Xia, L. Sekaric, and Y. A. Vlasov, "Mode conversion losses in silicon-on-insulator photonic wire based racetrack resonators", *Opt. Expr.*, vol. 14, no. 9, pp. 3872–3886, 2006.
- [4] T. T. Le and L. W. Cahill, "The modeling of MMI structures for signal processing applications", in *Integr. Opt. Dev. Mater. Technol. XII Proc. SPIE*, San Jose, USA, 2008, vol. 6896, pp. 68961G–68961G-7.
- [5] P. A. Besse, E. Gini, M. Bachmann, and H. Melchior, "New  $2 \times 2$  and  $1 \times 3$  multimode interference couplers with free selection of power splitting ratios", *IEEE J. Lightw. Technol.*, vol. 14, no. 10, pp. 2286–2293, 1996.
- [6] A. Yariv, "Critical coupling and its control in optical waveguide-resonator systems", *IEEE Photon. Technol. Lett.*, vol. 14, no. 4, pp. 483–485, 2002.
- [7] J. M. Choi, R. K. Lee, and A. Yariv, "Control of critical coupling in a ring resonator-fiber configuration: application to wavelength-selective switching, modulation, amplification, and oscillation", *Opt. Lett.*, vol. 26, no. 16, pp. 1236–1238, 2001.
- [8] M. Bachmann, P. A. Besse, and H. Melchior, "General self-imaging properties in  $N \times N$  multimode interference couplers including phase relations", *Appl. Opt.*, vol. 33, no. 18, pp. 3905–3911, 1994.
- [9] W. P. Huang, C. L. Xu, and S. K. Chaudhuri, "A finite-difference vector beam propagation method for three-dimensional waveguide structures", *IEEE Photon. Technol. Lett.*, vol. 4, no. 2, pp. 148–151, 1992.
- [10] W. P. Huang, C. L. Xu, W. Lui, and K. Yokoyama, "The perfectly matched layer (PML) boundary condition for the beam propagation method", *IEEE Photon. Technol. Lett.*, vol. 8, no. 5, pp. 649–651, 1996.

- [11] D. Dai and S. He, "Design of an ultrashort Si-nanowaveguide-based multimode interference coupler of arbitrary shape", *Appl. Opt.*, vol. 47, no. 19, pp. 38–44, 2008.
- [12] D. Dai and S. He, "Optimization of ultracompact polarization-insensitive multimode interference couplers based on Si nanowire waveguides", *IEEE Photon. Technol. Lett.*, vol. 18, no. 19, pp. 2017–2019, 2006.
- [13] E. Dulkeith *et al.*, "Group index and group velocity dispersion in silicon-on-insulator photonic wires", *Opt. Expr.*, vol. 14, no. 9, pp. 3853–3863, 2006.
- [14] J. I. Dadap *et al.*, "Nonlinear-optical phase modification in dispersion-engineered Si photonic wires", *Opt. Expr.*, vol. 16, no. 2, pp. 1280–1299, 2008.



**Trung-Thanh Le** received the B.E. and M.E. degrees in electronic and telecommunication engineering from the Hanoi University of Technology, Vietnam, in 2003 and 2005, respectively. Since 2003, he has been a lecturer at the National University of Transportation and Communications, Hanoi. He works towards the Ph.D. degree

in electronic engineering at the La Trobe University, Melbourne, Australia, since 2006. His research interests are

integrated optical devices and photonic signal processing, especially multimode interference based devices.

e-mail: thanh.latrobe@gmail.com

Department of Electronic Engineering

La Trobe University

Melbourne, Vic 3086, Australia



**Laurence W. Cahill** received the B.E., M.Sc. and Ph.D. degrees from the University of Melbourne, Australia. He is a Professor in the Department of Electronic Engineering, La Trobe University, Australia. He is a senior member of the IEEE, member of SPIE and a Fellow of Engineers Australia. His research interests lie in the area of

optical and photonics. Particular interests are semiconductor laser dynamics, high speed circuits for optical sources and receivers, mid infrared detectors, photonic switching, fibre sensors, computer aided design, image processing, and signal processing.

e-mail: L.Cahill@latrobe.edu.au

Department of Electronic Engineering

La Trobe University

Plenty road, Bundoora Campus

Melbourne, Vic 3086, Australia

# Linux Scheduler Improvement for Time Demanding Network Applications, Running on Communication Platform Systems

Marcin Hasse and Krzysztof Nowicki

**Abstract**—Communication platform systems as, e.g., advanced telecommunication computing architecture (ATCA) standard blades located in standardized chassis, provides high level communication services between system peripherals. Each ATCA blade brings dedicated functionality to the system but can as well exist as separated host responsible for servicing set of task. According to platform philosophy these parts of system can be quite independent against another solutions provided by competitors. Each system design can be different and can face with many computer systems design problems. One of the most difficult design problems to solve is system integration with a set of components running on different operating system levels. This paper presents Linux scheduler improvement possibility to make user space application classified as time demanding (required to be serviced by CPU in given amount of time) running in user space together with complicated kernel software structure in the system.

**Keywords**—communication platform systems, Linux, operating system, scheduler.

## 1. Introduction

Today's communication trends are consolidated to follow platform strategy. This strategy is to provide standard base solutions to be re-used over wide range of products. Advanced telecommunication computing architecture (ATCA) [1], [2] is a very good example which is successfully matching platform objectives. The communication architecture between subsystems, the major issue in platform implementation, seems to be prepared to manage restricted subsystems requirements. Standard base chassis with intelligent platform management interface (IPMI) and Ethernet communication makes a very friendly base for a big range of network products like switches and gateways as well as for computer base products like single board computers (SBC). Well organized communication between subsystems and advanced management opportunities makes ATCA platform very interesting solution for telecommunication market, especially because these systems are following restricted energy consumption and thermal norms.

From the other perspective, systems prepared to match ATCA standard have a big challenge to follow restricted norms and propose good enough performance for end users.

The law formed by Herb Grosh in 1965 [3] indicating that computer performance increases as the square of the cost. Regarding to this law SBC with more RAM memory and with bigger HDD would have a better performance, but ATCA system performance can not be limited only to regular PC specific costs. They need to be considered as well energy and thermal system assumptions which makes the cost more significant. This is causing that ATCA systems are designed with limited system resources mostly according only to design demands.

Platform strategy gives opportunity to application designers to choose ATCA hardware base on system demands. For example Ethernet line card hardware (based on network processor or other multicore processor) would be a good choice for IPsec gateway application.

Only one disadvantage of customize hardware to match ATCA platform standards is that blades (as line cards) can not easily be extended to additional system resources as RAM or flash memory. Application designers are responsible for achieving software goals with available resources starts from operating system and ends on specific application (as IPsec IKE [4], [5] for IPsec gateway example).

Linux is a most popular platform choice for ATCA blades used in network core: as gateways, routers, etc. It would be as well most reasonable choice for IPsec gateway Ethernet line card example. There are several Linux operating systems available with embedded system support and with ATCA blades board support packages (BSP) as MontaVista [6] or WindRiver [7]. Additional advantage of Linux OS is its open source nature and developers have access even to kernel sources. This is big opportunity to have more influence on system performance while developer can place program in the kernel level. Regarding Linux GPL [8] licence programs in kernel space suppose to be published as open source. This restriction creates a barrier for Linux commercial application providers – which are mostly offered as a user space programs. For example additional IPsec gateway functionalities as virtual router redundancy protocol (VRRP) [9] or simple network management protocol (SNMP) [10] can be taken from independent supplier as a user space application.

This paper indicates problems with limited system resources operated by Linux OS and common problems with user and kernel space applications working together in net-

work and real time environment. In Section 2 there will be Linux scheduler analyzed in order to present issue with time demand user space application working together with real time tasks in kernel level. Proposed scheduler improvement to make Linux more flexible if there are time demanding user space applications is described in Section 3. Measurement of improvement results plus comparison with standard scheduler, are described in Section 4.

## 2. Linux Scheduler against User Space Time Demanding Processes

In the Linux operating system there can be determined two kinds of threads [11]. First would be CPU bound, which spends a lot of time using central processing unit (CPU) and making computation. The second would be I/O bound most time waiting for a I/O operation to complete. Scheduler in Linux is designed to deal with both types of threads in the fair way, but there is no well known method to determine if thread should be classified as I/O bound or CPU bound. The reason why scheduler should treat I/O bound threads with bigger priority is slow nature of I/O. There is understandable requirement to service human input as fast as possible – most people simply do not like wait especially when they wanted to have something done by a computer. It takes a long time for service I/O so it is good if that kind of requests can be serviced as fast as possible.

Linux scheduler goals as efficiency and interactivity makes this mechanism more friendly for servers (most common usage of Linux these days) and for desktop (where Linux would like to be more important than today). Unfortunately, if something is more matching servers and desktops then it is probably less matching core network systems as, e.g., gateways.

In order to evaluate scheduler role in the system it would be efficient to determinate scheduler performance. Introduction this metric should allow checking if scheduler works properly for given set of margin conditions (different than for normal server or desktop usage conditions). In most cases performance determines the time required to finish the task. For process scheduler performance it would be time in which task (CPU or I/O) will be successfully serviced. In the other words performance  $P$  (for the process with priority  $X$ ), would be a process wait time until it will be serviced by CPU  $T_w$  and CPU execution slice time  $T_s$  with assumption that task could not be finished in  $Q$  CPU slices:

$$P(P_X) = \sum_{n=1}^Q [T_w(P_X) + T_s].$$

Waiting queue  $T_w$  time is dependent on several additional systems conditions as number of tasks  $N$  waiting for CPU and their priorities time slice  $T_{sX}$ :

$$T_w(P_X) = \sum_{n=1}^X (N \cdot T_{sX}).$$

While there will be several the same priority tasks, e.g.,  $S$  for scheduler it will service them in request order:

$$P(P_X) = \sum_{n=1}^Q \left\{ \sum_{k=1}^X (N \cdot T_{sX}) + T_s \right\}.$$

Priorities in Linux kernel 2.6 scheduler can be set between 0 and 139, where priorities between 0–99 determine kernel threads and 100–139 determine user threads. This thread priority is playing significant role when scheduler in kernel 2.6 assigns tasks into two queues: active and expired. Waked up thread is placed in active queue base on its priority. This means that when there are threads in systems with much different priorities, thread with bigger priority might be assigned again to active queue instead of expired queue. As long as there are threads in active queue as long threads from expiry queue will not get CPU time for execution. This might make situation while waked up threads stream can delay amount of time execution of tasks from expired queue. In the worst scenario this is possible even with only several CPU bound threads making situation in which low priority threads will be delayed more than several seconds.

In gateway example presented in Section 1 there is market driven possibility in which significant system applications are implemented to be executed in user space. As long as application providers are interested to not general public licence (GPL) it can not be implemented in kernel level. It is easy to imagine that set of user space applications can be executed in the system together with multiple kernel level tasks waked up quite often. Kernel priorities will take precedence over user space and will be serviced in active queue. In the same time expire queue threads will be still on hold.

The ATCA solutions on the market these days can give lots of communication opportunities to be used in professional systems. There is no communication connected issues any more. Separated parts of platform can exchange information base on standard backplane solutions offered by many suppliers. ATCA blades providers are proposing as well many systems working with energy save oriented CPUs like, e.g., ARM. For networking, these low performance cores are used to provide management opportunity for other CPUs like, e.g., network processors. Unfortunately, systems with good communication abilities might have some weak points in low performance management core areas. If system design assumes existence of many Linux kernel space threads (often waked up) together with critical for system user space applications, so less priority threads might wait to be serviced even several seconds.

### 2.1. Time Demanding User Space Processes

Common practice made by ATCA system providers is system integration on the application level. As long as stable kernel with support is offered by companies like WindRiver or Montavista, as long management software can be offered by many other suppliers. Only in networking there exist many areas in which 3th party applications can be used.

For example SNMP stack or Internet key exchange (IKE) support can be purchased from protocol specialized suppliers and integrated together with blade interfaces. There is a big advantage for that kind of solution, especially for companies specialized in restricted areas like, e.g., signaling. Thanks to integration possibilities these companies can provide final systems to the market even without specialized knowledge in all system functionalities. All they need to provide is integration of solutions with support from application suppliers.

In the group of networking applications to be used with integration model there are some “time demanding” examples. Advanced telecommunication systems used in core networking are often designed to provide redundancy opportunities. For example in the case of gateway failure the system is prepared to switch over to backup gateway. This redundancy can be serviced by VRRP protocol (see Fig. 1).

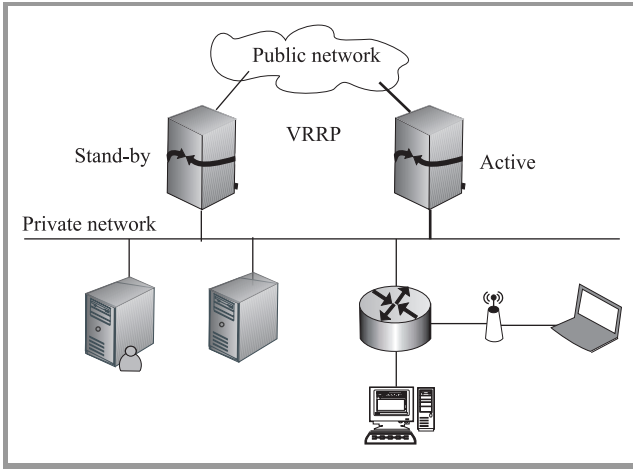


Fig. 1. Redundancy for network core nodes.

The VRRP protocol assumes continuous communication between active and backup gateway. In the case of communication lost for specified period of time, failover between gateways supposes to occur. In this VRRP example system is classified as unhealthy (dead), when packet exchange between gateways will not occur in given time.

One of the possible failover conditions would be user space VRRP application thread stockied in the expiry scheduler queue waiting until continuously waked up kernel threads will finally finish their jobs.

## 2.2. Critical Scenario Analyses

Linux scheduler is dealing with one run queue for each CPU in the system. Each run queue contains set of two priority arrays. All tasks begin in active priority array, and when they are executed on CPU there are moved to expired priority array and new time slice is calculated. Time slice describes time which given task will be able to spend on CPU before another task will be given a chance.

A change between active and expiry priority array will take place when there will be no tasks in the first active array. Linux 2.6 scheduler is designed to schedule always all the tasks with the biggest priority (see Fig. 2). If there are couples of tasks with the same priority then they will be scheduled with round robin algorithm.

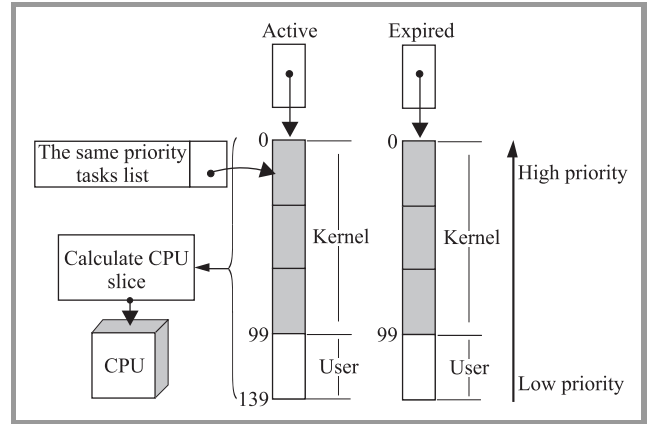


Fig. 2. Scheduler priority queue model for single CPU system.

In the Linux there can be user defined static values assigned to the priorities (*nice* from 20 to -19 by default 0). System is not intended to change static values to respect user input. To provide a difference between service I/O bound and CPU bound tasks scheduler uses dynamic priorities (0–139), which can award a bonus or depreciate task about 5 priority levels. Dynamic prioritization uses heuristic based on tracking how much time a task is sleeping against how long they are using CPU. Time  $T_{S\_AV}$  is never intended to be bigger than  $T_{max}$  and a bonus to bigger priority is given to tasks with bigger  $T_{S\_AV}$ . Priority can dynamically be changed based on average time  $T_{S\_AV}$  of CPU waiting on CPU (I/O bound). When task is waked up after  $T_S$  to be executed on CPU then

$$\forall_{T_{S\_AV} < T_{max}} T_{S\_AV} = T_{S\_AV} + T_S.$$

When task finishes using CPU after  $T_{CPU}$  then

$$\forall_{T_{S\_AV} < T_{max}} T_{S\_AV} = T_{S\_AV} - T_{CPU}.$$

Scheduler will not perform any heuristic priority changes for real time tasks (see Fig. 2 – priorities 0–100). Real time tasks are always executed with the current priority. For the rest of tasks bonus  $B$  (maximum  $B_{max}$ ) will be calculated in the following way:

$$B = NTJ \left( \frac{T_{S\_AV} B_{max}}{T_{max}} \right),$$

where  $NTJ - NS\_TO\_JIFFIES$  (see macro defined in sched.c [12]) depends on CPU frequency  $f$  [Hz],

$$NTJ(T) = \frac{T}{\frac{1000\,000\,000}{f \text{ [Hz]}}}.$$

When  $T_{S_{AV}}$  is high (I/O bound) then  $B$  might be 10 – task priority  $P$  will be increased about 5 and when  $T_{S_{AV}}$  is zero then  $B$  as well will be 0 – task priority  $P$  will be decreased about 5 levels.

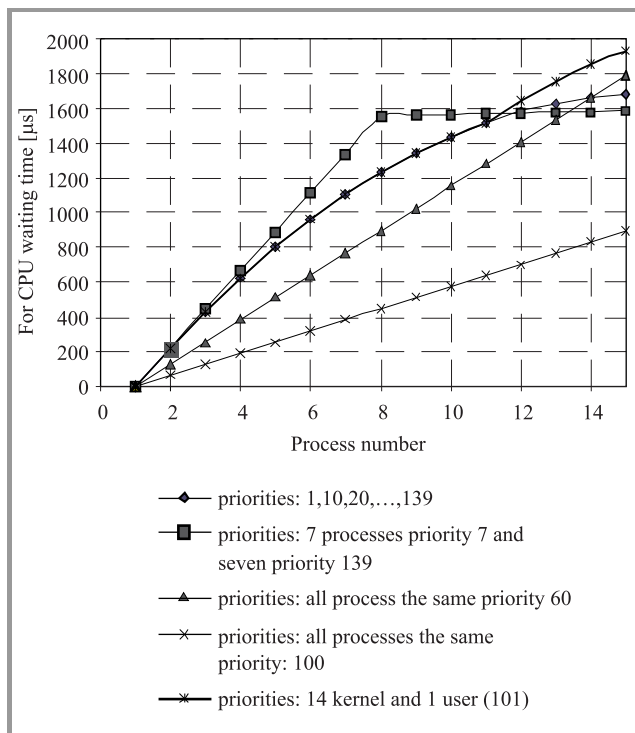
Priority is an essential metric for scheduler to calculate time slice. The lowest dynamic priority process will get the biggest time slice  $T_{CPU}$  (for given  $P_{\max}$  – maximal priority and  $P_{\max,U}$  – maximal user priority):

$$T_{CPU} = \max \left( T_{CPU\_DEF} \frac{(P_{\max} - P)}{\frac{P_{\max,U}}{2}}, T_{CPU\_min} \right),$$

$$T_{CPU\_DEF} = \frac{100 f \text{ [Hz]}}{1000},$$

$$T_{CPU\_min} = \max \left( \frac{5 f \text{ [Hz]}}{1000}, 1 \right).$$

Figure 3 presents a set of possible waiting for CPU times for 15 tasks with different priorities (CPU 800 MHz). Lower priorities tasks will receive less CPU time than task with bigger priorities. This chart presents data for single active queue without changing to expiry queue.

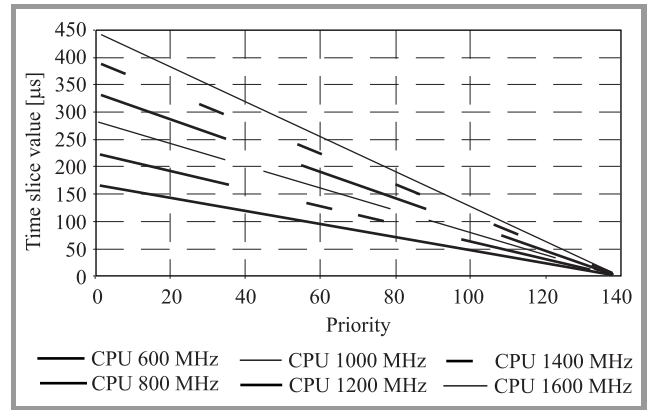


**Fig. 3.** Scheduler active queue tasks possible waiting for CPU time.

If system administrator assign the biggest possible nice priority to user space VRRP it is easy to prove that if there is many waked up processes in the active queue, then user space process will not be able to get CPU even after several milliseconds. It should be enough to set many processes in the kernel with high priorities. Delay in servicing VRRP process might be too big relative to its time demanding

behavior. If network node will not send frame notification that he is alive there might be failover procedure started.

Task with higher priority will be given with longer CPU time than task with lower priority. The CPU time slice will even be longer on machines with higher CPU frequency (see Fig. 4).



**Fig. 4.** Time slice estimation for different priorities on different CPUs.

To marginalize possibility of this situation there always can be said, that Linux kernel application supposes to be designed to avoid having important user space process stocked in the queue. In most cases it would be possible to work with design to make sure that database structures serviced in the kernel would have enough pointer references to avoid checking field by field. Unfortunately, close to this assumption exist many other possibilities (especially valid for ATCA) like, e.g., not enough memory to implement good enough data structures to avoid checking fields in the loop. Another common in the market reason is changing application assumptions when implementation is finished, that it is easier to find another solution to solve user space process stock issue than expensive redesign.

This set of explanations was accumulated in this section to assure about reasonability of researches presented in the next sections. Analyses of scheduler changes possibilities should always become first, before system designer decides to change base functionality of system kernel. Kernel level changes would make whole system less stable – unless validation in the field confirms that kernel patch works properly. On the basis of the following research results it should be much easier to decide if application should be redesigned or rather scheduler should be improved.

### 3. Scheduler Improvement for User Space Time Demanding Applications

Linux with its open source (OS) nature is giving this useful opportunity to provide changes even in the most critical parts of a code. This is allowed to change application as well as patch the kernel. Scheduler, as one of the most critical part of kernel, is already able to deal with user

processes base on heuristic method described in previous section.

To change a user space priority, there is average  $T_{S\_AV}$  process sleeping time metric introduced in Linux kernel 2.6. This metric is working good to determine I/O bound threads. To have time demanding user space application running with bigger priority there is a different heuristic metric needed.

**Scheduler metric to classify time demanding applications.** To create a functional metric for user space process (which can be classified as time demanding) there needs to be such process characteristic introduced. Base on this characteristic the new metric can be introduced to classify task priority to be changed.

Time demanding user space process:

- is awaked periodically for a specified amount of time;
- in most cases required to deal with I/O peripherals;
- its awake time can be different – depends on process functionality.

Usage of I/O peripherals can match many other processes, not necessarily time demanding and is not a good characteristic for metric. Much more useful seems to be periodic activity of time demanding applications. If scheduler could classify that application requests CPU access every defined amount of time (different for different processes), it can reassign bigger priority to application process.

In order to be more flexible in scheduler changes it would be good to use variables already implemented in the kernel. To classify task as time demanding the average waiting for CPU time  $T_{S\_AV}$  can not be easily used. For user space task this time depends on many conditions in kernel. For example  $T_{S\_AV}$  can be completely different while kernel threads are requested to make big amount of calculations. Base on priorities kernel threads will be given with CPU time slice before user processes (see Fig. 3).

Opposite possibility to detect time demanding tasks would be eliminate these, which are not matching characteristic. Scheduler on the beginning could give the same big priority to all the processes to make sure that all time slices will be the same. To notify Linux scheduler that process/task should get a CPU (normally based on I/O) there is kernel variable `need_resched = 1` used. If time demanding application would force `need_resched = 1` periodically then  $T_{S\_AV}$  should be enough to make task classification. There would be of course impact on whole system if scheduler would classify all tasks with the same priority at least for executing active, backup and again active queue. It should be enough to establish which process should be classified as time demanding. Unfortunately, this assumption could work when all of the processes would start the same time – which is bad assumption in the regular OS example.

Additional opportunity would be usage of average of  $T_{S\_AV}$  to eliminate sporadic activity of bigger priority tasks activity. Scheduler could be easily changed in order to save in

additional data structure  $K$  times  $T_{S\_AV}$  when given PID is executed on CPU:

$$\forall_{\text{schedule}(), \text{PID}} T_{AV}(K) = T_{S\_AV}.$$

Arithmetic average could easily be calculated on the basis of the data collected in created table:

$$T_{AAV} = \frac{\sum_{N=1}^K T_{AV}(N)}{K}.$$

This method could be successful as long as  $K$  would be estimated correctly and  $K$  average time calculation would be repeated couple of times. Additionally two average values would never be the same and there the range of error would need to be considered here as well:

$$1\_estimation : T_{AAV}(1)$$

$$2\_estimation : T_{AAV}(1) - X < T_{AAV}(2) < T_{AAV}(1) + X$$

$$n\_estimation : T_{AAV}(1) - X < T_{AAV}(n) < T_{AAV}(1) + X.$$

If it would be enough to classify that task matches time demanding characteristic after  $n = 2$  estimation, however probability that there is no mistake after  $n = 3$  estimation would be much bigger.

Scheduler could be designed to increase process priority about  $A$  when 2 estimation = TRUE and about  $B$  when 3 estimation = TRUE ( $B > A$ ).

## 4. Scheduler Improvement Measurement Results

Heuristic method in scheduler in kernel 2.6 assumes priority change about  $\pm 5$ . It is not too much, especially when several active kernel space tasks exist in active and expired queues. Figures 5 and 6 present changes in waiting for CPU time for 15 user space tasks with priorities from 100 to 114 and 100, 103, ..., 140.

Scheduler changes could provide more preemption than  $\pm 5$  change. Preemption patch [13] makes all tasks in-

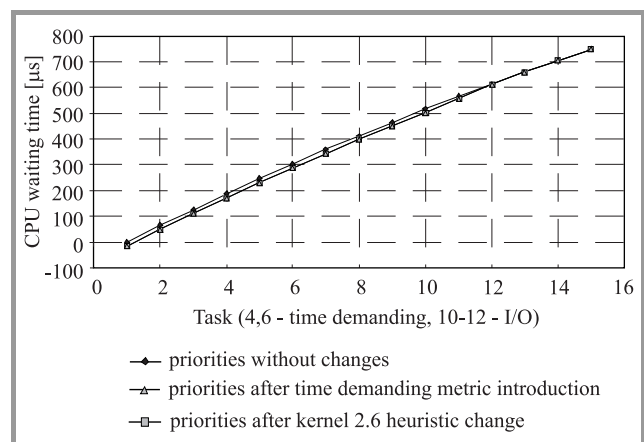


Fig. 5. Users space priority change – impact on waiting for CPU time – priorities 100–114.

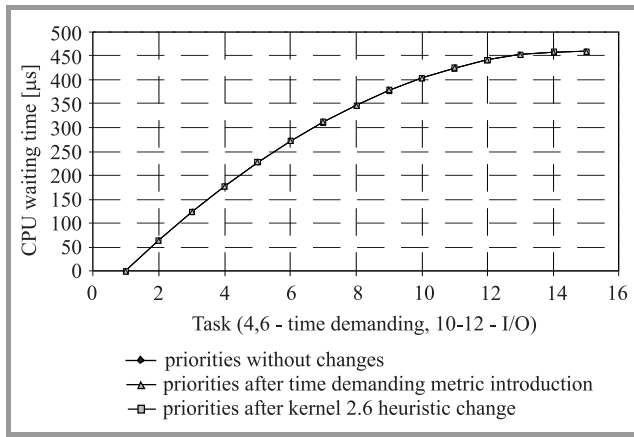


Fig. 6. Users space priority change – impact on waiting for CPU time – priorities 100, 103,..., 140.

cluding kernel soft-real-time available for priority change. For time demanding application executed in user space it would be more accurate to make opposite preemption and allow to change from user to kernel priority (from SCHED\_NORMAL to SCHED\_RR).

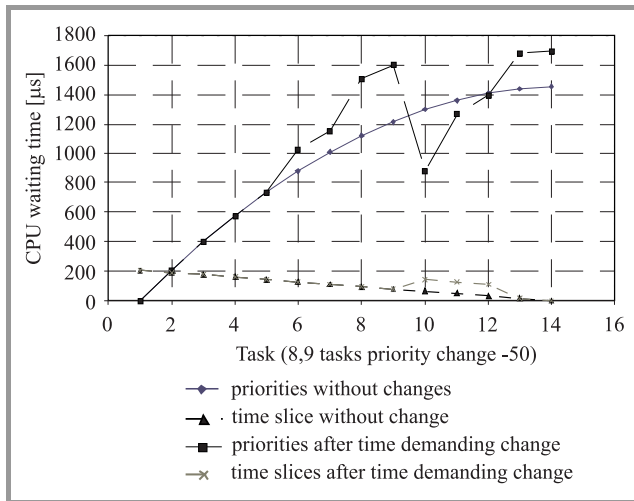


Fig. 7. Users space priority change – impact on waiting for CPU time and time slices – significant change -50.

Figure 7 presents total waiting for CPU and time slice values for priority change -50 (users pace moved to kernel).

**Scheduler metric efficiency experiment.** Time demanding processes classification metric bases on the average  $T_{S_{AV}}$  calculated after  $K$  measurement of  $T_{S_{AV}}$ . Average value is more valuable when it is calculated on the basis of more measurements. For scheduler it is not acceptable to make too many schedule() after priority change is done. For VRRP example if value  $K$  is determined incorrectly then scheduler could keep calculating which process should have priority changes while failover occurs. In the described metric method there is introduced value  $X$  which determines acceptable range to classify process request for CPU as periodic. Metric success basically depends on

correct  $X$  value, which should be not too big (to not classify accidental tasks) and not too small (to catch periodic nature even if there is major change in the queue for bigger priorities).

Table 1

Time  $T_{AAV}$  for different number of measurements

| $K$ | User1(120)         | User2(130)        | User3(134)        | User4(135)        |
|-----|--------------------|-------------------|-------------------|-------------------|
| 1   | 1352               | 1384              | 1400              | 1409.6            |
| 2   | 1420               | 1452              | 1468              | 1477.6            |
| 3   | <b>1431.466667</b> | <b>1463.46667</b> | <b>1479.46667</b> | <b>1489.06667</b> |
| 4   | 1431.6             | 1463.6            | 1479.6            | 1489.2            |
| 5   | <b>1420.48</b>     | <b>1452.48</b>    | <b>1468.48</b>    | <b>1478.08</b>    |
| 6   | 1416               | 1448              | 1464              | 1473.6            |
| 7   | 1406.857143        | 1438.85714        | 1454.85714        | 1464.45714        |
| 8   | 1405.4             | 1437.4            | 1453.4            | 1463              |
| 9   | 1398.577778        | 1432.35556        | 1447.64444        | 1457.06667        |
| 10  | 1399.52            | 1431.52           | 1447.52           | 1457.12           |

Table 1 and Fig. 8 describe possible average  $T_{S_{AV}}$  calculated for different  $K$ . This example consider only active queue with 10 kernel space tasks (priority 0–99) and 4 user space tasks (priority 100–139). Every schedule() CPU is given to the next process from active queue for a time slice calculated on the basis of priority. For every one from 10 experiments user space tasks in active queue have the same priorities while kernel can change to simulate difference of kernel tasks in a given amount of time.

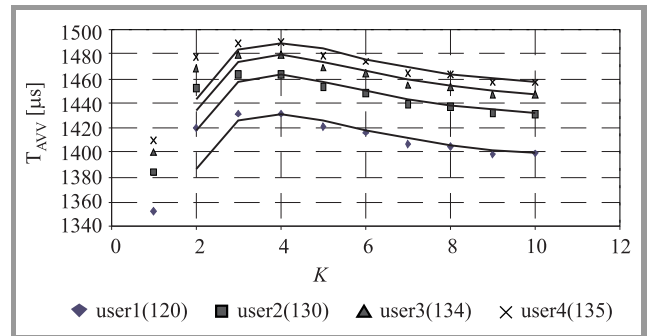


Fig. 8. Average  $T_{AVV}$  for user spaces processes and trend lines.

Measurement of  $K = 10$  active queue can provide information about average error described in Table 2.

For this example  $X = 80 \mu s$  would be valuable for time demanding process metric and would classify processes much better than  $X = 40$ .

Scheduler should be as well resistant to an average calculation errors. For that it can elect a process to increase

Table 2

Difference between average measurement for  $K = 10$

| User       | $T_{AAV} (1)$ | max $T_{AAV}$ | min $T_{AAV}$ | $X = \max - \min$ | $x = (\max - \min)/2$ |
|------------|---------------|---------------|---------------|-------------------|-----------------------|
| User1(120) | 1352          | 1431.6        | 1352          | 79.6              | 39.8                  |
| User2(130) | 1384          | 1463.6        | 1384          | 79.6              | 39.8                  |
| User3(134) | 1400          | 1479.6        | 1400          | 79.6              | 39.8                  |
| User4(135) | 1409.6        | 1489.2        | 1409.6        | 79.6              | 39.8                  |

```

X //average error range
C //allowed priority change for scheduler
PID //process ID
K //estimation
PRIO //process priority
TAVV[PID][K]

If schedule()
  K=K+1
  TAVV[PID][K] = TAVV
  //save average waiting time
  If (K==3)
    PRIO[PID]=PRIO[PID]+C
    //change priority
    Reschedule()
  end
  If (K==5)
    PRIO[PID]=PRIO[PID]+C*2
    //change priority
    Reschedule()
  end
  If (TAVV[PID][K] - TAVV[PID][K-1] > X)
    K=0
    //decline no time demanding processes
  end
  Reschedule()
end

```

Fig. 9. Scheduler estimation example pseudo code.

priority base on two values of  $K$ . Detailed algorithm is described in pseudo code on Fig. 9.

## 5. Summary

Time demanding user space application issue can be solved as many other computer science problems. To determinate if the cost of the solution is good enough to use it in the end user system a couple of numbers needs to be calculated together. Most important parts of the final grate would be the programming cost, improvement effect on real system, system stability after change.

Goal of this paper was to prove that such improvement in the kernel scheduler is possible and this or another idea can make time demanding user space application working more effective. According to measurement and calculation presented in previous section, Linux kernel scheduler can put more attention to the time demanding system activities. This can be done without breaking more important system rules. The scale of improvement depends on priority change level, which can be performed when process/task will be classified as time demanding. Presented solution shows as well that metric can depend on a set of additional parameters as classification range border or number of estimations. This leaves open door for system designers and developers and improvement, the base on several improvements can be parameterized for a dedicated system (e.g., ATCA SBC with a set of application running or ATCA line card with management application on it).

## Acknowledgment

Effort sponsored by the Ministry of Science and Higher Education, Poland, under grant PBZ-MNiSW-02-II/2007.

## References

- [1] ATCA – PICMG 3.0 R2.0: ECN 3.0-2.0-001 [Online]. Available: <http://www.picmg.org>
- [2] ATCA – Intelligent Platform Management Interface Specification Second Generation v2.0, Feb. 2006.
- [3] L. Null and J. Lobur, *The Essentials of Computer Organization and Architecture*. Sudbury: Jones & Bartlett Publ., 2006.
- [4] “Security Architecture for the Internet Protocol”, RFC 4301.
- [5] “Internet Key Exchange (IKEv2) Protocol”, RFC 4306.
- [6] Montavista Linux [Online]. Available: <http://www.montavista.com>
- [7] WindRiver [Online]. Available: <http://www.windriver.com>
- [8] Linux GPL [Online]. Available: <http://www.gnu.org>
- [9] “Virtual Router Redundancy Protocol (VRRP)”, RFC 3768.
- [10] “A Simple Network Management Protocol (SNMP)”, RFC 1157.
- [11] J. Aas, “Understanding the Linux 2.6.8.1 CPU scheduler”, SGI, 2005.
- [12] Linux kernel sources [Online]. Available: <http://kernel.org>
- [13] Linux kernel preemption project [Online]. Available: <http://kpreempt.sourceforge.net/>



**Marcin Hasse** received the M.Sc. degree in telecommunication from the Gdańsk University of Technology, Poland, in 2005. Currently he works for embedded computing leading company providing solutions for telecommunication market. His research interest and current work are related to operating system improvements for net-

working/telecommunication usage scenarios. He is an author of several publications in computer networking mechanisms improvements for end user services.

e-mail: [marcin@hasse.pl](mailto:marcin@hasse.pl)

Gdańsk University of Technology

G. Narutowicza st 11/12

80-952 Gdańsk, Poland



**Krzysztof Nowicki** received his M.Sc. and Ph.D. degrees in electronics and telecommunications from the Faculty of Electronics at the Gdańsk University of Technology, Poland, in 1979 and 1988, respectively. He is an author or co-author of more than 100 scientific papers and an author and co-author of five books. His scientific and re-

search interests include network architectures, analysis of communication systems, network security problems, modeling and performance analysis of cable and wireless communication systems, analysis and design of protocols for high speed LANs.

e-mail: [krzysztof.nowicki@eti.pg.gda.pl](mailto:krzysztof.nowicki@eti.pg.gda.pl)

Gdańsk University of Technology

G. Narutowicza st 11/12

80-952 Gdańsk, Poland

# Simple Dynamic Threshold Decryption Based on CRT and RSA

Bartosz Nakielski and Jacek Pomykała

**Abstract**—In the paper we present a simple threshold decryption system based on the RSA cryptosystem. Our model avoids the application of the Shamir secret sharing protocol and is based only on the Chinese remainder theorem. The flexibility in the threshold level is attained due to the suitable preparation of the input data. The second part of the article describes a modification of the basic model, which admits the sender's impact on the choice of the real receiver's group.

**Keywords**—CRT, RSA, threshold decryption.

## 1. Introduction

Threshold cryptography is one of the most important directions in modern cryptology. The basic idea is the division of the private key (used to decrypt or sign the messages) into shares, such that at least the given number of them (called the threshold level) is necessary to its reconstruction. This allows to distribute the trust or responsibility over the group members who are involved in the decryption or signing process, respectively. On the other hand, the distributed data or services allow to increase their availability, reliability or security.

The potential draw-back of the threshold cryptosystems (particularly in the encryption systems) is the lack of flexibility in the threshold level aspect. However in many applications the messages have different “priorities” or importance. In this connection we should require that some data should have higher threshold level than the others. In the classical approach this implies the necessity of generation of several polynomial threshold sharing protocols each responsible for the different threshold level. Recently there were some attempts to partially solve this problem (see, e.g., [1], [2], [3]) in the digital signature context.

As concerns the threshold decryption systems, very interesting solution was presented by H. Ghodosi, J. Pieprzyk, R. Safavi-Naini [4]. They apply the RSA (Rivest, Shamir, Adleman) cryptosystem [5] together with the Chinese remainder theorem (CRT) and the Shamir secret sharing protocol [6] to obtain the flexibility of the threshold level in dynamic group decryption process.

The Shamir protocol allowed the sender to share the “session key” among the members of the corresponding decryption group. In [7] the application of the above model for the databases systems was presented.

In this paper we were able to avoid the application of Shamir protocol completely, while still keeping the possibility to vary the threshold level together with the encrypted messages. In view of the additional “formatting”

conditions concerning the encrypted data, our model is based only on the RSA cryptosystem and Chinese remainder theorem.

## 2. Mathematical Background

### 2.1. Chinese Remainder Theorem

Given the pairwise coprime positive integers  $n_1, n_2, \dots, n_k$  and any integers  $a_1, a_2, \dots, a_k$  one can compute the integer  $a$  satisfying the following conditions:

$$a \equiv a_i \pmod{n_i} \text{ (for } i = 1, 2, \dots, k).$$

It can be obtained explicitly from the formula below:

$$a = \left( \sum_{i=1}^k a_i z_i y_i \right) \pmod{n},$$

where:

$$n = \prod_{i=1}^k n_i,$$

$$z_j = \frac{n}{n_j} = \prod_{i=1}^{j-1} n_i \times \prod_{i=j+1}^k n_i,$$

$$y_j = z_j^{-1} \pmod{n_j}.$$

### 2.2. RSA Cryptosystem

The RSA cryptosystem may be applied for the data encryption process as well as to the digital signatures. Its security is based on the factorization problem (for positive integers), which is believed to be computationally hard.

#### Parameters of RSA encryption scheme

Public key –  $(e, N)$  and private key  $d$ ,

$N = p \cdot q$  ( $p, q$  are prime numbers),

$\phi(N) = (p-1) \cdot (q-1)$ ,

$e$  – a number coprime to  $\phi(N)$ ,

$d$  – a number satisfying the condition:

$$e \cdot d \equiv 1 \pmod{\phi(N)}.$$

To encrypt the message  $m$  we compute the cryptogram  $c = m^e \pmod{N}$ . To decrypt the ciphertext  $c$  we compute the value  $m = c^d \pmod{N}$ .

More information concerning the RSA cryptosystems may be found in [5] and [8].

### 2.3. Dynamic threshold decryption

The idea of the threshold decryption cryptosystem is based on the splitting of the decryption key into several parts called the shares, which are applied for the reconstruction of the plaintext from the given ciphertext. The shares attached to the group members result in the fact that the decryption process has to be done collectively.

The typical threshold decryption systems use Lagrange interpolation formula to reconstruct the secret (being the free coefficient of the corresponding polynomial) from the shares being its values in positive integers. The degree of the polynomial defines the minimal number of shares needed to reconstruct the secret value. Thus the change of the threshold level causes the requirement of a new random polynomial to be generated and the corresponding shares to be distributed among the decryption group members. This makes the traditional approach to such systems completely impractical especially when the dynamic groups are considered. The solution proposed in [4] makes the sender responsible for the corresponding polynomial choice and pointing out the group of receivers of a given message, by means of some kind of “session keys”. As a result the flexibility of the threshold level (depending on the message) is admissible. Moreover the impact of the sender on the choice of the “decryption” group is achieved. The more general dynamic decryption group model admitting the distribution of the same shares among the distinct members (c.f. [9], [10]) or hierarchical model (c.f. [11]) could be also considered within the similar framework.

In this paper we present the simple threshold decryption protocol which avoids the application of the Shamir secret sharing protocol, still keeping the flexibility of the corresponding threshold level in the decryption process. It was possible due to the suitable preparation (representation) of the data, according to the assumed threshold level, before its encryption by the public keys of the decryption group members. In the decryption phase we use the corresponding private keys and the shares of the plaintext to reconstruct the original message.

## 3. Model

### 3.1. Notation

Let  $G = \{P_1, \dots, P_n\}$  be the group of users equipped with RSA keys  $(d_i, (e_i, N_i))$ , respectively.

Let us assume that the RSA moduli  $N_i$  are localized in the intervals:

$$(*) \quad N_i \in (2^{i-1}N_0, 2^iN_0) \text{ for } i = 1, 2, \dots, n$$

and  $K$  be a fixed number (security parameter) satisfying the inequality:

$$(**) \quad 2 \log_2 \log_2 N_0 < K < \frac{\log_2 N_0}{10}.$$

The communication among the group  $G$  goes through the group message board (*GMB*), where all the decrypted fragments of the message are published. The access to

the *GMB* requires RSA keys so only the members of  $G$  can read and write in *GMB*.

According to the application of CRT we denote:

$$N = \prod_{i=1}^n N_i,$$

$$N^j = \frac{N}{N_j} = \prod_{i=1}^{j-1} N_i \cdot \prod_{i=j+1}^n N_i,$$

$$Y_j = (N^j)^{-1} \bmod N.$$

### 3.2. Data Preparation

Let  $\lfloor x \rfloor$  stand for the largest integer not exceeding  $x$ , while  $\lceil x \rceil$  stand for the smallest integer not less than  $x$ . We define:

$$l_1 = l_1(t) = \lfloor \log_2 \prod_{i=n-t+2}^n N_i \rfloor,$$

$$l_2 = l_2(t) = \lfloor \log_2 \prod_{i=1}^t N_i \rfloor.$$

Assume for the moment that  $l_1 + 4K < l_2$  (see Lemma 1). By the CRT any message  $M$  of the length contained in the interval  $(l_1 + K, l_1 + 4K) \subset (l_1, l_2)$  is represented uniquely by the values of the residue classes:  $M_{ij} \bmod N_{ij}$  for  $j = 1, 2, \dots, t$ .

If  $M$  has the length  $k$  greater or equal to  $l_1 + 4K$  we can divide it on the suitable messages of length  $l_1 + K$  and at most one message of length less than  $l_1 + K$ . For the message  $M$  of length  $k < l_1 + K$  we shall apply the message padding procedure (described, e.g., in [8] in the framework of hash functions). Namely we require  $M$  to be represented by  $M_t$  according to the following steps:

1. Select a random number  $l \in (l_1 + 3K, l_1 + 4K)$ .
2. Add  $l - k - \lceil \log_2(l_1 + K) \rceil$  random bits to the message  $M$  (at the left side).
3. Add  $\lceil \log_2(l_1 + K) \rceil$  bits (at the right – hand side) denoting the length of the original message  $M$  (this part will contain a few zeros, since the length of  $M$  requires only  $\lceil \log_2 k \rceil$  bits).

After the padding phase the message  $M_t$  is built of three parts:

| Random bits                             | Message $M$ | Bits denoting the length of $M$ |
|-----------------------------------------|-------------|---------------------------------|
| $l - k - \lceil \log_2(l_1 + K) \rceil$ | $k$         | $\lceil \log_2(l_1 + K) \rceil$ |

**Lemma 1.** Let  $n \geq 3$ ,  $1 \leq t \leq n$  and  $N_0 > \max(2^{n^2}, 2^{500})$ . Assume that the RSA moduli satisfy the condition  $(*)$  (see Subsection 3.1) and let  $K$  satisfying  $(**)$  be fixed. Then  $l_1 + 4K < l_2$  and  $M_t$  contains at least  $K$  random bits.

**Proof.** By  $(*)$  and  $(**)$  we obtain that  $l_2 - l_1 \geq \log_2 N_0 - (t-1)(n-t+1) = \log_2 N_0 - \frac{n^2-1}{4} \geq \log_2 N_0 - \frac{n^2}{4}$ . Therefore by the lower bound for  $N_0$  and the upper bound for  $K$  we obtain that  $l_1 + 4K < l_2$ . Moreover, by the lower

bound (\*\*) for  $K$  the number of random bits in  $M_t$  is at least  $l - k - \log_2(l_1 + K) > l_1 + 3K - (l_1 + K) - K > K$  as required. ■

From the above we see that each fragment of  $M$  of length  $l_1 + K$  has the  $K$ -bit margin against the possible decryption of the message by any group of at most  $t - 1$  members of  $G$ . On the other hand, the fragment of  $M$  of length  $< l_1 + K$  has, after the padding procedure, the corresponding margin of size  $K$  according to minimum  $K$  random bits in  $M_t$ . Finally, let us also remark that the length of  $M_t$  is chosen randomly in the interval of length  $K$ .

## 4. Encryption and Decryption Algorithms

### 4.1. Encryption

To send the encrypted message  $M$  to the group  $G$  the following steps are performed by the sender:

1. Select the threshold level  $t$  ( $t \leq n$ ).
2. Create the  $M_t$  from  $M$  according to the description in Subsection 3.2.
3. Using the public keys belonging to group members compute 
$$c_i \equiv M_t^{e_i} \bmod N_i \text{ for } i = 1, 2, \dots, n.$$
4. With the aid of CRT compute  $C$  such that  $C \equiv c_i \bmod N_i$  for  $i = 1, 2, \dots, n$ .
5. Send the cryptogram  $(C, t)$  to the group  $G$ .

*Remark 1.* For the sake of complexity the values of  $N_i$  and  $Y_i$  (see Subsection 2.1) used in the Step 4 above should be precomputed and published for the users in advance.

### 4.2. Decryption

Decryption of the given ciphertext runs as follows:

1. Group members who decide to decrypt the message compute 
$$m_i = C^{d_i} \bmod N_i \text{ and publish the triple } ((C, t), m_i) \text{ on the } GMB.$$
2. When  $t$  triples occur on  $GMB$  any member can combine the fragments and reconstruct the plaintext  $M_t$  (and then the original message  $M$ ).

### 4.3. Modification

In the above model all the users have the same rights in the decryption process. However by the slight modification in the protocol the sender can have an impact on the choice of the members (say from some set  $B \subset G$ ) participating

in the decryption process. The modified protocol runs as follows:

1. Select the threshold level  $t$  ( $t \leq |B|$ ).
2. Create the  $M_t$  from  $M$  following the procedure described in Subsection 3.2.
3. Using the public keys belonging to group members compute 
$$c_i \equiv M_t^{e_i} \bmod N_i \text{ for } i \in B,$$
4. With the aid of CRT compute  $C$  such that  $C \equiv c_i \bmod N_i$  for  $i \in B$ .
5. Send the cryptogram  $(C, t)$  to the group  $B$ .

If the message is declared only for users of  $B$  (and nobody else should read it) they should distribute the corresponding values among the group  $B$  (instead of publishing them in  $GMB$ ).

## 5. Conclusion

In the paper we presented a threshold decryption protocol for the dynamic group based on the RSA cryptosystem and CRT, with the full flexibility of the threshold level. Our model avoids the application of the classical Shamir secret sharing protocol.

Instead we use some kind of formatting data technique which allows to replace Shamir protocol by the CRT method. The required ingredient was the application of the well known padding methodology in our context. The slight modification of the principal model admits the impact of the sender for the choice of the real receiver's group. The presented model can be extended for the more general framework of the threshold decryption systems (c.f. [10], [11]).

## Appendix – an Example

In order to simplify the description we present the suitable example omitting the data preparation phase in the protocol.

1. Let  $G = \{P_1, P_2, P_3\}$ .
2. Values of the private and public keys belonging to the members of  $G$ : 
$$\begin{array}{lll} e_1 = 17 & d_1 = 1289 & N_1 = 23 \cdot 167 = 3841 \\ e_2 = 11 & d_2 = 3459 & N_2 = 59 \cdot 83 = 4897 \\ e_3 = 13 & d_3 = 4501 & N_3 = 47 \cdot 107 = 5029 \end{array}$$
3. The user  $S$  encrypts and sends the message  $M = 452009$ , with the threshold level  $t = 2$  for the group  $G$ .
4. Encryption: 
$$\begin{array}{l} c_1 = 452009^{17} \bmod 3841 = 2053 \\ c_2 = 452009^{11} \bmod 4897 = 2197 \\ c_3 = 452009^{13} \bmod 5029 = 3845 \end{array}$$

Parameters of the CRT:

$$N^1 = 4897 \cdot 5029 = 24627013$$

$$N^2 = 3841 \cdot 5029 = 19316389$$

$$N^3 = 3841 \cdot 4897 = 18809377$$

$$(N^1)^{-1} \bmod N_1 = 24627013^{-1} \bmod 3841 = 174$$

$$(N^2)^{-1} \bmod N_2 = 19316389^{-1} \bmod 4897 = 1973$$

$$(N^3)^{-1} \bmod N_3 = 18809377^{-1} \bmod 5029 = 2775$$

$$C = (2053 \cdot 24627013 \cdot 174 + 2197 \cdot 19316389 \cdot 1973 + 3845 \cdot 18809377 \cdot 2775) \bmod 94592356933 = 79682507303$$

5. The user  $S$  sends the ciphertext  $(79682507303, 2)$  to the group  $G$ .

6. Users  $P_1$  and  $P_2$  decrypt the message:

$$m_1 = C^{d_1} \bmod N_1 = 79682507303^{1289} \bmod 3841 = 2612$$

$$m_2 = C^{d_2} \bmod N_2 = 79682507303^{3459} \bmod 4897 = 1485$$

$$\begin{aligned} M_t &= 2612 \cdot 4897 \cdot (4897^{-1} \bmod 3841) \\ &\quad + 1485 \cdot 3841 \cdot (3841^{-1} \bmod 4897) \\ &= (2612 \cdot 4897 \cdot 3139 \\ &\quad + 1485 \cdot 3841 \cdot 895) \bmod 18809377 \\ &= 452009 \end{aligned}$$

## References

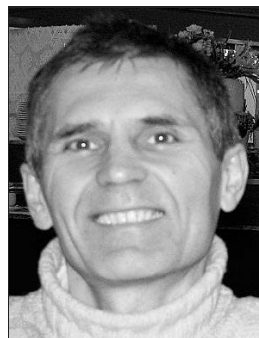
- [1] B. Nakielski, J. Pomykała, and J. A. Pomykała, "Multi-threshold signature", *J. Telecommun. Inform. Technol.*, no. 1, pp. 51–55, 2008.
- [2] J. Pomykała and T. Warchoń, "Threshold signatures in dynamic groups", in *Proc. Fut. Gener. Commun. Netw. 2007*, Jeju-Island, Korea, 2007, pp. 32–37.
- [3] J. Pomykała and T. Warchoń, "Dynamic multi-threshold signatures without the trusted dealer", *Int. J. Multimed. Ubiquit. Eng.*, vol. 3, pp. 31–42, July 2008.
- [4] H. Ghodosi, J. Pieprzyk, and R. Safavi-Naini, "Dynamic threshold cryptosystems: a new scheme in group oriented cryptography", in *Proc. Pragocrypt'96*, Prague, Czech Republic, 1996, pp. 370–379.
- [5] R. L. Rivest, A. Shamir, and L. M. Adleman, "A method for obtaining digital signatures and public-key cryptosystems", *Commun. ACM*, vol. 21, no. 2, pp. 120–126, 1978.
- [6] A. Shamir, "How to share a secret", *Commun. ACM*, vol. 22, pp. 612–613, 1979.
- [7] B. Nakielski, J. Pomykała, and J. A. Pomykała, "Wykorzystanie deszyfrowania progowego w bazach danych", *Biul. WAT*, vol. LVII, no. 4, pp. 183–196, 2008 (in Polish).
- [8] B. Schneier, *Applied Cryptography*. New York: Wiley, 1996.
- [9] R. Di Pietro, L. V. Mancini, and G. Zanin, *Efficient and Adaptive Threshold Signatures for Ad-Hoc Networks*. Electronic Notes in Theoretical Computer Science. Amsterdam: Elsevier, 2007, vol. 171, pp. 93–105.
- [10] J. Pomykała and B. Żrałek, "Threshold flexible signature scheme in dynamic groups", in *Proc. ACS Conf.*, Międzyzdroje, Poland, 2008.
- [11] T. Tassa, "Hierarchical threshold secret sharing", *J. Cryptol.*, vol. 20, pp. 237–264, 2007.



**Bartosz Nakielski** was born in Warsaw, Poland, in 1979. He received his M.Sc. in mathematics from Department of Mathematics, Mechanics and Informatics of Warsaw University. His thesis was titled "Arithmetical aspects of digital signatures". He was working as a certificate authority administrator in Information Security Department

in Social Insurance Institution (years 2004–2007). Since January 2008 he works in Security Department in National Bank of Poland.

e-mail: barteknakielski@aster.pl



**Jacek Pomykała** works at the Faculty of Mathematics, Informatics and Mechanics of Warsaw University, Poland, since 1985. He received the Ph.D. and D.Sc. degrees in 1986 and 1997, respectively. He has published over 20 papers mainly in mathematical journals. He is also one of the authors of the book concerning the information systems and cryptography.

He had many research visits (two long term visits) in Europe, America, Asia and has been an invited speaker in many international conferences in mathematics and computer science.

e-mail: pomykala@mimuw.edu.pl

Institute of Mathematics

Faculty of Mathematics, Informatics and Mechanics

University of Warsaw

Banacha st 2

02-097 Warsaw, Poland

# Generating Pseudorandom S-Boxes – a Method of Improving the Security of Cryptosystems Based on Block Ciphers

Piotr Mroczkowski

**Abstract**—The paper presents a general framework for improving the security of the cryptosystem based on the symmetric block cipher. The main idea is based on possibility of changing substitution boxes (called S-boxes) in encryption/decryption algorithm. In order to make it possible, it is necessary to generate identical boxes by an encryption and decryption party. This is the main reason, why deterministic methods of generating substitution boxes based on the pseudorandom sequences will be presented.

**Keywords**—block cipher, cryptosystem, permutation box (P-box), (pseudo)random bit generator, substitution box (S-box).

## 1. Introduction

The confidentiality of the information, transmitted via contemporary telecommunications systems, is ensured by cryptographic devices. In such devices a cryptosystem is implemented, which provides confidentiality using cryptographic ciphers. The cryptosystem (Fig. 1) is defined as a quintuple  $(PT, CT, K, E_K, D_K)$ , where:  $PT$  - plaintext,  $CT$  - ciphertext,  $K$  - key,  $E_K$  - encryption algorithm,  $D_K$  - decryption algorithm, such that:

$$CT = E_K(PT),$$

$$PT = D_K(CT) = D_K(E_K(PT)).$$

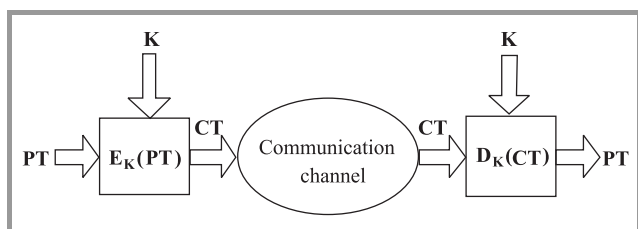


Fig. 1. Cryptosystem.

As an encryption/decryption algorithm the block cipher can be used. A lot of modern block ciphers are built using Shannon's concept [1], which uses two basic transformations: confusion and diffusion. Such product cipher uses S-boxes that provide confusions and P-boxes that provide diffusions and spread out the output bits to different

S-boxes of the next round. Simple substitution and transposition transformations individually do not provide a very high level of security. However, by combining these transformations it is possible to obtain strong product cipher, which guarantees resistance against linear and differential cryptanalysis. The strength of such cipher mainly comes from the property of S-boxes. It is well known that weaknesses of substitution boxes may be compensated for by the increased number of rounds. Another way to improve security of such cryptosystem is possibility of changing S-boxes in the product cipher.

The S-box  $S: Z_2^n \rightarrow Z_2^m$  is a nonlinear transformation which transforms the  $n$ -binary sequence to the  $m$ -binary sequence. They can be constant (generated at the stage of specification and used during the lifetime of algorithm, e.g., Data Encryption Standard, Advanced Encryption Standard) or variable (generated before a session on the basis of data, which are available to the encryption/decryption party and used during the session). The idea of improving the security of cryptosystem based on the replacement S-boxes in the encryption/decryption algorithm by new generated S-boxes. This causes, that for the same plaintext and main key, we obtain different ciphertext. We can say that we have another block cipher. The changing of S-boxes prevents from receiving enough information to execute the effective cryptanalytical attack.

So the main question is how to change S-boxes without delivering them to the cryptographic devices. The answer is really simple. The encryption/decryption party has to generate identical substitution boxes, which will be used in the encryption/decryption process. Therefore deterministic methods of generating S-boxes using pseudorandom sequences are proposed. The encryption and decryption party has to have the same parameters of the generation process, for example, the seed and the algebraic module, so that it can generate identical sequences. Using these sequences and presented in Section 2 methods, it is possible to generate identical S-boxes, which can be change in the block cipher.

S-boxes, which can be applied in the block cipher, should fulfill the cryptographical criterion like balancedness, non-linearity, strict avalanche criterion. The article proposes two algorithms of generation good cryptographically S-boxes using pseudorandom sequences. The generated S-boxes

will be checked towards their nonlinearity (called nl) and distance to strict avalanche criterion (called dSAC).

## 2. The S-Box Generation Method Using Binary Sequence

The analysis of opportunities of replacement substitution boxes in block ciphers was preceded by analysis of generation them using binary sequence. With this view two constructions are proposed which allow generating truth table of the S-box  $S: Z_2^n \rightarrow Z_2^m$  from binary sequence [2].

### Construction 1

Let  $Z = (z_0, z_1, z_2, \dots, z_{r-1})$ ,  $z_i \in Z_2$ , be a binary sequence. Using this sequence we construct  $m$ -binary vectors:  $\underline{z}_0 = (z_0, \dots, z_{m-1})$ ,  $\underline{z}_1 = (z_m, \dots, z_{2m-1})$ ,  $\dots$ ,  $\underline{z}_k = (z_{km}, \dots, z_{(k+1)m-1})$ ,  $\underline{z}_i \in Z_2^m$ . The first nonzero vector  $\underline{z}_l = (z_{lm}, \dots, z_{(l+1)m-1})$  is the first row in the generated truth table. The next vectors  $\underline{z}_{l+1}, \underline{z}_{l+2}, \dots$  are checked if they are different from the vectors in the truth table. In the case of the positive verification such a vector is the next row in truth table, otherwise the vector is omitted and the next vector is checked. The process is continued until the truth table is full.

---

### Algorithm 1: Algorithm of S-box $S: Z_2^n \rightarrow Z_2^m$ generation using Construction 1

---

**Input:**  $Z = (z_0, z_1, z_2, \dots, z_{r-1})$ ,  
 $z_i \in Z_2$  – the binary sequence,  
**Output:**  $Sb = [s_0, s_1, \dots, s_{2^n-1}]$  – the S-box table.  
 $l := 0$ ;  
**repeat**  
     $Sb[0] := 0$ ;  
    **for**  $j := 0$  to  $m-1$  **do**  
         $Sb[0] := Sb[0] + z_l \cdot 2^j$ ;  
     $l := l + 1$ ;  
**until**  $Sb[0] \neq 0$ ;  
**for**  $i = 1$  to  $2^n - 1$  **do**  
    **repeat**  
         $Sb[i] := 0$ ;  
        **for**  $j := 0$  to  $m-1$  **do**  
             $Sb[i] := Sb[i] + z_l \cdot 2^j$ ;  
         $l := l + 1$ ;  
        **if**  $(Sb[i] \in \{s_0 \dots s_{i-1}\})$  **then**  $spr := 1$ ,  
        otherwise  $spr := 0$ ;  
    **until**  $spr = 0$ ;

---

### Construction 2

Let:  $Z^0 = (z_0^0, z_1^0, \dots, z_{r-1}^0)$ ,  $z_i^0 \in Z_2$ ,  
 $Z^1 = (z_0^1, z_1^1, \dots, z_{r-1}^1)$ ,  $z_i^1 \in Z_2, \dots$ ,  
 $Z^{m-1} = (z_0^{m-1}, z_1^{m-1}, \dots, z_{r-1}^{m-1})$ ,  $z_i^{m-1} \in Z_2$ ,  
be  $m$  random sequences.

Using them we construct  $m$ -binary vectors in the following way:  $\underline{z}_0 = (z_0^0, z_1^0, \dots, z_{r-1}^{m-1})$ ,  $\underline{z}_1 = (z_1^0, z_1^1, \dots, z_1^{m-1})$ ,  $\dots$ ,  $\underline{z}_k = (z_k^0, z_k^1, \dots, z_k^{m-1})$ ,  $\underline{z}_i \in Z_2^m$ . The first nonzero vector  $\underline{z}_l = (z_l^0, z_l^1, \dots, z_l^{m-1})$  is the first row in the generated truth table. The next vectors  $\underline{z}_{l+1}, \underline{z}_{l+2}, \dots$  are checked if they are different from the vectors in the truth table. In the case of the positive verification such vector is the next row in the truth table, otherwise the vector is omitted and the next vector is checked. The process is continued until the truth table is full.

---

### Algorithm 2: Algorithm of S-box $S: Z_2^n \rightarrow Z_2^m$ generation using Construction 2

---

**Input:**  $Z^0 = (z_0^0, z_1^0, \dots, z_{r-1}^0)$ ,  
 $Z^1 = (z_1^1, z_1^1, \dots, z_{r-1}^1), \dots$ ,  
 $Z^{m-1} = (z_0^{m-1}, z_1^{m-1}, \dots, z_{r-1}^{m-1})$ ,  
 $z_i^j \in Z_2$  – binary sequences,  
**Output:**  $Sb = [s_0, s_1, \dots, s_{2^n-1}]$  – the S-box table.  
 $l := 0$ ;  
**repeat**  
     $Sb[0] := 0$ ;  
    **for**  $j := 0$  to  $m-1$  **do**  
         $Sb[0] := Sb[0] + z_l^j \cdot 2^j$ ;  
     $l := l + 1$ ;  
**until**  $Sb[0] \neq 0$ ;  
**for**  $i = 1$  to  $2^n - 1$  **do**  
    **repeat**  
         $Sb[i] := 0$ ;  
        **for**  $j := 0$  to  $m-1$  **do**  
             $Sb[i] := Sb[i] + z_l^j \cdot 2^j$ ;  
         $l := l + 1$ ;  
        **if**  $(Sb[i] \in \{s_0 \dots s_{i-1}\})$  **then**  $spr := 1$ ,  
        otherwise  $spr := 0$ ;  
    **until**  $spr = 0$ ;

---

The presented above methods were implemented. Generated S-boxes were verified towards their nonlinearity and strict avalanche criterion.

## 3. Generation S-Boxes Using Random Binary Sequence

The random numbers play a crucial part in cryptography. They are used in many cryptographical applications and devices and may be generated exclusively by hardware binary sequence generator, for example SGCL-1 [3], which was designed in Military Communication Institute. Using this generator and proposed method 1000000 S-boxes  $S: Z_2^8 \rightarrow Z_2^8$  were generated. They were tested towards their nonlinearity and strict avalanche criterion. The number of

S-boxes depending on nonlinearity and distance to SAC was shown in Tables 1 and 2, and was illustrated in Figs. 2–5.

Table 1  
The results of the nonlinearity test of S-boxes generated using random binary sequence

| S-box                         | nl(S) | Construction 1 | Construction 2 |
|-------------------------------|-------|----------------|----------------|
| $S : Z_2^8 \rightarrow Z_2^8$ | 70    | 0              | 1              |
|                               | 72    | 1              | 0              |
|                               | 74    | 1              | 0              |
|                               | 76    | 4              | 3              |
|                               | 78    | 25             | 25             |
|                               | 80    | 153            | 175            |
|                               | 82    | 699            | 666            |
|                               | 84    | 2849           | 2874           |
|                               | 86    | 11420          | 11380          |
|                               | 88    | 41955          | 41527          |
|                               | 90    | 132211         | 131664         |
|                               | 92    | 313810         | 314461         |
|                               | 94    | 385570         | 386342         |
|                               | 96    | 109767         | 109350         |
|                               | 98    | 1535           | 1532           |

Table 2  
The results of the distance to SAC test of S-boxes generated using random binary sequence

| S-box                         | dSAC(S) | Construction 1 | Construction 2 |
|-------------------------------|---------|----------------|----------------|
| $S : Z_2^8 \rightarrow Z_2^8$ | 32      | 6359           | 6303           |
|                               | 36      | 228686         | 228574         |
|                               | 40      | 459965         | 460667         |
|                               | 44      | 227993         | 227715         |
|                               | 48      | 61804          | 61624          |
|                               | 52      | 12644          | 12608          |
|                               | 56      | 2190           | 2144           |
|                               | 60      | 307            | 318            |
|                               | 64      | 44             | 41             |
|                               | 68      | 6              | 5              |
|                               | 72      | 1              | 1              |
|                               | 76      | 1              | 0              |

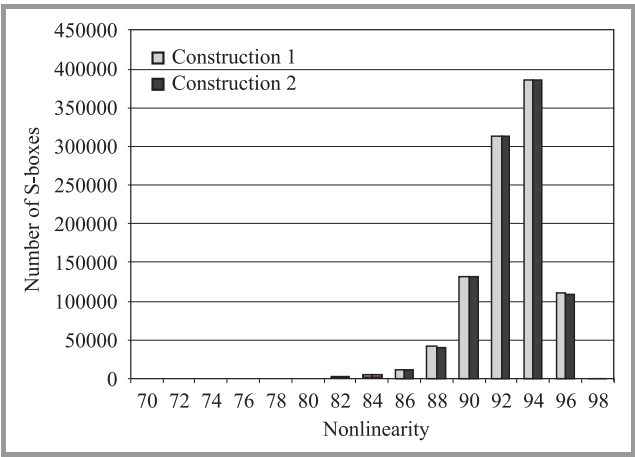


Fig. 2. The number of S-boxes generated using random binary sequence depending on the nonlinearity.

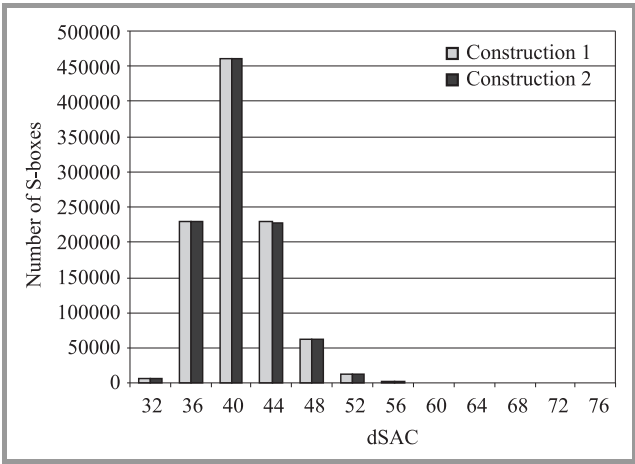


Fig. 3. The number of S-boxes generated using random binary sequence depending on the distance to SAC.

The maximum value of the nonlinearity is 98 and it is achieved in 0.1535% (Construction 1) and 0.1532% (Construction 2) generated S-boxes. The nonlinearity, which accomplishes the maximum number of S-boxes is equal 94 and it is achieved in 38.5570% (Construction 1) and 38.6342 (Construction 2) S-boxes.

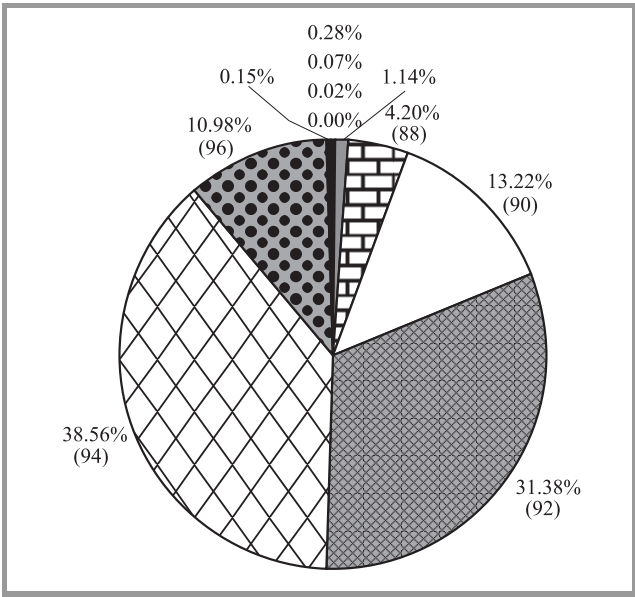
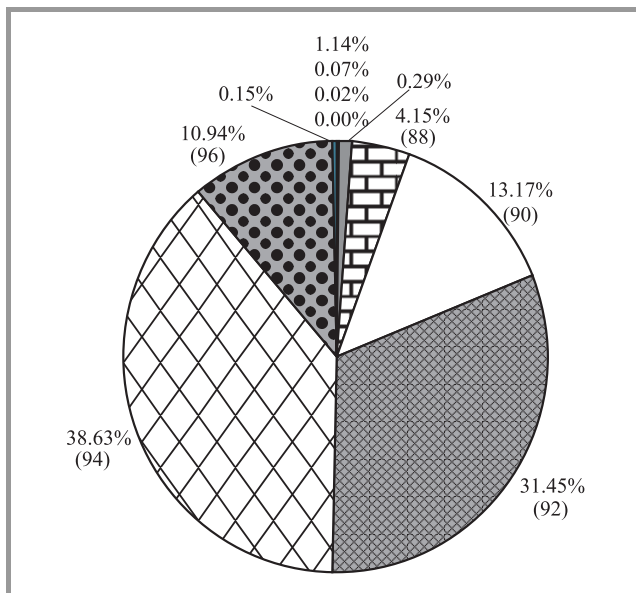


Fig. 4. The percent estimation of the number of S-boxes generated using random binary sequence depending on the nonlinearity (Construction 1). The value in parentheses specifies the nonlinearity of S-boxes.

The minimum value of the dSAC is 32 and it is achieved in 0.6359% (Construction 1) and 0.6303% (Construction 2) generated S-boxes. The dSAC, which accomplishes the maximum number of S-boxes is equal 40 and it is achieved in 45.9965% (Construction 1) and 46.0667 (Construction 2) S-boxes.

If we assume that “good” S-box should have nonlinearity at the level of 90 and dSAC at more than 48 then the probability, that generated S-box will satisfy above-mentioned



**Fig. 5.** The percent estimation of the number of S-boxes generated using random binary sequence depending on the nonlinearity (Construction 2). The value in parentheses specifies the nonlinearity of S-boxes.

assumptions, is very high and is equal 0.9286 for both constructions.

#### 4. Generation S-Boxes Using Pseudorandom Bit Sequence

The pseudorandom numbers play an important part in cryptography, too. They are used in many cryptographic applications and devices and may be generated by pseudorandom generators, for example, Legendre's generator [4], [5], BBS generator [4], inverse generator [4] and so on. The pseudorandom sequence generators are composed as linear or nonlinear. In cryptography nonlinear pseudorandom generators are used because of characteristics like:

- deterministic method of generating sequences – this feature allows to generate the same sequence by independent computation party;
- improvement of quality generated sequences;
- nonlinearity of the generated sequence – very important feature, as S-boxes should have high nonlinearity, so they should be generated using nonlinear sequences.

The Legendre's generator is defined in the following way. Let  $p$  be an odd prime number then  $X_n = 1$  if  $n$  is quadratic residue modulo  $p$  and  $X_n = 0$  otherwise, where  $n$  takes following natural value. The Legendre's sequences satisfy Golomb's postulates [6].

Using Legendre's generator and proposed constructions 1000000 S-boxes  $S : Z_2^8 \rightarrow Z_2^8$  were generated. They were tested towards their nonlinearity and strict avalanche criterion.

The number of S-boxes depending on nonlinearity and distance to SAC was showed in Tables 3 and 4, and illustrated in Figs. 6–9.

Table 3

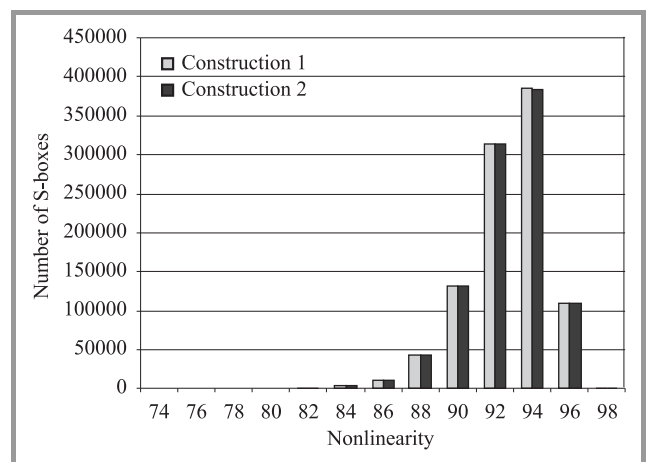
The results of the nonlinearity test of S-boxes generated using pseudorandom Legendre's sequence

| S-box                         | nl(S) | Construction 1 | Construction 2 |
|-------------------------------|-------|----------------|----------------|
| $S : Z_2^8 \rightarrow Z_2^8$ | 74    | 2              | 0              |
|                               | 76    | 6              | 3              |
|                               | 78    | 28             | 39             |
|                               | 80    | 156            | 143            |
|                               | 82    | 666            | 643            |
|                               | 84    | 2816           | 2942           |
|                               | 86    | 11295          | 11430          |
|                               | 88    | 41727          | 41761          |
|                               | 90    | 131692         | 132081         |
|                               | 92    | 314333         | 314254         |
|                               | 94    | 385979         | 384790         |
|                               | 96    | 109769         | 110420         |
|                               | 98    | 1531           | 1494           |

Table 4

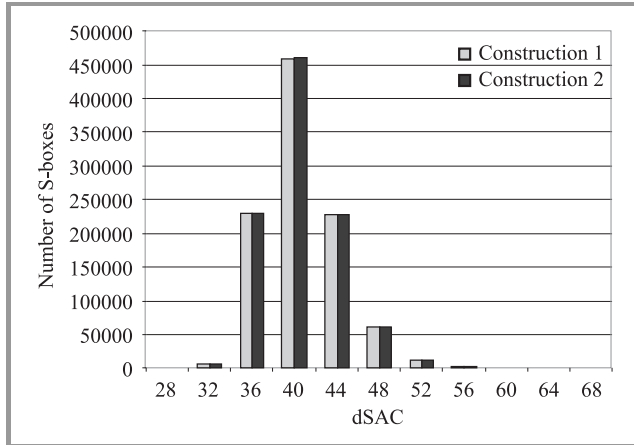
The results of the distance to SAC test of S-boxes generated using pseudorandom Legendre's sequence

| S-box                         | dSAC(S) | Construction 1 | Construction 2 |
|-------------------------------|---------|----------------|----------------|
| $S : Z_2^8 \rightarrow Z_2^8$ | 28      | 1              | 1              |
|                               | 32      | 6254           | 6169           |
|                               | 36      | 229105         | 228828         |
|                               | 40      | 459813         | 460664         |
|                               | 44      | 228003         | 228002         |
|                               | 48      | 61710          | 61344          |
|                               | 52      | 12628          | 12547          |
|                               | 56      | 2095           | 2079           |
|                               | 60      | 341            | 315            |
|                               | 64      | 42             | 48             |
|                               | 68      | 8              | 3              |



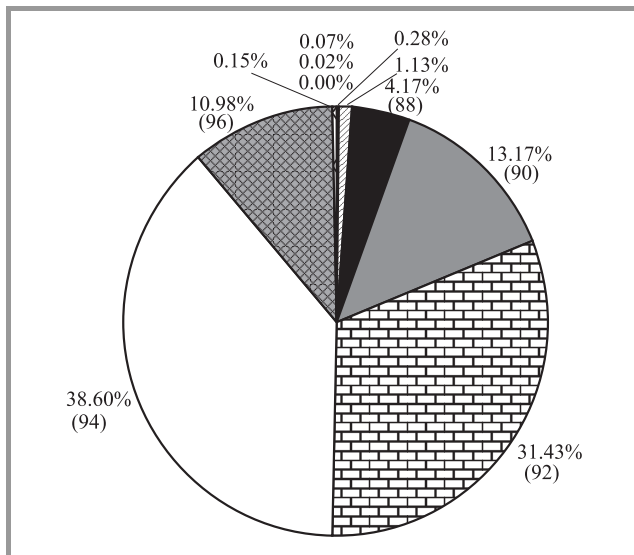
**Fig. 6.** The number of S-boxes generated using pseudorandom Legendre's sequence depending on the nonlinearity.

The maximum value of the nonlinearity is 98 and it is achieved in 0.1531% (Construction 1) and 0.1494% (Construction 2) generated S-boxes. The nonlinearity, which accomplishes the maximum number of S-boxes is equal 94 and achieves it 38.5979% (Construction 1) and 38.4790% (Construction 2) S-boxes.



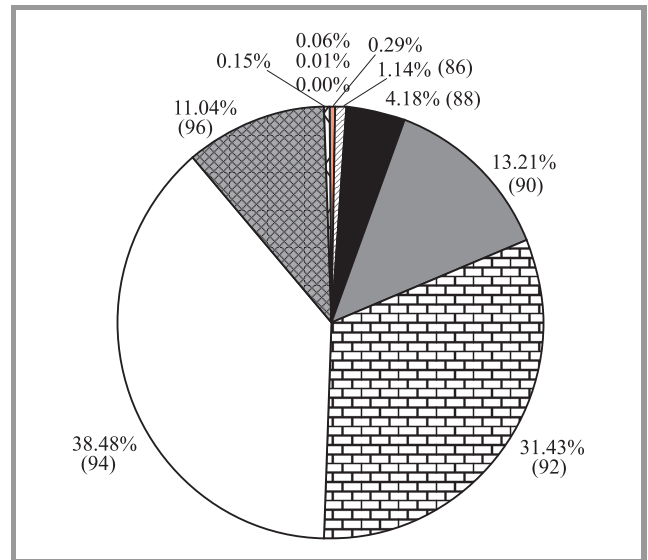
**Fig. 7.** The number of S-boxes generated using pseudorandom Legendre's sequence depending on the distance to SAC.

The minimum value of the distance to the strict avalanche criterion is 28 and it is achieved in 0.0001% (both constructions) generated S-boxes. The dSAC, which accomplishes the maximum number of S-boxes is equal 40 and it is achieved in 45.9813% (Construction 1) and 46.0664% (Construction 2) S-boxes.



**Fig. 8.** The percent estimation of the number of S-boxes generated using pseudorandom Legendre's sequence depending on the nonlinearity (Construction 1). The value in parentheses specifies the nonlinearity of S-boxes.

If we assume that "good" S-box should have nonlinearity at least the level of 90 and dSAC at more than 48 then the probability, that generated S-box will satisfy above-men-



**Fig. 9.** The percent estimation of the number of S-boxes generated using pseudorandom Legendre's sequence depending on the nonlinearity (Construction 2). The value in parentheses specifies the nonlinearity of S-boxes.

tioned assumptions, is very high and is equal 0.9290 (Construction 1) and 0.9289 (Construction 2).

## 5. The Method of Improving the Security of the Cryptosystem Based on the Block Cipher

The idea of improving the security of cryptosystem based on block cipher depends on replacement of substitution boxes in the encryption/decryption algorithm. So that, it would be possible, the encryption and decryption party have to generate identical S-boxes for the cryptographic session. These possibilities are given by the method of generating S-boxes from pseudorandom sequences, because the encryption/decryption party can generate identical bits streams and as a result thanks to proposed method – identical S-boxes. With this view we can apply the following method based on the Diffie-Hellman protocol [6] to arrange the seed:

- 1) the encryption/decryption party has to have identical numbers  $g, m$ ;
- 2) the encryption party chooses a number:  $1 < x < m - 1$  and calculates:  $X = g^x \mod m$ ;
- 3) the decryption party chooses a number:  $1 < y < m - 1$  and calculates:  $Y = g^y \mod m$ ;
- 4) the encryption/decryption party exchanges numbers  $X$  and  $Y$ ;
- 5) the encryption/decryption party calculates the seed:

$$s_{encr} = Y^x \mod m;$$

$$s_{decr} = X^y \mod m,$$

$$s = s_{encr} = s_{decr} = g^{xy} \mod m.$$

The encryption/decryption party using the seed  $s$  and module  $m$  generates Legendre's sequence and then substitution boxes. The generated S-boxes should characterize high nonlinearity (at least  $nl_{\min}$ ) and low dSAC (at most  $dSAC_{\max}$ ), so the algorithm of generation of such S-boxes is as follow.

---

**Algorithm 3: Algorithm of S-box table generation**


---

**Input:**  $nl_{\min}$  – the minimal value of the nonlinearity of the generated S-boxes;  
 $dSAC_{\max}$  – the maximal value of the dSAC of the generated S-boxes;

**Output:**  $TSB = [SB_0, SB_1, \dots, SB_{r-1}]$  – the table of S-boxes.

**for**  $i := 0$  **to**  $r - 1$  **do**  
  **repeat**  
     $SB_i := SbGen$ ;  
    //  $SbGen$  – a S-box generation algorithm using Constr. 1 or Constr. 2;  
     $nl := Licz\_NL(SB_i)$ ;  
    //  $Licz\_NL$  – a nonlinearity of the S-box calculation algorithm;  
     $dsac := Licz\_dSAC(SB_i)$ ;  
    //  $Licz\_dSAC$  – an dSAC of the S-box calculation algorithm;  
  **until** ( $nl \geq nl_{\min}$  and  $sac \leq dSAC_{\max}$ );

---

The exchanging of the S-boxes in block ciphers causes that a different block cipher is applied in the cryptosystem. This makes impossible to collect sufficient amount of information to carry back the cryptanalytical attack and in consequence raises the security of cryptosystem.

## 6. Conclusions

The sequences generated by the Legendre's generator satisfy statistical tests described in Federal Information Processing Standards Publications (FIPS 140-1), have statistical characteristic like random sequences and are strongly nonlinear. For this reason they carry out experiments of

making substitution boxes using the proposed method and Legendre's sequences. The researches of nonlinearity and dSAC of generated S-boxes show that, giving the maximum nonlinearity and minimum dSAC up, it is possible to generate "good" S-boxes. It gives the possibility of making S-boxes used in block ciphers and makes them replaceable. This treatment secures cryptosystems against many cryptographical attacks, especially differential and algebraic cryptanalysis.

## References

- [1] C. E. Shannon, "Communication theory of secrecy system", *Bell Syst. Techn. J.*, no. 28, pp. 656–715, 1949.
- [2] P. Mroczkowski and A. Paszkiewicz, "About the designing method of strong cryptographically boolean functions", in *Proc. XI KKKiOI ENIGMA 2007 Conf.*, Warsaw, Poland, 2007.
- [3] M. Leśniewicz, "Sprzętowy generator ciągów losowych do zastosowań kryptograficznych", in *Proc. Symp. XVII KST 2001*, Bydgoszcz, Poland, 2001 (in Polish).
- [4] T. W. Cusik, C. Ding, and A. Renevall, *Stream Ciphers and Number Theory*. North Holland: Mathematical Library, 1998.
- [5] I. D. Damagard, "On the randomness of Legendre and Jacobi sequences", in *Advanced in Cryptology – Eurocrypt'88*. Berlin: Springer-Verlag, 1988.
- [6] A. Menezes, P. van Oorschot, and S. Vanstone, *Handbook of Applied Cryptography*. Boca Raton: CRC Press, 1997.



**Piotr Mroczkowski** received the M.Sc. degree in 2000 and the Ph.D. degree in 2008 from the Military University of Technology in Warsaw, Poland. He is an Assistant Professor at the Military Communication Institute. He is interested in cryptology and computer science.

e-mail: p.mroczkowski@wil.waw.pl  
 Military Communication Institute  
 Warszawska st 22A  
 05-130 Zegrze, Poland

# Information for Authors

**Journal of Telecommunications and Information Technology (JTIT)** is published quarterly. It comprises original contributions, dealing with a wide range of topics related to telecommunications and information technology. **All papers are subject to peer review.** Topics presented in the JTIT report primary and/or experimental research results, which advance the base of scientific and technological knowledge about telecommunications and information technology.

JTIT is dedicated to publishing research results which advance the level of current research or add to the understanding of problems related to modulation and signal design, wireless communications, optical communications and photonic systems, voice communications devices, image and signal processing, transmission systems, network architecture, coding and communication theory, as well as information technology.

Suitable research-related papers should hold the potential to advance the technological base of telecommunications and information technology. Tutorial and review papers are published only by invitation.

**Manuscript.** TEX and LATEX are preferable, standard Microsoft Word format (.doc) is acceptable. The author's JTIT LATEX style file is available:

<http://www.itl.waw.pl/publ/jtit/doc/jtitaut.zip>

Papers published should contain up to 10 printed pages in LATEX author's style (Word processor one printed page corresponds approximately to 6000 characters).

The manuscript should include an abstract about 150–200 words long and the relevant keywords. The abstract should contain statement of the problem, assumptions and methodology, results and conclusion or discussion on the importance of the results. Abstracts must not include mathematical expressions or bibliographic references.

Keywords should not repeat the title of the manuscript. About four keywords or phrases in alphabetical order should be used, separated by commas.

The original files accompanied with pdf file should be submitted by e-mail: [redakcja@itl.waw.pl](mailto:redakcja@itl.waw.pl)

**Figures, tables and photographs.** Original figures should be submitted. Drawings in Corel Draw and PostScript formats are preferred. Figure captions should be placed below the figures and can not be included as a part of the figure. Each figure should be submitted as a separated graphic file, in .cdr, .eps, .ps, .png or .tif format. Tables and figures should be numbered consecutively with Arabic numerals.

Each photograph with minimum 300 dpi resolution should be delivered in electronic formats (TIFF, JPG or PNG) as a separated file.

**References.** All references should be marked in the text by Arabic numerals in square brackets and listed at the end of the paper in order of their appearance in the text, including exclusively publications cited inside. Samples of correct formats for various types of references are presented below:

- [1] Y. Namihiro, "Relationship between nonlinear effective area and mode field diameter for dispersion shifted fibres", *Electron. Lett.*, vol. 30, no. 3, pp. 262–264, 1994.
- [2] C. Kittel, *Introduction to Solid State Physics*. New York: Wiley, 1986.
- [3] S. Demri and E. Orłowska, "Informational representability: Abstract models versus concrete models", in *Fuzzy Sets, Logics and Knowledge-Based Reasoning*, D. Dubois and H. Prade, Eds. Dordrecht: Kluwer, 1999, pp. 301–314.

**Biographies and photographs of authors.** A brief professional author's biography of up to 200 words and a photo of each author should be included with the manuscript.

**Galley proofs.** Authors should return proofs as a list of corrections as soon as possible. In other cases, the article will be proof-read against manuscript by the editor and printed without the author's corrections. Remarks to the errata should be provided within one week after receiving the offprint.

**Copyright.** Manuscript submitted to JTIT should not be published or simultaneously submitted for publication elsewhere. By submitting a manuscript, the author(s) agree to automatically transfer the copyright for their article to the publisher, if and when the article is accepted for publication. The copyright comprises the exclusive rights to reproduce and distribute the article, including reprints and all translation rights. No part of the present JTIT should not be reproduced in any form nor transmitted or translated into a machine language without prior written consent of the publisher.

For copyright form see:

<http://www.itl.waw.pl/publ/jtit/wskazowki.html>

A copy of the JTIT is provided to each author of paper published.

*(Contents Continued from Front Cover)*

**Linux Scheduler Improvement for Time Demanding Network Applications, Running on Communication Platform Systems**

*M. Hasse and K. Nowicki*

*Paper*

**63**

**Simple Dynamic Threshold Decryption Based on CRT and RSA**

*B. Nakielski and J. Pomykała*

*Paper*

**70**

**Generating Pseudorandom S-Boxes – a Method of Improving the Security of Cryptosystems Based on Block Ciphers**

*P. Mroczkowski*

*Paper*

**74**

**Editorial Office**

National Institute  
of Telecommunications  
Szachowa st 1  
04-894 Warsaw, Poland

tel. +48 22 512 81 83  
fax: +48 22 512 84 00  
e-mail: [redakcja@itl.waw.pl](mailto:redakcja@itl.waw.pl)  
<http://www.itl.waw.pl/jtit>