

INTERNATIONAL JOURNAL OF
COMPUTER NETWORKS (IJCN)

ISSN : 1985-4129

Publication Frequency: 6 Issues / Year



CSC PUBLISHERS
<http://www.cscjournals.org>

INTERNATIONAL JOURNAL OF COMPUTER NETWORKS (IJCN)

VOLUME 4, ISSUE 4, 2012

**EDITED BY
DR. NABEEL TAHIR**

ISSN (Online): 1985-4129

International Journal of Computer Networks (IJCN) is published both in traditional paper form and in Internet. This journal is published at the website <http://www.cscjournals.org>, maintained by Computer Science Journals (CSC Journals), Malaysia.

IJCN Journal is a part of CSC Publishers

Computer Science Journals

<http://www.cscjournals.org>

INTERNATIONAL JOURNAL OF COMPUTER NETWORKS (IJCN)

Book: Volume 4, Issue 4, October 2012

Publishing Date: 25-10-2012

ISSN (Online): 1985-4129

This work is subjected to copyright. All rights are reserved whether the whole or part of the material is concerned, specifically the rights of translation, reprinting, re-use of illustrations, recitation, broadcasting, reproduction on microfilms or in any other way, and storage in data banks. Duplication of this publication of parts thereof is permitted only under the provision of the copyright law 1965, in its current version, and permission of use must always be obtained from CSC Publishers.

IJCN Journal is a part of CSC Publishers

<http://www.cscjournals.org>

© IJCN Journal

Published in Malaysia

Typesetting: Camera-ready by author, data conversion by CSC Publishing Services – CSC Journals, Malaysia

CSC Publishers, 2012

EDITORIAL PREFACE

The International Journal of Computer Networks (IJCN) is an effective medium to interchange high quality theoretical and applied research in the field of computer networks from theoretical research to application development. This is the fourth issue of volume four of IJCN. The Journal is published bi-monthly, with papers being peer reviewed to high international standards. IJCN emphasizes on efficient and effective image technologies, and provides a central for a deeper understanding in the discipline by encouraging the quantitative comparison and performance evaluation of the emerging components of computer networks. Some of the important topics are ad-hoc wireless networks, congestion and flow control, cooperative networks, delay tolerant networks, mobile satellite networks, multicast and broadcast networks, multimedia networks, network architectures and protocols etc.

The initial efforts helped to shape the editorial policy and to sharpen the focus of the journal. Starting with volume 5, 2013, IJCN aims to appear with more focused issues. Besides normal publications, IJCN intend to organized special issues on more focused topics. Each special issue will have a designated editor (editors) – either member of the editorial board or another recognized specialist in the respective field.

IJCN give an opportunity to scientists, researchers, engineers and vendors to share the ideas, identify problems, investigate relevant issues, share common interests, explore new approaches, and initiate possible collaborative research and system development. This journal is helpful for the researchers and R&D engineers, scientists all those persons who are involve in computer networks in any shape.

Highly professional scholars give their efforts, valuable time, expertise and motivation to IJCN as Editorial board members. All submissions are evaluated by the International Editorial Board. The International Editorial Board ensures that significant developments in computer networks from around the world are reflected in the IJCN publications.

IJCN editors understand that how much it is important for authors and researchers to have their work published with a minimum delay after submission of their papers. They also strongly believe that the direct communication between the editors and authors are important for the welfare, quality and wellbeing of the journal and its readers. Therefore, all activities from paper submission to paper publication are controlled through electronic systems that include electronic submission, editorial panel and review system that ensures rapid decision with least delays in the publication processes.

To build its international reputation, we are disseminating the publication information through Google Books, Google Scholar, Directory of Open Access Journals (DOAJ), Open J Gate, ScientificCommons, Docstoc and many more. Our International Editors are working on establishing ISI listing and a good impact factor for IJCN. We would like to remind you that the success of our journal depends directly on the number of quality articles submitted for review. Accordingly, we would like to request your participation by submitting quality manuscripts for review and encouraging your colleagues to submit quality manuscripts for review. One of the great benefits we can provide to our prospective authors is the mentoring nature of our review process. IJCN provides authors with high quality, helpful reviews that are shaped to assist authors in improving their manuscripts.

Editorial Board Members

International Journal of Computer Networks (IJCN)

EDITORIAL BOARD

EDITOR-in-CHIEF (EiC)

Dr. Min Song

University of Toledo, Ohio (United States of America)

ASSOCIATE EDITORS (AEiCs)

Dr. Qun Li

The College of William and Mary
United States of America

Dr. Sachin Shetty

Tennessee State University
United States of America

Dr. Liran Ma

Michigan Technological University
United States of America

Dr. Benyuan Liu

University of Massachusetts Lowell
United States of America

Assistant Professor Tommaso Melodia

University at Buffalo
United States of America

EDITORIAL BOARD MEMBERS (EBMs)

Dr. Wei Cheng

George Washington University
United States of America

Dr. Yu Cai

Michigan Technological University
United States of America

Dr. Ravi Prakash Ramachandran

Rowan University
United States of America

Dr. Bin Wu

University of Waterloo
Canada

Dr. Jian Ren

Michigan State University
United States of America

Dr. Guangming Song
Southeast University
China

Dr. Jiang Li
Howard University
China

Dr. Fang Liu
University of Texas at Pan American
United States of America

Dr. Enyue Lu
Salisbury University
United States of America

Dr. Chunsheng Xin
Norfolk State University
United States of America

Dr. Imad Jawhar
United Arab Emirates University
United Arab Emirates

Dr. Yong Cui
Tsinghua University
China

Dr. Zhong Zhou
University of Connecticut
United States of America

Associate Professor Cunqing Hua
Zhejiang University
China

Dr. Manish Wadhwa
South University
United States of America

Associate Professor Paulo de Figueiredo Pires
Federal University of Rio de Janeiro
Brazil

Associate Professor Vijay Devabhaktuni
University of Toledo
United States of America

Dr. Mukaddim Pathan
CSIRO-Commonwealth Scientific and Industrial Research Organization
Australia

Dr. Bo Yang
Shanghai Jiao Tong University
China

Assistant Professor Yi Gu

University of Tennessee at Martin
United States of America

Assistant Professor Tarek Guesmi

University of Nizwa
Oman

Dr Yan Sun

Washington State University
United States of America

Associate Professor Flavia C. Delicato

Federal University of Rio de Janeiro
Brazil

Dr. Rik Sarkar

Free University of Berlin
Germany

Associate Professor Mohamed Younis

University of Maryland, Baltimore County
United States of America

Dr. Jinhua Guo

University of Michigan
United States of America

Associate Professor Habib M. Ammari

University of Michigan Dearborn
United States of America

TABLE OF CONTENTS

Volume 4, Issue 4, October 2012

Pages

- | | |
|-----------|--|
| 72 - 81 | On the Malware Front
<i>Robert Kooij, H.J. van der Molen</i> |
| 82 - 97 | On the Speedup/Delay Trade-Off in Distributed Simulations
<i>Alessandra Pieroni, Giuseppe Iazeolla</i> |
| 98 - 103 | Reliability Evaluation of Stochastic-Flow Network Under Quickest Path and System Capacity Constraints
<i>Moatamad R. Hassan</i> |
| 104 - 134 | Analytical Models for Dimensioning of OFDMA-based Cellular Networks Carrying VoIP and Best-Effort Traffic
<i>Bruno Baynat</i> |
| 135 - 155 | Synergies of Radio Frequency and Free Space Optics Communication: New Hybrid Solutions for Next Generation Wireless Mesh Networks
<i>Mohammad Ahmad, Ala Al-Fuqaha, Osama Awwad, Bilal Khan</i> |

On the Malware Front

R.E. Kooij

*Delft University of Technology &
TNO (Netherlands Organization for Applied Scientific Research)
Brassersplein 2, 2612 CT Delft, the Netherlands*

robert.kooij@tno.nl

H.J. van der Molen

*Wageningen University
P.O. box 2, 6700 AA Wageningen, the Netherlands*

henk-jan.van.der.molen@hswageningen.nl

Abstract

The purpose of this article is to extend related research on the spread of malware in networks and to assess the security impact of certain measures against the spread of malware. We examine the influence of heterogeneous infection and curing rates for a Susceptible-Infected-Susceptible (SIS) model, that is used to describe the spread of malware on the Internet. The topology structure considered is the regular graph, which represents homogeneous network structures. We present a new method to calculate the steady state of heterogeneous populations, for the general case with m subpopulations. Using this method, we give the explicit conditions under which the malware persists in the network. Under the condition that all infection rates in the heterogeneous model are equal, we provide a logistic solution for the model.

Next we give calculation examples which are based on the assumption of two subpopulations and explore this method in more detail, proving that the method produces valid outcomes and that the basic reproduction numbers R for each subpopulation are the only factors determining the steady state situation. The value of R depends on the effectiveness of attacking malware and the defending countermeasures.

We then consider some special cases for subpopulations using this method. In the first case the protection against malware is assumed to be absent within one subpopulation. The calculations show that it pays off for the subpopulations with the best protection to help other, less protected subpopulations.

The second case describes the effect of diversification against malware, when one subpopulation does not share the vulnerabilities with the rest of the population to become infected with malware and propagate that malware. The results show that diversification is an effective countermeasure against the propagation of malware. Based on the market share of the software, we demonstrate how to calculate the 'resistance' of different compartments against malware.

Using statistical data, we finally show that dividing a population in two subpopulations increases the accuracy of the model. Based on this data, we also show that the use of security software does not correlate very well with the number of reported infections.

Keywords: Virus spread, epidemic threshold, heterogeneous networks, diversification.

1. INTRODUCTION

In our current society the Internet represents an enormous societal and economic value. Unfortunately where there is value, crime is soon to follow and on the Internet now many cybercriminals are active and malware is ubiquitous. The term "malware" is defined as a piece of software with a harmful payload, which needs (vulnerabilities in) a specific software package to propagate from an infected system to other systems. In 2010, the average rate of malware in email traffic was 1 in every 284 emails and the average number of malicious web sites blocked each day rose to 3,188. Almost 90% of these blocked sites are legitimate sites, which were compromised [1]. When infected computers spread the infection to other computers, the number of malware sources explodes in a short time. When confronted with such malware avalanches, relying on prevention alone is not realistic any more. It is necessary to identify the factors that

International Journal of Computer Networks (IJCN), Volume (4) : Issue (4) : 2012 72

control the spread of malware on the Internet, predict how many computers will be infected and how effective countermeasures are.

The purpose of this article is to extend related research on the spread of malware in networks and to assess the security impact of certain measures against the spread of malware. Our malware spreading model is based upon the Susceptible-Infected-Susceptible (SIS) infection model, which arose in mathematical biology, which is often used to model the spread of computer viruses [2], [3], [4]. The SIS model assumes that a node in the network is in one of two states: infected and therefore infectious, or healthy and therefore susceptible to infection. The SIS model usually assumes instantaneous state transitions. Thus, as soon as a node becomes infected, it becomes infectious and likewise, as soon as a node is cured it is susceptible to re-infection.

In epidemiological theory, a crucial notion is the epidemic threshold τ_c , see for instance [2], [3], [5], [6], [7], [8]. If it is assumed that the infection rate along each link is β while the curing rate for each node is δ then the effective spreading rate of the virus can be defined as $\tau = \beta/\delta$. The epidemic threshold can be defined as follows: for effective spreading rates below τ_c the virus contamination in the network dies out, while for effective spreading rates above τ_c the virus is prevalent. In the case of persistence we will refer to the prevailing state as a steady state. The epidemic threshold is related to the so-called basic reproduction number R , see [6]. In fact, the epidemic threshold τ_c corresponds to the case $R=1$, with virus extinction for $R \leq 1$ and virus prevalence for $R > 1$.

Between 1999 and 2009 many articles considered more modeling aspects for homogeneous populations like incubation periods, variable infection rate, a curing process that takes a certain amount of time, adaptive networks and so on, see [6], [7], [9], [10].

In 2009, ref. [11] derived analytical results for the epidemic threshold in the case of heterogeneous curing rates for a specific class of graphs. It is assumed in [11] that the infection rate at every link is the same, namely β . The aim of this paper is to generalize the results from [11] by also considering heterogeneous infection rates.

The rest of the paper is organized as follows. In Section 2 we derive and analyze the spread of viruses in regular graphs in case of m subpopulations, with curing rate δ_i and infection rate β_i , for $i=1..m$. In Section 3, we discuss the specific case of regular graphs with 2 subpopulations. In the subsequent sections we consider some special cases; in Section 4 we look at the effect of a population without defense and in Section 5 we look at the impact of diversification. In Section 6 we discuss some statistics obtained through Eurostat in order to determine the relation between security software deployed and the percentage of infected computers. We summarize our results in Section 7.

2. VIRUS SPREAD ON REGULAR GRAPHS WITH M SUBPOPULATIONS

In this section, we derive the threshold for the spread of viruses and the steady state of m subpopulations on regular graphs, each with their own curing rate and infection rate. We assume that each node in the connected regular graphs has exactly k neighbors. Denote n_i as the fraction of nodes in subpopulation i , with $i = 1..m$. Obviously, it holds that $\sum_{i=1}^m n_i = 1$. For every node in subpopulation i we denote the curing rate as δ_i , and the infection rate of all incoming links as β_i , with $i = 1..m$. Our assumptions imply that we are considering bi-directional links where the infection rate in the two directions in general is not equal. The latter condition also reflects the assumption that the rate of infection is determined by the node itself, for instance by the type of software it is running. We will come back to this assumption later on in the paper.

It is important to note that our assumptions imply complete symmetry, each node sees the same fraction of nodes from every subpopulation. So every node has a fraction n_1 of neighbors from subpopulation #1, a fraction n_2 of neighbors from subpopulation #2 and so on. Therefore, the number of subpopulations should not exceed the number of direct neighbors, or $m \leq k$.

For subpopulation i at time t , we denote the number of infected nodes as $X_i(t)$ and the fraction of infected nodes as $v_i(t)$. Then, the probability that a randomly chosen node within subpopulation i is infected in the total population with N nodes is $v_i(t) \equiv \frac{X_i(t)}{Nn_i}$.

The rate at which the probability of infection for nodes in subpopulation i changes is due to two processes: susceptible nodes becoming infected and infected nodes being cured. The curing rate for an infection probability v_i is $\delta_i v_i$. The rate at which the probability v_i grows is proportional to the probability of a node in subpopulation i being susceptible, i.e. $1-v_i$. For every susceptible node the rate of infection is the product of the infection rate per node in that subpopulation (β_i) and the probability that on a given link the susceptible node connects to an infected node is $k \sum_{j=1}^m n_j v_j$. Therefore, the following system of differential equations describes the time evolution of $v_i(t)$ with $i, j = 1..m$: $\frac{dv_i}{dt} = \beta_i k \left(\sum_{j=1}^m n_j v_j \right) (1 - v_i) - \delta_i v_i$ (1)

Note that for $\delta_1 = \dots = \delta_m$, and for $\beta_1 = \dots = \beta_m$, the system of equations (1) reduces to a single differential equation, describing the general solution for a homogeneous population with $v = \sum_{j=1}^m n_j v_j$. (2)

For the general case with different curing and infection rates, it is impossible to obtain an explicit solution for the system of equations (1).

Theorem 1. If the effective spreading rate $\tau = \sum_{i=1}^m \frac{\beta_i n_i}{\delta_i}$ for a system of m differential equations in Eq. (1), then the epidemic threshold satisfies $\tau_c = \frac{1}{k}$.

Proof. We will use a Lyapunov function [12] to show that, under the condition $\sum_{i=1}^m \frac{\beta_i n_i}{\delta_i} \leq \frac{1}{k}$, the origin is a global attractor for $\{v_1 \geq 0, v_2 \geq 0, \dots, v_m \geq 0\}$, hence, that the virus dies out.

Let $V = \sum_{i=1}^m \frac{n_i v_i}{\delta_i}$, then we have $\frac{dV}{dt} = \left(k \sum_{i=1}^m \frac{\beta_i n_i}{\delta_i} (1 - v_i) - 1 \right) \sum_{i=1}^m n_i v_i$. (3)

Because $v_i \geq 0$, it follows that $1 - v_i \leq 1$. Therefore Eq. (3) implies that $\frac{dV}{dt} \leq k \sum_{i=1}^m \frac{\beta_i n_i}{\delta_i} - 1 \sum_{i=1}^m n_i v_i$.

Hence under the condition $\sum_{i=1}^m \frac{\beta_i n_i}{\delta_i} \leq \frac{1}{k}$ it holds that $\frac{dV}{dt} \leq 0$. The claim follows directly by applying Lyapunov's stability theorem.

Next, we consider the case $\sum_{i=1}^m \frac{\beta_i n_i}{\delta_i} > \frac{1}{k}$. We first note that any trajectory of the system (1) can never leave the box $B = \{(v_1, \dots, v_m) \mid 0 \leq v_1 \leq 1, \dots, 0 \leq v_m \leq 1\}$. This follows from $\frac{dv_1}{dt} \big|_{v_1=0} = \beta_1 k \sum_{i=1}^m n_i v_i \geq 0$ and similar inequalities at the borders of the box B . From the construction of the above Lyapunov function V , we can see that for $\sum_{i=1}^m \frac{\beta_i n_i}{\delta_i} > \frac{1}{k}$ and for $(v_1, \dots, v_m)$ in B and sufficiently close to the origin, $\frac{dV}{dt} > 0$. This implies that the origin has an unstable manifold in B . Therefore, since any trajectory of system (1) can never leave the box B , system (1) has an attractor as the ω -limit set and the virus survives. This finishes the proof of the theorem.

If the graphs considered are limited to connected regular graphs where each node has exactly k neighbors, then calculations can be simplified by introducing the basic reproduction numbers $R_i = \frac{\beta_i k}{\delta_i}$. Under the condition $\forall i = 1..m \rightarrow \delta_i > 0$, Eqs. (1) and (2) lead to:

$$\frac{dv_i}{dt} = \delta_i R_i v (1 - v_i) - v_i \quad (4)$$

Theorem 2. For a system of m differential equations in Eq. (4), the steady state of (v) can be calculated by solving a polynomial equation of order m .

Proof. Solving Eq. (4) leads to: $\frac{dv_i}{dt} = 0 \rightarrow v_i = \frac{R_i v}{R_i v + 1}$ (5)

If $v \neq 0$ then Eq. (2) and (5) lead to: $\sum_{i=1}^m \frac{n_i R_i}{R_i v + 1} = 1$ (6)

Thus, when $\sum_{i=1}^m (n_i R_i) \leq 1 \rightarrow v \leq 0$; $\sum_{i=1}^m (n_i R_i) > 1 \rightarrow v > 0$. (7)

Eq. (7) confirms the epidemic threshold found in Theorem 1. Next, we use Eq. (6) to calculate the steady state: $h(v) = \prod_{i=1}^m (R_i v + 1) - \sum_{i=1}^m n_i R_i \prod_{j=1, j \neq i}^m (R_j v + 1) = 0$. (8)

So $h(v)$ is a polynomial equation of order m with these preconditions: $\{n_1, \dots, n_m \in \langle 0..1 \rangle \wedge R_1, \dots, R_m > 0\}$. Solving $h(v)$ provides the steady state of v and by filling in v in Eq. (5) the steady state of every v_j . This finishes the proof of the theorem.

Theorem 3. *If $\beta_1 = \dots = \beta_m$, then Eq. (1) has $m - 1$ solutions in the form of hyper-planes passing through the origin. The intersection of the hyper-planes is also a solution of Eq. (1) and its dynamics are described by a logistic equation.*

Proof. Assuming $v_j = \lambda_j v_1$ ($i=2..m$), where the λ_j 's are constants, it follows that

$$\left. \frac{dv_j}{dt} - \lambda_j \frac{dv_1}{dt} \right|_{v_j = \lambda_j v_1} \equiv -v_1 (c_2 \lambda_j^2 + c_1 \lambda_j + c_0) = -v_1 f(\lambda_j) \quad \text{with}$$

$$c_2 = n_j k (v_1 (\beta_j - \beta_1) + \beta_1); c_1 = n_1 v_1 k (\beta_j - \beta_1) + n_1 \beta_1 k - n_j \beta_j k + \delta_j - \delta_1; c_0 = -n_1 \beta_j k. \quad (9)$$

Hence, if $\beta_1 = \dots = \beta_m$, then c_0 , c_1 and c_2 are constants. Then, because $c_2 > 0$ and $c_0 < 0$, it follows that $f(\lambda_j)$ has exactly one positive root λ_j^* , for $i=2..m$. Therefore, the hyper-planes $v_j = \lambda_j^* v_1$ ($i=2..m$) are solutions of Eq. (1), when $\beta_1 = \dots = \beta_m$. Using the first equation in Eq. (1), we can show that on the intersection of the $m - 1$ hyper-planes, the dynamics are described by a logistic equation:

$$\frac{dv_1}{dt} = \beta_1 k v_1 \left(\sum_{j=1}^m n_j \lambda_j^* \right) (1 - v_1) - \delta_1 v_1, \quad (10)$$

where $\lambda_1^* = 1$.

This concludes the proof of the theorem.

3. CALCULATION METHOD IN DETAIL FOR TWO SUBPOPULATIONS

In this section, the least complex heterogeneous situation is explored in more detail. By filling in $m=2$ in Eq. (8) it follows that:

$$h(v) = R_1 R_2 v^2 + (R_1 + R_2 - R_1 R_2) v + 1 - n_1 R_1 - n_2 R_2 = 0. \quad (11)$$

The method always yields a solution for v , since in Eq. (11) the discriminant $d > 0$. For two subpopulations with parameters $\{n_1, n_2, R_1, R_2\}$ d is calculated as:

$$d = R_1^2 R_2^2 + (R_1 - R_2)^2 + 2 R_1 R_2 (R_1 - R_2) (n_1 - n_2). \quad (12)$$

We know from Theorem 1 that for $n_1 R_1 + n_2 R_2 \leq 1$ system (4) with $m=2$ the virus dies out, i.e. $v=0$ is the global attractor. Hence we only consider the case $n_1 R_1 + n_2 R_2 > 1$. Then, because $h(0) < 0$ and $h(1) > 0$, it follows that there is a unique solution $0 < v < 1$ for Eq. (11). A simple calculation shows that the only solution satisfies $v = \frac{R_1 R_2 - R_1 - R_2 + \sqrt{d}}{2 R_1 R_2}$.

Next, we will show that for Eq. (4) with $m=2$, under the condition $n_1 R_1 + n_2 R_2 > 1$, the steady state of v corresponds to a stable equilibrium point of the system in Eq. (4).

It follows from the analysis above that for $n_1 R_1 + n_2 R_2 > 1$ system (4) with $m=2$ has an equilibrium point located in the region $A = \{(v_1, v_2) \mid 0 < v_1 < 1, 0 < v_2 < 1\}$.

We know from the proof of Theorem 1 that for $n_1 R_1 + n_2 R_2 > 1$ the origin has an unstable manifold entering A , while trajectories of the system can never leave the region A . Therefore, by application of the Poincaré-Bendixson theorem [12] on A , the ω -limit set for system (4) for $m=2$, can be either an equilibrium point or an isolated periodic orbit. To rule out the existence of periodic orbits for system (4) with $m=2$, we can use the Bendixson-Dulac criterion, see [12]. In fact, because periodic orbits cannot intersect $v_1=0$ or $v_2=0$, we can use the Dulac function

$$D(v_1, v_2) = \frac{1}{v_1 v_2} \text{ which leads to}$$

$$\frac{\partial(D \frac{dv_1}{dt})}{\partial v_1} + \frac{\partial(D \frac{dv_2}{dt})}{\partial v_2} = -(\delta_1 n_1 R_1 v_1^2 v_2 + \delta_1 n_2 R_1 v_2^2 + \delta_2 n_1 R_2 v_1^2 + \delta_2 n_2 R_2 v_1^2 v_2) D < 0. \quad (13)$$

Therefore, the system in Eq. (4) with $m=2$ has no periodic orbits and hence the equilibrium point, corresponding with the steady state v , is globally stable.

Next, we will assess the impact of the fraction nodes of type 1, i.e. n_1 , on the value of v , i.e. the fraction of infected nodes. From Eq. (9) it is easy to verify that

$$\frac{\partial v}{\partial n_1} = \frac{R_1 - R_2}{\sqrt{d}}. \quad (14)$$

Hence v is always a monotonic function of n_1 , unless $R_1=R_2$, which corresponds to the homogeneous case. Figure 1 depicts several possibilities for the case $R_2=2$.

For $n_1=0$, the homogeneous case were all nodes belong to type 2, the fraction of infected nodes v equals $1-1/R_2=0.5$. For the case $R_1>R_2$, according to Eq. (14), v increases monotonically with n_1 hence for this case the homogeneous case $n_1=0$ gives the least number of infected nodes.

Clearly, $R_1=R_2$ is the homogeneous case with $v=0.5$.

If $1<R_1<R_2$ then v decreases monotonically with n_1 , while for $n_1=1$ it holds that $v>0$. Hence for this case the homogeneous case $n_1=1$ gives the least number of infected nodes. Finally, if $R_1<1$ then v decreases monotonically with n_1 while for $n_1=1$ it holds that $v\leq 0$. Hence for this case the optimal situation, where the virus dies out, occurs from $n_1 = \frac{R_2-1}{R_2-R_1}$ onwards.

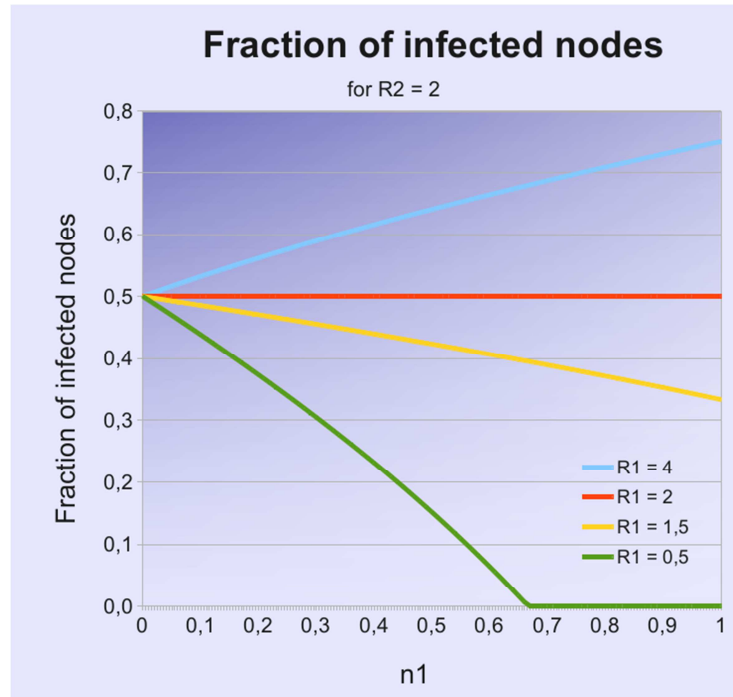


FIGURE 1: The Steady State of v as a function of n_1

We would like to stress once more, that our model is more general than previous models, see e.g. [2], [4], [7, [8], [11], because none of these models can deal with heterogeneous spreading rates. As a simple example, consider the case with two malware populations, with $n_1 = 0.7$, $\beta_1 = 0.4$, $\delta_1 = 1$ and $n_2 = 0.3$, $\beta_2 = 0.05$, $\delta_2 = 1$, where every node has four neighbors, i.e. $k = 4$. If we would want to apply the results of [2], [4], [7, [8], [11], and we would use for the spreading rate β the mean of β_1 and β_2 , then the effective spreading rate would become 0.225, which is below the

epidemic threshold $1/k = 0.25$. However, applying our, more accurate model, we arrive at the conclusion that the malware persists, because $n_1 \beta_1 + n_2 \beta_2 = 0.295 > 0.25 = 1/k$.

4. SPECIAL CASE #1: MALWARE RESERVE

In this section we consider the case that subpopulation #2 has no defense against malware, i.e. $\delta_2=0$, or equivalently, $R_2=\infty$. Then it follows from Eq. (1) that in steady state, the whole subpopulation #2 is infected, i.e. in steady state $v_2=1$ holds.

Using Eqs. (2), (5) and taking $\lim R_2 \rightarrow \infty$, it is easy to show that the steady state fraction of

infected nodes v satisfies $v = \frac{R_1 - 1 + \sqrt{R_1^2 + 2(n_1 - n_2)R_1 + 1}}{2R_1}$.

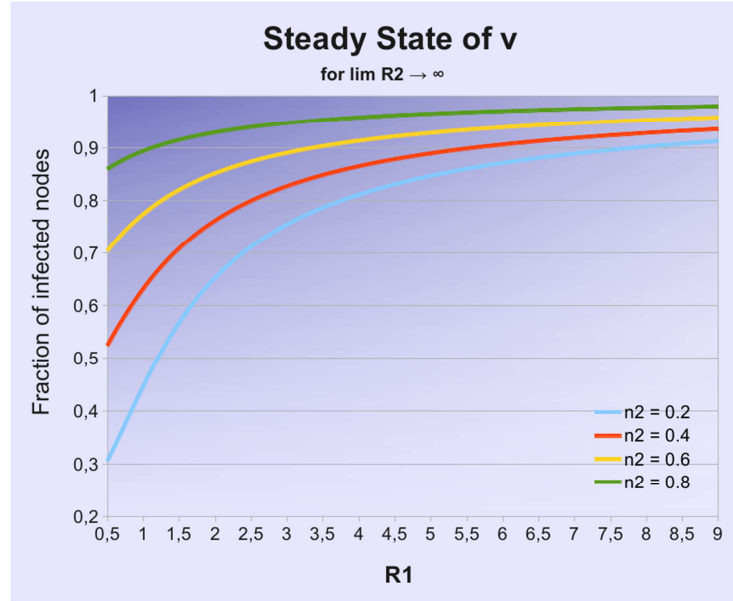


FIGURE 2: Steady state v for the case $\delta_2=0$

Figure 2 shows the steady state v as a function of R_1 , for different values of n_2 , the fraction of nodes that belong to the malware reserve. We observe that the whole population benefits if the size of the malware reserve is decreased. We also see that even if subpopulation #1 has adequate protection against malware (i.e. $R_1 < 1$), they still become infected because of the lack of security for subpopulation #2. For instance, if $R_1=0.5$ and $n_2=0.2$, then in steady state 13% of subpopulation #1 is infected. This is easily verified from Figure 2 and the equality $v=n_1v_1+n_2$, which holds for $\delta_2=0$.

5. SPECIAL CASE #3: DIVERSIFICATION AGAINST MALWARE

In this section we consider the case that subpopulation #2 is immune for malware infections, i.e. $\beta_2=0$, or equivalently, $R_2=0$. Then it follows from Eq. (1) that in steady state, the whole subpopulation #2 is uninfected, i.e. in steady state $v_2=0$ holds.

Using Eq. (12) with $R_2=0$ it follows that the steady state fraction of infected nodes v satisfies

$$v_1 = 1 - \frac{1}{n_1 R_1}; v_2 = 0 \rightarrow v = n_1 - \frac{1}{R_1}. \quad (15)$$

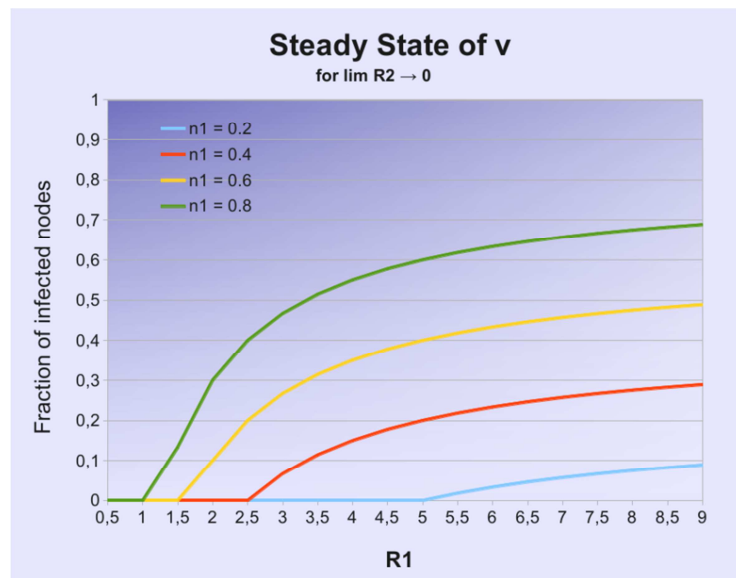
FIGURE 3: Steady state v for the case $\beta_2=0$

Figure 3 shows the steady state v as a function of R_1 , for different values of n_1 , the fraction of nodes not immune to the malware. We observe that the whole population benefits if the size of the immune population increases. We also see that in case of an immune subpopulation, there exists a threshold value for the basis reproduction number R_1 , below which the virus dies out. It is clear from Eq. (15) that this threshold value satisfies $R_1 = \frac{1}{n_1}$.

This threshold value for R is the minimum value necessary to “sustain” an infection level above zero in that compartment. The lower the threshold value, the more malware is able to match it. In Table 1 the threshold value (R) was calculated for popular software using market shares [13]. In this example, the software for which the most malware is expected is MS Windows, MS Office, MS Internet Explorer and Mozilla Firefox – in that order. However, market share statistics vary too much to draw explicit conclusions based on the numbers presented.

6. CORRELATING SECURITY MEASURES WITH MODEL PARAMETERS

Like all models, the SIS model is an approximation of reality. It should be applied with care and respect for its limitations and premises. One of its limitations is that it is necessary to assume that the population is completely symmetrical, i.e. the different nodes are distributed evenly in the network.

Platform	Software Platform	Market Share 3Q10	Threshold R	Remark
Webclient	Windows	0,8821	1,13	
Webclient	MacOS	0,0682	14,66	
Webclient	Linux	0,0108	92,59	
Webclient	Symbian	0,0021	476,19	
Webclient	Blackberry	0,0045	222,22	
Webclient	Other	0,0323		
Webbrowser	IE	0,4622	2,16	
Webbrowser	Firefox	0,2992	3,34	
Webbrowser	Chrome	0,1240	8,06	
Webbrowser	Safari	0,0555	18,02	
Webbrowser	Opera	0,0193	51,81	
Webbrowser	Other	0,0398		
Office Suite	MS Office	0,8800	1,14	(Dutch market share only)
Office Suite	OpenOffice	0,0800	12,5	(Dutch market share only)
Office Suite	Wordperfect	0,0090	111,11	(Dutch market share only)
Office Suite	Other	0,0310		

TABLE 1: Calculation Example of Threshold Values for popular Software

In its simplest, homogeneous form, the value of the single set of parameters of the SIS-model (β , δ , R) are determined by the effectiveness of security measures taken by the defenders and the attempts of the attackers to outsmart them (see Table 2).

	Cyber Defense (Reduce R)	Cyber Attack (Increase R)
β	Prevent malware infection: - Intrusion prevention system, firewall, heuristic AV software for on-access scanning "white-list" software, USB media, e-mail content, web content to download - Configuration: restrict user rights, harden systems, sandboxing, good passwords frequently changed,... - Separate compartments: network, software, user rights, encrypted files - Good procedures for changes / updates - Invest in knowledge and awareness - Better, less vulnerable software - Preventive security audits - Legal software for employees @home	Increase risk of infection by malware: - IP / MAC address spoofing - Multiple attack patterns in malware - Web site offers customized malware - Domain Generation Algorithms (DGA) - Social engineering, imitation of legitimate software, e.g. AV software - Malware on trusted sites - Sharing / stealing knowledge, source code - Commercial and freeware Malware kits - Fuzz testing of software for vulnerabilities - Testing malware with security software - Stealth malware, encryption, code obfuscation - Massive and rapid spread of malware (reverse engineering of patches?) - Targeted malware, APT ("precision ammo")
δ	Improve disinfection (detection+correction): - Multiple AV packages for scheduled scans - Intrusion Detection System, logging - Management procedures for incidents and changes, including an Incident Response Plan - Invest in knowledge and awareness - Postmortem security audits - Periodically re-install clean software image on all PC's	Reduce loss of infected computers: - Root kits, anti-virtualization techniques, disable security software and update mechanisms - Encryption, remove trace data, multiple layered code obfuscation - Malware self-activation / self removal under certain conditions - Malware updates faster than AV - Patching of infected computers (!) - Continuity plan for botnet, e.g. rotating web servers, integrate infected computer in >1 botnets, bullet proof hosting of C&C servers - Imitation behavior of legitimate software

TABLE 2: The Battle Between Cyber Attack and Cyber Defense

The infectivity of different occurrences of malware may vary widely, depending on the knowledge of the attacker and the purpose of the malware. For this moment we neglect the differences between malware samples, because we like to focus on the protective measures of the defenders. The infectivity of all malware is then considered to be equal.

When individuals or organizations use different security measures, the parameter (R) is likely to differ. For instance, if one organization prohibits the user to install software, this will reduce the risk of infection considerably.

Correctly estimating the corresponding value of the parameter (R) from the operational security measures is difficult. The Eurostat Newsrelease of Feb 8, 2011 presents a list of statistics on Internet security of the EU countries in 2010.

One of the statistics was the percentage of individuals who reported that they caught a computer infection resulting in loss of information or time using the Internet in the 12 months prior to the survey. Another statistic was the percentage of individuals who used the Internet in the last 12 months and stated that they used IT security software to protect their private computer and data.

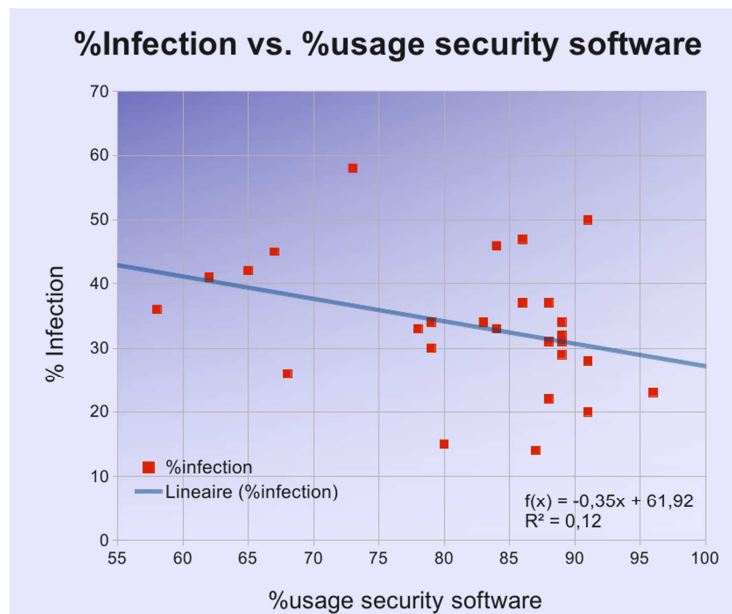


FIGURE 4: Correlating the usage of security software and percentage of infections

After deleting the incomplete data entry of Romania, both statistics were plotting against each other using an XY diagram (see Figure 4). When we treat this population as homogeneous, the lowest value of 14 percent infections and highest value of 58 are far off the mean value of 33.45. The usage of security software vs. the (resulting) percentage of infections seems a clearcut case, but the data shows little correlation - the trend line shows an R^2 value of 0.12. With a total variance of 100.54 for 29 countries, this model's accuracy is low.

We can decrease the total variance by introducing subpopulations with their own mean values. Sorting the list of countries using the percentage of infection, we can divide the population in a group of leaders and a group of laggards. We have determined that the minimal variance of the total population is 41.29 when the group of leaders consists of 22 countries (mean value is 29.14) and the group of laggards consists of 7 countries (mean value is 47.0).

Interestingly, the higher accuracy has little influence on the correlation per subpopulation between the usage of security software vs. the resulting percentage of infections; in fact, for the group of laggards, the correlation line shows that the higher the usage of security software is, the higher the resulting percentage of infections becomes. A possible explanation is that security software is necessary to detect malware infections and that the laggards are more often tricked in using bogus anti-virus software, which in fact is malware. However, even the correlation in the leaders group has dropped to 0.06, so the usage of security software seems to be a poor predictor of the resulting percentage of infections.

7. CONCLUSIONS

We have introduced a new method to calculate the steady state for heterogeneous populations. Based on analysis of this method, we think that a heterogeneous model can be accurately matched with a logistic function. We can also predict that either a minimal or maximum value of the infected fraction of the population occurs when the heterogeneous population becomes homogeneous. Thus, when the security level of the least secure group increases or the fraction of this group decreases, the whole population benefits from this.

The analysis also reveals that a minimum occurs when the population ceases to be a mono-culture, i.e. not every node shares the same vulnerabilities for malware. More diversity is an effective measure against the propagation of malware. Although all separate compartments can attract malware, the existing measures become more effective and the total level of infections is less than in a similar population which all use the same hard- and software. The opposite is also true: the bigger a mono-culture is, the less infectious malware has to be to persist.

Finally, we like to suggest some directions for future research. First of all, since it has been proven difficult to correlate model parameters of the (heterogeneous) SIS model to the use of security software, future research can be directed to derive realistic parameter values for other (clusters of) security measures. Secondly, our presented model assumes a static topology, whereas in real-life, computer networks are changing in time. In line with [10], where homogeneous virus spread for dynamic networks is considered, we suggest to generalize our results for heterogeneous malware populations to dynamic networks. Finally, although our heterogeneous model incorporates different infection and curing rates per subpopulation (β_i , δ_i , R_i), our results depend on the assumption of a form of complete symmetry in the network, i.e. every node is connected to the same fractions of nodes from each subpopulation. More research is needed to reveal how the results are influenced when this symmetry assumption is dropped.

8. REFERENCES

- [1] MessageLabs Intelligence. "2010 Annual Security Report", December 7, 2010 <http://www.inteco.es/file/27gHxrzWsYyeyRTFYq8MuQ> [2012-10-05]
- [2] J.O. Kephart and S.R. White. "Direct-graph epidemiological models of computer viruses", Proc. IEEE Computer Society Symposium on Research in Security and Privacy, pp. 343-359, 1991.
- [3] R. Pastor-Satorras and A. Vespignani. "Epidemic Spreading in Scale-Free Networks", Physical Review Letters, Vol. 86, No. 14, April, 3200-3203, 2001.
- [4] A. Ganesh, L. Massoulié and D. Towsley. "The Effect of Network Topology on the Spread of Epidemics", Proc. IEEE INFOCOM.05, Miami, 2005.
- [5] N.T.J. Bailey. "The Mathematical Theory of Infectious Diseases and its Applications", London: Charlin Griffin & Company, 2nd ed., 1975.
- [6] D.K. Daley and J. Gani. "Epidemic modelling: An Introduction", Cambridge University Press, 1999.
- [7] Y. Wang, D. Chakrabarti, C. Wang, and C. Faloutsos. "Epidemic spreading in real networks: An eigenvalue viewpoint", IEEE Computer Society, 22nd International Symposium on Reliable Distributed Systems (SRDS'03), pages 25—34, Los Alamitos, CA, USA, 2003.
- [8] P. Van Mieghem, J. Omic, and R.E. Kooij. "Virus spread in networks". IEEE/ACM Transactions on Networking, 17(1), 1-14, 2009.
- [9] Y. Wang and C. Wang. "Modeling the Effects of Timing Parameters on Virus Propagation". ACM Workshop on Rapid Malcode, Washington, DC, Oct. 27, 2003.
- [10] T. Gross, C. Dommar D'Lima and B. Blasius. "Epidemic dynamics on an adaptive network", Physical Review Letters 96, 208701–4, 2006.
- [11] J. Omic, R.E. Kooij and P. Van Mieghem. "Heterogenous protection in regular and complete bipartite networks", Proc. of Networking 2009, Aachen Germany, 11-15 May, 2009.
- [12] J. Guckenheimer and P. Holmes. "Nonlinear oscillations, dynamical systems, and bifurcations of vector fields", New York: Springer, 1983
- [13] See for market share used (OS, Browser and Office software) [2012-05-20]: <http://marketshare.hitslink.com/operating-system-market-share.aspx?qprid=8>; <http://marketshare.hitslink.com/browser-market-share.aspx?qprid=0&qpcustommd=0&qptimeframe=M&qpsp=155>; www.webmasterpro.de/portal/news/2010/02/05/international-openoffice-market-shares.html

On the Speedup/Delay Trade-Off in Distributed Simulations

Alessandra Pieroni

*Enterprise Engineering Department
University of Rome TorVergata
Rome, Italy*

alessandra.pieroni@uniroma2.it

Giuseppe Iazeolla

*Engineering Department
University of Rome TorVergata
Rome, Italy*

giuseppe.iazeolla@uniroma2.it

Abstract

Assume a local simulator (LS) of a given system is available and we wish to turn it into a distributed simulator (DS). In the DS case, the LS is partitioned into segments called federates, each federate being run by a separate host. Before implementing the DS (i.e., at design-time) we wonder: will the DS execution time be shorter than LS one? In some cases the DS may run slower than the equivalent LS. To answer this question we are to consider that the execution time of a distributed simulation system depends on 3 interacting factors: 1) the speedup (or run-time gain) resulting from the partitioning of the local simulator into federates. 2) The network delays in the federate synchronization messages exchange. 3) The network delays in the federatedata messages exchange. The combination of such factors makes very hard predicting the benefits of the LS-to-DS transformation. In this paper, a LS/DS decision procedure to support the LS/DS decision process at design-time. The procedure is guided by a performance model of the DS. The use of the High Level Architecture (HLA) distributed simulation standard is assumed.

Keywords: distributed simulation, parallel speedup, computer networks delay.

1. INTRODUCTION

A simulation model can be seen as consisting of a set of sub-models. In local simulation (LS), a single model exists that simulates the entire system and is run by a single host. In distributed simulation (DS), various sub-models (called federates) simulate distinct parts of the system and are run by separated hosts connected via a LAN, MAN or WAN computer network or a composition thereof.

Predicting at *design-time* the convenience of implementing the DS version of the LS can be of interest. Indeed, the development of a DS system is a complex and expensive task, since of the cost of achieving the necessary know-how of the distributed simulation standard [1], the cost of the extra-lines of code to develop for each federate [2], the cost of the hosts, the computer networks, and the number of design alternatives to face (in terms of simulator partitioning, host capabilities, network potentialities and so on).

This paper introduces a method to support the evaluation of the DS convenience before implementation. The method investigates the effects of three interacting factors:

1) The speedup (or run-time gain) resulting from partitioning the local simulator into federates, spread across various hosts that operate in parallel; 2) the synch-communication overhead due to network delays in the exchange of synchronization messages among federates; 3) the data-communication overhead due to network delays in the exchange of data messages among federates.

The two communication overheads lower down the run-time gain obtained with the speedup.

A LS/DS decision procedure is proposed to choose (at design-time) whetherto remain on the LS version of the simulator orcarry out the implementation of its DS version. The procedure is guided by a performance model (PM) of the DS. The model can be used to perform what-if analysis and sensitivity analysis to observe how changing one or the other factor may affect the DS execution time. The PM assumes the DS is based on the HLA protocol standard and middleware [3].

The paper is organized as follows: Sect.2 presents the problem statement. Sect.3 illustrates the PM. Sect.4 illustrates the PM implementation in the OMNet++ simulation language and its use in the LS/DS decision procedure. Sect.5 presents the paper contribution with respect to existing literature and finally Sect.6 gives concluding remarks.

2. PROBLEM STATEMENT

Assume a local simulator (LS) of a given system Σ is available, and that we wish to turn it into a distributed simulator (DS).

In the DS case, the LS is partitioned into segments called federates, each federate being run by a separate host. Fig.1 shows the two federate case, with N_s denoting the network for the exchange of synch messages and N_D the one for data messages.

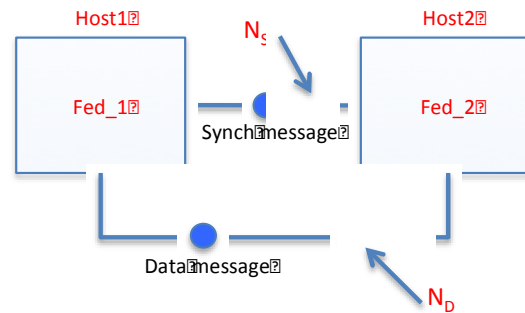


FIGURE 1:DS system with two federates.

Before implementing the DS (i.e., at design-time) we wonder: will the DS execution time be shorter than LS one? In some cases the DS may run slower than the equivalent LS. To answer this question a Performance Model (PM) of ak -federate system is introduced in Sect.3.

The following terminology will be used throughout the paper:

- Σ = System to be simulated
- $LS(\Sigma)$ = Local Simulator of Σ
- T_{LS} = LS execution time
- $DS(\Sigma)$ = Distributed Simulator of Σ
- T_{DS} = DS execution time
- $PM(DS(\Sigma))$ = Performance Model of $DS(\Sigma)$ to predict the execution time T_{DS} .

The question is: when does $DS(\Sigma)$ run faster than $LS(\Sigma)$? In other words, when does $T_{DS} < T_{LS}$ hold?

There are 3 conflicting factors that determine the T_{DS} value:

- *Speedup*: the run-time gain resulting from partitioning LS into federates spread across many hosts that operate in parallel.

Thanks to the speedup one may obtain $T_{DS} < T_{LS}$, the speedup being defined by $S = T_{LS} / T_{DS}$, a positive speedup meaning $S > 1$. Let us call S the **no-delay speedup**, for reasons that will be soon clear.

- *Synch communication overhead*: all DS simulations must incorporate techniques to coordinate the execution of federates across the many hosts by synchronization messages. Such messages travel along a synch network N_S (that may be a LAN, a MAN or a WAN, or a composition thereof) whose delay ΔN_S may yield a $T'_{DS} > T_{DS}$ thus reducing the no-delay speedup S to a **synch-delay speedup** $S' < S$ with $S' = T_{LS} / T'_{DS}$.
- *Data communication overhead*: the federates also need to exchange data-packets by way of data messages. Such messages travel along a data network N_D (that may or may not coincide with N_S), whose delay ΔN_D may yield a $T''_{DS} > T'_{DS}$ thus reducing the synch-delay speedup S' to a **synch&data-delay speedup** $S'' < S'$ with $S'' = T_{LS} / T''_{DS}$.

The question above then becomes:

When does T''_{DS} turn out to be lower than T_{LS} ($T''_{DS} < T_{LS}$), thus still yielding a positive speedup $S'' > 1$?

In other words, when can the no-delay speedup win over the synchronization and data communication overheads?

Next section tries to answer such a question.

2.1 The speedup/communication overhead trade-off

As with most parallel computations, to obtain a positive speedup the portion of LS that can be parallelized must be large relative to the portion that is inherently serial. Let us denote by $S(N)$ the maximum speedup that can be achieved using N processors, and by Q the fraction of computation that is inherently serial. According to Amdahl's law [4,5] even with an arbitrarily large number of processors ($N \rightarrow \infty$), $S(N)$ can be no larger than the inverse of the inherently serial portion Q of LS.

$$S(N) = \frac{1}{Q + \frac{1-Q}{N}} \quad (1)$$

Thus, one requirement for the DS code to achieve positive speedups is that the fraction Q should be small.

An appropriate partitioning of LS into a set of federates should then be found at design-time that improves S while maintaining the synch and data overheads low. In other words, a partitioning that yields a high computation-to-communication ratio (i.e., a large amount of computation between communications).

On this basis, an LS/DS decision procedure can be foreseen (Fig.2) to decide whether to remain on the LS version of the simulation system or carry out the implementation of its DS version.

In other words, assume an $LS(\Sigma)$ has been developed and that its T_{LS} is not satisfactory. A search for an appropriate partitioning of $LS(\Sigma)$ into federates and for an appropriate choice of the N_S and N_D networks has to be performed by the iterative use of the $PM(DS(\Sigma))$, to obtain a $T''_{DS} < T_{LS}$.

At each iteration, if the T''_{DS} predicted by the PM is sufficiently lower than T_{LS} , the decision to implement the $DS(\Sigma)$ can be taken.

Otherwise, one may either try a new tentative partitioning or try alternative networks N_S and N_D of improved capabilities. In case no partitioning nor network improvements can be found, one may decide not to implement the $DS(\Sigma)$.

An example use of the PM in the LS/DS decision procedure is illustrated in Sect.4. The PM cannot be evaluated by analytic methods and thus its evaluation is simulation-based. The coding of the PM is done in the OMNet++ simulation language [6] and an example coding is provided in Sect.4.

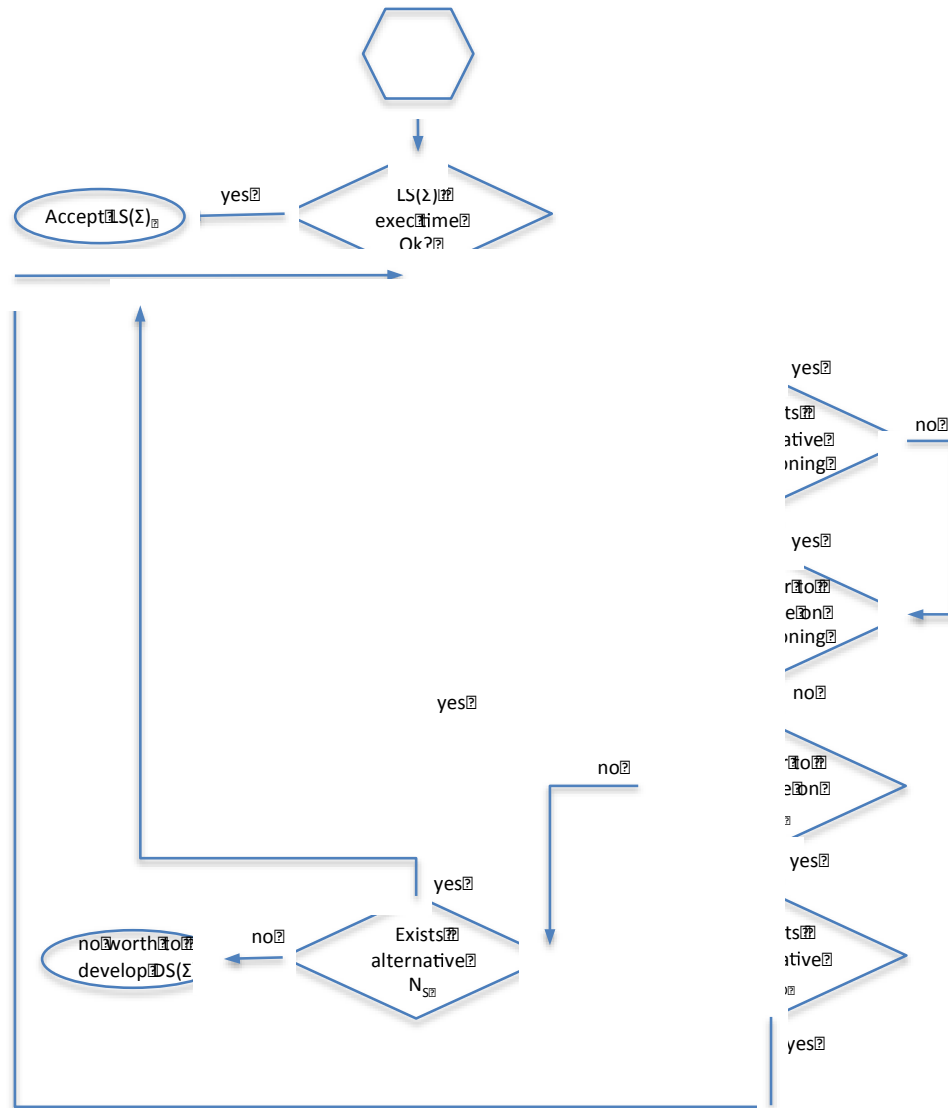


FIGURE 2: The LS/DS decision procedure

3. RELATED WORK

A number of existing contributions can be found in literature that address the prediction of execution times of simulation systems, see e.g., [5, 15, 16, 17, 18, 19, 20].

Contributions in [15, 16, 18, 19, 20] however deal with parallel simulation case rather than distributed case, as is this paper. Parallel simulations are run by a set of processors that are tightly connected by use of dedicated networks rather than by a computer network, as is this paper case (a set of processors connected by a computer network, such as a LAN, a MAN or a WAN, or a composition thereof). Moreover, parallel simulations are coordinated by an ad hoc network-operating system rather than by DS middlewares (e.g. HLA) as is this paper case.

In other words, our paper cannot take advantage from results of the parallel simulation literature. On the other hand, looking at the distributed simulation literature, the only work that, to our knowledge, one can refer to is [17], which however only deals with the performance of DS shared data-access algorithms, a topic that is not of interest to our paper, which instead is interested to the evaluation of the whole DS execution time (T_{DS}).

There are two essentially different ways for evaluating the T_{DS} . One way is to base the analysis on the execution of a DS run. The other way (this paper way) is to use the LS version of the simulator to derive parameters for predicting the T_{DS} . While the first method is potentially more accurate, its main disadvantage is that it requires the existence of the DS program, hence it cannot really predict T_{DS} , and it can only be used to evaluate the T_{DS} of a given DS implementation.

In a previous work [12], the trace information generated during the LS run has been obtained and it will be used in this paper now to derive parameters to give to the performance model (PM) of the DS for the T_{DS} prediction.

The PM, this paper now introduces, is able to separately investigate the effects of the model partitioning, and also investigate separately the effects of the N_S delay and of the N_D delay. Besides being important for the LS/DS decision procedure, the knowledge of the effects of the two communication overheads is of importance to evaluate the representativeness of the $DS(\Sigma)$ at design-time.

Indeed, depending on the nature of system Σ , there are situations in which the data and synch message delays are not critical and thus a communication network of any capability can be used. On other situations, instead, the system Σ can be of such a nature that the synch and data delays become very critical for the representativeness of the system simulator. In other words, the $DS(\Sigma)$ loses the capability of realistically representing the original Σ in case the N_S and N_D networks are not sufficiently well performing.

4. THE PERFORMANCE MODEL OF $DS(\Sigma)$

It is assumed the reader is familiar with the structure of an HLA federation, based on the so-called *Run Time Infrastructure* (RTI) [7]. The RTI is the software that allows the federates to execute together. In Fig.3 the interface between the RTI and the federates is illustrated [8]. The federates do not talk to each other directly. They are instead connected to the RTI and communicate with each other using services provided by the RTI. The RTI offers to each federate an interface called *RTI Ambassador*.

Each federate on the other hand presents an interface called *Federate Ambassador* to the RTI.

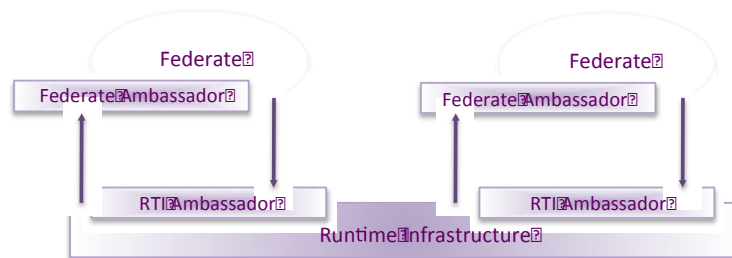


FIGURE 3: HLA federation structure

In the following we shall denote by:

- LEX the local execution internal to a federate, in other words, the standard simulation operations such as event processing, event routines, scheduling of local events, etc.

- HLAR the execution internal to a federate of an RTI service, e.g., an invocation of a time advance request.
- HLA-Ex the execution internal to a federate of a service request coming from the Federate Ambassador.

Assume we deal a federation consisting of k federates. The $PM(DS(\Sigma))$ model will consist of one sub-model for each of the k federates, and of two network models, one for the N_S and one for the N_D network.

The single federate sub-model is illustrated in Sect.3.1, and consists of a non-conventional EQN (Extended Queueing Network) model, that we denote NC_EQN , that includes both conventional EQN nodes and a number of *AND/OR* logic gates necessary to bring-in the logic of the HLA standard. The network model, instead, is a conventional EQN model and illustrated in Sect.3.2, to be used (with appropriate parameters) for both N_S and N_D .

3.1 The federation performance model

To answer the Fig.2 question “Exists alternative partitioning?” we shall assume that the LS is partitioned into a k -federates DS and shall evaluate the PM of such a partitioning.

The NC_EQN model of a k -federates DS is shown in Fig.4, where the details of the PM of only one Federate (Fed_i) are illustrated. One may partition the LS code into the k -portions of the DS code in various ways. As we shall better see in Sect.4, the effect of the partitioning choice is reflected in the value given to parameter p_{SYNC} at model parameterization time.

In Fig.4 the interactions are shown between Fed_i and all remaining federates (in the *publish/subscribe* RTI assumption). The set of all remaining federates is denoted by using the bold **Fed_x** notation. Therefore, the x_i (or ix) notation will be used in the subscript of various components in Fig.4. For example, gate AND_{x_i} relates to the synch-messages exchanged between **Fed_x** and Fed_i . Consequently, we are to figure out number $k-1$ such *AND* gates in the illustration. The same can be said for all other *AND* and *OR* gates with a bold **x** in the subscript.

As visible in Fig.4, each Fed_i sends (receives) messages to (from) **Fed_x** through the N_S and N_D networks. The entire federation PM will thus consist of a set of Fed_i sub-models (as in Fig.4) that interact between themselves through the N_S and N_D nodes as in Fig.5, that shows how messages from various federates are enqueued in front of N_S or N_D to be served, i.e. forwarded to the destination federates.

The Fed_i model in Fig.4 includes:

- a time-consuming node (*Fed_i Host CPU*), that synthetically represents the host that runs the federate. Such a node service-time parameters vary with the serviced job class (C_i or C_{RC} see later).
- a set of non-time consuming nodes, namely:
 - *AND* nodes that perform *AND*-logic operations.
 - *OR* nodes that perform *OR*-logic operations.
 - *SPLIT* nodes, that split an incoming job class into two or more outgoing classes.
 - *Classifier* nodes that, basing on the class of the input job, forward the job in one or the other direction.
 - *Router* nodes that perform probabilistic routing of the incoming jobs.
 - a *Merge* node that merges two job classes.

The computation performed by the federation starts by launching the RTI interface and by initializing the HLA components local to each federate. Such initial computations are performed

only once and do not substantially affect the federation execution time and thus are omitted from the modeling. They are synthetically represented by the INIT job on the top of Fig.4.

The INIT job enters the $Split_0$ node and yields a main thread for each federate belonging to the Federation. The main thread C_i for Fed_i and its follow-on is detailed in Fig.4.

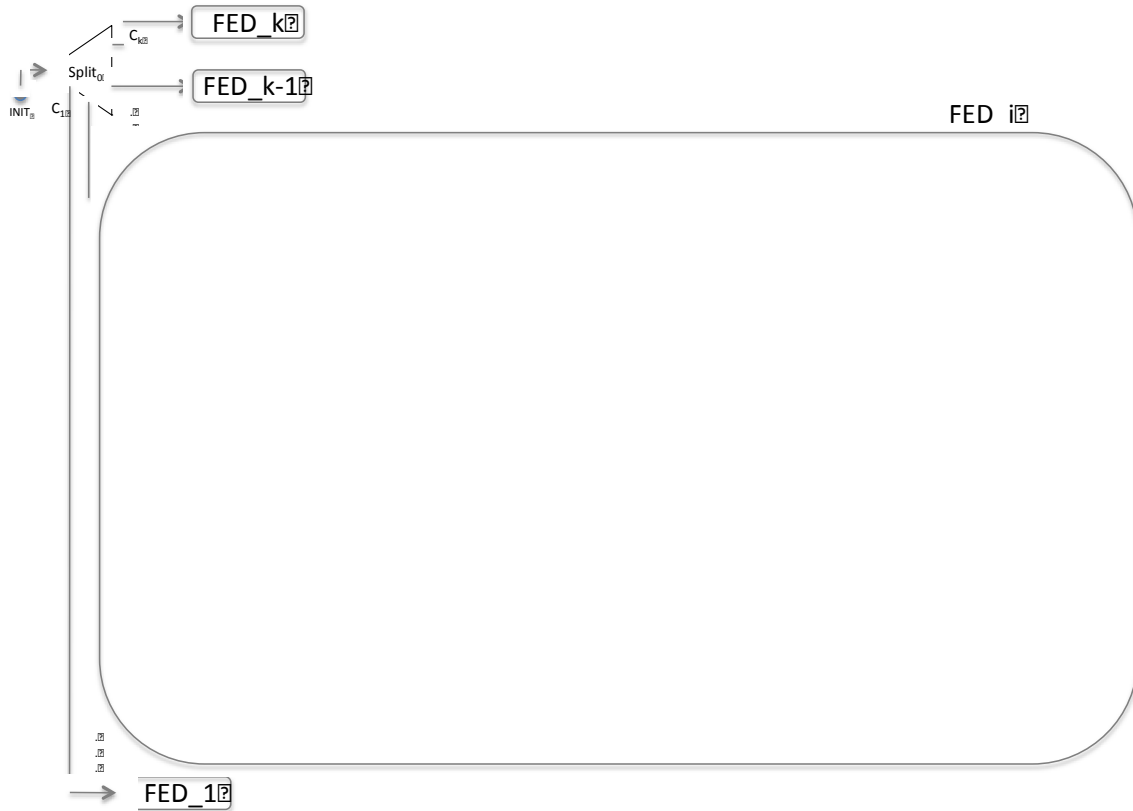


FIGURE 4: View of the federation PM with details of the i -th federate.

It is assumed that the *conservative-time management* HLA option is used, in other words, no federate advances logical time except when it can be guaranteed not to receive any events in its past. If we also assume *zero lookahead* (actually HLA does not accept a null value for lookahead and thus a very small value is given to this parameter) there is guarantee that federates do not receive events in the past and thus that they are fully synchronized.

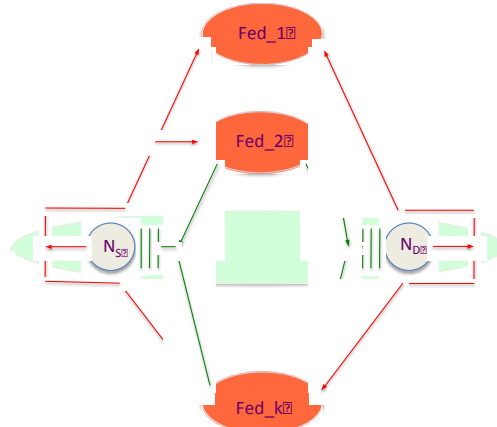


FIGURE 5: View of the federation PM including the communication networks.

For such choices, the federates will not process events in parallel and parallelism will only be found when federates include intrinsically parallel portions of LS. If this holds, a positive speedup will be obtained when transforming the LS into its DS version.

The computation performed by Fed_i is carried out by jobs of various classes that circulate in its PM, namely:

- Class C_i jobs
- Class C_i^D jobs
- Class C_{HLA} jobs
- Class C_{RL} jobs
- Class C_{RL}^D jobs
- Class C_{RC} jobs
- Class C_{RC}^D jobs

The only jobs that consume CPU execution time are C_i and C_{RC} .

The class C_i job¹ simulates the so-called *federate main thread* [8], performing LEX and HLA-R computations.

The class C_{RC} job¹ simulates the so-called *federate RTI callback* [8], performing HLA-R-Ex computations.

The class C_i^D job is a job derived from C_i and holding the data payload to be forwarded to **Fed_x** through network N_D , when the RTI-Ack arrives from **Fed_x** (see the AND_{ix}^D node). A class C_{HLA} job is a job derived from C_i and holding the synch-message to be forwarded to **Fed_x** through network N_S . A class C_{RL} job represents the so-called *federate request listener thread*, waiting for synch-messages from **Fed_x** (see the AND_{xi} node). A class C_{RL}^D job is the federate request listener thread, waiting for data messages from **Fed_x** HLA-R-Ex computations. A class C_{RC}^D job is the federate request callback thread holding the data payload coming from **Fed_x** and to be used by the C_i job class.

The main thread C_i enters the $Split_1$ node and yields three outgoing jobs: C_i itself again, C_{RL} and C_{RL}^D . The job of class C_i enters the CPU processing queue from $Split_1$ and circulates in the model (in a way that we shall soon illustrate), so iteratively re-entering the CPU processing queue coming from the AND_{MTi} node. The job of class C_{RL} , instead, enters the OR_{xi} node and from here the AND_{xi} and waits for a synch-message from **Fed_x** to generate a C_{RC} job, which through the $Split_3$ produces both a new C_{RL} job (that waits for future synch messages) and the C_{RC} itself again that enters the CPU processing queue. The same logic applies to the C_{RL}^D job coming from $Split_1$, which enters the OR_{xi} and the AND_{xi}^D nodes waiting for a data-message from **Fed_x**. The C_{RC}^D job outgoing $Split_4$ does not enter the CPU processing queue directly but merges itself with the C_i circulating main thread through the merge node M_i . As said above, the C_i job entering the CPU performs LEX and HLA-R computations, while the C_{RC} job performs HLA-R-Ex computations.

The job leaving the CPU can be a C_{RC} or a C_i job.

- In case the job leaving the CPU is a C_{RC} job, the $Classifier_i$ node forwards it to the router R_3 , which sends the job to the AND_{MTi} node in case the *synchronous HLA service*² is invoked [8].

¹ The service time *parameters* for such a job class (distribution, mean $E(t_{CPU})$ and variance) can be obtained [11] basing on the model of the software run by the Fed_i CPU, and on the CPU capacity.

² In other words, when the federate needs to wait for a RTI callback (C_{RC}), in the case of invocation of a *Time Advance Request Service*, or of a *Next Message Request Available* service.

Otherwise, the C_{RC} job has no effects and is absorbed by the sink node. If directed to the AND_{MTi} node, the C_{RC} job gives consensus to the circulation of the main thread C_i , which thus re-enters the CPU processing queue.

- In case instead, the job leaving the CPU is a C_i job, the $Classifier_1$ directs it to the R_1 router, which sends the job to $Classifier_2$ if the simulation is not ended ($1-p_{QUIT}$). Here, if C_i contains a data-message, a C^D_i job is produced which enters the AND^D_{ix} node, and waits for the RTI-Ack from **Fed_x** in order to be forwarded to **Fed_x** through network N_D . If instead, the outcome from $Classifier_2$ is a no-data message $\bar{C}_i$ this enters the $Split_2$ node and yields a C_{HLA} job (holding a synch-message to be forwarded to **Fed_x** through network N_S) and again a circulating main thread C_i , which (in case a synchronous HLA service is invoked (p_{SYNC})) reaches the aforementioned AND_{MTi} node to iterate the main thread circulation. In case, instead, of *no-synchronous HLA service*³ ($1-p_{SYNC}$), the C_i job does not need the AND_{MTi} consensus to iterate the main thread circulation, and returns directly to the CPU processing queue.

In summary, synchronization and data messages that Fed_i exchanges with other federates **Fed_x** are enqueued in front of the Fed_i Host CPU to be processed.

Considered that in the *publish/subscribe* assumption Fed_i interacts with all remaining $k-1$ federates, the message flow arriving into the queue of the Fed_i Host CPU scales-up with the dimension k of the federation.

Another element that may increase the message flow into the CPU queue is the use of lookahead. Indeed, the frequency of the synchronization messages exchanged between federates per wall clock time-unit may be affected by the value of the lookahead parameter set by the user.

Such a parameter, however, assumes significant values only in some kind of distributed simulation models. So, in many cases, the federate PM needs not to model the rise of synch messages due to lookahead. This is the case of the Fig.4 model, which however can be easily extended to include lookahead synch messages generators, if needed.

Let us conclude this Section by pointing out that in building the Fed_i model we did not make any mention of the simulated system Σ . This is since the federate model we introduce in the paper is independent from Σ , i.e. it is valid for any Σ . In other words, the paper model can be used for any HLA-based simulation. Only its parameters may depend on Σ , as better seen in Sect.4.

3.2 The network performance model

A further model is necessary to answer the second Fig.2 question of the LS/DS decision procedure: "Exists alternative N_S (or N_D)?". The needed model is the model of the computer network connecting the federation hosts. By use of such a model, the "Exists alternative N_S (or N_D)?" question can be answered by making what-if and sensitivity analysis of the various network components (LANs, GWs, WAN, etc.) of both N_S and N_D .

As said above, the entire federation PM consists of a set of Fed_i sub-models (as in Fig.4) and of the N_S and N_D network models for communication between federates (as in Fig.5). Such networks are used in common by all federates and thus synch and data-messages will be enqueued in front of N_S and N_D as shown in Fig.5. Network N_S will thus introduce a ΔN_S delay for the synch-messages and similarly the N_D a ΔN_D delay for the data-messages.

³ In other words, in case of invocation of *Send Interaction* service.

In other words, when Fed_i sends a synch-message to **Fed_x** through N_S, the message reaches **Fed_x** after a time that can be calculated by evaluating the ΔN_S introduced by N_S.

The evaluation of ΔN_S requires knowledge of the detailed model of N_S (and similarly for N_D). The generic network architecture we shall assume is illustrated in Fig.6 and consists of:

- A set of LANs (LAN₁, ..., LAN_k) where LAN_i is the LAN to which the Fed_i host is connected.
- A set of GATEWAYS (GW₁, ..., GW_k) where GW_i is the gateway that connects LAN_i to the WAN.
- The WAN communication backbone.

Fig.7 gives the EQN performance model of such a network assuming the TCP/IP protocol is used.

The interaction between the Fed_i Host and the **Fed_x** Hosts is based on message exchanges carried out by packet flows over the various components of the network with the WAN being a X.25 packet switching network. The packet flow involves several technologies: The LAN₁ through LAN_k technologies (e.g.: Ethernet, Token Ring, etc.) the Gateways technology and the X.25 WAN technology.

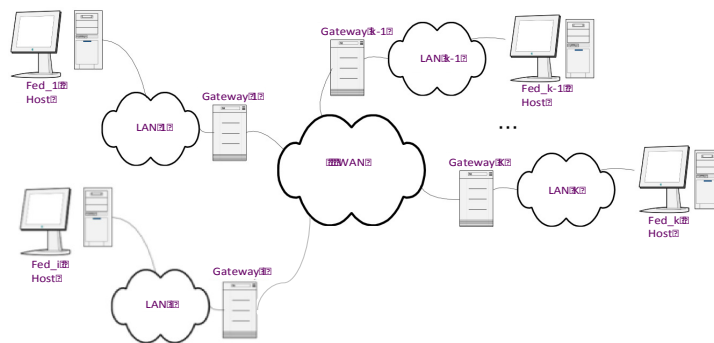


FIGURE 6: View of the network architecture.

The communications between Fed_i Host and **Fed_x** Hosts are based on three basic mechanisms (m1, m2, m3):

- (m1) protocol conversion, from the transport level protocol TCP, to the network level protocol IP, to the data-link level and physical level protocols (and vice versa), in either direction from Fed_i Host to **Fed_x** Hosts, with the IP to X.25 protocol conversion (and vice versa) at the gateway level,
- (m2) packet fragmentation and re-assembly at many protocol conversion interfaces,
- (m3) window-type flow control procedure operated at transport level by protocol TCP for a fixed window size of value C (for the sake of simplicity no varying window sizes are considered, nor the use of congestion-avoidance algorithms).

In the Fed_i-to-**Fed_x** flow, illustrated in Fig.7, the packets are originated by the Fed_i Host application level in TCP format and then translated into IP format by the Fed_i Host network level to enter LAN_i. From LAN_i they exit in LLC/MAC802.5 format to enter the GW_i fragmentation section (FRAG) that fragments them into X.25 format to be transferred by the transfer section (TRANS) to the WAN. Vice versa for the **GW_x**, where X.25 packets are re-assembled by its reassembly section (REA) into LLC/MAC802.3 format to be forwarded to the LAN_x by the Transfer section (TRANS).

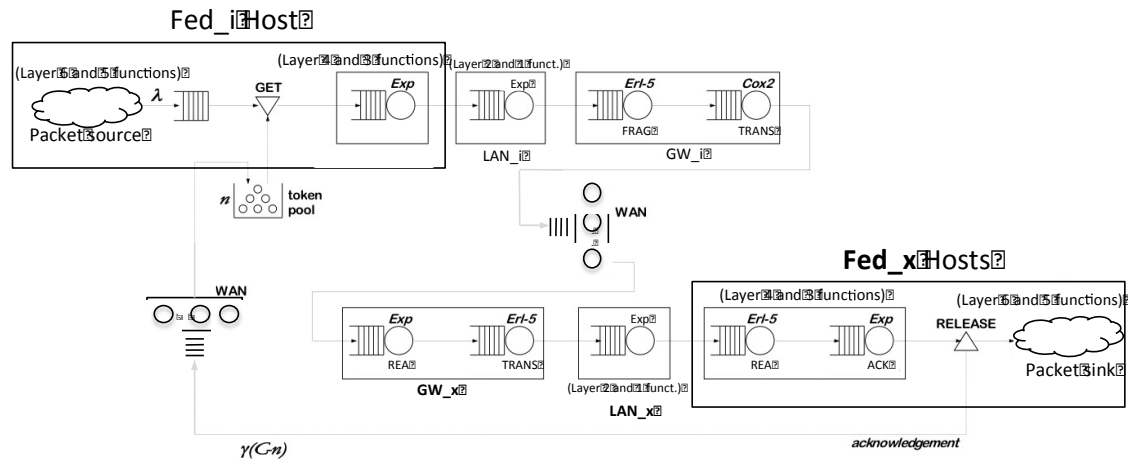


FIGURE 7: View of the network performance model.

LAN_x transfers such frames to **Fed_x** Hosts, which in turns re-assembles them into IP packets and then into TCP frames, in the re-assembly section (REA). The received TCP frames are finally passed to the application level and are acknowledged by the sending of an ACK packet back to Fed_i Host, by the acknowledger section (ACK). The token pool [9] is introduced to represent the window-type flow control procedure implemented by the TCP between the source and the sink (see later).

In summary, the considered network consists of various subsystems each of different complexity, each sub-system in turn consisting of various subsystems of various complexities. Each LAN, for example, is in itself a complex network (not illustrated here), and similarly the WAN.

Producing a model of the network with all details of each LAN, all details of the WAN, etc., could yield so many components and details to make the model very difficult to handle, and its evaluation very time-consuming.

In order to obtain a tractable model, a hierarchical hybrid approach [11] can be foreseen. To this scope, three abstraction levels are introduced:

Level-1 abstraction: At this level the separable sub-systems are identified according to decomposability theory [10], and studied in isolation. Assume the LANs are separable sub-systems. In this case they can be preliminarily studied separately from the rest of the network, then evaluated to obtain their end-to-end delay, and finally substituted in the network model by equivalent service centers whose service times are the end-to-end delays obtained above. If separation is possible, each LAN model (that normally consists of a very large number of service centers) is collapsed into a single equivalent center. The decomposability conditions for the LANs, can be verified⁴, and are respected in the considered model.

In conclusion, at this level sub-systems LAN₁, through LAN_k are separately evaluated to obtain their equivalent end-to-end delay and are each replaced by a single equivalent center, as illustrated in Fig.7. The evaluation gives the distributions (e.g.: exponential in the Fig.7 case) of the LAN equivalent service time and its *parameters* (mean, variance, etc.), calculated basing on the model of the software run by the LAN and the capacity of the hardware [11].

⁴ Generally speaking, the decomposability condition holds when the events that are injected from the external systems into the separable sub-system (i.e. from Fed_i Host into LAN_i and from GW_x into LAN_x) are very rare with respect to the events that take place internally to the sub-system. This can be also empirically verified by comparing the average service rates of the external systems with the ones internal to the LANs. In the specific case, the formers are orders of magnitude smaller than the latters.

Note that the evaluation of the equivalent service time of each LAN may take into consideration the fact that there might exist many Hosts on the some LAN and that some of them might not be part of the federation.

To complete *Level-1* network model the GW_i, the **GW_x**, the WAN and the Hosts are also to be modeled.

The GW_i can be shown to consist of two stages (of Erl-5 and Cox-2 distribution, as in the illustration), with *parameters* (mean, variance, etc.) again calculated basing on the model of the software run by the gateway and its hardware capacity [11]. A similar work is done for **GW_x**, which is shown to consist of two stages (Exp and Erl-5) and relating *parameters*. As far as the WAN is concerned, this can be shown to be globally modeled by an Exp multi-server center, as illustrated in Fig.7 and relating *parameter*.

Finally, the two Hosts are modeled as being each divided into two sections to represent the division of work between the Layer 6 and 5 OSI protocol functions and the Layer 4 and 3 functions.

Level-2 abstraction: At this level, the window-type flow control procedure operated at transport level by the TCP protocol is modeled on the simplified network obtained at Level.1.

In order to represent such a flow control between the exit of the first section of Fed_i Host and the entrance of the first section of **Fed_x** Hosts, the so-called passive queue [9] is used, consisting of a token pool with the GET and RELEASE nodes. For a window size C, the pool consists of C tokens, and so up to C consecutive TCP frames can get a token at the GET node and be admitted. Non-admitted packets are enqueued in front of the GET node. On the other hand, each leaving packet releases its token at the RELEASE node, thus allowing another packet to enter. When data transfer takes place in the opposite direction, the GET node with its queue takes the place of the RELEASE node, and vice versa.

The Level-2 model is however still too complex to be evaluated in closed form, and thus its evaluation is made by simulation. The evaluation will yield the acknowledgement throughput [11] (or number of returned ACKs per time unit), denoted as $\gamma(C-n)$, where C is the chosen window size, n the number of acknowledged packets and (C-n) the number of still unacknowledged ones in the network.

Level-3 abstraction: At this level the entire network N_s (or N_D) is replaced by a single equivalent center (see Fig.8) whose service rate is the ACK throughput $\gamma(C-n)$ calculated at abstraction level 2. In other words, the entire network N_s (or N_D) is now seen as a single server system with arrival rate λ (the packets arrival rate from the user application in Fed_i Host) and mean service time *parameters* of value E(t_s) depending on the $\gamma(C-n)$ throughput [11], namely:

$$E(t_s) = \begin{cases} 1/\chi(i), & 0 \leq i \leq C \\ 1/\chi(C), & i > C \end{cases} \quad (2)$$

with *i* the number of packets in the queue, including the server. Such a model is of the M/M/1 type with state-dependent service time (i.e. dependent on the number n of packets in the center), which can be evaluated according to standard procedures⁵[9]. Its response time gives the network ΔN_s

⁵The Poisson assumption (M) for the arrival process of mean λ is a reasonable assumption for the packets flow from the client user application. The exponential assumption (M) for the network service time is a pessimistic assumption that introduces a security factor, which can however be replaced by a general service time assumption (G) by introducing the coxian approximation.

(or ΔN_D) delay to be used for N_S (or N_D) in the Fig.5 PM (in other words, the N_S or N_D equivalent service times to be used in the federation PM).

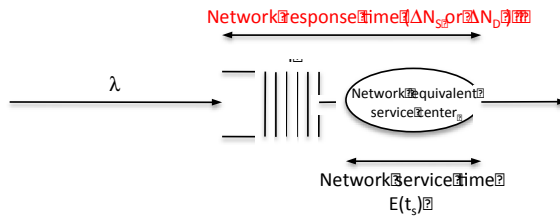


FIGURE 8: Synthetic model of N_S (or N_D)

As said at the beginning of this Section, by use of the N_S (or N_D) network model, what-if and sensitivity analysis can be performed of various network components (LANs, GWs, WAN, etc.) or of various functions (window size C) to answer the Fig.2 question “exists alternative N_S (or N_D)?” of the LS/DS decision procedure.

5. The OMNet++ version of the PM(DS(Σ)) and model parameterization

To perform an example prediction of the DS(Σ) execution time (T_{DS}) to be used in the Fig.2 decision procedure, we developed the OMNet++ simulation version of the Fig.5 model for a $k=2$ federates case (Fed_1 and Fed_2). Only the Fed_1 part (Fig.4) and the N_S and N_D nodes are shown in Fig.9. As said above, the model structure is valid for any system Σ and only its parameters, illustrated in Tab.1, (i.e., the CPU service time, the N_D and N_S service times, the p_{QUIT} and p_{SYNC} routing probabilities) may change with Σ .

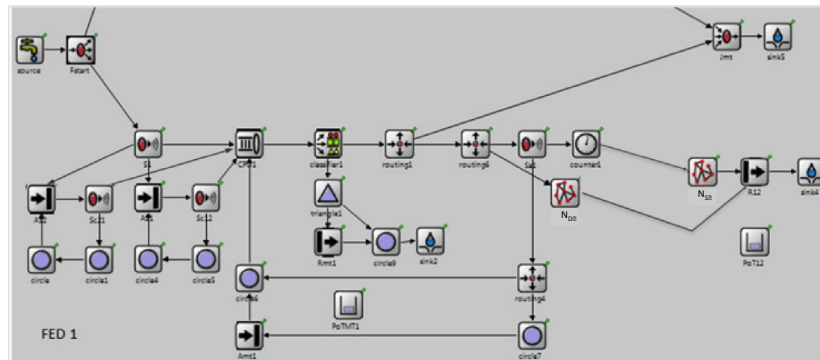


FIGURE 9: OMNet++ simulation version of the PM(DS(Σ)) Fed_i.

The derived parameters for a Σ example case [12] are illustrated in Tab.1.

	Distribution	Parameters
Fed_i Host CPU service time t_{CPU} ($i=1, 2$)	positive truncated-Normal	$E(t_{CPU}) = 10\text{ms}$ (Scen.A) $E(t_{CPU}) = 500\text{ms}$ (Scen.B) $\sigma^2(t_{CPU}) = 1$
N_S, N_D service time t_s	k -Pareto, $k=4$	$E(t_s) = 21\text{ms}$
Routing parameters	p_{QUIT}	0,001 (Fed_1); 0,001 (Fed_2)
	p_{SYNC}	0,82 (Fed_1); 0,74 (Fed_2)

TABLE 1: Model parameters for a two-federate DS.

As can be seen from the table, there exist three types of parameters: the Fed_i Host CPU parameters, the N_S and N_D parameters and the routing parameters (p_{QUIT} and p_{SYNC}).

The Fed_i Host CPU service time parameters vary with the job class (C_i or C_{RC}) and are derived from the CPU capacity and the Fed_i software run by the CPU, as seen in Sect.3.1. For the sake of simplicity, in this example a common mean $E(t_{CPU})$ of 10ms or 500ms (for Scenarios A and B respectively, see later) is chosen for both classes.

The parameters for the N_D and N_S networks are instead derived from the software run by the network components and their capacity, as seen in Sect.3.2.

The routing parameters p_{QUIT} and p_{SYNC} , finally, can be derived from measurements on $LS(\Sigma)$, in particular, by counting the number of events $n_{IntEvents}$, $n_{disEvents}$, $n_{disToIntEvents}$ which respectively denote the number of local events (internal events), the number of events sent from a potential Fed_1 to a potential Fed_2, and the number of events received from a potential Fed_2. Such counting can be easily performed collecting the number of LS events in a simulation experiment for a given hypothetical LS partitioning into two federates. Indeed, it is possible to be convinced [12] that under the *conservative time-management* assumption, one may write:

$$p_{QUIT} = 1 / n_{cycles}, \quad (3)$$

where n_{cycles} is the number of local-HLA processing cycles. Value n_{cycles} can be estimated by the number of events locally processed within the model partition. More specifically,

$$n_{cycles} = n_{IntEvents} + n_{disToIntEvents} \quad (4)$$

Similarly, under the same assumption, one may write:

$$p_{SYNC} = n_{IntEvents} / (n_{IntEvents} + n_{disEvents}) \quad (5)$$

Basing on the Fed_i Host CPU parameters, the N_S and N_D parameters and the routing parameters, the OMNet++ code simulation model has been run to obtain the T''_{DS} predictions shown in Tab.2. This was carried-out [13,14] in two scenarios A and B: Scenario A being one in which the fraction Q of inherently serial computation was high and Scenario B in which Q was low.

The first column in Tab.2 reports the local simulator execution time T_{LS} . The second column reports the distributed simulator execution time T''_{DS} predicted by OMNet++ simulator of the PM, and the third column the times of the real distributed simulator DS (that was implemented in Java+HLA). Such a column thus provides a validation of the PM results, and shows how the predicted results adequately match the real ones. Note that in Scenario B the execution times are in minutes while in Scenario A they are in seconds. This is since Scenario B is built in a way to yield a high computation-to-communication ratio. In other words, a large amount of computation between communications.

	T_{LS}	PM results (OMNet++ predictions)	PM validation (real DS measurements)
A (high Q)	0.7s	$T''_{DS} = 8.3s$	$T''_{DS} = 8.2s$
B (low Q)	33 min	$T''_{DS} = 12.5 \text{ min}$	$T''_{DS} = 12.0 \text{ min}$

TABLE 2: Execution-time results.

Tab.2 also shows how in the Scenario B the distributed simulator outperforms the local one. Indeed, in such a Scenario the DS execution time (T''_{DS}) is much lower than the LS time (T_{LS}).

Finally by using the expression $S'' = T_{LS}/T''_{DS}$, the results in Tab.2 were used to obtain the speedup results shown in Tab.3.

	PM results
A: High Q	$S = 0.08$
B: low Q	$S = 2.64$

TABLE 3: Speedup results.

This table shows that a quite good speedup ($S'' = 2.64$) is obtained in the B Scenario. In other words, in this case the run-time gain obtained by the parallel execution on two hosts compensates for the data and synch communication overheads. In the scenario A, instead, the parallelism does not yield a sufficient run-time gain to compensate for the overheads, and the resulting speedup ($S=0.08$) is practically irrelevant.

The Tab.2 and 3 results are used by the decision procedure of Fig.2 to decide at design-time whether to remain on the LS version of the simulator or implement its DS version. In case the T''_{DS} execution time are not considered “ok” (see Fig.2), one may either try a new tentative partitioning (to modify the p_{SYNC} parameters, see Sect.3.1) or try alternative networks N_S and N_D of improved capabilities (to modify the $E(t_S)$ parameters, see Sect.3.2). In case no partitioning nor network improvements can be found, one may decide not to implement the DS(Σ).

6. CONCLUSION

The execution time of a Distributed Simulator (DS) depends on 3 interacting factors: the speedup, the synch-communication overhead and the data-communication overhead, due to network delays.

The combination of such 3 factors makes it very hard to predict the advantage of transforming a local version of the simulator (LS) into a distributed version (DS).

A LS/DS decision procedure has been proposed to decide at design-time whether to remain on the LS version of the simulator or carry out the implementation of its DS version. The procedure is guided by a performance model (PM) of the DS. The PM assumes the DS is based on the HLA protocol standard and middleware. The model can be used both to support the LS/DS decision process and to evaluate the representativeness of the DS(Σ) at design-time.

ACKNOWLEDGMENTS

Work partially supported by funds from the FIRB project on “Software frameworks and technologies for distributed simulation”, from the FIRB project on “Performance evaluation of complex systems”, from the University of Rome TorVergata research on “Performance modeling of service-oriented architectures” and from the CERTIA Research Center.

REFERENCES

- [1] D. Gianni, A. D’Ambrogio and G. Iazeolla. “A Layered Architecture for the Model-driven Development of Distributed Simulators”, Proceedings of the First International Conference on Simulation Tools (SIMUTOOLS’08), Marseille, France, pp. 1-9, 2008.
- [2] A. D’Ambrogio, D. Gianni, G. Iazeolla: “A Software Architecture to ease the development of Distributed Simulation Systems”, Simulation-Transaction of the Society for Modeling and Simulation International, Vol. 87, n.9, pp. 819-836, 2011.
- [3] IEEE Std 1516. “IEEE Standard for Modeling and Simulation (M&S) High Level Architecture (HLA) frameworks and rules”, 2000.

- [4] R.M. Fujimoto. "Parallel and Distributed Simulation Systems", John Wiley & Sons 1999.
- [5] A. Park. "Parallel Discrete Event Simulation", College of Computing. vol. PhD: Georgia Institute of Technology, 2008
- [6] OMNeT++ Discrete event simulation v.4.0. User Manual, <http://www.omnetpp.org>.
- [7] Pitch. The Certified Runtime Infrastructure for HLA 1516 – User's guide. <http://www.pitch.se>, 2005.
- [8] F. Kuhl, R. Weatherly, J. Dahmann. "Creating Computer Simulation Systems", Prentice-Hall, 1999.
- [9] S.S. Lavenberg. "Computer Performance Modeling Handbook", Academic Press, New York, 1983.
- [10] P.J. Courtois. "Decomposability: Queueing and Computer System and Applications", Academic Press, 1997.
- [11] A. D'Ambrogio, G. Iazeolla. "Steps towards the Automatic Production of Performance Models of Web-Applications", Computer Networks, n.41, pp 29-39, Elsevier Science, 2003.
- [12] D. Gianni, G. Iazeolla, A. D'Ambrogio. "A methodology to predict the performance of distributed simulation", PADS10 the 24th ACM/IEEE/SCS Workshop on Principles of Advanced and distributed simulation Atlanta May 17-19, 2010.
- [13] G. Iazeolla, A. Gentili, F. Ceracchi. "Performance prediction of distributed simulations", Technical Report RI.02.10, Software Engineering Lab, Dept. Computer Science, University of Roma TorVergata , 2010.
- [14] G. Iazeolla, M. Piccari. "The Speedup in distributed simulation", Technical Report RI.03.10, Software Engineering Lab, Dept. Computer Science, University of Roma TorVergata, 2010.
- [15] L. Chu-Cheow, L. Yoke-Hean, et al. "Performance prediction tools for parallel discrete-event simulation", Proceedings of the thirteenth workshop on Parallel and distributed simulation Atlanta, Georgia, United States: IEEE Computer Society, 1999.
- [16] J. Liu, D. Nicol, et al. "A Performance prediction of a parallel simulator", Proceedings of the thirteenth workshop on Parallel and distributed simulation Atlanta, Georgia, United States: IEEE Computer Society, 1999.
- [17] R. Ewald, D. Chen, et al. "Performance Analysis of Shared Data Access Algorithms for Distributed Simulation of Multi-Agent Systems", Proceedings of the 20th Workshop on Principles of Advanced and Distributed Simulation: IEEE Computer Society, 2006.
- [18] R. Ewald, J. Himmelsbach, et al. "A Simulation Approach to Facilitate Parallel and Distributed Discrete-Event Simulator Development", Proceedings of the 10th IEEE international symposium on Distributed Simulation and Real-Time Applications: IEEE Computer Society, 2006.
- [19] K. S. Perumalla, R.M. Fujimoto, et al. "Performance prediction of large-scale parallel discrete event models of physical systems", Proceedings of the 37th conference on Winter Simulation Orlando, Florida: Winter Simulation Conference, 2005.
- [20] P. Teo, S. J. Turner, Z. Juhasz. "Optimistic Protocol Analysis in a Performance Analyzer and Prediction Tool", PADS '05 - Proceedings of the 19th Workshop on Principles of Advanced and Distributed Simulation Pages 49 – 58, 2005.

Reliability Evaluation of Stochastic-Flow Network under Quickest Path and System Capacity Constraints

M. R. Hassan

*Computer Science Branch, Department of Mathematics,
Faculty of Science, South Valley University,
Aswan, Egypt.*

m_r_hassan73@yahoo.com

Abstract

This paper proposes an algorithm to evaluate the probability that d units of data can be sent from the source node to the sink node through a valid group of paths on a network. The conditions of transmission are such that the transmission time of each path belonging to this valid group of paths does not exceed the specified upper bound T , and that the maximal capacity of that path is not less than the specified lower bound C_s (the required system capacity). Such a probability, which is called the system reliability, is denoted by $R_{(d,C_s,T)}$. Based on minimal paths, the algorithm generates all the lower boundary points for (d,C_s,T) , and the system reliability can then be calculated in terms of all the lower boundary points for (d,C_s,T) by applying the inclusion-exclusion rule.

Keywords: Time and capacity constraints, minimal path, stochastic-flow network, system reliability.

1. INTRODUCTION

The system reliability of a flow network R_d is the probability that the maximum flow of the network is not less than a given demand d (which is known as a single commodity). For the case when each arc has several capacities and may fail, [1] presented an algorithm to evaluate R_d in terms of minimal pathsets. In [2], Lin presents an algorithm to evaluate R_d for the case when both the arc and the node have several capacities and may fail. A flow network with two commodities has been studied in [3].

The system reliability of a flow network $R_{(d,C)}$ is defined as the probability that d units of flow can be transmitted from the source node to the sink node, such that the total transmission cost is less than or equal to C ; this can be computed in terms of minimal path vectors to level (d,C) (named (d,C) -MPs). In [4], Lin presented an algorithm to generate all (d,C) -MPs of such a system for each level (d,C) in terms of minimal path sets, considering the fact that each arc has several capacities and may fail. For the case when each node and arc having a designated capacity has a different lower level owing to various partial and complete failures, Lin [5] proposed an efficient algorithm, based on minimal paths, that generates all lower boundary points for (d,C) . The system reliability can then be calculated in terms of all lower boundary points for (d,C) by applying the inclusion-exclusion rule. The system reliability $R_{(d,C)}$ of a multicommodity flow network has been studied in [6] and [7].

The system reliability of a flow network $R_{(d,T)}$ is defined as the probability that d units of data can be sent from the source to the sink through a stochastic-flow network within T units of time. Based on minimal paths, Lin presented an algorithm to calculate $R_{(d,T)}$ [8].

The idea of considering path capacity and the required system capacity C_s in the reliability evaluation is referred to in [9]. The authors stated that a system is good if and only if it is possible to successfully transmit the required capacity from the source node to the sink node.

In this paper, we will extend the idea of using C_s to a flow network. We suppose that the system has a limited capacity value C_s , and we want to send d units of data within T units of time. Thus, we have a new measure of the system reliability of a flow network denoted by $R_{(d,C_s,T)}$. This new measure is defined as the probability that d units of data can be sent from the source node to the sink node through a valid group of paths on the network within the transmission time (T) under the required system capacity (C_s) constraints.

2. NOTATIONS and ASSUMPTIONS

2.1 Notations

$G(A, N, C)$ A stochastic-flow network with a set of arcs $A = \{a_i | 1 \leq i \leq n\}$, a set of nodes N , and

$C = \{C_1, C_2, \dots, C_n\}$ with C_i (an integer) being the maximum capacity of each arc a_i .

X Capacity vector; $X = (x_1, x_2, \dots, x_n)$.

MPs Minimal paths.

mp_j Minimal path no. j ; $j = 1, 2, \dots, m$.

l_i The lead time of arc a_i .

C_s The required system capacity.

$R_{(d, C_s, T)}$ The system reliability for a given demand d under the constraints of T and C_s .

2.2 Assumptions

- 1- The capacity of each component a_i is an integer-valued random variable that takes values $0 < 1 < 2 < \dots < M^i$ according to a given distribution.
- 2- The flow in G must satisfy the so-called flow-conservation law.
- 3- The capacities of different components are statistically independent.

3. AN ALGORITHM for COMPUTING $R_{(d, C_s, T)}$

3.1 Definition of lower boundary points for (d, C_s, T) .

If X is a minimal capacity vector such that the network can send d units of data from the source to the sink within T units of time under a system capacity C_s , then X is a lower boundary point for (d, C_s, T) .

3.2 Generate all Lower Boundary Points for (d, C_s, T) .

In the following steps, for each minimal path $mp_j = \{a_{j1}, a_{j2}, \dots, a_{jn}\}$, we show how to find the minimal capacity vector $X^j = (x_1, x_2, \dots, x_b, \dots, x_n)$ such that the network sends d units of data within T units of time under a maximum system capacity C_s .

1. For all mp_j , examine the path capacities C_{mpj} as

$$C_{mpj} = \min\{C_i | a_i \in mp_j\}, j = 1, 2, \dots, m. \quad \dots(1)$$
2. For all mp_j , calculate the transmission time of the path T_j as

$$T_j = \sum_{i=1}^n \{l_i | a_i \in mp_j\} + \lceil d/C_s \rceil \quad \dots(2)$$

3. Determine the valid group paths, $V_{mp} = \{mp_j | C_{mpj} \geq C_s \text{ and } T_j \leq T, j = 1, 2, \dots, m\}$.
4. Generate the system capacity vector $X^j = (x_1, x_2, \dots, x_b, \dots, x_n)$ for each mp_j that belongs to V_{mp} as follows:

$$x_i = \begin{cases} C_s & \text{if } a_i \in mp_j \\ 0 & \text{otherwise} \end{cases} \quad \dots(3)$$

where x_i is an element of X^j .

Lemma 1. If X is a lower boundary point for (d, C_s, T) , then the system capacity under X is greater than or equal to C_s , and the minimum transmission time under X is less than T .

Lemma 2. The set generated by the algorithm 3.2— $X^1, X^2, \dots, X^q$ —is the set of lower boundary points for (d, C_s, T) .

3.3 Evaluate $R_{(d, C_s, T)}$

If $X^1, X^2, \dots, X^q$ are the collection of all (d, C_s, T) -mp, then the system reliability $R_{d, C_s, T}$ is defined by

$$R_{(d, C_s, T)} = \Pr\left\{\bigcup_{i=1}^q \{Y | Y \geq X^i\}\right\} \quad \dots(4)$$

where $\Pr\{Y\} = \Pr\{y_1\} \cdot \Pr\{y_2\} \cdot \dots \cdot \Pr\{y_n\}$. We will use the inclusion-exclusion rule presented in [10] to calculate $R_{d, C_s, T}$ as follows:

If $A_1 = \{Y|Y \geq X^1\}$, $A_2 = \{Y|Y \geq X^2\}$, ..., $A_q = \{Y|Y \geq X^q\}$, then apply the inclusion-exclusion rule to calculate $R_{d,Cs,T}$ using the following relationship:

$$R_{(d,Cs,T)} = \sum_i \Pr\{A_i\} - \sum_{i \neq j} \Pr\{A_i \cap A_j\} + \sum_{i \neq j \neq k} \Pr\{A_i \cap A_j \cap A_k\} - \dots + (-1)^{q-1} \Pr\{A_1 \cap A_2 \cap \dots \cap A_q\} \quad \dots(5)$$

4. AN ILLUSTRATIVE EXAMPLE

Here, we use the network in Fig. 1 that was studied in [8]. This network has five nodes and eight arcs, which are numbered from a_1 to a_8 . The capacity and lead time of each arc are shown in Table 1.

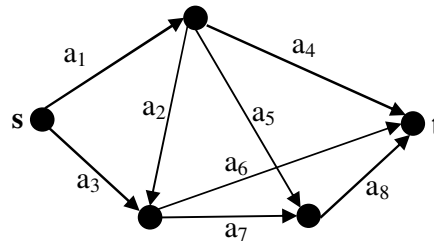


FIGURE 1: Computer network

There are six minimal paths: $mp_1 = \{a_1, a_2\}$, $mp_2 = \{a_1, a_5, a_8\}$, $mp_3 = \{a_1, a_2, a_6\}$, $mp_4 = \{a_1, a_2, a_7, a_8\}$, $mp_5 = \{a_3, a_6\}$, and $mp_6 = \{a_3, a_7, a_8\}$. Given $d = 8$ and $T = 9$, Tables 2, 3, and 4 summarize the values of C_{mpj} and T_j for each path mp_j for the different values of C_s , using the algorithm 3.2. Also, below each table we show the value of V_{mp} and the corresponding X vectors, as well as the system reliability $R_{d,Cs,T}$.

Arc	Capacity	Probability	Lead time	Arc	Capacity	Probability	Lead time
a_1	3	0.80	2	a_6	4	0.60	2
	2	0.10			3	0.20	
	1	0.05			2	0.10	
	0	0.05			1	0.05	
a_2	2	0.80	1	a_7	5	0.55	2
	1	0.10			4	0.10	
	0	0.05			3	0.10	
	0	0.05			2	0.10	
a_3	1	0.85	3	a_8	1	0.10	1
	0	0.10			0	0.05	
	0	0.05			0	0.05	
a_4	1	0.90	3	a_8	3	0.80	1
	0	0.10			2	0.10	
a_5	1	0.90	1	a_8	1	0.05	1
	0	0.10			0	0.05	

Table 1: Arc data for Fig. 1

mp_j	C_{mpj}	Transmission Time T_j	mp_j	C_{mpj}	Transmission Time T_j
mp_1	1	13	mp_4	3	14
mp_2	1	12	mp_5	2	13
mp_3	3	13	mp_6	2	14

Table 2: Values of C_{mpj} and T_j when $C_s = 1$

According to Table 2, there is no path satisfies constraints of T and C_s . So, $V_{mp} = \emptyset$, i.e., the $X^1, X^2, \dots, X^6$, does not exist. Therefore, $R_{(8,1,9)} = 0$.

mp_j	C_{mpj}	Transmission Time T_j	mp_j	C_{mpj}	Transmission Time T_j
mp_1	1	9	mp_4	3	10
mp_2	1	8	mp_5	2	9
mp_3	3	9	mp_6	2	10

Table 3: Values of C_{mpj} and T_j when $C_s = 2$

According to Table 3, $V_{mp} = \{mp_3, mp_5\}$ because C_{mp3} is greater than C_s and C_{mp5} is equal to C_s . Furthermore, both T_3 and T_5 are equal to T . Thus, we have $X^3 = (2\ 2\ 0\ 0\ 0\ 2\ 0\ 0)$ and $X^5 = (0\ 0\ 2\ 0\ 0\ 2\ 0\ 0)$. Let $A_3 = \{Y|Y \geq X^3\}$ and $A_5 = \{Y|Y \geq X^5\}$. Then, the system reliability $R_{(8,2,9)} = \Pr\{A_3 \cup A_5\} = 0.89145$, using the inclusion-exclusion rule, where

$$\begin{aligned}
\Pr\{A_3\} &= \Pr\{Y \geq (2\ 2\ 0\ 0\ 0\ 2\ 0\ 0)\} \\
&= \Pr\{x_1 \geq 2\} \times \Pr\{x_2 \geq 2\} \times \Pr\{x_3 \geq 0\} \times \Pr\{x_4 \geq 0\} \times \Pr\{x_5 \geq 0\} \times \Pr\{x_6 \geq 2\} \times \Pr\{x_7 \geq 0\} \times \Pr\{x_8 \geq 0\} \\
&= 0.90 \times 0.90 \times 1 \times 1 \times 1 \times 0.90 \times 1 \times 1 = 0.729 \\
\Pr\{A_5\} &= \Pr\{Y \geq (0\ 0\ 2\ 0\ 0\ 2\ 0\ 0)\} \\
&= \Pr\{x_1 \geq 0\} \times \Pr\{x_2 \geq 0\} \times \Pr\{x_3 \geq 2\} \times \Pr\{x_4 \geq 0\} \times \Pr\{x_5 \geq 0\} \times \Pr\{x_6 \geq 2\} \times \Pr\{x_7 \geq 0\} \times \Pr\{x_8 \geq 0\} \\
&= 1 \times 1 \times 0.95 \times 1 \times 1 \times 0.90 \times 1 \times 1 = 0.855 \\
\Pr\{A_3 \cap A_5\} &= \Pr\{Y \geq (2\ 2\ 2\ 0\ 0\ 2\ 0\ 0)\} \\
&= \Pr\{x_1 \geq 2\} \times \Pr\{x_2 \geq 2\} \times \Pr\{x_3 \geq 2\} \times \Pr\{x_4 \geq 0\} \times \Pr\{x_5 \geq 0\} \times \Pr\{x_6 \geq 2\} \times \Pr\{x_7 \geq 0\} \times \Pr\{x_8 \geq 0\} \\
&= 0.90 \times 0.90 \times 0.95 \times 1 \times 1 \times 0.90 \times 1 \times 1 = 0.69255
\end{aligned}$$

mp_j	C_{mpj}	Transmission Time T_j	mp_j	C_{mpj}	Transmission Time T_j
mp_1	1	8	mp_4	3	9
mp_2	1	7	mp_5	2	8
mp_3	3	8	mp_6	2	9

Table 4: Values of C_{mpj} and T_j when $C_s = 3$

It is clear that $V_{mp} = \{mp_3, mp_4\}$ because both C_{mp3} and C_{mp4} are equal to C_s . Also, T_3 is less than T and T_4 is equal to T .

Then, we have $X^3 = (3\ 3\ 0\ 0\ 0\ 3\ 0\ 0)$ and $X^4 = (3\ 3\ 0\ 0\ 0\ 0\ 3\ 3)$. Let $A_3 = \{Y|Y \geq X^3\}$ and $A_4 = \{Y|Y \geq X^4\}$.

Then, the system reliability is $R_{(8,3,9)} = \Pr\{A_3 \cup A_4\} = 0.5888$, using the inclusion-exclusion rule, where

$$\Pr\{A_3\} = \Pr\{Y \geq (3\ 3\ 0\ 0\ 0\ 3\ 0\ 0)\}$$

$$= \Pr\{x_1 \geq 3\} \times \Pr\{x_2 \geq 3\} \times \Pr\{x_3 \geq 0\} \times \Pr\{x_4 \geq 0\} \times \Pr\{x_5 \geq 0\} \times \Pr\{x_6 \geq 3\} \times \Pr\{x_7 \geq 0\} \times \Pr\{x_8 \geq 0\}$$

$$= 0.80 \times 0.80 \times 1 \times 1 \times 1 \times 0.80 \times 1 \times 1 = 0.512$$

$$\Pr\{A_4\} = \Pr\{Y \geq (3\ 3\ 0\ 0\ 0\ 0\ 3\ 3)\}$$

$$= \Pr\{x_1 \geq 3\} \times \Pr\{x_2 \geq 3\} \times \Pr\{x_3 \geq 0\} \times \Pr\{x_4 \geq 0\} \times \Pr\{x_5 \geq 0\} \times \Pr\{x_6 \geq 0\} \times \Pr\{x_7 \geq 3\} \times \Pr\{x_8 \geq 3\}$$

$$= 0.80 \times 0.80 \times 1 \times 1 \times 1 \times 1 \times 0.75 \times 0.80 = 0.384$$

$$\Pr\{A_3 \cap A_4\} = \Pr\{Y \geq (3\ 3\ 0\ 0\ 0\ 3\ 3\ 3)\}$$

$$= \Pr\{x_1 \geq 3\} \times \Pr\{x_2 \geq 3\} \times \Pr\{x_3 \geq 0\} \times \Pr\{x_4 \geq 0\} \times \Pr\{x_5 \geq 0\} \times \Pr\{x_6 \geq 3\} \times \Pr\{x_7 \geq 3\} \times \Pr\{x_8 \geq 3\}$$

$$= 0.80 \times 0.80 \times 1 \times 1 \times 1 \times 0.80 \times 0.75 \times 0.80 = 0.3072$$

5. DISCUSSION

Algorithm 3.2 needs $O(mn)$ time to generate all lower boundary points for (d, C_s, T) in the worst case, where n is the number of arcs and m is the number of minimal paths. Algorithm 3.3 needs $O(m^2n)$ time to evaluate the system reliability in the worst case [8], using the inclusion-exclusion rule. Therefore, the total time needed by the algorithm is $O(mn) + O(m^2n)$ to calculate the system reliability $R_{(d, C_s, T)}$ in the worst case. In comparison with the algorithm presented in [8] to evaluate reliability under the constraint T , the presented algorithm in this paper needs the same time ($O(mn) + O(m^2n)$) to evaluate reliability under C_s and T constraints.

6. CONCLUSIONS and FUTURE WORK

A new definition of the system reliability of a flow network to a given demand d has been presented, which takes into account both the required system capacity (C_s) and the transmission time (T). In addition, an algorithm has been presented for the calculation of $R_{(d, C_s, T)}$. The algorithm is based on the use of minimal paths to generate all lower boundary points for (d, C_s, T) , and to then calculate the system reliability $R_{(d, C_s, T)}$ using the inclusion-exclusion rule.

Finally, we have illustrated the use of the proposed algorithm by calculating the reliability of a flow network for a given network taken from the literature.

The algorithm has proved to be efficient and may be used to compute the system reliability of a multicommodity flow network.

6. REFERENCES

- [1] JS. Lin, CC. Jane and J. Yuan, "On reliability evaluation of a capacitated-flow network in terms of minimal path sets", *Networks*, 25:131-138, 1995.
- [2] YK. Lin, "A simple algorithm for reliability evaluation of a stochastic-flow network with node failure", *Computers and Operations Research*, vol. 28(13):1277-1285, 2001.
- [3] YK. Lin, "Two-commodity reliability evaluation for a stochastic-flow network with node failure", *Computers & Operations Research*, 29(13):1927-1939, 2002.
- [4] JS. Lin, "Reliability evaluation of capacitated-flow networks with budget constraints", *IEEE Transactions*, 30(12):1175-1180, 1998.
- [5] YK. Lin, "Reliability of a Stochastic-Flow Network With Unreliable Branches and Nodes", Under Budget Constraints, *IEEE Transactions on Reliability*, 53(3):381-387, 2004.
- [6] YK. Lin, "Reliability evaluation for an information network with node failure under cost constraint", *IEEE Transactions on Systems, Man, and Cybernetics, Part A: System and Humans*, 37(2):180-188, 2007.

- [7] YK. Lin, "System Reliability of a Limited-Flow Network in Multicommodity Case", *IEEE Transactions on Reliability*, 56(1):17-25, 2007.
- [8] YK. Lin, "Extend the quickest path problem to the system reliability evaluation for a stochastic-flow network", *Computers and Operation Research*, 30:567-575, 2003.
- [9] K. K. Aqqarwal, Y. C. Chopra, and J.S. Bajwa, "Capacity consideration in reliability analysis of communication system", *IEEE Transactions on Reliability*, 31(2):177-181, 1982.
- [10] J. Xue, "On Multistate System Analysis", *IEEE Transactions on Reliability*, 34(4):329-337, 1985.

Analytical Models for Dimensioning of OFDMA-based Cellular Networks Carrying VoIP and Best-Effort Traffic

Bruno Baynat

*LIP6 - UPMC Sorbonne University - CNRS
4, place Jussieu
75005 Paris, France*

bruno.baynat@lip6.fr

Abstract

The last years have seen an exponentially growing interest for mobile telecommunication services. As a consequence, a great diversity of applications is expected to be supported by cellular networks. To answer this ever increasing demand, the ITU-R defined the requirements that the fourth generation (4G) of mobile standards must fulfill. Today, two especially promising candidates for 4G stand out: WiMAX and LTE. However, 4G cellular networks are still far from being implemented, and the high deployment costs render over-provisioning out of question. We thus propose in this paper accurate and convenient analytical models well-suited for the complex dimensioning of these promising access networks. Our main interest is WiMAX, yet, we show how our models can be easily used to consider LTE cells since both technologies are based on OFDMA. Generic Markovian models are developed specifically for three service classes defined in the WiMAX standard: UGS, ertPS and BE, respectively corresponding to VoIP, VoIP with silence suppression and best-effort traffic. First, we consider cells carrying either UGS, ertPS or BE traffic. Three methods to combine the previous models are then proposed to assume both UGS and BE traffic in the studied cell. Finally, we provide a way to easily integrate the ertPS traffic and obtain a UGS/ertPS/BE model able to account for multiple traffic profiles in each service class while keeping an instantaneous resolution. The proposed models are compared in depth with realistic simulations that show their accuracy. Lastly, we demonstrate through different examples how our models can be used to answer dimensioning issues which would be intractable with simulations.

Keywords: performance evaluation, analytical models, OFDMA, 4G, cell dimensioning, service integration.

1 INTRODUCTION

The fourth generation (4G) of mobile networks is coming to answer the ever increasing demand. Two main candidates for 4G are emerging: WiMAX (Worldwide Interoperability for Microwave Access) and 3GPP LTE (Long Term Evolution). They both propose air interfaces based on OFDMA.

WiMAX leans on the IEEE 802.16 family of standards. The first operative version of IEEE 802.16 is 802.16-2004 (fixed/nomadic WiMAX) [2]. It was followed by a ratification of amendment IEEE 802.16e (mobile WiMAX) in 2005 [3]. A new standard, 802.16m, is currently under definition to provide even higher efficiency. In addition, the consortium WiMAX Forum was found to specify profiles (technology options are chosen among those proposed by the IEEE standard), define an end-to-end architecture (IEEE does not go beyond physical and MAC layer), and certify products (through inter-operability tests).

As for LTE, it has first been introduced in 3GPP Release 8 as a set of improvements to UMTS (Universal Mobile Telecommunications System), a widespread third generation mobile technology. An enhanced version of the LTE technology, named LTE Advanced, is under development in 3GPP Release 10 to achieve even better performance.

A great number of services such as voice, video and web are to be offered by 4G mobile networks. To this aim, several service classes have been defined in the WiMAX standard corresponding to specific QoS needs. Among them are UGS (Unsolicited Granted Service), ertPS (enhanced real-time Polling

Service) and BE (Best Effort). UGS corresponds to applications reserving a part of the resource to obtain a constant bitrate (mostly VoIP without silence suppression). ertPS has been especially designed to carry VoIP with silence suppression traffic. Finally, BE carries elastic traffic generated by web applications. Unlike WiMAX, LTE standards do not specify service classes however different mechanisms are proposed to achieve similar QoS.

Most manufacturers and operators are still under trial phases. As deployment of 4G cellular networks is under way, the need arises for fast and efficient tools used for network design and performance evaluation and able to account for these different services.

Literature on performance evaluation of cellular networks with service integration is constituted of two sets of papers: i) packet-level simulations that precisely implement system details and scheduling schemes; ii) analytical models and optimization algorithms that derive performance metrics at user-level.

Among the latter set of papers, Borst and Hegde presented in [8] an analytical framework for wireless networks supporting a combination of streaming and elastic traffic. The authors proposed to handle the streaming connections first because of their priority over elastic traffic. Then, they used the quasi-stationary assumption first formulated by Delcoigne and al. in [12] to account for the elastic traffic without exponentially increasing the resolution complexity of their modeling. Indeed, this assumption enables to exploit the different time scales of streaming and elastic flows through astute averages to obtain the performance of the elastic flows. Note however that their study do not provide closed-form expressions of the service rates which in some cases can require high computation times or even turn intractable. In addition, this approach does not allow to observe the impact of QoS degradations on the performance parameters.

In [13], Dirani, Tarhini and Chahed designed a simple Markov chain for performance evaluation of a mobile network with one dimension corresponding to streaming and the other to elastic traffic. They too made use of the quasi-stationary assumption from [12] to simplify the resolution of their bidimensional model. However, the variations of the radio channel were not taken into account in their study. To answer this problem, Tarhini and Chahed introduced in [30] an extension of the previous model. They included in the states of the Markov chain the current channel conditions of the streaming and elastic users. They also added transitions corresponding to the probabilities that the channel conditions of a user change before the end of a connection. The resulting multi-dimensional model requires time-consuming numerical resolution, and thus, prevents in-depth dimensioning.

Niyato and Hossain presented in [27] a queuing model for bandwidth allocation in a WiMAX cell. To account for multiple services with specific QoS requirements independently from each other, they introduced a complete partitionning of the available resource among the different types of traffic using linear programming. Yet, setting fixed thresholds for each services can lead to a huge waste of bandwidth in cellular networks where radio conditions and demands of users can widely fluctuate in short amounts of time. Also, they do not account for the adaptive slot scheduling specific to OFDMA-based cellular networks.

Not specific to service integration, generic analytical models for performance evaluation of cellular networks have been proposed in [7], [6], [24]. They are mostly based on multi-class processor-sharing queues with each class corresponding to users having similar radio conditions and subsequently equal data rates. These models implicitly consider that users can only switch class between two successive data transfers. However, in broadband systems like WiMAX and LTE, radio conditions and thus data rates of a particular user can change frequently during a data transfer. In addition, the capacity of a cell may change as a result of varying radio conditions of users.

We presented in [15], [14] novel analytical models dedicated to BE traffic that take into account frame structure, precise slot sharing-based scheduling and channel quality variation of broadband wireless systems. Unlike existing models [7], [6], [24], ours are adapted to the specifics of OFDMA systems. They also offer instantaneous resolution even in multi-traffic cases: closed-form expressions are provided for all performance parameters. Moreover, our approach makes it possible to consider the so-called "outage" situation. A user experiences an outage, if at a given time radio conditions are so bad that it cannot transfer any data and is thus not scheduled.

In this paper, we extend our models to consider UGS and ertPS traffic in addition to BE traffic. We first only consider either UGS, ertPS or BE traffic in the studied cell and provide extensions of the resulting models to take into account multi-profile traffic. Then, we propose methods to combine these models into a UGS/BE model considering a cell with both UGS and BE traffic. Finally, we integrate ertPS traffic into this model and obtain a UGS/ertPS/BE model able to account for each service class while keeping an instantaneous resolution.

To avoid confusion between WiMAX and LTE, we focus on WiMAX throughout this paper. However, we detail how the models presented here can be used for performance evaluation of an LTE cell.

The paper is organized as follows. In Section 2, the modeling assumptions are listed and the specific details on both WiMAX and LTE networks needed to understand our analytic framework are provided. Sections 3, 4 and 5 present our analytical models for UGS, ertPS and BE service classes respectively. As a first step, methods to combine these models into an UGS/BE model are introduced and compared in Section 6. Then, in Section 7, we integrate ertPS traffic into our model and validate the resulting UGS/ertPS/BE model through comparisons with simulations. Lastly, Section 8, provides examples of WiMAX dimensioning processes using this model.

2 MODELING ASSUMPTIONS

Our analytical models stand on several assumptions related to the system, the channel and the traffic. The validity of these assumptions has been thoroughly discussed in our study of BE traffic [15]. To avoid any possible confusion, between WiMAX and LTE specifics, we initially only consider the WiMAX technology. The assumptions shared by the models mentioned in this paper are presented and, wherever required, related particulars of WiMAX systems are specified. In addition, various notations are introduced. Lastly, we explain how these assumptions can be adapted to consider an LTE cell.

2.1 WiMAX Modeling

2.1.1 System

A WiMAX time division duplex (TDD) frame is divided in slots using Orthogonal Frequency Division Multiplexing (OFDM). A slot occupies space both in the time and frequency domains and is the smallest unit of resource that can be allocated to a mobile. A frame is comprised of two parts: one is dedicated to uplink and the other to downlink. Besides, a portion of the frame is used for overhead (e.g., UL_MAP and DL_MAP). The duration T_F of this TDD frame is equal to 5 ms [3].

1. We consider a single WiMAX cell and focus in this paper on the downlink part which is a critical portion of asymmetric data traffic. However, our models can also be used for the dimensioning of the uplink part in a similar way.
2. We assume that there is a mean number of slots available for data transmission in the downlink part of each TDD frame denoted by $\bar{N}_S$. This number is a mean value because the size of the downlink part can vary with the downlink to uplink bandwidth ratio which can be adjusted dynamically over time. The size of the overhead, increasing with the number of multiplexed transmissions per frame, can also affect the number of available downlink slots. However, we consider that the small variations of the amount of overhead are not significant in regard to the size of the downlink part.
3. In the case of UGS (respectively ertPS) traffic, we consider that there is a limit V_{max} (resp. W_{max}) to the number of simultaneous calls accepted in the cell. On the contrary, regarding BE traffic, we assume that all mobiles can simultaneously be in active transfer. As a consequence, any BE connection demand will be accepted and no blocking can occur.

Note that this last assumption implicitly states that no admission control of BE connections is implemented in the system. However, the BE model can be easily modified to account for a system with an

admission controller limiting the number of simultaneous active BE transfers. Indeed, as detailed in [15], we just have to truncate the BE Markov chain accordingly.

2.1.2 Channel

One of the important features of IEEE 802.16e is link adaptation: different modulation and coding schemes (MCS) allows a dynamic adaptation of the transmission to the radio conditions. Many subcarrier permutations defining how the pilot and data subcarriers should be distributed over the subchannels are proposed in the standard. As the number of data subcarriers per slot is the same for all permutation schemes [11], the number of bits carried by a slot for a given MCS is constant. The selection of appropriate MCS is carried out according to the value of signal to interference plus noise ratio (SINR). In case of outage, i.e., if the SINR is too low, no data can be transmitted without error. We denote the radio channel states as: MCS_k , $1 \leq k \leq K$, where K is the number of MCS. By extension, MCS_0 represents the outage state. The number of bits transmitted per slot by a mobile using MCS_k is denoted by m_k . For the particular case of outage, $m_0 = 0$.

The radio link quality in broadband wireless networks like WIMAX is highly variable. As such, the MCS used by a given mobile can change very often.

4. We assume that each mobile sends a feedback channel estimation on a frame by frame basis, and thus, the base station can change its MCS every frame. Since we do not make any distinction between users and consider all mobiles as statistically identical, we associate a probability p_k with each coding scheme MCS_k , and assume that, at each time-step T_F , any mobile has a probability p_k to use MCS_k . Table 1 presents examples of MCS and their associated probabilities.

MCS	bits per slot	probability
Outage	$m_0 = 0$	$p_0 = 0.02$
QPSK-1/2	$m_1 = 48$	$p_1 = 0.12$
QPSK-3/4	$m_2 = 72$	$p_2 = 0.31$
16QAM-1/2	$m_3 = 96$	$p_3 = 0.08$
16QAM-3/4	$m_4 = 144$	$p_4 = 0.47$

Table 1: Channel Parameters.

As a result, our analytical model only depends upon stationary probabilities of using the different MCS and does not explicitly take into account the radio channel dynamics. These probabilities can be accurately obtained from famous statistical fading models such as Rayleigh or Rician models as shown in numerous publications including [20], [25], [31]. In addition, other methods can be considered. By example, in [19] a spatial model is used while in [26] a semi-analytical approach is proposed based on an interpolation of simulation results.

Finally, note that the robustness of this assumption to temporal channel correlation has been validated through extensive simulations considering radio channels with memory as shown in section 7.2.

2.1.3 Traffic

The traffic model is based on the following assumptions.

5. We assume that there is a fixed number N of mobiles sharing the available bandwidth of the cell. The numbers of mobiles generating UGS, ertPS or BE traffic present in the cell are denoted, respectively, by N_{ugs} , N_{ertps} and N_{be} .

Note that operators find finite population models more suitable for the dimensioning of a cell. Indeed, they have means to estimate the number of users they will have to serve in a cell and, as such, consider

those models more appropriate. However, our models can be easily adapted to Poisson arrivals should an infinite population assumption be considered [15].

6. Each of the N mobiles is assumed to generate an infinite length ON/OFF traffic. In the case of UGS and ertPS traffics, an ON period corresponds to a call and is characterized by its duration. In the case of BE traffic, an ON period corresponds to the download of an element (e.g., a web page including all embedded objects). As opposed to UGS and ertPS ON periods, the downloading duration depends on the system load and the radio link quality, so BE ON periods must be characterized by their size. Lastly, in each case, an OFF period corresponds to an idle time independent of the system load and, as such, is characterized by its duration.
7. We assume that UGS and ertPS ON durations, BE ON sizes and each OFF duration are exponentially distributed. We denote by $\bar{t}_{on}^{ugs}$ and $\bar{t}_{on}^{ertps}$ the average durations of UGS and ertPS ON periods (in seconds), by $\bar{x}_{on}^{be}$ the average size of ON data volumes (in bits) and by $\bar{t}_{off}^{ugs}$, $\bar{t}_{off}^{ertps}$ or $\bar{t}_{off}^{be}$ the average durations of OFF periods (in seconds).

Memoryless BE traffic distributions are strong assumptions that have been validated by numerous theoretical results. Several works on insensitivity (e.g., [5], [7], [18]) have shown (for systems fairly similar to the one we are studying) that the average performance parameters are insensitive to the distribution of ON and OFF periods. Moreover, note that we compare in Section 7.2 the results of our analytical model with those of simulations considering a truncated Pareto ON size distribution. These comparisons tend to prove that insensitivity still holds or is at least a very good approximation. Thus, memoryless distributions are the most obvious choice to model BE traffic.

8. We consider absolute priorities between each service class. As such, at each frame, the available slots are first allocated to UGS, then to ertPS and at last to BE connections.

No specific inter-class scheduling is suggested in the WiMAX standard. However, the QoS requirements of each service class suggest that UGS traffic should always be served first, followed by ertPS traffic and finally by BE traffic [23], [28], [29].

2.2 WiMAX to LTE Modeling

We detail here how a 3GPP LTE cell can be easily considered instead of a WiMAX cell as both technologies are based on OFDMA. In addition, we introduce the LTE particulars required to apprehend these modifications whenever needed.

2.2.1 System

LTE frames are organized differently than WiMAX frames. Each LTE frame lasts 10 ms and is divided into 10 subframes [1]. Contrary to WiMAX, a new scheduling is not done at each frame but at each of these subframes. As a consequence, we need, in our modeling, to consider different time intervals between two consecutive schedulings:

- WiMAX: frame duration, $T_F = 5$ ms.
- LTE: subframe duration, $T_F = 1$ ms.

An LTE subframe comprises of resource blocks which, similarly to WiMAX slots, occupy space in both time and frequency domains. Although both LTE and WiMAX use OFDMA as their multiple access scheme, they specify different smallest data allocation units. Indeed, the smallest data unit that can be allocated to an LTE user is formed by a pair of resource blocks (i.e., two consecutive resource blocks in either the frequency or the time domain). So, $\bar{N}_S$, the mean quantity of downlink resource available during one time interval must be adjusted as follows:

- WiMAX: $\bar{N}_S$ is the mean number of downlink slots in a frame.
- LTE: $\bar{N}_S$ is the mean number of downlink resource block pairs in a subframe.

2.2.2 Channel

WiMAX and LTE systems share the same adaptive modulation and coding mechanism so the MCS_k and p_k parameters stay the same. But, the smallest data allocation unit differs for each technology and, as such, the m_k now represent different values whether we consider WiMAX or LTE:

- WiMAX: numbers of transmitted bits per slot using MCS_k .
- LTE: numbers of transmitted bits per resource block pair using MCS_k ,

$$m_k = 2 N_{re} m_{re k}, \quad (1)$$

where N_{re} is the number of downlink resource elements in a resource block, which depends on the antenna settings, and $m_{re k}$ is the number of transmitted bits per resource element using MCS_k .

2.2.3 Traffic

The traffic generated by the mobiles present in the cell is not affected by the considered mobile technology. Thus, no adjustments are required to our traffic modeling assumptions.

Lastly, unlike WiMAX standards, LTE does not specify service classes. The traffic separation is possible instead by defining Evolved Packet System bearers (bearers for short) which provide differential treatment for traffic with differing QoS requirements. The QoS parameters of the bearers enable to consider the same scheduling assumption than we do for our WiMAX models. As such, we keep those assumptions for our modeling of LTE.

3 UGS MODEL

The UGS (*Unsolicited Grant Service*) service class has been designed to support real-time applications periodically generating fixed-size data packets (e.g., VoIP without silence suppression). In this section, we provide the models we use to characterize the UGS traffic of a WiMAX cell. Let us highlight that even though they are based on the famous Engset model [16], we still need to adapt to the specifics of OFDMA and our channel model. Indeed, it would not be possible to formulate crucial performance parameters such as the mean throughput achieved by the mobiles and the mean utilization of the frames otherwise.

We consider in these models that only UGS mobiles (i.e., mobiles generating only UGS traffic) are present in the cell we study. How to model cells with traffic of several service classes is addressed in Sections 6 and 7.

3.1 Mono-Traffic Model

An UGS call corresponds to an utilization of the resource in circuit mode. The reserved bit rate associated to each UGS connection is called Guaranteed Bit Rate (GBR). In a first phase, no distinctions between users are made: all mobiles are considered statistically identical. As such, we assume that the N_{ugs} users are generating infinite-length ON/OFF constant bit rate traffics with the same traffic profile $(GBR, \bar{t}_{on}^{ugs}, \bar{t}_{off}^{ugs})$.

The amount of resource, i.e., the number of slots, needed at each frame by an UGS connection to achieve its GBR varies with the MCS it uses. In order to prevent the losses caused by outage periods, we assume that an UGS connection reserves a slightly greater bit rate than its GBR, called Delivered Bit Rate (DBR):

$$DBR = \frac{GBR}{1 - p_0}. \quad (2)$$

We model this system by a continuous-time Markov chain (CTMC) where each state v , represents the total number of concurrent UGS calls, regardless of the coding scheme they use. The maximum number of UGS calls accepted being $V_{ugs} = \min(N_{ugs}, V_{max})$, this chain is thus made of $V_{ugs} + 1$ states as shown in Fig 1.

- A transition out of a generic state v to state $v + 1$ occurs when a mobile in OFF period starts its call. This “arrival” transition corresponds to one mobile among the $(N_{ugs} - v)$ in OFF period, ending its idle time, and is performed with a rate $(N_{ugs} - v)\lambda_{ugs}$, where λ_{ugs} is defined as the inverse of the idle time between two calls: $\lambda_{ugs} = \frac{1}{t_{off}^{ugs}}$.
- On the opposite, a transition out of a generic state v to state $v - 1$ occurs when a mobile in ON period finishes its call. This “departure” transition corresponds to one mobile among the v in ON period, ending its call, and is performed with a rate $v\mu_{ugs}$, where μ_{ugs} is defined as the inverse of the average ON duration: $\mu_{ugs} = \frac{1}{t_{on}^{ugs}}$.

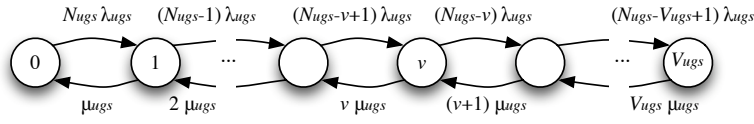


Figure 1: Mono-traffic UGS CTMC.

This results in the famous Engset model [16] which steady state probabilities $\pi_{ugs}(v)$ of having v current calls are derived as:

$$\pi_{ugs}(v) = \frac{\rho_{ugs}^v}{v!} \frac{N_{ugs}!}{(N_{ugs} - v)!} \pi_{ugs}(0), \quad (3)$$

with

$$\rho_{ugs} = \frac{\lambda_{ugs}}{\mu_{ugs}}. \quad (4)$$

and $\pi_{ugs}(0)$ obtained by normalization.

We then deduce the performance parameters as follows. First, the probability of rejecting a call P_{rej} is expressed:

$$P_{rej} = \frac{\pi_{ugs}(V_{ugs})(N_{ugs} - V_{ugs})}{\sum_{v=0}^{V_{ugs}} \pi_{ugs}(v)(N_{ugs} - v)}. \quad (5)$$

We compute $\bar{Q}_{ugs}$, the mean number of current UGS calls as:

$$\bar{Q}_{ugs} = \sum_{v=1}^{V_{ugs}} v \pi_{ugs}(v). \quad (6)$$

To attain its DBR, a mobile using MCS_k needs g_k slots:

$$g_k = \frac{DBR T_F}{m_k}. \quad (7)$$

Obviously, no slots are allocated to a mobile in outage so $g_0 = 0$. The available resource being limited, an UGS mobile does not always achieve its GBR if V_{max} is too big. Thus, we also derive $\bar{X}_{ugs}$, the instantaneous throughput obtained by an UGS mobile:

$$\bar{X}_{ugs} = \sum_{v=1}^{V_{ugs}} \frac{\pi_{ugs}(v)}{1 - \pi_{ugs}(0)} \sum_{\substack{(v_0, \dots, v_K) = (0, \dots, 0) \\ v_0 + \dots + v_K = v \\ v_0 \neq v}}^{(v, \dots, v)} p(v_0, \dots, v_K) \frac{\bar{N}_S}{\max\left(\sum_{k=1}^K v_k g_k, \bar{N}_S\right)} GBR, \quad (8)$$

where $p(v_0, \dots, v_K)$ is the probability that the v mobiles are distributed among the K MCS as $(v_0, \dots, v_K)$ (v_k being the number of mobiles using MCS_k):

$$p(v_0, \dots, v_K) = \binom{v}{v_0, \dots, v_K} \left(\prod_{k=0}^K p_k^{v_k} \right), \quad (9)$$

with $\binom{v}{v_0, \dots, v_K}$ the multinomial coefficient, and $\left(\frac{\bar{N}_S}{\max(\sum_{k=1}^K v_k g_k, \bar{N}_S)} GBR \right)$ corresponds to the throughput achieved by a mobile when the v connections are distributed in $(v_0, \dots, v_K)$. So, to obtain $\bar{X}_{ugs}$ we average this throughput for every possible distributions and remove the case when no UGS connections are active (i.e., state $v = 0$). Note that when there is no degradation, this expression leads to $\bar{X}_{ugs} = GBR$.

Finally, $\bar{U}_{ugs}$, the average utilization of the TDD frame by UGS connections is expressed as:

$$\bar{U}_{ugs} = \sum_{v=1}^{V_{ugs}} \pi_{ugs}(v) \sum_{\substack{(v_0, \dots, v_K) = (0, \dots, 0) \\ v_0 + \dots + v_K = v \\ v_0 \neq v}}^{(v, \dots, v)} p(v_0, \dots, v_K) \frac{\sum_{k=1}^K v_k g_k}{\max\left(\sum_{k=1}^K v_k g_k, \bar{N}_S\right)}. \quad (10)$$

Note that when V_{ugs} is small enough to guarantee that UGS calls are never degraded, i.e., $V_{ugs} \leq \frac{\bar{N}_S}{g_1}$, this expression can be greatly simplified:

$$\bar{U}_{ugs} = \sum_{v=1}^{V_{ugs}} \frac{\bar{g}(v)}{\bar{N}_S} \pi_{ugs}(v), \quad (11)$$

where $\bar{g}(v)$ is the mean number of slots needed by v UGS calls to obtain their DBR:

$$\bar{g}(v) = v \sum_{k=1}^K p_k g_k. \quad (12)$$

3.2 Multi-Traffic Extension

We now relax the assumption that all users have the same traffic profile. To do so, we distribute the mobiles among R traffic profiles defined by $(GBR_r, \bar{t}_{on}^r, \bar{t}_{off}^r)$. Thus, the mobiles of a given profile r generate an infinite-length ON/OFF traffic, with a guaranteed bit rate of GBR_r bits per second, an average ON duration of $\bar{t}_{on}^r$ seconds and an average OFF duration of $\bar{t}_{off}^r$ seconds. We consider that there is a fixed number N_{ugs}^r of mobiles belonging to each profile in the cell. So, there are $N_{ugs} = \sum_{r=1}^R N_{ugs}^r$ users in the cell with different traffic profiles.

Note that for the sake of clarity, the ugs indexes are removed from the $\bar{t}_{on}^r$ and $\bar{t}_{off}^r$ notations in this section.

Similarly to the mono-traffic model, we define DBR_r , the bit rate demanded by a class- r mobile:

$$DBR_r = \frac{GBR_r}{1 - p_0}. \quad (13)$$

To model this system, we use the multi-class extension of the Engset model [22]. The associated CTMC contains as much dimensions as considered traffic profiles and each of its states is characterized by a specific R-tuple $(v_1, \dots, v_R)$ where v_r is the number of active connections of class r .

- A transition out of a generic state $(v_1, \dots, v_r, \dots, v_R)$ to state $(v_1, \dots, v_r + 1, \dots, v_R)$ occurs when a class- r mobile in OFF period starts its call. This "arrival" transition is performed with a rate $\lambda_r(v_1, \dots, v_R) = (N_r - v_r) \lambda_r$, where λ_r is defined as the inverse of a class- r average idle time: $\lambda_r = \frac{1}{\bar{t}_{off}^r}$.

- A transition out of a generic state $(v_1, \dots, v_r, \dots, v_R)$ to state $(v_1, \dots, v_r - 1, \dots, v_R)$ occurs when a class- r mobile in ON period ends its call. This “departure” transition is performed with a rate $\mu_r(v_1, \dots, v_R) = v_r \mu_r$, where μ_r is defined as the inverse of a class- r average call duration: $\mu_r = \frac{1}{\bar{t}_{on}^r}$.

We also assume that V_{max} , the limit on the maximum number of concurrent UGS calls, is observed regardless of the classes they belong to. So, we define V_{ugs}^r , the maximum number of possible class- r simultaneous calls as:

$$V_{ugs}^r = \min(N_{ugs}^r, V_{max}). \quad (14)$$

However, nothing prevents from considering more complex limitations (e.g., privileging certain traffic profiles over others). Indeed, we then only need to adapt the possible states of the CTMC and the resolution of the model stays the same [21].

Fig. 2 presents this model's CTMC when considering only two traffic profiles ($R = 2$) and $V_{max} \leq \min(N_{ugs}^1, N_{ugs}^2)$.

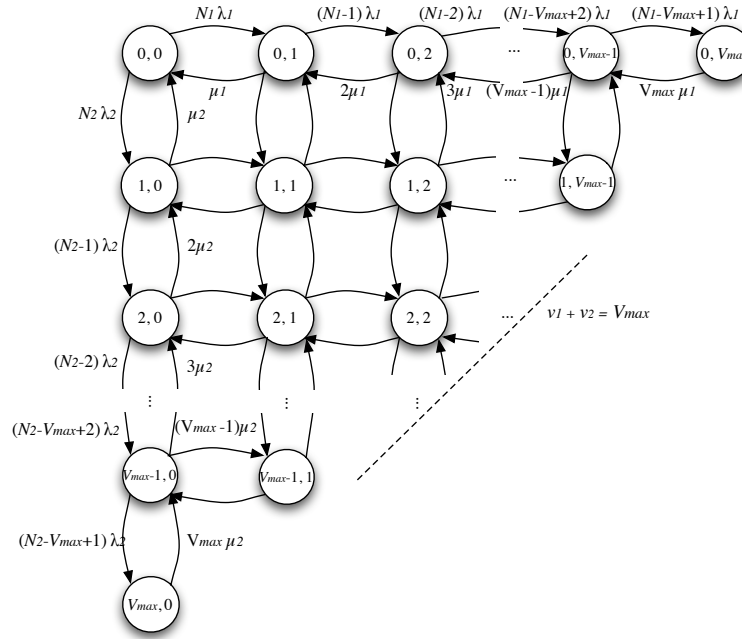


Figure 2: 2-dimensional multi-traffic UGS CTMC.

The steady state probabilities $\pi_{ugs}(v_1, \dots, v_R)$ of having $(v_1, \dots, v_R)$ concurrent UGS calls are given by:

$$\pi_{ugs}(v_1, \dots, v_R) = \left(\prod_{r=1}^R \frac{\rho_r^{v_r}}{v_r!} \frac{N_{ugs}^r!}{(N_{ugs}^r - v_r)!} \right) \pi_{ugs}(0, \dots, 0), \quad (15)$$

with

$$\rho_r = \frac{\lambda_r}{\mu_r}, \quad (16)$$

and $\pi_{ugs}(0, \dots, 0)$ obtained by normalization.

The performance parameters are derived from the steady-state probabilities as follows. The probability

of rejecting a class- r call P_{rej}^r is given by:

$$P_{rej}^r = \frac{\sum_{\substack{(V_{ugs}^1, \dots, V_{ugs}^R) \\ (v_1, \dots, v_R) = (0, \dots, 0) \\ v_1 + \dots + v_R = V_{max}}} \pi_{ugs}(v_1, \dots, v_R)(N_{ugs}^r - v_r)}{\sum_{\substack{(V_{ugs}^1, \dots, V_{ugs}^R) \\ (v_1, \dots, v_R) = (0, \dots, 0) \\ v_1 + \dots + v_R \leq V_{max}}} \pi_{ugs}(v_1, \dots, v_R)(N_{ugs}^r - v_r)}. \quad (17)$$

Obviously, this probability is null when $N_{ugs} \leq V_{max}$ as there can be no blocking in this case.

We can compute $\bar{Q}_{ugs}^r$, the mean number of concurrent UGS calls belonging to class r as:

$$\bar{Q}_{ugs}^r = \sum_{\substack{(V_{ugs}^1, \dots, V_{ugs}^R) \\ (v_1, \dots, v_R) = (0, \dots, 0) \\ v_1 + \dots + v_R \leq V_{max}}} v_r \pi_{ugs}(v_1, \dots, v_R). \quad (18)$$

Finally, by first defining $\bar{g}(v_1, \dots, v_R)$, the mean number of slots needed by $(v_1, \dots, v_R)$ UGS mobiles to achieve their respective GBR_r :

$$\bar{g}(v_1, \dots, v_R) = \sum_{r=1}^R v_r \sum_{k=1}^K p_k \frac{DBR_r T_F}{m_k}, \quad (19)$$

we can then express $\bar{X}_{ugs}^r$, the instantaneous throughput achieved by class- r mobiles as:

$$\bar{X}_{ugs}^r = \sum_{\substack{(V_{ugs}^1, \dots, V_{ugs}^R) \\ (v_1, \dots, v_R) = (0, \dots, 0) \\ v_1 + \dots + v_R \leq V_{max}}} \frac{\pi_{ugs}(v_1, \dots, v_R)}{1 - p_{v_r=0}} \frac{\bar{N}_S}{\max(\bar{g}(v_1, \dots, v_R), \bar{N}_S)} GBR_r, \quad (20)$$

with $p_{v_r=0}$ the probability that no class- r mobile is active:

$$p_{v_r=0} = \sum_{\substack{(V_{ugs}^1, \dots, V_{ugs}^R) \\ (v_1, \dots, v_R) = (0, \dots, 0) \\ v_1 + \dots + v_R \leq V_{max} \\ v_r = 0}} \pi_{ugs}(v_1, \dots, v_R), \quad (21)$$

and compute $\bar{U}_{ugs}$, the average utilization of the TDD frame as:

$$\bar{U}_{ugs} = \sum_{\substack{(V_{ugs}^1, \dots, V_{ugs}^R) \\ (v_1, \dots, v_R) = (0, \dots, 0) \\ v_1 + \dots + v_R \leq V_{max}}} \frac{\bar{g}(v_1, \dots, v_R)}{\max(\bar{g}(v_1, \dots, v_R), \bar{N}_S)} \pi_{ugs}(v_1, \dots, v_R). \quad (22)$$

4 ERTPS MODEL

The ertPS (*enhanced real-time Polling Service*) service class has been especially added to the WiMAX standards in order to carry traffic from VoIP with silence suppression. As such, an ertPS call only occupies the resource during the talk spurts of the conversation. We show, in this section, how our UGS model can be easily adapted to take into account the impact of silence suppression on the cell capacity. To this aim, we now consider a cell carrying only ertPS traffic.

A telephonic conversation can be seen as a succession of ON periods (talk spurts) and OFF periods (silences). As shown in [9], [10], the durations of these periods can be accurately modeled by exponential

distributions with mean values $\bar{t}_{talk}$ et $\bar{t}_{sil}$ respectively. Note that [10] recommends to set these values to $\bar{t}_{talk} = 1.2$ s et $\bar{t}_{sil} = 1.8$ s.

Knowing this, we can account for the effect of silence suppression in a mono-traffic scenario by adjusting the expression of g_k , the number of slots needed by a mobile using MCS_k to attain its DBR, as:

$$g_k = \frac{\bar{t}_{talk}}{\bar{t}_{talk} + \bar{t}_{sil}} \frac{DBR T_F}{m_k}. \quad (23)$$

Similarly, in a multi-traffic scenario, we just need to modify the expression of $\bar{g}(v_1, \dots, v_R)$, the mean number of slots needed by $(v_1, \dots, v_R)$ mobiles using to achieve their respective GBR_r , as:

$$\bar{g}(v_1, \dots, v_R) = \sum_{r=1}^R v_r \sum_{k=1}^K p_k \frac{\bar{t}_{talk}}{\bar{t}_{talk} + \bar{t}_{sil}} \frac{DBR_r T_F}{m_k}, \quad (24)$$

to consider the bandwidth saved by the silence suppression.

The rest of our UGS modeling approach still stands whether we consider UGS or ertPS traffic. So, the computing of the ertPS steady state probabilities and performance parameters is the same as in the UGS model.

5 BE MODELS

The BE (*Best Effort*) service class of the WiMAX standard has been planned to carry the traffic of applications without QoS guarantee (e.g., web applications). In this section, we provide an overview of the analytical models we use to perform the traffic analysis of this service class. In all the models presented here, we consider that only BE mobiles (i.e., mobiles generating only BE traffic) are present in the studied cell. Models for cells carrying traffic belonging to several service classes are detailed in Sections 6 and 7.

Our main concern in this paper is to show how these BE models can be used as stepping stones for the study of cells carrying traffics of several service classes. As a consequence, we only introduce the various parameters needed for the dimensioning procedure and do not detail their expressions. However, note that these models have already been fully explained, discussed and validated in [15], [14].

5.1 Scheduling

Several scheduling schemes can be considered. In [15], we focused on three traditional schemes:

- The slot sharing fairness scheduling equally divides the slots of each frame between all active users that are not in outage.
- The instantaneous throughput fairness scheduling shares the resource in order to provide the same instantaneous throughput to all active users not in outage.
- The opportunistic scheduling gives all the resources to active users having the highest transmission bit rate, i.e., the better MCS.

Lastly, in [14], we proposed an alternative scheduling called throttling which forces an upper bound on the users' throughputs, the Maximum Sustained Traffic Rate (MSTR):

- The throttling scheduling tries to allocate at each frame the right number of slots to each active mobile in order to achieve its MSTR. If a mobile is in outage it does not receive any slot and its throughput is degraded. If at a given time the total number of available slots is not enough to satisfy the MSTR of all active users (not in outage), they all see their throughputs equally degraded.

5.2 Mono-Traffic Model

As a first step, we do not make any distinction between users and consider all mobiles as statistically identical. Thus, we consider that the N_{be} users are generating infinite-length ON/OFF BE traffics with the same traffic profile $(\bar{x}_{on}^{be}, \bar{t}_{off}^{be})$.

We model this system by a continuous-time Markov chain (CTMC) where each state n , represents the total number of concurrent active BE connections, regardless of the coding scheme they use. So, the resulting CTMC is made of $N_{be} + 1$ states as depicted in Fig 3.

- A transition out of a generic state n to a state $n + 1$ occurs when a mobile in OFF period starts its transfer. This “arrival” transition corresponds to one mobile among the $(N_{be} - n)$ in OFF period, ending its reading, and is performed with a rate $(N_{be} - n)\lambda_{be}$, where λ_{be} is defined as the inverse of the average reading time: $\lambda_{be} = \frac{1}{\bar{t}_{off}^{be}}$.
- A transition out of a generic state n to a state $n - 1$ occurs when a mobile in ON period completes its transfer. This “departure” transition is performed with a generic rate $\mu_{be}(n)$ corresponding to the total departure rate of the frame when n mobiles are active.

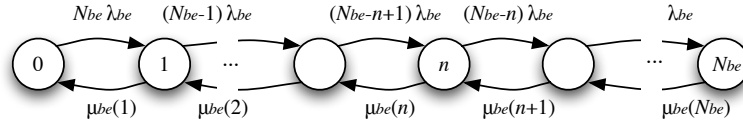


Figure 3: Mono-traffic BE CTMC with state-dependent departure rates.

Obviously, the main difficulty of the model resides in estimating the aggregate departure rates $\mu_{be}(n)$.

If we consider either the instantaneous throughput fairness, the slot sharing fairness or the opportunistic policy, they are expressed as follows [15]:

$$\mu_{be}(n) = \frac{\bar{m}(n) \bar{N}_S}{\bar{x}_{on}^{be} T_F}, \quad (25)$$

where $\bar{m}(n)$ is the average number of bits transmitted per slot when there are n concurrent active transfers. These parameters are strongly dependent on the scheduling policy. As a consequence, we provide their different expressions for each policy.

With the slot sharing policy:

$$\bar{m}(n) = \sum_{\substack{(n_0, \dots, n_K) = (0, \dots, 0) \\ n_0 + \dots + n_K = n \\ n_0 \neq n}}^{(n, \dots, n)} \frac{n!}{n - n_0} \sum_{k=1}^K m_k n_k \prod_{k=0}^K \frac{p_k^{n_k}}{n_k!}. \quad (26)$$

With the instantaneous throughput fairness policy:

$$\bar{m}(n) = \sum_{\substack{(n_0, \dots, n_K) = (0, \dots, 0) \\ n_0 + \dots + n_K = n \\ n_0 \neq n}}^{(n, \dots, n)} \frac{(n - n_0) n! \prod_{k=0}^K \frac{p_k^{n_k}}{n_k!}}{\sum_{k=1}^K \frac{n_k}{m_k}}. \quad (27)$$

With the opportunistic policy:

$$\bar{m}(n) = \sum_{k=1}^K m_k \left(1 - \sum_{j=k+1}^K p_j\right)^n \left(1 - \left(1 - \frac{p_k}{\sum_{j=0}^k p_j}\right)^n\right). \quad (28)$$

If we consider the throttling policy, the departure rates $\mu_{be}(n)$ become:

$$\mu_{be}(n) = \frac{\bar{N}_S}{\max(n\bar{g}, \bar{N}_S)} n \frac{MSTR}{\bar{x}_{on}^{be}}. \quad (29)$$

with $\bar{g}$, the average number of slots per frame needed by a mobile to obtain its MSTR:

$$\bar{g} = T_F MSTR \sum_{k=1}^K \frac{p_k}{(1 - p_0)m_k}. \quad (30)$$

Once the departure rates $\mu_{be}(n)$ have been determined, the steady state probabilities $\pi_{be}(n)$ of having n concurrent transfers in the cell, can easily be derived from the birth-and-death structure of the Markov chain:

$$\pi_{be}(n) = \left(\prod_{i=1}^n \frac{(N_{be} - i + 1)\lambda_{be}}{\mu_{be}(i)} \right) \pi_{be}(0), \quad (31)$$

where $\pi_{be}(0)$ is obtained by normalization.

Note that the $\bar{x}_{on}^{be}$ and $\bar{t}_{off}^{be}$ traffic parameters are only involved in this last expression through their ratio. As a consequence, we define the intensity ρ_{be} of the traffic generated by a mobile:

$$\rho_{be} = \frac{\bar{x}_{on}^{be}}{\bar{t}_{off}^{be}}. \quad (32)$$

The following performance parameters of the system can be obtained from the steady state probabilities. The average number of active users $\bar{Q}_{be}$ is expressed as:

$$\bar{Q}_{be} = \sum_{n=1}^{N_{be}} n \pi_{be}(n), \quad (33)$$

and $\bar{D}$, the mean number of departures (i.e., mobiles completing their transfer) per unit of time, is obtained as:

$$\bar{D}_{be} = \sum_{n=1}^{N_{be}} \mu_{be}(n) \pi_{be}(n). \quad (34)$$

From Little's law, we can thus derive the average duration $\bar{t}_{on}^{be}$ of an ON period (duration of an active transfer):

$$\bar{t}_{on}^{be} = \frac{\bar{Q}_{be}}{\bar{D}_{be}}. \quad (35)$$

and compute the average throughput $\bar{X}_{be}$ obtained by one BE mobile in active transfer as:

$$\bar{X}_{be} = \frac{\bar{x}_{on}^{be}}{\bar{t}_{on}^{be}}. \quad (36)$$

Finally, we can express the average utilization $\bar{U}_{be}$ of the TDD frame. This last parameter depends on the scheduling policy. Indeed, with the instantaneous throughput fairness, the slot sharing fairness or

the opportunistic policy, the cell is considered fully utilized as long as there is at least one active mobile not in outage:

$$\bar{U}_{be} = \sum_{n=1}^{N_{be}} (1 - p_0^n) \pi_{be}(n). \quad (37)$$

However, if we consider the throttling policy, $\bar{U}_{be}$ is now expressed as the weighted sum of the ratios between the mean number of slots needed by the n mobiles to reach their MSTR and the mean number of slots they obtain:

$$\bar{U}_{be} = \sum_{n=1}^{N_{be}} \frac{n\bar{g}}{\max(n\bar{g}, \bar{N}_S)} \pi_{be}(n). \quad (38)$$

5.3 Multi-Traffic Extension

Now, we relax the assumption that all users have the same traffic profile. To this aim, we associate to each mobile one of the R traffic profiles, $(\bar{x}_{on}^r, \bar{t}_{off}^r)$. The mobiles of a given profile r thus generate an infinite-length ON/OFF traffic, with an average ON size of $\bar{x}_{on}^r$ bits and an average reading time of $\bar{t}_{off}^r$ seconds. We assume that there is a fixed number N_{be}^r of mobiles belonging to each profile in the cell. As a consequence, there are $N_{be} = \sum_{r=1}^R N_{be}^r$ users in the cell with different traffic profiles.

Similarly to the notations in Section 3.2, the be indexes are removed from the $\bar{x}_{on}^r$, $\bar{t}_{on}^r$ and $\bar{t}_{off}^r$ notations in this section.

To compute the performance parameters, we first transform this system into an equivalent one where all profiles of traffic have the same average ON size $\bar{x}_{on}$, and different average OFF durations $\bar{t}_{off}^r$, such that [15]:

$$\frac{\bar{x}_{on}}{\bar{t}_{off}^r} = \frac{\bar{x}_{on}^r}{\bar{t}_{off}^r}. \quad (39)$$

With this transformation, the equivalent system can be described as a multi-class closed queueing network with two stations as shown by Fig. 4:

1. An *Infinite Servers* (IS) station that models mobiles in OFF periods. This station has profile-dependent service rates $\lambda_r = \frac{1}{\bar{t}_{off}^r}$;
2. A *Processor Sharing* (PS) station that models active mobiles. This station has profile-independent service rates $\mu_{be}(n)$ that in turn depend on the total number active mobiles (whatever their profiles). They are given by the same relations than the departure rates of the mono-traffic model (see relations 25 and 29)

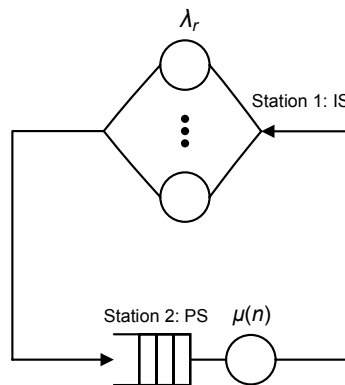


Figure 4: Closed-queueing network.

A direct extension of the BCMP theorem [4] for stations with state-dependent rates can now be applied to this closed queueing network. The detailed steady state probabilities are expressed as follows:

$$\pi_{be}(\vec{n}_1, \vec{n}_2) = \frac{1}{G} f_1(\vec{n}_1) f_2(\vec{n}_2), \quad (40)$$

where $\vec{n}_i = (n_{i1}, \dots, n_{iR})$, n_{ir} is the number of profile- r mobiles present in station i ,

$$f_1(\vec{n}_1) = \frac{1}{n_{11}! \dots n_{1R}!} \frac{1}{(\lambda_1)^{n_{11}} \dots (\lambda_R)^{n_{1R}}} \quad (41)$$

$$f_2(\vec{n}_2) = \frac{(n_{21} + \dots + n_{2R})!}{n_{21}! \dots n_{2R}!} \frac{1}{\prod_{k=1}^{n_2} \mu_{be}(k)}, \quad (42)$$

and G is the normalization constant:

$$G = \sum_{\vec{n}_1 + \vec{n}_2 = \vec{N}_{be}} f_1(\vec{n}_1) f_2(\vec{n}_2). \quad (43)$$

All the performance parameters of interest can be derived from the steady state probabilities as follows. The average number of profile- r active mobiles, $\bar{Q}_r$, is given by:

$$\bar{Q}_r = \sum_{\vec{n}_1 + \vec{n}_2 = \vec{N}_{be}} n_{2r} \pi_{be}(\vec{n}_1, \vec{n}_2), \quad (44)$$

and the average number of profile- r mobiles completing their download by unit of time, $\bar{D}_r$, can be expressed as:

$$\bar{D}_r = \sum_{\vec{n}_1 + \vec{n}_2 = \vec{N}_{be}} \mu(n_2) \pi_{be}(\vec{n}_1, \vec{n}_2), \quad (45)$$

with $n_2 = \sum_{r=1}^R n_{2r}$.

The average download duration of profile- r mobiles, $\bar{t}_{on}^r$, is obtained from Little's law:

$$\bar{t}_{on}^r = \frac{\bar{Q}_r}{\bar{D}_r}. \quad (46)$$

And we can then calculate the average throughput obtained by customers of profile r during their transfer, denoted by $\bar{X}_r$, as:

$$\bar{X}_r = \frac{\bar{x}_{on}^r}{\bar{t}_{on}^r} \quad (47)$$

Finally, the utilization $\bar{U}_{be}$ of the TDD frame is expressed differently whether we consider the instantaneous throughput fairness, the slot sharing fairness, the opportunistic policy:

$$\bar{U}_{be} = \sum_{\vec{n}_1 + \vec{n}_2 = \vec{N}_{be}} (1 - p_0^{n_2}) \pi_{be}(\vec{n}_1, \vec{n}_2), \quad (48)$$

or the throttling policy:

$$\bar{U}_{be} = \sum_{\vec{n}_1 + \vec{n}_2 = \vec{N}_{be}} \frac{n_2 \bar{g}}{\max(n_2 \bar{g}, N_S)} \pi_{be}(\vec{n}_1, \vec{n}_2). \quad (49)$$

Again, fully detailed explanations on the multi-traffic model and the different relations are available in [14]. Finally, let us add that [14] also includes a method to consider the throttling policy and traffic profiles with different MSTR values.

6 UGS/BE MODEL

Using the previous models, we can consider a cell carrying either UGS or BE traffic. However, we now desire to take into account a cell with both kinds of traffic without increasing the resolution complexity of the resulting model.

To do so, we first propose three methods to combine both mono-traffic UGS and BE models. Then, we compare results brought by these methods and conclude on the best one to use. Lastly, we explain how to extend the resulting UGS/BE model to multiple traffic profiles in both service classes.

6.1 Combining the Mono-Traffic Models

As a first step, we only consider one UGS traffic profile and one BE traffic profile.

6.1.1 General approach

The WiMAX standard does not recommend any scheduling between service classes. However, it is common sense and widely admitted in the literature that UGS calls preempt BE traffic [23], [28], [29]. Indeed, UGS connections reserve a part of the resource in each frame and BE mobiles share whatever is left to them. This has two major consequences: i) our UGS model is sufficient to characterize the UGS traffic since UGS calls are not affected by the presence of the BE traffic; ii) the performance evaluation of the BE mobiles are strongly dependent on the portion of resource left to them by the UGS traffic at each frame. As such, the part of the resource, i.e., the mean number of slots, left to the BE connections needs to be evaluated and accounted for in the characterization of the BE traffic.

Our general approach to combine both UGS and BE models consists in the following steps:

1. We compute the steady state probabilities $\pi_{ugs}(v)$ of having v active UGS connections by only using the UGS model.
2. We then estimate the mean number of slots occupied by these v UGS calls and deduce the mean number of slots available to BE connections when there are v active UGS calls.
3. For each possible value of v , we solve a CTMC of the BE model with the $\bar{N}_S$ parameter set to the corresponding mean number of slots available. We obtain the steady state probabilities $\pi_{be}(n|v)$ of having n active BE connections provided that there are v concurrent UGS calls.
4. We express the steady state probabilities $\pi(v, n)$ of having simultaneously both v active UGS and n active BE connections as:

$$\pi(v, n) = \pi_{ugs}(v)\pi_{be}(n|v). \quad (50)$$

Three methods to combine the UGS and BE models are proposed, each one corresponding to a different level of precision in evaluating the portion of the resource left to the BE mobiles.

6.1.2 DTL (detailed) method

As its name indicates, this method is the most precise of the three. Indeed, for each possible value v of simultaneous active UGS connections, we specifically consider each possible distributions $(v_0, \dots, v_K)$ of the v connections among the $K + 1$ MCS (including outage). And, for each of these distributions, we compute $\bar{N}_S^{be}(v_0, \dots, v_K)$, the corresponding mean number of slots left to the BE connections:

$$\bar{N}_S^{be}(v_0, \dots, v_K) = \bar{N}_S - \bar{N}_S^{ugs}(v_0, \dots, v_K), \quad (51)$$

with

$$\bar{N}_S^{ugs}(v_0, \dots, v_K) = \min \left(\sum_{k=1}^K v_k g_k, \bar{N}_S \right). \quad (52)$$

Then, for each $\bar{N}_S^{be}(v_0, \dots, v_K)$, we solve a corresponding BE CTMC and obtain the steady state probabilities $\pi_{be}(n|(v_0, \dots, v_K))$ of having n active BE transfers provided that there are $(v_0, \dots, v_K)$ active UGS connections.

From these, we first deduce the probabilities $\pi_{be}(n|v)$:

$$\pi_{be}(n|v) = \sum_{\substack{(v_0, \dots, v_K) = (0, \dots, 0) \\ v_0 + \dots + v_K = v}}^{(v, \dots, v)} p(v_0, \dots, v_K) \cdot \pi_{be}(n|(v_0, \dots, v_K)), \quad (53)$$

then, the probabilities $\pi(v, n)$ with relation 50.

Note that for the distributions where no slots remain for BE traffic ($\bar{N}_S^{be}(v_0, \dots, v_K) = 0$), we set:

$$\pi_{be}(n|(v_0, \dots, v_K)) = \begin{cases} 1 & \text{if } n = N_{be} \\ 0 & \text{else} \end{cases} \quad (54)$$

because we consider that BE connections keep on being initiated but none of them can end without available resource.

Lastly, let us highlight that this method, illustrated by Fig.5(a), requires to solve as much BE CTMC as there are possible values of $\bar{N}_S^{be}(v_0, \dots, v_K)$.

6.1.3 AVG (averaged) method

This method has been proposed to significantly reduce the number of BE CTMC to solve. We now only evaluate $\bar{N}_S^{ugs}(v)$, the mean numbers of slots occupied by v UGS calls by averaging the $\bar{N}_S^{ugs}(v_0, \dots, v_K)$ as follows:

$$\bar{N}_S^{ugs}(v) = \sum_{\substack{(v_0, \dots, v_K) = (0, \dots, 0) \\ v_0 + \dots + v_K = v \\ v_0 \neq v}}^{(v, \dots, v)} p(v_0, \dots, v_K) \bar{N}_S^{ugs}(v_0, \dots, v_K). \quad (55)$$

So, this time, to obtain the $\pi_{be}(n|v)$ probabilities, we only solve one CTMC for each possible value of v , considering that $\bar{N}_S^{be}(v) = \bar{N}_S - \bar{N}_S^{ugs}(v)$ slots remain available to BE transfers when there are v active UGS connections. Thus, with this method (see Fig. 5(b)), the number of BE CTMC to solve is reduced to $V_{ugs} + 1$.

Let us highlight that this method is not based on the quasi-stationary assumption introduced in [12] by Delcoigne and al. Indeed, here, we do not need to make any assumption on the time scales of UGS and BE traffics.

6.1.4 AGG (aggregated) method

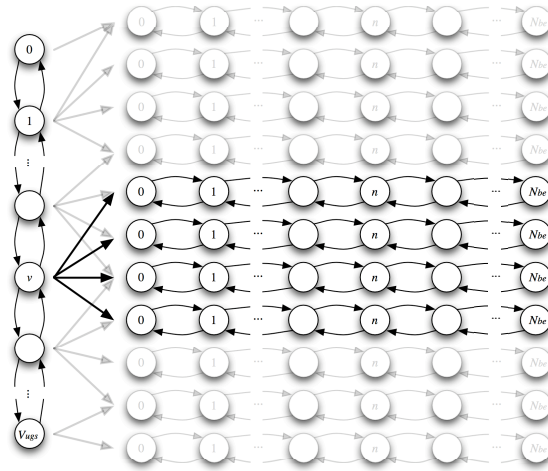
This last method, as shown in Fig. 5(c), is the most straightforward of the three: only one BE CTMC has to be solved. We first compute $\bar{N}_S^{ugs}$, the mean number of slots occupied by UGS connections as:

$$\bar{N}_S^{ugs} = \sum_{v=1}^{V_{ugs}} \pi_{ugs}(v) \bar{N}_S^{ugs}(v). \quad (56)$$

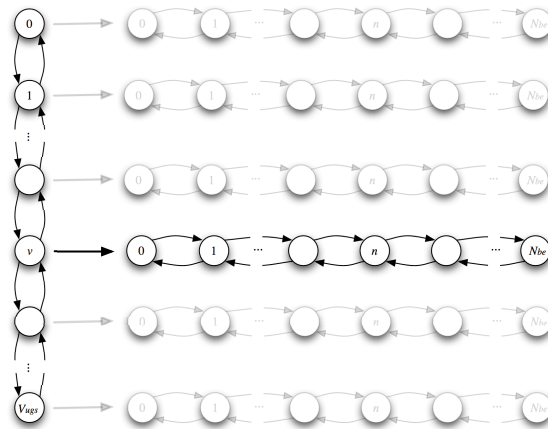
Then we only have to solve the BE CTMC corresponding to $\bar{N}_S^{be} = \bar{N}_S - \bar{N}_S^{ugs}$ slots available to the BE mobiles to obtain the probabilities $\pi_{be}(n)$ and the various performance parameters.

6.1.5 Performance parameters

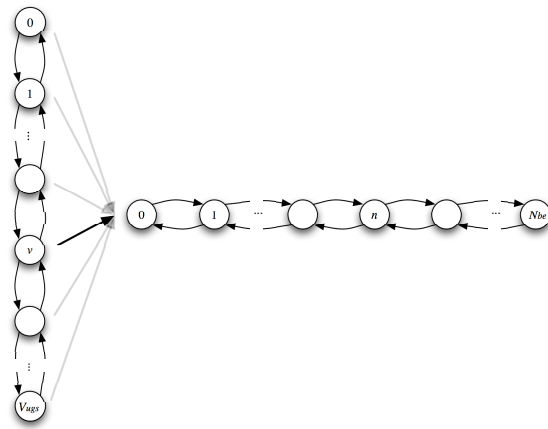
UGS



(a) DTL method



(b) AVG method



(c) AGG method

Figure 5: Three methods to combine the UGS and BE models. (UGS CTMC is vertical, BE ones horizontal.)

As stated earlier, the UGS traffic is not affected by the presence of the BE traffic. So, the UGS performance parameters are computed as detailed in Section 3.1.

BE DTL/AVG

When using either the DTL or AVG method, the BE performance parameters are derived from the steady state probabilities $\pi(v, n)$ as follows:

$\bar{Q}_{be}$, the mean number of active BE users is given by:

$$\bar{Q}_{be} = \sum_{n=1}^{N_{be}} n \sum_{v=0}^{V_{ugs}} \pi(v, n). \quad (57)$$

$\bar{D}_{be}$, the mean number of BE departures per unit of time, depends on the number of slots left to the BE mobiles. As such, its expression differs whether we consider the DTL or AGG method. So, for DTL, we consider all the possible distributions of the v UGS connections among the K MCS:

$$\bar{D}_{be} = \sum_{v=0}^{V_{ugs}} \sum_{\substack{(v_0, \dots, v_K) = (0, \dots, 0) \\ v_0 + \dots + v_K = v}}^{(v, \dots, v)} p(v_0, \dots, v_K) \sum_{n=1}^{N_{be}} \bar{N}_S^{be}(v_0, \dots, v_K) \mu_{be}(n) \pi(v, n), \quad (58)$$

whereas for AVG we only use the $\bar{N}_S^{be}(v)$:

$$\bar{D}_{be} = \sum_{v=0}^{V_{ugs}} \sum_{n=1}^{N_{be}} \bar{N}_S^{be}(v) \mu_{be}(n) \pi(v, n). \quad (59)$$

In both cases, we obtain from Little's law $\bar{t}_{on}^{be}$, the average duration of a BE transfer:

$$\bar{t}_{on}^{be} = \frac{\bar{Q}_{be}}{\bar{D}_{be}}, \quad (60)$$

and deduce $\bar{X}_{be}$, the average throughput achieved by a BE mobile:

$$\bar{X}_{be} = \frac{\bar{x}_{on}^{be}}{\bar{t}_{on}^{be}}. \quad (61)$$

Finally, the average utilization of the TDD frame by BE transfers, $\bar{U}_{be}$, needs to be adapted to both methods but also to the considered BE scheduling. With DTL and the slot sharing, the instantaneous throughput fairness or the opportunistic policy:

$$\bar{U}_{be} = \sum_{v=0}^{V_{ugs}} \sum_{\substack{(v_0, \dots, v_K) = (0, \dots, 0) \\ v_0 + \dots + v_K = v}}^{(v, \dots, v)} p(v_0, \dots, v_K) \frac{\bar{N}_S^{be}(v_0, \dots, v_K)}{\bar{N}_S} \sum_{n=1}^{N_{be}} (1 - p_0^n) \pi(v, n). \quad (62)$$

With DTL and the throttling policy:

$$\bar{U}_{be} = \sum_{v=0}^{V_{ugs}} \sum_{\substack{(v_0, \dots, v_K) = (0, \dots, 0) \\ v_0 + \dots + v_K = v}}^{(v, \dots, v)} p(v_0, \dots, v_K) \cdot \frac{\bar{N}_S^{be}(v_0, \dots, v_K)}{\bar{N}_S} \sum_{n=1}^{N_{be}} \frac{n\bar{g}}{\max(n\bar{g}, \bar{N}_S^{be}(v_0, \dots, v_K))} \pi(v, n). \quad (63)$$

With AVG and the slot sharing, the instantaneous throughput fairness or the opportunistic policy:

$$\bar{U}_{be} = \sum_{v=0}^{V_{ugs}} \frac{\bar{N}_S^{be}(v)}{\bar{N}_S} \sum_{n=1}^{N_{be}} (1 - p_0^n) \pi(v, n). \quad (64)$$

With AVG and the throttling policy:

$$\bar{U}_{be} = \sum_{v=0}^{V_{ugs}} \frac{\bar{N}_S^{be}(v)}{\bar{N}_S} \sum_{n=1}^{N_{be}} \frac{n\bar{g}}{\max(n\bar{g}, \bar{N}_S^{be}(v))} \pi(v, n). \quad (65)$$

BE AGG

As for the AGG method, the BE performance parameters are obtained using the $\pi_{be}(n)$ as explained in Section 5.2. We just have to replace $\bar{N}_S$ by $\bar{N}_S^{be}$ in the various expressions.

Only the average utilization of the resource by BE traffic, $\bar{U}_{be}$, needs to be adapted following the considered BE scheduling. With the slot sharing, the instantaneous throughput fairness or the opportunistic policy:

$$\bar{U}_{be} = \frac{\bar{N}_S^{be}}{\bar{N}_S} \sum_{n=1}^{N_{be}} (1 - p_0^n) \pi_{be}(n). \quad (66)$$

With the throttling policy:

$$\bar{U}_{be} = \frac{\bar{N}_S^{be}}{\bar{N}_S} \sum_{n=1}^{N_{be}} \frac{n\bar{g}}{\max(n\bar{g}, \bar{N}_S^{be})} \pi_{be}(n). \quad (67)$$

6.2 Comparison

Here, we compare results obtained with the DTL, AVG and AGG methods. The channel, cell and traffic parameters are summarized in Tables 1 and 2.

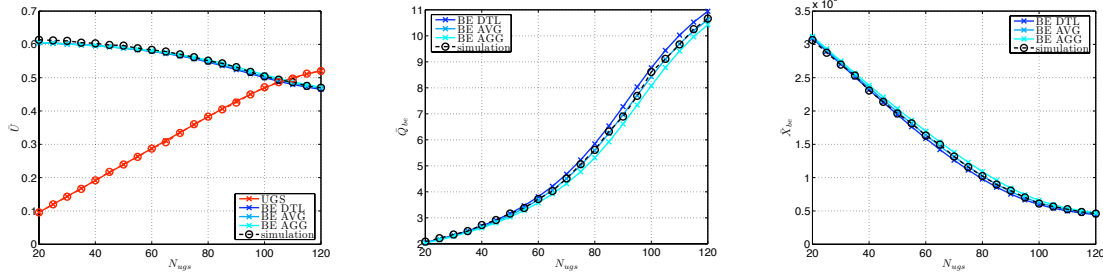
Parameter		Value
Number of slots per frame, N_S		450
Frame duration, T_F		5 ms
Limit on UGS calls, V_{max}		40
BE scheduling		slot fairness
UGS	Number of UGS mobiles, N_{ugs}	20 to 120
	Guaranteed bit rate, GBR	128 Kbps
	Mean ON duration, t_{on}^{ugs}	60 s
	Mean OFF duration, t_{off}^{ugs}	120 s
BE	Number of BE mobiles, N_{be}	40
	Mean ON size, $\bar{x}_{on}^{be}$	1 Mbits
	Mean OFF duration, t_{off}^{be}	6 s

Table 2: Cell and traffic parameters.

The duration T_F of one TDD frame of WiMAX is 5 ms. We consider $\bar{N}_S = 450$ slots per frame available for downlink. This value roughly corresponds to a system bandwidth of 10 MHz, a downlink/uplink ratio of 2/3, a Fast Fourier Transform size equal to 2048, a PUSC subcarrier permutation and an average protocol overhead length of 4 symbols.

We observe the behaviors of $\bar{U}$, $\bar{Q}_{be}$ and $\bar{X}_{be}$ parameters given by each method while considering a cell with an increasing number of UGS mobiles, N_{ugs} and a fixed number of BE mobiles N_{be} . The maximum number of concurrent UGS calls accepted in the cell is set to $V_{max} = 40$ and BE connections are scheduled using the slot fairness policy.

Fig. 6(a) presents the evolution of the average utilization of the TDD frame by UGS and BE traffics when the UGS traffic load increases. Obviously, the portions of the frames occupied by UGS connections increases with the number of UGS mobiles in the cell and the BE transfers get less and less resource.



(a) Average utilization, $\bar{U}$ of the TDD frame by UGS and BE connections. (b) Average number of concurrent BE transfers, $\bar{Q}_{be}$. (c) Instantaneous throughput of a BE mobile, $\bar{X}_{be}$.

Figure 6: Comparison of the UGS/BE methods: Customary BE traffic parameters when UGS traffic increases.

So, as shown in Fig. 6(b) and 6(c) respectively, BE mobiles stay in ON periods longer and achieve smaller throughputs.

Moreover, these figures also provide simulation results. Indeed, as a first step to validate our approach, we have developed a simulator that integrates an ON/OFF traffic generator, a wireless channel for each user and a centralized scheduler allocating radio resources, i.e., slots, to active users on a frame by frame basis. The simulator allocates the resources to active users at each frame serving first the UGS connections and then the BE connections. At the beginning of a frame, a new MCS is drawn for each active mobile according to the channel probabilities p_k . Then, the scheduler allocates slots to those mobiles depending on their MCS and their service class. In the simulations presented here, a mobile generates an exponentially distributed ON/OFF traffic and is subject to a memoryless channel. Note however that results of more realistic simulations are provided in section 7.2 to show the robustness of our analytical modeling toward the channel model and the distribution of BE ON sizes.

In all the figures, we can see that the results obtained with each of the three UGS/BE methods and the simulations are always very close. Furthermore, we have compared these methods and the simulations in all sorts of configurations (different traffic loads, different bit rates reserved by UGS calls, different BE schedulings, etc.). Each time, the results they gave proved to be almost equivalent: the difference between them remained below 3%.

This leads us to conclude that as straightforward as the AGG method may appear, its precision is sufficient to efficiently combine the UGS and BE models.

6.3 Multi-Traffic Extensions

As stated above, the 3 methods give similar results so we now only consider the AGG method. The simplicity of this method enables us to extend the resulting UGS/BE model to multiple traffic profiles in both service classes as follows.

6.3.1 Multi-traffic UGS

Again, UGS traffic preempts BE traffic. So, to consider multiple UGS traffic profiles, we just use the multi-traffic extension introduced in Section 3.2. The only modification to the UGS/BE model resides in the expression of $\bar{N}_S^{ugs}$:

$$\bar{N}_S^{ugs} = \sum_{\substack{(v_1, \dots, v_R) = (0, \dots, 0) \\ v_1 + \dots + v_R \leq V_{max}}}^{(V_{ugs}^1, \dots, V_{ugs}^R)} \pi_{ugs}(v_1, \dots, v_R) \sum_{v=1}^{V_{ugs}} \min(\bar{N}_S^{ugs}(v_1, \dots, v_R), \bar{N}_S), \quad (68)$$

where $\bar{N}_S^{ugs}(v_1, \dots, v_R)$ is the mean number of slots occupied by UGS connections distributed in $(v_1, \dots, v_R)$ among the R traffic profiles:

$$\bar{N}_S^{ugs}(v_1, \dots, v_R) = \sum_{r=1}^R \sum_{\substack{(j_0, \dots, j_K) = (0, \dots, 0) \\ j_0 + \dots + j_K = v_r}}^{(v_r, \dots, v_r)} p(j_0, \dots, j_K) \min \left(\sum_{k=1}^K j_k g_k, \bar{N}_S \right). \quad (69)$$

The successive minima in these expressions enables to ensure that we never count more slots than $\bar{N}_S$ in the averaging of the numbers of slots occupied by UGS traffic.

Lastly, note that when V_{max} , the limit on the maximum simultaneous UGS calls allowed in the cell is small enough to ensure that no degrading will occur, the expression of $\bar{N}_S^{ugs}(v_1, \dots, v_R)$ can be drastically simplified as:

$$\bar{N}_S^{ugs}(v_1, \dots, v_R) = \bar{g}(v_1, \dots, v_R). \quad (70)$$

6.3.2 Multi-traffic BE

To consider multiple BE traffic profiles, we use the multi-traffic extension presented in Section 5.3 while replacing $\bar{N}_S$ by $\bar{N}_S - \bar{N}_S^{ugs}$ in the different expressions. We then obtain the steady states $\pi_{be}(n_1, \dots, n_R)$ and the performance parameters in the same way.

7 UGS/ERTPS/BE MODEL

In this section, we first explain how to integrate the traffic of the ertPS service class in our multiclass modeling. Then, we validate the resulting UGS/ertPS/BE model through comparison with simulations.

7.1 Adding ertPS

To integrate the ertPS service class in our UGS/BE model, we follow a similar approach than in Section 6.1.1. Again, no specific scheduling is suggested in the WiMAX standard. However, the QoS needs characterizing each service class lead to consider that ertPS connections preempt BE connections but are preempted by UGS ones [23], [28], [29].

So, our UGS/ertPS/BE model consists in the cascading resolution of the three models, each corresponding to the characterization of the traffic of a given service class. The 3-steps resolution is as follows:

1. We first solve the UGS model to characterize the UGS traffic and compute $\bar{N}_S^{ugs}$, the mean number of slots occupied by the UGS connections.
2. We then solve the ertPS model. Although, this time, we only consider $\bar{N}_S - \bar{N}_S^{ugs}$ available slots in the cell. Similarly to the previous step, we compute the mean number $\bar{N}_S^{ertps}$ of slots occupied by the ertPS connections.
3. Finally, we solve the BE model to obtain the performance parameters of the BE service class using the AGG method as explained in Section 6.1.4. But, we here consider only $\bar{N}_S^{be} = \bar{N}_S - \bar{N}_S^{ugs} - \bar{N}_S^{ertps}$ available slots for the BE connections.

7.2 Validation

To validate our UGS/ertPS/BE model, we now compare its results with those of simulations. To this aim, we use a simulator which implements an ON/OFF traffic generator and a wireless channel for each user. Besides, a centralized scheduler allocates the slots to the active mobiles on a frame by frame basis according to their current MCS and service class. At each frame, the scheduler first serves the UGS connections, then the ertPS connections and lastly the BE connections.

In the simulations, the durations of UGS and ertPS ON and OFF periods are exponentially distributed. In addition, ertPS ON periods are decomposed in talk spurts and silences as recommended in [10].

BE OFF durations are also exponentially distributed. However, contrary to our model, the BE ON sizes are drawn according to a truncated Pareto distribution. Indeed, the truncated Pareto distribution is well known to fit the reality of WEB data traffic. The mean value of the truncated Pareto distribution is given by:

$$\bar{x}_{on} = \frac{\alpha L}{\alpha - 1} [1 - (L/H)^{\alpha-1}], \quad (71)$$

where α is the shape parameter, L is the minimum value of Pareto variable and H is the cutoff value for truncated Pareto distribution. For the sake of comparison, we keep the same the mean BE ON size $\bar{x}_{on}$. We set the value of $H = 100 \bar{x}_{on}$ and consider $\alpha = 1.2$ as suggested in [17]. Finally, we deduce the value of L using relation (71).

The wireless channel parameters are summarized in Table 1. In our analytical model, the channel model is assumed to be memoryless, i.e., MCS are independently drawn from frame to frame for each user. In order to show the robustness of this assumption, we model the channel variations of a simulated mobile by a finite state Markov chain (FSMC) corresponding to a slowly varying Nakagami-m fading channel as proposed in [27]. Each state of the FSMC corresponds to one of the 5 considered MCS (outage included). The state transition matrix C associated to this FSMC is as follows:

$$C = \begin{pmatrix} 0.020 & 0.980 & 0 & 0 & 0 \\ 0.163 & 0.120 & 0.717 & 0 & 0 \\ 0 & 0.277 & 0.620 & 0.103 & 0 \\ 0 & 0 & 0.398 & 0.080 & 0.522 \\ 0 & 0 & 0 & 0.089 & 0.911 \end{pmatrix}, \quad (72)$$

where $C_{i,j}$ is the probability that the MCS of a mobile change from MCS_i to MCS_j . The transitions occur only between adjacent states as we assume in the simulations that the channel is slowly fading. Let us highlight that the steady state probabilities of this FSMC are equal to the p_k probabilities used in the analytical model.

We have repeated comparisons between results from both our analytical models and simulations while considering all sorts of scenarios (e.g., different numbers of traffic profiles in each service classes, different BE scheduling, etc.). Here we present the results of two representative scenarios. In both cases, we observe the behaviors of $\bar{Q}$, $\bar{U}$ and $\bar{X}$ parameters of each service class obtained with our analytical models and compare them with the results of simulations.

7.2.1 First validation scenario

The cell and traffic parameters constituting this first scenario are detailed in Table 3. We assume a total number, N , of mobiles present in the cell ranging from 20 to 180. 50% of these mobiles generates UGS traffic and 30% rtPS traffic. The remaining 20% constitutes the population of BE mobiles in the cell and are equally distributed into two traffic profiles each representing 10% of the total population of mobiles. Finally, note that we consider the throttling policy to allocate slots among the BE connections and that we associate a different $MSTR$ to each BE traffic profile.

Fig. 7(a) presents the average numbers, $\bar{Q}$, of concurrent active connections in the cell. Obviously as the traffic load increases so does these numbers since more and more connections share the limited amount of resource.

The average utilization, $\bar{U}$, of the resource by UGS, ertPS and BE traffics is depicted in Fig. 7(b) and the instantaneous throughput per mobile, $\bar{X}$, is illustrated in Fig. 7(c). At first ($N < 50$), there is always enough resource to satisfy all demands. The parts of the resource occupied by each service class keep on increasing and all connections get their desired throughputs. However, when there are more mobiles in the cell ($N > 50$), the BE mobiles are the first to suffer the lack of resource. As such, their utilization of the frames and their throughputs dive. Finally, observe that when even more mobiles are present in the cell ($N > 150$), even ertPS connections start to see their performances deteriorate. This is explained by the fact that UGS connections are always served first, followed by ertPS connections and then by BE connections.

Parameter		Value
Number of slots per frame, N_S		450
Frame duration, T_F		5 ms
Limit on UGS calls, V_{max}		50
Limit on ertPS transfers, W_{max}		50
BE scheduling		throttling
Number of mobiles in the cell, N		20 to 180
UGS	Number of UGS mobiles, N_{ugs}	50% of N
	Guaranteed bit rate, GBR	200 Kbps
	Mean ON duration, t_{on}^{ugs}	60 s
	Mean OFF duration, t_{off}^{ugs}	120 s
ertPS	Number of ertPS mobiles, N_{ertps}	30% of N
	Guaranteed bit rate, GBR	400 Kbps
	Mean ON duration, t_{on}^{ertps}	90 s
	Mean OFF duration, t_{off}^{ertps}	120 s
Traffic profile		1 2
Number of mobiles per profile		$N_1 = N_2$
BE	Number of BE mobiles, N_{be}	20% of N
	Maximum bit rate, $MSTR$	1024 Kbps 2048 Kbps
	Mean ON size, $\bar{x}_{on}^{be}$	3 Mbits 3 Mbits
	Mean OFF duration, t_{off}^{be}	6 s 6 s

Table 3: First validation scenario.

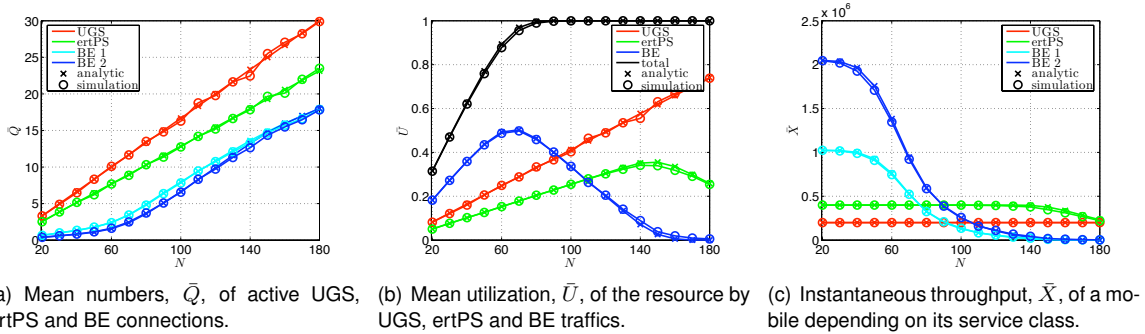


Figure 7: First validation scenario: Customary traffic parameters when traffic increases in the three service classes.

It is obvious from the curves depicted in the three figures that the results given by our analytical model match those of simulations. Indeed, the difference between them is less than 4% in most cases and less than 9% in the worst case. This shows the robustness of our model toward the distribution of BE ON sizes and the channel model.

7.2.2 Second validation scenario

The cell and traffic parameters of the second scenario are presented in Table 4. This time we consider a fixed population of $N = 60$ mobiles in the cell and observe the impact of $\bar{N}_S$ the mean number of available slots on the performance. The numbers of mobiles of each service class are as follows: 24 mobiles generate UGS traffic, 12 ertPS traffic and 24 BE traffic. Also, the UGS and BE connections are equally distributed into two traffic profiles. Lastly, note that the slots are allocated to BE connections

according to the throughput fairness policy.

Parameter		Value	
Number of slots per frame, $\bar{N}_S$		200 to 1000	
Frame duration, T_F		5 ms	
Limit on UGS calls, V_{max}		24	
Limit on ertPS transfers, W_{max}		12	
BE scheduling		throughput fairness	
Number of mobiles in the cell, N		60	
Traffic profile		1	2
UGS	Number of UGS mobiles, N_{ugs}	12	12
	Guaranteed bit rate, GBR	100 Kbps	200 Kbps
	Mean ON duration, $\bar{t}_{on}^{ugs}$	80 s	40 s
	Mean OFF duration, $\bar{t}_{off}^{ugs}$	120 s	120 s
ertPS	Number of ertPS mobiles, N_{ertps}	12	
	Guaranteed bit rate, GBR	300 Kbps	
	Mean ON duration, $\bar{t}_{on}^{ertps}$	80 s	
	Mean OFF duration, $\bar{t}_{off}^{ertps}$	120 s	
BE	Number of BE mobiles, N_{be}	12	12
	Mean ON size, $\bar{x}_{on}^{be}$	1 Mbits	4 Mbits
	Mean OFF duration, $\bar{t}_{off}^{be}$	5 s	5 s

Table 4: Second validation scenario.

Fig. 8(a), 8(b) and 8(c) respectively show the evolution of the $\bar{Q}$, $\bar{U}$ and $\bar{X}$ performance parameters of each traffic profile when $\bar{N}_S$ increases. The mean numbers of active UGS and ertPS mobiles are constant as their ON and OFF durations do not depend on the number of available slots.

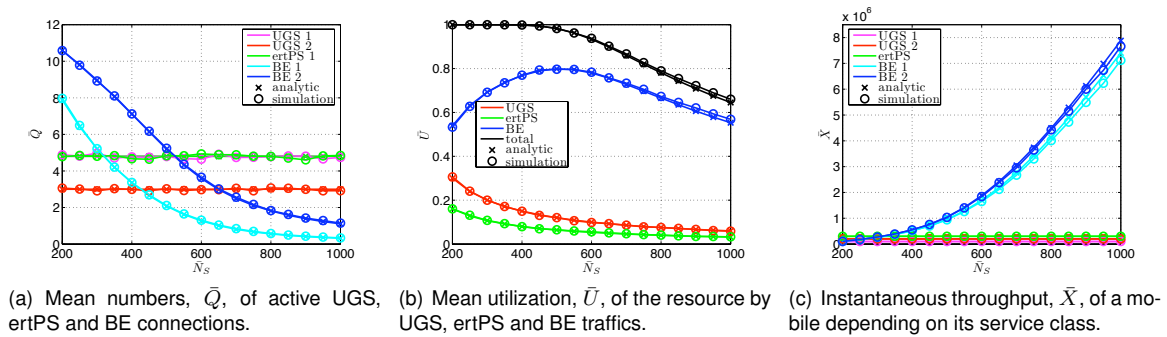


Figure 8: Second validation scenario: Customary traffic parameters when $\bar{N}_S$ the number of available slots increases.

On the contrary, the mean number of active BE connections drop when $\bar{N}_S$ increases. Indeed, more available slots mean faster BE transfers. The UGS and ertPS calls need a given average amount of resource to attain their respective GBR. This constant amount corresponds to a decreasing proportion of the frame as the value of $\bar{N}_S$ rises. When $\bar{N}_S$ is still small ($\bar{N}_S < 500$ slots), the BE mobiles need all the remaining slots of the frames. As such, their mean utilization of the frame grows as more and more slots are left to them by UGS and ertPS calls. On the opposite, when $\bar{N}_S$ is big enough ($\bar{N}_S \geq 500$ slots), the frames are not fully occupied anymore and the proportion of the frame needed for the BE transfers

decreases. Finally, the throughput of BE connections keep on increasing with the number of available slots because they are now scheduled according to the throughput fairness policy.

The results of our analytical model and those of simulations still match each other with the same precision. Here, the difference between them is less than 3% in most cases and less than 7% in the worst case. Once again, the curves show the strong robustness of our model toward the distribution of BE data volumes and the channel model. Indeed, the results of simulations considering a truncated Pareto distribution of BE ON sizes and a slowly fading channel match very closely those provided by our analytical model.

Besides, note that the results presented in this section are representative of the results obtained for the numerous scenarios we have considered. Indeed, each time, the simulations results validated our analytical models with similar accuracy.

At last, let us highlight that the simulations required very long computation durations (nearly a day) whereas the results of the analytic models were obtained instantaneously.

8 DIMENSIONING

In this section, we provide dimensioning examples to demonstrate possible applications of our models. We study the case of an operator wishing to dimension a WiMAX cell. This operator considers subscribers distributed in two profiles, business and domestic, corresponding to their utilization of the cell as detailed in Table 5.

Parameter		Value	
Number of slots per frame, N_S		450	
Frame duration, T_F		5 ms	
Limit on UGS calls, V_{max}		60	
Limit on ertPS transfers, W_{max}		30	
BE scheduling		opportunistic	
Subscriber profile		Business	Domestic
Number of subscribers in the cell		10 to 150	
Proportion of subscribers		50%	50%
UGS	Guaranteed bit rate, GBR	128 Kbps	64 Kbps
	Mean ON duration, t_{on}^{ugs}	120 s	60 s
	Mean OFF duration, t_{off}^{ugs}	180 s	180 s
ertPS	Guaranteed bit rate, GBR	256 Kbps	no domestic ertPS traffic
	Mean ON duration, t_{on}^{ertps}	60 s	
	Mean OFF duration, t_{off}^{ertps}	180 s	
BE	Mean ON size, $\bar{x}_{on}^{be}$	varying	
	Mean OFF duration, t_{off}^{be}	varying	

Table 5: Cell and traffic parameters.

The channel parameters are provided in Table 1.

Business subscribers are assumed to generate important UGS, ertPS and BE traffics. As such, a business subscriber is represented in our model by three mobiles, one for each service class, and the traffic he generates is the sum of the traffic generated by the three mobiles. On the opposite, domestic subscribers generate a lower amount of UGS, an equivalent amount of BE traffics and no ertPS traffic at all. So, a domestic subscriber is only represented in our model by two mobiles: one mobile for UGS and one for BE. In the following examples, we consider the same proportion of business and domestic

subscribers present in the cell.

Let us highlight that results can be obtained for any other possible configuration (i.e., any mono or multi-profile traffic scenario) by using the according models.

8.1 Performance graphs

Here, we study contour graphs in which a performance parameter is drawn as a function of the parameters to dimension, e.g., the total number of subscribers in the cell, the number of available slots, $\bar{N}_S$, traffic intensities, etc.

The mean radio resource utilization of the WiMAX cell $\bar{U}$, and the average throughput of a BE connection $\bar{X}_{be}$ are illustrated in Fig. 9(a) and 9(b).

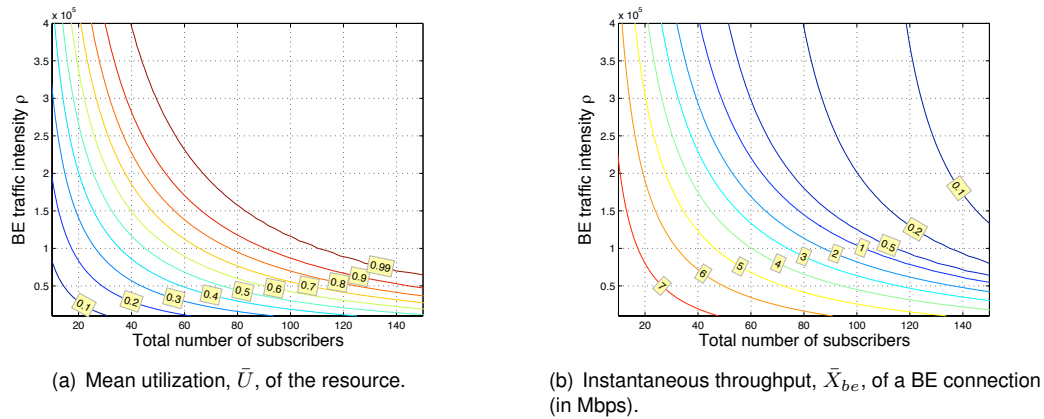


Figure 9: $\bar{U}$ and $\bar{X}_{be}$ traffic parameters when the population of subscribers in the cell and the BE traffic intensity increase.

These parameters are presented as functions of the total number of subscribers in the cell and ρ_{be} , the intensity of the traffic generated by the BE mobiles. ($\rho_{be} = \frac{\bar{x}_{on}^{be}}{\bar{t}_{off}^{be}}$ as described in Section 5.2.) To determine $\bar{U}$ or $\bar{X}_{be}$ for a given number of subscribers in the cell and a given BE traffic intensity, we first locate the point at these coordinates in the corresponding graph. This point is located between two contour lines. The line with the higher value gives an upper bound on the parameter value and the line with the smaller value gives a lower bound. For example, 40 subscribers in the cell and a BE traffic intensity of 200 Kbps (e.g., $\bar{x}_{on}^{be} = 2$ MB and $\bar{t}_{off}^{be} = 10$ s) lead to $0.7 < \bar{U} < 0.8$ and $3 \text{ Mbps} < \bar{X}_{be} < 4 \text{ Mbps}$.

As depicted on these graphs, the mean utilization $\bar{U}$ of the frame rises and the instantaneous throughput of the BE connections $\bar{X}_{be}$ drops when the number of subscribers and/or the BE traffic intensity increase.

Similarly, Fig. 10(a) and 10(b) show the same performance parameters but this time as functions of both the number of available slots and the total number of subscribers in the cell. ρ_{be} is set to 200 Kbps. We can observe that low frame utilizations and important BE throughputs are achieved when high amounts of slots are available and only small numbers of subscribers share the resource.

Each graph is the result of several thousands of input parameter sets. Obviously, any simulation tool or even any multidimensional Markov chain requiring numerical resolution, would have precluded the drawing of such graphs.

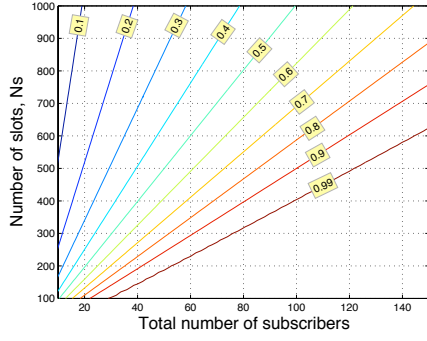
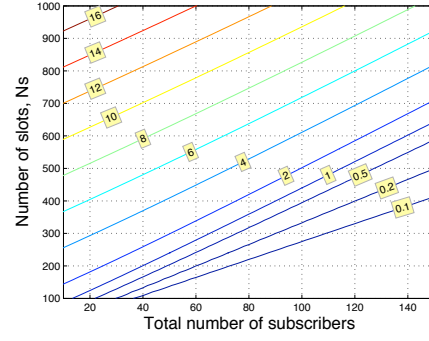
(a) Mean utilization, $\bar{U}$, of the resource.(b) Instantaneous throughput, $\bar{X}_{be}$, of a BE connection (in Mbps).

Figure 10: $\bar{U}$ and $\bar{X}_{be}$ traffic parameters when the population of subscribers in the cell and the number N_S of available downlink slots per frame increase.

8.2 Dimensioning study

Here, we show how our model can be advantageously used for dimensioning issues. Two examples, each corresponding to a certain QoS criterion, are presented in Fig. 11(a) and 11(b). To draw these two graphs, we computed the considered performance parameter ($\bar{U}$ or $\bar{X}_{be}$) for each possible number of subscribers and ρ_{be} pair while increasing the number N_S of available slots until the chosen QoS criterion could not be guaranteed anymore. Note that this straightforward method is only rendered possible due to the instantaneous resolution aspect of our model.

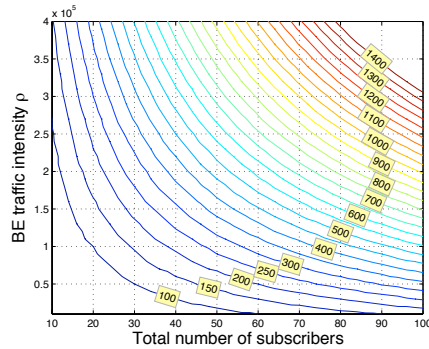
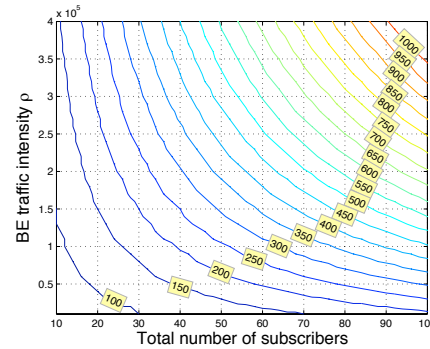
(a) Maximum number of slots, N_S^{max} to guarantee $\bar{U} \geq 0.85$.(b) Minimum number of slots, N_S^{min} to guarantee $\bar{X}_{be} \geq 1024$ Kbps.

Figure 11: Dimensioning of the number N_S of available downlink slots per frame.

In Fig. 11(a) we find the maximum number $\bar{N}_S^{max}$ of available slots guaranteeing an average radio utilization over 85%. This kind of criterion allows operators to avoid uneconomical over-provisioning of the network resources in regard to the traffic load of their customers. To obtain the optimal value of $\bar{N}_S^{max}$ associated with a number of subscribers and a given value of ρ_{be} , we look for the point at the corresponding coordinates in the graph. This point is located between two contour lines, and the one with the lower value gives the value of $\bar{N}_S^{max}$.

The QoS criterion chosen as a second example is the throughput per BE connection. We observed BE traffic throughputs because BE connections are served in last. As such, guaranteeing a minimum throughput for BE connections leads to ensure that UGS and rtPS connections obtain their reserved bit rate. We decided on 1024 Kbps as the arbitrary value of the minimum throughput of a BE connection. Now, we want to find the minimum number $\bar{N}_S^{min}$ of available slots guaranteeing this minimum throughput

threshold. In Fig. 11(b), a given point is located between two contour lines. The line with the higher value gives $\bar{N}_S^{min}$.

The graphs of Fig. 11(a) and 11(b) can be jointly used to satisfy multiple QoS criteria. For example, if we consider a WiMAX cell with 50 subscribers and a BE traffic intensity of 200 Kbps, Fig. 11(a) gives $\bar{N}_S^{max} = 500$ slots, and Fig. 11(b) gives $\bar{N}_S^{min} = 350$ slots. The combination of these two results recommends to have a number of slots $\bar{N}_S \in [350; 500]$ to guarantee both a reasonable resource utilization and an acceptable minimum throughput to the subscribers. This roughly corresponds to a system bandwidth between 8 and 10.5 MHz if we consider a downlink/uplink ratio of 2/3, a Fast Fourier Transform size equal to 2048, a PUSC subcarrier permutation and an average protocol overhead length of 4 symbols.

9 CONCLUSION

As deployment of 4G networks is underway, need arises for operators and manufacturers to develop dimensioning tools. In this paper, we have designed useful analytical models for the three service classes defined in the WiMAX standard: UGS, ertPS and BE, respectively corresponding to VoIP, VoIP with silence suppression and best-effort traffic. Also, we have explained how these models could be easily adapted for LTE systems.

As a first step, we presented models considering cells carrying either UGS, ertPS or BE traffic. For each model, we proposed multi-traffic extensions in order to account for users generating traffics of different intensities. Besides, we provided closed-form expressions of all the required performance parameters.

Then, we have detailed three methods to combine the UGS and BE models. The results obtained with each method proved to be very close. We deduced that averaging the numbers of slots occupied by the traffics of the different service classes was sufficient to combine the models. The resulting UGS/ertPS/BE model, based on the cascading resolution of the UGS, ertPS and BE models, is able to instantaneously evaluate the performance parameters of each service class. This multi-class model can even take into account users of the same service class generating traffics of different profiles with a minimal increase of its resolution complexity. Therefore it renders possible efficient and advanced dimensioning studies as shown in this paper.

We used extensive simulations to validate our analytical approach. To show the robustness of our traffic and channel assumptions, we compared our model's results to simulations considering a slowly fading channel and a truncated Pareto distribution of BE ON sizes. The simulation results showed the accuracy and robustness of our analytical modeling: for all considered scenarios, maximum relative errors never exceeded 9%.

Lastly, let us highlight that we focused in this paper on VoIP and best-effort traffic. However, another type of traffic is expected to be supported by 4G networks: streaming. The streaming connections need guarantees on both their minimum bit rate and the maximum latency of their packets. This last constraint renders the modeling of streaming traffic very difficult. In our next work, we will tackle the challenging task of developing a model dedicated to streaming in OFDMA-based cellular network. Ultimately, our aim will be to integrate this model in our cascade and thus be able to instantaneously obtain results for all types of traffic.

REFERENCES

- [1] 3GPP LTE: Homepage with specifications - <http://www.3gpp.org/LTE>.
- [2] IEEE Standard for local and metropolitan area networks - Part 16: Air Interface for Fixed Broadband Wireless Access Systems, 2004.
- [3] Draft IEEE std 802.16e/D9. IEEE Standard for local and metropolitan area networks - Part 16: Air Interface for Fixed Broadband Wireless Access Systems., 2005.
- [4] F. Baskett, K. Chandy, R. Muntz, and F. Palacios. Open, closed, and mixed networks of queues with different classes of customers. *Journal of the Association of Computing Machinery*, 22(2):248–260, April 1975.
- [5] A. Berger and Y. Kogan. Dimensioning bandwidth for elastic traffic in high-speed data networks. *IEEE/ACM Transactions on Networking*, 8(5):643–654, October 2000.
- [6] T. Bonald and A. Proutiere. Wireless downlink channels: User performance and cell dimensioning. In *ACM Mobicom*, 2003.
- [7] S. Borst. User-level performance of channel-aware scheduling algorithms in wireless data networks. In *IEEE Infocom*, 2003.
- [8] S. Borst and N. Hedge. Integration of Streaming and Elastic Traffic in Wireless Networks. In *IEEE Infocom*, 2007.
- [9] P. T. Brady. A technique for investigating ON/OFF patterns for speech. In *Bell Labs Syst. Tech. J.*, vol. 44, no. 1, pages 1–22, 1965.
- [10] P. T. Brady. A model for generating ON/OFF speech patterns in two-way conversations. In *Bell Labs Syst. Tech. J.*, vol. 48, no. 9, pages 2445–2472, 1969.
- [11] K.-C. Chen and J. R. B. de Marca. *Mobile WiMAX*. Wiley, February 2008.
- [12] F. Delcoigne, A. Proutière, and G. Rœgini. Modelling integration of streaming and data traffic. In *Performance Evaluation*, 2004.
- [13] M. Dirani, C. Tarhini, and T. Chahed. Cross-layer modeling of capacity in wireless networks Application to UMTS-HSDPA, IEEE802.11 WLAN and IEEE802.16 WiMAX. In *Computer Communications*, 2007.
- [14] S. Doirieux, B. Baynat, M. Maqbool, and M. Coupechoux. An Analytical Model for WiMAX Networks with Multiple Traffic Profiles and Throttling Policy. In *Proc. of the 7th International Symposium on Modeling and Optimization in Mobile, Ad Hoc, and Wireless Networks*, June 2009.
- [15] S. Doirieux, B. Baynat, M. Maqbool, and M. Coupechoux. An Efficient Analytical Model for the Dimensioning of WiMAX Networks Supporting Multi-profile Best Effort Traffic. *Computer Communications Journal*, to appear, 2010.
- [16] T. O. Engset. On the calculation of switches in an automatic telephone system. In *Tore Olaus Engset: The man behind the formula*, 1998.
- [17] A. Feldmann, A. C. Gilbert, P. Huang, and W. Willinger. Dynamics of IP traffic: A study of the role of variability and the impact of control. In *Computer Communication Review*, October 1999.
- [18] D. Heyman, T. Lakshman, and A. Neidhardt. New method for analyzing feedback protocols with applications to engineering web traffic over the internet. In *Proc. of the ACM Sigmetrics*, June 1997.
- [19] C. Hoymann. Analysis and performance evaluation of the OFDM-based metropolitan area network IEEE 802.16. In *Computer Networks*, 2005.

- [20] M. Ibrahim, K. Khawam, A. E. Samhat, and S. Tohme. Analytical framework for dimensioning hierarchical WiMax-WiFi networks. In *Computer Networks*, 2009.
- [21] V. B. Iversen. *Teletraffic engineering and network planning*. Technical University of Denmark, 2006.
- [22] A. Jensen. Truncated multidimensional distributions. *The Life and Works of A.K. Erlang*, pages 58–70, 1948.
- [23] C.-H. Jiang and T.-C. Tsai. Token bucket based CAC and packet scheduling for IEEE 802.16 broadband wireless access networks. In *Proc. of the 3rd IEEE Consumer Communications and Networking Conference*, January 2006.
- [24] S. Liu and J. Virtamo. Performance Analysis of Wireless Data Systems with a Finite Population of Mobile Users. In *19th ITC*, 2005.
- [25] J. Lu and M. Ma. A cross-layer elastic CAC and holistic opportunistic scheduling for QoS support in WiMAX. In *Computer Networks*, 2010.
- [26] M. Maqbool, M. Coupechoux, P. Godlewski, S. Doirieux, B. Baynat, and V. Capdevielle. Dimensioning Methodology for OFDMA Networks. In *Proc. of the Wireless World Research Forum (WWRF22)*, 2009.
- [27] D. Niyato and E. Hossain. A queuing-theoretic and optimization-based model for radio resource management in IEEE 802.16 broadband networks. *IEEE ToC (vol. 55)*, 2006.
- [28] A. Sayenko, O. Alanen, J. Karhula, and T. Hamalainen. Ensuring the QoS requirements in 802.16 scheduling. In *Proc. of the 9th ACM International Symposium on Modeling Analysis and Simulation of Wireless and Mobile Systems*, 2006.
- [29] V. Singh and V. Sharma. Efficient and fair scheduling of uplink and downlink in IEEE 802.16 OFDMA networks. In *Proc. of the IEEE Wireless Communications and Networking Conference*, 2006.
- [30] C. Tarhini and T. Chahed. On Mobility of Voice-Like and Data Traffic in IEEE802.16e. In *Global Telecommunications*, 2008.
- [31] Z. Yumei and S. Yu. Scheduling Algorithm with Quality of Service Support in IEEE 802.16 Networks. In *Computer Networks*, 2009.

Synergies of Radio Frequency and Free Space Optics Communication: New Hybrid Solutions for Next Generation Wireless Mesh Networks

Mohammad Ahmad
*Mobile Services Department
Salesforce.com
San Francisco, 94105, USA*

mahmad@salesforce.com

Ala Al-Fuqaha
*Computer Science Department
Western Michigan University
Kalamazoo, 49008, USA*

alfuqaha@cs.wmich.edu

Osama Awwad
*Computer Science Department
Western Michigan University
Kalamazoo, 49008, USA*

oaawwad@cs.wmich.edu

Bilal Khan
*John Jay College
City University of New York
New York, 10019, USA*

bkhan@jjay.cuny.edu

Abstract

In this article, we describe the limitations of radio frequency (RF) and Free Space Optics (FSO) technologies, and show that a hybrid approach that uses both communication media in a suitably coordinated manner is capable of addressing the shortcomings of each. The nature of the required coordination between RF and FSO communication technologies, and the merits derived thereof, are the subject of our research. We report on an initial hardware implementation of a prototype Hybrid RF/FSO node. The properties of the node are used to inform the development of an Integer Linear Programming (ILP) model for this coordination process in RF/FSO networks. We show that by making suitable choices of beam-widths and power levels, the proposed model can be used to design robust hybrid RF/FSO communication infrastructures that minimize power consumption, while satisfying specified joint throughput and end-to-end delay requirements.

Keywords: Hybrid RF/FSO, Wireless Optical, RF, FSO, Topology Control, QoS, Linear Programming.

1. INTRODUCTION

Over the past years, the fields of computer and telecommunication networks have experienced tremendous growth, as many new communication technologies have been developed to help address scalability challenges, arising from requirements in the face of widespread adoption. These requirements include: higher bandwidth, security and privacy guarantees, low end to end delay and responsiveness, high network connectivity and coverage, etc.

In the history of the intellectual development of ideas, there is frequent academic bias towards “pure solutions”—in this case, communication systems that use just one core communication technology. Unfortunately, no single communication technology has been found, to date, which can satisfy all these requirements. For instance, in ad-hoc networks, researchers recently started

proposing hybrid models to address the scalability problem in mobile ad hoc networks. These solutions use a sparse network of fixed, wired base stations within an ad hoc network to improve the network properties. In such hybrid networks, nodes can communicate either in a multi-hop fashion using both wireless and wired links. Figure 1 demonstrates an example of such a hybrid ad hoc network. Certainly, a model in mixed nodes (some fixed and some mobile) and mixed edges (some wireless and some wired), is less aesthetically appealing from the vantage point of mathematical uniformity. However, issues of aesthetics aside, the hybrid model has been shown to successfully trade off homogeneity for better performance [2, 4]. Most (but not all) nodes are mobile, and in exchange, nodes sometimes (but not always) have access to an orthogonal communication medium, which they can use to improve network properties. Employing two or more technologies simultaneously, and carefully coordinating their use, in effect leverages the advantages of each while circumventing their weaknesses.

The main shortcoming of the aforementioned hybrid wireless/wired networks is that they have high deployment costs [6]. In establishing telecommunications infrastructures in hostile environments (e.g. disaster recovery and battlefield settings), establishing a wired network is not feasible because of financial costs, setup time, and infrastructure vulnerability. Simply put, a fiber can be cut anywhere along its length, rendering the link inoperative. Since the endpoints of the fiber are static, so is the fiber—and thus the wired communication layer of the network is much more vulnerable to physical attack. How might such vulnerabilities be addressed? First, one might seek to make the links impossible to cut; the only way to break a link is to attack the endpoint node. And second, one might make it easier to move the link endpoint. This is effectively what is achieved by Free Space Optics (FSO) technology.

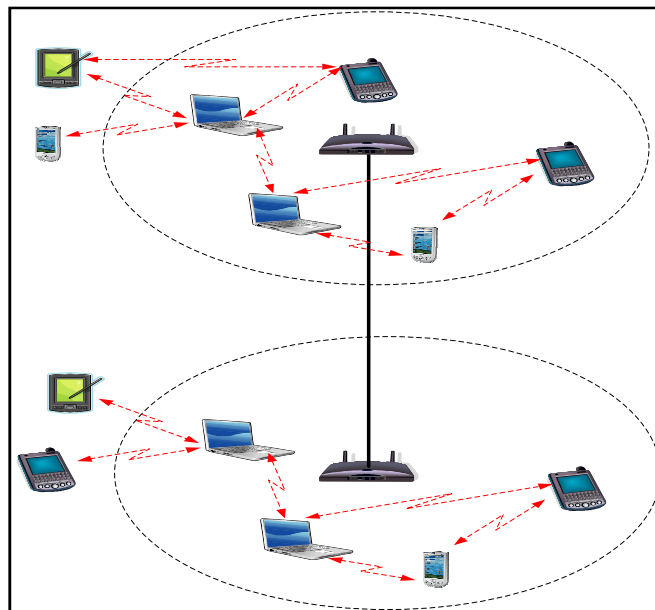


FIGURE 1:Hybrid solution in ad hoc networks.

Historically, FSO was proposed and developed for wireless communications to address QoS, interference, and security limitations of RF. FSO (like wired networks) has the potential for much higher reliable link data rates compared to present RF technology. In addition, because FSO uses directed optical transmissions in which channel beam-width is adjustable, inter-FSO communication interference can be controlled. Finally, the ability to use narrow beam-widths sidesteps the reliance on radial broadcasting in RF; this provides FSO with an advantage over RF in terms of security against eavesdropping [8].

Here consider the coordinated use of Free Space Optics (FSO) and Radio Frequency (RF) communications. At present, most wireless networks are deployed strictly in the RF domain, since RF channels provide natural support for radial broadcast operations. The downside of RF channels is that they introduce many limiting externalities that make providing scalable quality of service (QoS) support difficult, if not intractable [10]. These well-known technical challenges include bandwidth scarcity, high interference, and vulnerabilities to eavesdropping—all of which can be mitigated through the carefully coordinated incorporation of FSO communications. The benefits of FSO do not come without a price, however. Most notable of these drawbacks is the need to maintain line of sight (LOS) between the transmitter and the receiver during the course of communication. In addition, FSO link availability can be degraded by adverse weather conditions like fog, rain, snow, and haze. A hybrid approach that uses both RF and FSO communication must be designed with the strengths and weaknesses of each of the two channel types in mind. One channel type can serve as the “backup” channel and facilitate control of the other type [12].

We believe that the hybrid RF/FSO approach is especially well suited to three application areas: (i) battlefield environments, (ii) intelligent transportation systems (ITS), and (iii) telemetry and telesurgery. In battlefield environments, the hybrid system enables Free Space Optical (FSO) communications bandwidth without giving up RF reliability and “adverse-weather” performance [40]. The RF/FSO technology has many applications in next generation military networks [1] such as:

- Ultra high capacity cross-links between satellites and potentially space-to-air or space-to-ground platforms.
- Airborne networks.
- Air-to-ground links to increase the high-rate RF links currently used for communication.

ITS can benefit from FR/FSO channels in providing a queuing system with alternating service rates [41]. ITS objectives include the development of effective roadside-equipment (RSE) to roadside equipment communication. Unfortunately, the broadcast nature and the low data rates of RF channels make the technology unsuitable for dense deployments. In contrast, the directionality and high data rates of FSO make them an attractive alternative. Telemetry and telesurgery seek to provide of emergency medical care to people in remote areas and/or harsh environments, such as war zones, Polar Regions, or space stations—even if a doctor is not locally available. In this kind of application, the surgeons operate using robotic arms, based on a visual feedback from miniature video cameras at the patient location. Receiving such video streams requires reliable and high data rate connections, which are better achieved using a hybrid RF/ FSO [42].

The remainder of this paper is organized as follows: In Section 2, we describe the challenges faced in FSO and RF technologies when considered in isolation, as well as the particular difficulties of a hybrid model. In Section 3, we describe a working hardware prototype of a hybrid RF/FSO node. In Section 4, we formulate the optimization problem of topology control for hybrid RF/FSO networks. Finally, in Section 5, we present overall assessments and conclusions, closing the planned future trajectory of our research efforts.

2. THE CHALLENGES OF RF, FSO, AND HYBRIDS

2.1 Free Space Optics (FSO) Technology

FSO communications is a line-of-sight technology that uses invisible beams of light to transmit and receive voice, video and data information through the air with a rate up to 1.25 Gbps. FSO operates in the infrared (IR) spectral range—akin to many fiber optics solutions—by using wavelengths near the visible spectrum around 850 and 1550nm (corresponding to 200THz portion) of the spectrum [3, 5]. The Federal Communication Commission (FCC) does not presently require a permit for FSO transmitters, since FSO operates in an unregulated region of the spectrum [43].

FSO signals are transmitted by an infrared laser or LED, and received by high sensitivity receivers at the remote link side [14, 44]. Typically, FSO systems are based on communications links between pairs of free space optical transceivers (transmitter and receiver) that provides full duplex communication links. The optical units use lenses or telescopes, which narrow the light beam and project it toward the receivers. The transmitted light is picked up at the receiver side by a lens, which is connected to a sensitive detector via optical fiber [5, 11, 12, 45]. FSO transceivers can be positioned on a rooftop, at the corner of a building, or indoors behind a building as (See Figure 2). Presently FSO transmission ranges vary between 700 feet to about a mile in a clear, dry atmosphere [11, 46].

In FSO, channel impairments can be caused by fiber attenuation, insertion losses, chromatic dispersion, PMD, or fiber nonlinearities. The attenuation in fiber can be described as the change in signal power over the course of the transmission, and can be defined as: $dP/dz = -\alpha P$, where α is the power attenuation coefficient per unit length [47]. In general, the attenuation of an optical fiber measures the amount of light lost between input and output. Total attenuation is the sum of all losses [48]. Moreover, fiber loss can be caused by the fiber splices and fiber connectors. The fiber splices can be fused or joined together by some mechanical means, with typical attenuation being 0.01–0.1 dB per fused splice and slightly above 0.1 dB per mechanical splice [47]. Chromatic Dispersion occurs when optical pulses are spread out into a much broader temporal distribution than actual optical fiber channel [47, 49]. Another cause of impairments is Polarization Mode Dispersion (PMD). It happens when the polarization unit vector, representing the state of polarization (SOP) of the electric field vector, does not remain stable in practical optical fibers; rather, it changes in a random manner along the fiber because of its fluctuating birefringence [47, 50]. At last, Fiber nonlinearities present a new realm of obstacle that must be overcome. Fiber nonlinearities arise from two basic mechanisms: The most detrimental mechanism arises from the refractive index of glass being dependent on the optical power going through the material. The second mechanism that causes fiber nonlinearities is the scattering phenomena [51].

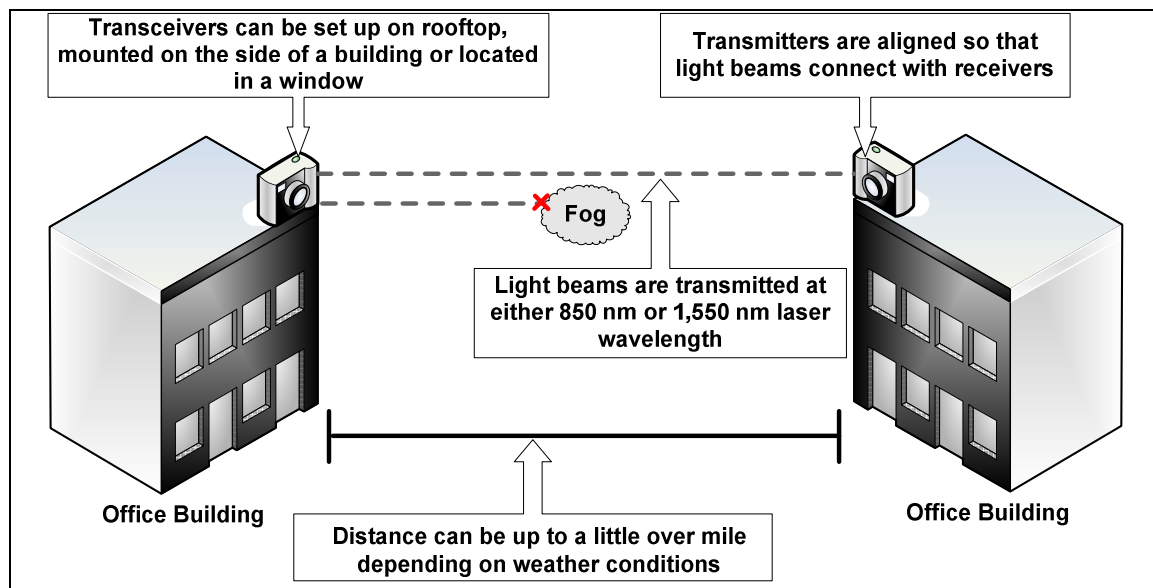


FIGURE 2:FSO connections with line of sight access.

FSO has several advantages over existing wireless or wired communication technologies. As a communication technology enjoys the higher capacity advantages of fiber, so greater data access can be provided. The spatial confinement of each FSO laser beam allows for beams to operate independently, and provides virtually unlimited frequency reuse. FSO is extremely secure, since an adversary cannot intercept communications without placing itself in direct path between two FSO

points, which is difficult and often detectable (as energy loss) by the legitimate endpoints [1]. Finally, FSO technology supports much more rapid deployment compared to other access technology [9]. In fact, FSO offers significant cost savings since there is no need to make infrastructure changes (e.g. digging up streets to lay cables), since air acts as the transmission medium instead of fiber optics. The installation of FSO systems is cheaper than fiber optics, DSL or cable modem services, and many FSO installations can be completed in a few days [6, 9].

The source of FSO's advantages is also a source of limitations: the air transmission medium. Whenever the medium's condition is unstable or unpredictable, it becomes difficult to manage free space optics transmissions, since volatility in the medium cause disruptions in "line of sight", and thus impact system availability and capacity. The main causes of such disruptions include:

- **Fog**, composed of water droplets, can completely obstruct the passage of light by a combination of absorption, scattering and reflection. This can reduce the power density of the transmitted beam, and therefore decrease FSO link transmission range.
- **Absorption** decreases the power density of the link and degrades system availability.
- **Scattering** occurs when photons collide with the airborne particles. Even in the absence of energy loss, this results in directional redistribution of beam energy, which, in turn, causes multipath effects.
- **Scintillation**, occurs when light propagates through space which contains sharp temperature variations across different pockets of air. Scintillation causes periodic fluctuations of the signal amplitude, which can lead to a sharp increase in the bit-error-rate of the FSO link.

There are additional secondary challenges faced by FSO systems such as physical obstructions, building sway, rain, snow, and eye safety [1, 3, 5, 11].

While FSO technology has been in use by the military and NASA for almost 30 years, it is still a relatively new technology in civilian telecommunication applications. Recently, it has been considered as a candidate replacement for fiber optics, in solutions seeking to develop a high-speed wireless network access [9, 11].

2.2 Radio Frequency (RF) Technology

RF refers to a number of different technologies, products and industries. Broadly speaking, RF refers to the portion of the frequency spectrum in which radio waves can be generated (resp. received) by driving (resp. detecting) alternating current through an antenna. In RF systems, data is transmitted through air by digital radio signals at a given frequency, typically in the 3kHz to 300GHz range. Such communication can be used to maintain bi-directional, online radio connection between a mobile telephone and an antenna host for example.

RF technology has the advantage of being wireless, and thus requires no cable deployment expenses to provide communication access. This instant accessibility also provides a timesaving in the deployment process. Also, the real-time accuracy of RF is high (>99%) and can be considered a general advantage.

On the other hand, the development of technologies based on RF communications is hindered by regulation in the assignment of radio spectrum bands by the FCC and National Telecommunications and Information Administration (NTIA). As of today, the FCC has allocated frequencies between 9 kHz and 275 GHz, with the highest bands reserved for radio astronomy and satellites [3, 52].

2.3 RF vs.FSO

RF and FSO are two different wireless technologies and there is no direct *competition* between them. In what follows, we summarize the main differences between them:

- A) *Throughput.* The major advantage of FSO is its high throughput compared to RF. At present, FSO can provide a throughput of 1.25Gbps for links of up to a kilometer or more. In contrast, wireless RF technologies such as 802.11a provide link throughput of tens of Mbps with distances of tens of meters. Although Ultra Wide Band (UWB) technology can provide a throughput of hundreds of Mbps, it usually drops to levels lower than even 802.11a at long ranges [11, 14, 15].
- B) *Deployment Costs.* FSO links have low cost of deployment compared to fibers, and system installation can be completed within few days [23].
- C) *Reliability.* FSO link traffic can experience disruptions due to atmospheric turbulence, while RF links are largely immune to such effects [17]. Low deployment cost makes FSO easy to maintain in the event of equipment failures, and thus suitable as components in a high reliability backbone for military applications and MANs.
- D) *Security.* RF link transmissions can be intercepted or detected by 3rd parties. This makes RF unsecure without addition time and energy expenditures [18] in the form of encryption algorithms. In comparison, FSO technology is extremely secure, even without encryption overlays because it uses a directional narrow beam point-to-point line of sight transmission, whose interception is easy to detect. A laser FSO at 1550nm offers an excellent wireless transmission solution and provides the highest possible level of physical layer security [19, 20].
- E) *Availability.* Link availability for RF is frequently above 99.999%. Unfortunately, performance level for FSO may vary from one area to another depending on the atmospheric conditions [36]. Effective FSO transmission range can thus vary from more than 700 feet up to about a mile in a clear, dry atmosphere [8], though ensuring carrier-class availability often requires that link length be limited to 200–500 meters [6].
- F) *Frequency Bands.* RF systems operate in frequency bands between 3kHz and 300GHz. Some frequencies in this spectrum range are regulated and assigned by the FCC [21]. In contrast, FSO systems operate in frequencies around 200THz, which presently requires no RF spectrum licensing [22].
- G) *Eye Safety.* RF signals are relatively harmless to the human eye. In comparison, FSO signals can damage the cornea and lens of the eye, particularly if transmissions are in the visible or near infrared spectrum regions¹ (100-1,200nm).
- H) *Calibration:* FSO systems require beam pointing, acquisition, and tracking, since optical transceiver must be aligned to each other for communication to take place. The precision tolerances mandated in the alignment process depend on the link power budget and the design of transmitters and receivers [1]. In contrast, in RF systems, there is no special alignment process required between transmitter and receiver, since transmission is radial and does not operate via line of sight [18].

2.4 Major Technical Challenges in Hybrid RF/FSO

By hybridizing RF and FSO technologies intelligently, we seek to address the needs of current applications in a manner that supersedes what is achievable using a single technology alone. Doing so, however, requires surmounting certain technical challenges. In this section we describe the three major technical challenges: (A) Pointing, Acquisition and Tracking, (B) Routing and Path

¹This is a compelling reason for selecting the longer 1550nm wavelength as the basis for FSO Systems.

Protection, and (C) Topology Control. We will also present currently proposed solutions for each of these challenges.

A) *Pointing Acquisition and Tracking*

One challenge faced in real FSO deployments is the necessity to maintain Line Of Sight (LOS) between sender and receiver during transmission using the FSO channel. This problem is called the “Pointing-Acquisition-Tracking” (PAT) process. We do not consider PAT to be a limitation or disadvantage, since several researchers have developed adequate solutions to this problem using different techniques. Nevertheless, PAT is at present, an active research area. In considering hybridized RF/FSO schemes where nodes are mobile, the PAT process gets even more difficult.

Derenick et al. proposed a hierarchical PAT system in which they used a vision-based system to maintain LOS [14]. The system assumes prior knowledge regarding the initial position of each FSO node and its partner receiver. In addition, both sender and receiver need to refer to a relative “origin” to interpret position and orientation updates. The authors accomplished this using a GPS system and inclinometer sensors. The alignment process is carried out using feedback from a high zoom camera system. The dependency on a GPS system makes this approach restricted for large-scale outdoor applications.

Akella et al. the authors proposed an omnidirectional spherical FSO transceiver to maintain LOS in hybrid RF/FSO MANETs [24] as illustrated in Figure 3. The surface of this sphere is covered with FSO transceivers, which contain LEDs and photodetectors. When two mobile FSO nodes in this design move away from each other the existing LOS link will be lost and a new one will be established. The omnidirectional FSO transceiver is a promising approach to the alignment problem in mobile environments. Unfortunately, current hardware implementations are not fast enough to switch the beam to another FSO channel without significant breaks in connectivity—especially when the environment is highly dynamic one (e.g. a tactical setting). Omnidirectional FSO transceivers may also sacrifice some covertness if movement is discontinuous, and the process of repairing broken optical channels is not done carefully.

Nichols proposed a framework using a dynamic RF/FSO staged acquisition technique [21]. The proposed framework decomposes PAT into three major stages as illustrated in Figure 5: A very wide RF signal is sent by the sender in order to identify the location of the FSO receiver. After locating the FSO receiver object, a wide FSO beam is sent by the sender to establish the FSO link. Finally, the sender and receiver narrow the FSO beam width down, to meet the QoS requirement. This approach also seems a promising solution to solve PAT problem in dynamic environment. Unfortunately, there has been no implementation or field test of this proposal² to date.

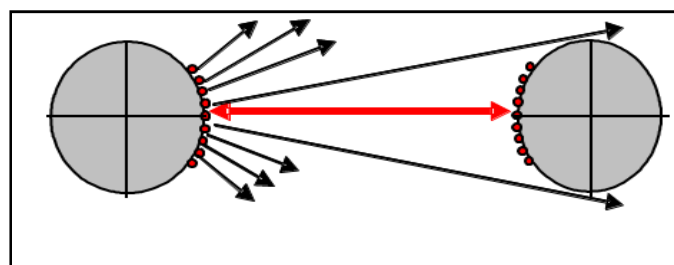


FIGURE 3:FSO connections with line of sight access.

² No hardware implementation of the proposed system exists as of the writing of this article.

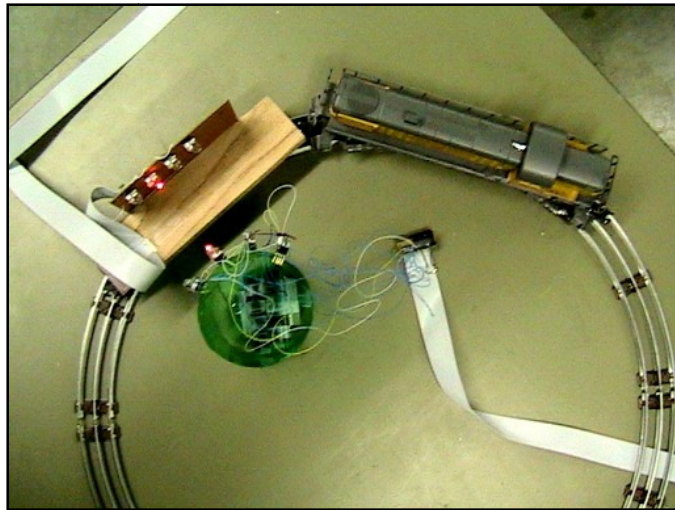


FIGURE 4:3D-Omnidirectional spherical FSO transceiver [24]

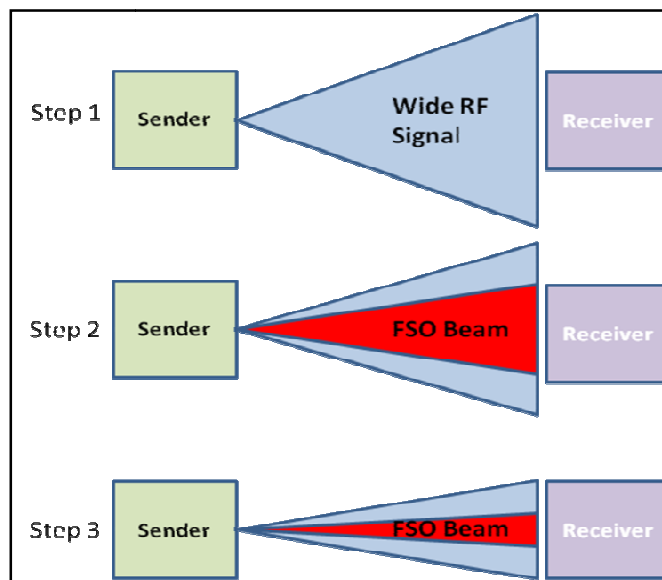


FIGURE 5:Dynamic RF/FSO staged acquisition technique [21]

B) Routing and Path Protection in Hybrid RF/FSO

The disparate (and time-varying) natures of the RF and FSO channels make routing a difficult problem in hybrid RF/FSO networks. Kashyap and Shayman introduced a routing framework for hybrid RF/FSO networks that relies on the assumption that RF links should have lower backbone traffic demand compared with FSO links [17]. Based on this assumption, they introduced a concept called “critically index” which determines the fraction of each traffic profile entry. A path is then computed for each traffic profile entry using the shortest widest path (SWP) algorithm, which is implemented as an extension of the OSPF routing protocol [25].

The fact RF links are more reliable than FSO channels led several researchers to consider using the RF channel for path protection in hybrid RF/FSO networks. For instance the Kashyap et al

proposed a joint topology control and routing framework where the RF links serve to provide instantaneous backup to traffic in hybrid RF/FSO networks when FSO links are degraded [27]. Their proposal does not treat the RF channel as “just an additional wavelength”, since the failure of the RF channel is not a single point failure as the case in the wired optical links³. Because of this fact, the problem cannot be solved using classical path protection schemes developed for optical networks (e.g. see Ramamurthy and Mukherjee [26]) and requires the development of new strategies.

C) Topology Control

The disparate (and time-varying) natures of the RF and FSO channels also make topology control a difficult problem in hybrid RF/FSO networks. Topology control in traditional pure-RF networks has been extensively studied. Examples of objectives considered have been, for example, how to adjust the power levels to have connected network yet use the minimum possible total power; or, how to reduce interference to meet some specific QoS requirements, [28, 29, 30] etc. Very limited research has been carried out in the area of topology control for hybrid RF/FSO.

Baskaran et al. studied the ability to provide topology reconfiguration according to changes in links capacities and traffic demands in RF/FSO networks [31]. The authors developed a heuristic for finding a topology configuration with the minimum packet drop rate. The heuristic cost of the packet drop rate includes factors reflecting both the link congestion cost and the packet drops that occur while the topology reconfiguration process is underway.

In section 4, we will present an integer linear program (ILP) model to solve the topology control problem, based on adapting and adjusting the transmission power and the beam-width of individual nodes according to Quality of Service (QoS) requirements. Our main focus will be to meet QoS requirements in terms of end-to-end delay and throughput jointly.

3. HYBRID RF/FSO PROTOTYPE

In our research lab, we were able to implement a hybrid RF/FSO prototype. To accomplish this objective, the development work was carried out in the following steps:

- 1) Implement a real system to transmit and receive data over one FSO channel.
- 2) Enhance the FSO transceiver to be able to operate over three FSO channels.
- 3) Implement an RF transceiver.
- 4) Improve the RF transceiver to be able to operate over three RF channels.

Develop and propose a hybrid system that combines systems produced in steps (2) and (4) to provide RF and FSO data access.

3.1 Single-channel FSO System Design

The central element of an FSO design is to have a line of sight data transmission. FSO technology pioneers such as LightPointe or fSONA offer a wide variety of FSO transceivers. In this section we design and implement a simple prototype FSO system that has the capability to send and receive data over a single FSO channel. We found that a simple infrared circuit can achieve this adequately, since FSO technology operates in the IR spectrum. The main difference between the industrial FSO transceivers and our infrared circuit is the transmission range. In industrial FSO laser transceivers, the transmission range can reach up to 1 mile in clear atmospheric conditions, while it is limited to just a few meters in the infrared. However, coverage distance is not a critical issue at this point since our goal is a prototype that will be used to test the signals transmitted and received.

³ Take for example, the interference in RF area, which can hit all RF channels in that interference range and not only a single RF channel.

The requirements and the IC components involved in our implementation of this system are as follows:

- A PC workstation
- Serial port interfacing Cable
- A software or a program that can transmit and receive data at the same time over the serial cable
- UART MAX 232 Dallas
- MCP 2120
- A Timer or a crystal oscillator TLC555
- MAX 3120
- Infrared LED and photodiode

Serial Port Interfacing

The system was designed to read and write data from and to a PC workstation. One of the most common ways to interface circuits with PCs is through serial port communication. In our prototype, a simple JAVA application was developed to send and receive data over serial ports using the RS-232 cable.

UART MAX232 Dallas

The MAX232 chip is simple a driver/receiver that converts TTL/CMOS input levels into TIA/EIA-232-F levels and vice versa. In our design, it acts as a bridge to carry data to or from the RS232 cable to or from the IR transceivers [32].

MCP 2120

This device is an infrared encoder/decoder, which is placed between a UART and an infrared (IR) optical transceiver where the data received from a standard UART is encoded (modulated) and then sent as electrical pulses to the IR Transceiver. When the IR Transceiver receives data, it sends it as electrical pulses that are decoded (demodulated) by the MCP2120. This modulation and demodulation method is performed in accordance with the IrDA standard. The MCP supports both hardware and software baud rate selection. In this design, we selected hardware baud rate by tying the pins BAUD2, BAUD1 and BAUD0 to the values 1,0,0 respectively. By doing so, the baud rate by the MCP 2120 was $F_{osc}/64$ where F_{osc} is the frequency generated by an external timer or a crystal oscillator. It is necessary to match the baud rates at the MCP2120 and the RS 232 serial cable so that the system works properly [33].

A TLC555 timer or a crystal oscillator

Since the MCP 2120 encoder requires an external timer or oscillator to generate the required baud rate, a TLC555 timer was used to supply the MCP2120 with the appropriate frequency [34]. The TLC555 timer was set up to generate a frequency of 6037.66 KHz which made the MCP 2120 provide a baud rate of $(F_{osc}/64) = (6037.66 \text{ KHz}/64) = 9433 \text{ bps}$. Also, a 1.84 MHz crystal oscillator was used instead of the TLC555 timer to supply the MCP2120. The generated baud rate was $1.84 \text{ MHz}/64 = 28.8 \text{ Kbps}$.

MAX 3120

This is an 8-pin chip that operates as an IR transceiver. With this chip and few more components—an LED, a photodiode, capacitor, and a few current-setting resistors—the Infrared application design is complete and ready to send and receive data over one channel [35].

Infrared LED and Photodiode

We used one IR LED transmitter and one photodiode detector to transmit and receive infrared data over one channel. The chosen LED has a wavelength peak emission of 950nm and the photodiode has a spectral bandwidth range of 620nm to 980nm.

As illustrated in Figure 6 and Figure 7, the system was built so that it can transmit and receive data at the same time. The same PC workstation can transmitted signals and receive them. To do so, we developed a simple JAVA program that uses threads to continue listening for the user commands while it keeps reading any data that may have been received by the serial port.

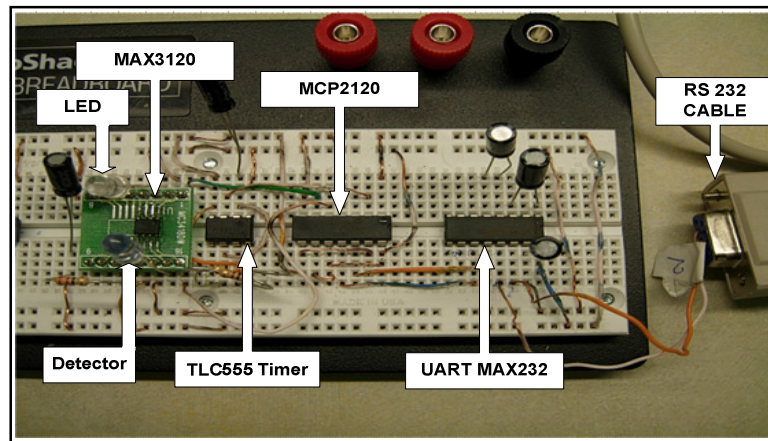


FIGURE 6:IR transceiver circuit for 1-channel communication.

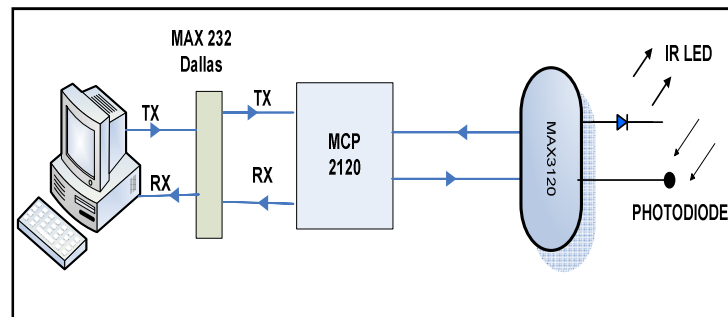


FIGURE 7:Block diagram of the 1-channel IR transceiver.

The system was tested, by sending and receiving English, encoded as a stream of 8-bit ASCII characters. The transmission was not 100% accurate, especially when any object obstructed the line of sight between the LED and the photodiode. In fact, this was expected because FSO and infrared are line of sight technologies and, therefore, data transmission is sensitive to alignment/obstructions between the transmitter and receiver⁴.

3.2 Multi-channel FSO System Design

The motivation behind transmitting infrared data over three channels is to improve system performance by providing more redundant channels for transmission, and thus higher reliability and availability of the communication system. In this design we added a PIC microcontroller to do FSO channel duplication. The PIC can take input on one pin and duplicate it to any of the three independent channel circuits, thereby providing the ability to do smart multiplexing/de-multiplexing to select based on the best quality of service channel or other global coordination logistics. The three channel circuits must each operate on a different wavelength to ensure concurrent multi-channel transmission can really provide an increase in reliability.

⁴ Other reasons for the data loss might come from the small difference in the baud rate. The serial cable has a set of predefined baud rates at which it transmits the signals. These rates were not exactly matching the baud rates at the MCP2120 encoder/decoder.

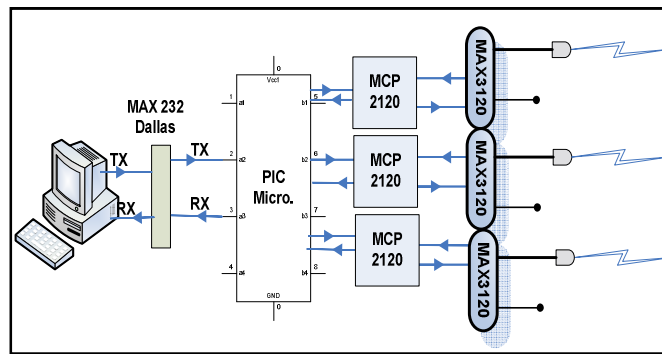


FIGURE 8:Block diagram of 3-channel IR transceivers.

This system is illustrated in Figure 8. As seen, it is composed of one UART, three infrared encoders/decoders and three transceivers. Also, there are three LEDs and three photodiodes, each of which operates in different wavelengths.

3.3 Single-channel RF System Design

RF systems can be built over different architectures and implementations, depending on the underlying technology used (e.g. Bluetooth, ZigBee, OFDM or any other circuit that works on the RF frequency band). The one-channel RF transceiver we implemented works on the 434 MHz frequency band; it has been tested by sending and receiving English ASCII characters using a test jig that is very similar to the unit tests of the FSO circuit. The components required to build this transceiver were:

- A PC workstation
- Serial port interfacing Cable
- A software or a program that can transmit and receive data at the same time over the serial cable
- UART MAX 232 Dallas
- MCP 2120
- TLP434A Transmitter
- RLP434A Receiver
- Two 400 MHz RF antennas

The first four components listed above were used in the RF system identically to their manner of use in the FSO transceiver—namely to send data to and receive data from the workstation. In place of IR LED and photodiode, two RF modules used to send and receive the RF signals: these were the TLP434A and RLP434A, provided by LAIPAC [36]. The TLP434A module is an Ultra Small Wireless Transmitter that is ideal for remote control projects or data transfers to a remote object via the 434MHz frequency band. The RLP434A module is a Compact Radio Receiver (RF) that works directly with the TLP434A transmitter over the 433.92Mhz frequency band. Each of the two modules requires a 400 MHz antenna for transmitting and receiving the signals [37]. The implementation is illustrated in Figures 9 and 10.

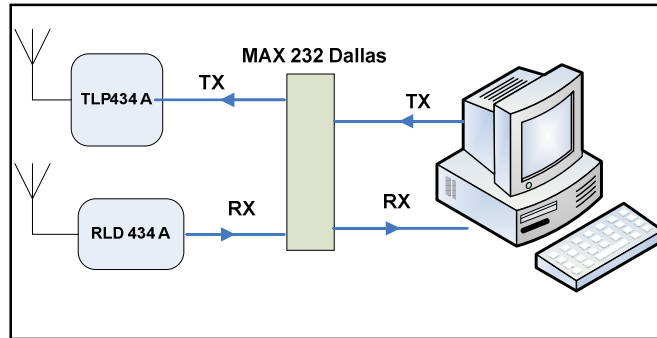


FIGURE 9:One channel RF transceiver block diagram

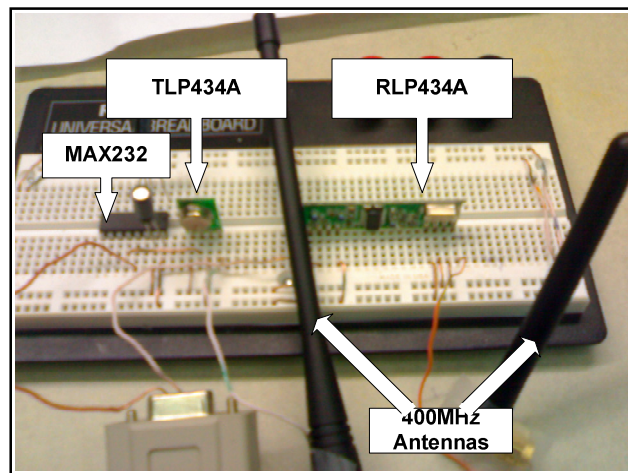


FIGURE 10:Single RF transceiver system

3.4 Multi-channel RF System Design

The single RF channel system was extended to provide transmission over three different RF channels, each based on a different frequency band. The motivation behind this was, once again, to improve the reliability and availability of the channel. This revised system used the three LAIPAC RF modules: TLP434A/RLP434A, RF900DV and RF2400DV to provide frequency bands of 434 MHz, 900 MHz and 2.4 GHz successively [38, 39]. A PIC microcontroller was also used in this design to distribute the single input data channel into three and to offer the ability to do smart multiplexing/de-multiplexing based on the best quality of service channel or other global coordination logistics.

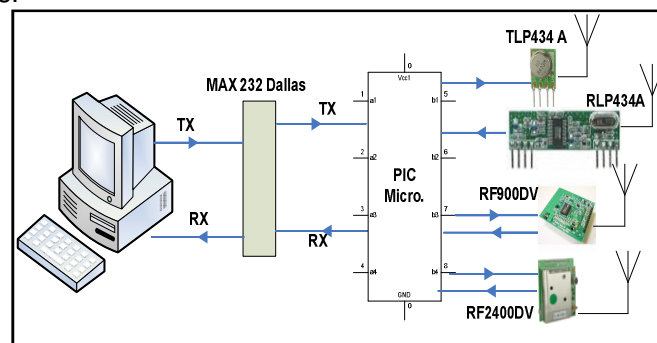


FIGURE 11:Three RF transceivers system block diagram

The proposed hybrid FSO/RF system illustrated in Figure 12 combines the two orthogonal architectures: FSO multi-channel (see Section 3.2 above) and RF multi-channel (see Section 3.4 above). Such a hybrid system offers FSO connectivity over three different wavelengths and reliable RF wireless communication over three different frequency bands. This prototype may not clearly state how hybrid RF/FSO networks exploit the strengths of each technology. But, such system increases the overall bandwidth via FSO, as well as high availability in dense fog or in bad atmospheric conditions where the RF transceiver can be used as a backup channel. The architecture provides a highly reliable link since it concurrently operates multi-channel communication links over FSO and RF, over different wavelengths and frequency bands.

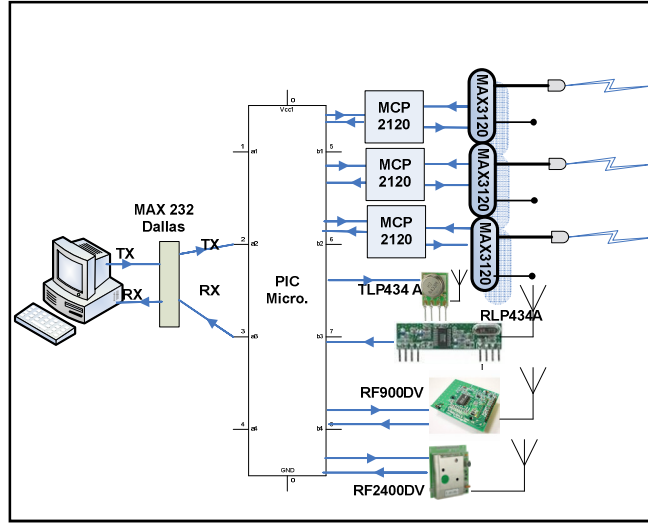


FIGURE 12: A hybrid FSO/RF system architecture

4. ILP FORMULATION OF THE TOPOLOGY CONTROL PROBLEM IN HYBRID RF/FSO NETWORKS

In this section we formulate the topology control problem for hybrid RF/FSO wireless mesh networks, as an integer linear programming (ILP) problem. Informally, we seek to construct a robust topology by minimizing the transmission power, adapting the beam-width, and selecting different channels—all in a manner that allows us to meet specified throughput and end-to-end delay requirements.

We make the following assumptions: First, that the network topology is a mesh with directed links. Second, that each node is both an RF and an FSO transceiver. Finally, we assume that RF transceivers are omnidirectional, while FSO transceivers are directional.

We begin with our **resources** and their **limitations**: A set of mobile nodes V . For each node $i \in V$, we have its location, and the number of RF and FSO transceivers present. Let T_i be the set of transceivers at node i . For each transceiver t in T_i we have: its maximum capacity C_{MAX} , its sensitivity S , its diameter D , its maximum angular beam width, and its maximum transmission power level. We consider a discrete set of possible transmission powers P , and a discrete set of possible (angular) beam openings Φ .

Next, we state our **constraints**: A set of source-destination pairs SD , whose QoS is to be guaranteed. For each $(s, d) \in SD$, we have the maximum delay, $H_{(s,d)}$. Maximum delay, and the minimum acceptable throughput $Th_{(s,d)}$.

In the preprocessing step, we construct a possible network topology $G = (V, E)$. Our objective is to select an optimal construction based on our ILP formulation. The links in a given graph can be enumerated as:

- $l_{i,j,t}^{p,\theta_t,\theta_r}$: For a given transmission power p , transmitter beam-width θ_t , and receiver beam-width θ_r , $l_{i,j,t}^{p,\theta_t,\theta_r} = 1$ if there is a $link(i,j,t)$ available from node i to node j using transceiver t , otherwise $l_{i,j,t}^{p,\theta_t,\theta_r} = 0$. The $link(i,j,t)$ is said to be available if node j is inside the coverage area of node i using transceiver t . This can be verified easily by calculating the transmitter's maximum range, and determining the coverage area based on a scaled sector shape (for FSO channels), or a scaled circle shape area (for RF channels). If t is an FSO transceiver, then for $link(i,j,t)$ to be considered available it must be verified that there is a line of sight between the transmitter and the receiver.
- $BER_{i,j,t}^{p,\theta_t,\theta_r}$: For a given transmission power p , transmitter beam-opening θ_t , and receiver beam-opening θ_r , $BER_{i,j,t}^{p,\theta_t,\theta_r}$ represents the bit error rate on $link(i,j,t)$.
- $B_{i,j,t}^{p,\theta_t,\theta_r}$: For a given transmission power p , transmitter beam-opening θ_t , and receiver beam-opening θ_r , $B_{i,j,t}^{p,\theta_t,\theta_r}$ represents the bandwidth of $link(i,j,t)$:

$$B_{i,j,t}^{p,\theta_t,\theta_r} = \frac{B_{\max_{i,t}}}{\sum l_{i,j,t}^{p,\theta_t,\theta_r}}$$

where $B_{\max_{i,t}}$ is the bandwidth of transceiver t at node i .

We are now ready to formalize the ingredients of the ILP, namely the variables, objective function, and constraints.

We introduce the following Boolean **variables**:

- $l_{i,j,t}^{s,d}$: Boolean variable, $l_{i,j,t}^{s,d} = 1$ if the path of (s,d) connection pair uses $link(i,j,t)$; otherwise $l_{i,j,t}^{s,d} = 0$.
- $g_{i,j,t}^{p,\theta_t,\theta_r}$: Boolean variable selector, $g_{i,j,t}^{p,\theta_t,\theta_r} = 1$ if $l_{i,j,t}^{p,\theta_t,\theta_r}$ is selected to construct the topology; otherwise $g_{i,j,t}^{p,\theta_t,\theta_r} = 0$.
- $x_{i,t}^p$: Boolean power indicator variable. $x_{i,t}^p = 1$ if transceiver t at node i is transmitting using power p ; otherwise $x_{i,t}^p = 0$.

The **objective function** is to minimize transmission power by all nodes in the network.

$$\text{Min} \left(\sum_{i,t,p} p \cdot x_{i,t}^p \right) \quad (1)$$

Finally, we describe the constraints of the ILP. These fall into various classes: Routing constraints, bandwidth constraints, throughput constraints, power constraints, beam opening constraints, alignment constraints, and selector constraints.

Routing Constraints. To ensure that each (s,d) pair is routed:

$$\sum_t \sum_j l_{i,j,t}^{(s,d)} - \sum_t \sum_j l_{j,i,t}^{(s,d)} = \begin{cases} 1 & \text{if } s=i \\ -1 & \text{if } d=i \\ 0 & \text{otherwise} \end{cases} \quad \forall i \in V \text{ and } (s,d) \in SD \quad (2)$$

and that a single route is assigned for a given (s,d) pair:

$$l_{i,j,t}^{(s,d)} \leq \sum_{(p,\theta_i,\theta_r)} l_{i,j,t}^{p,\theta_i,\theta_r} \cdot g_{i,j,t}^{p,\theta_i,\theta_r} \quad \forall i,j \in V, t \in T_i, (s,d) \in SD \quad (3)$$

Delay Constraint. To ensure that the number of hops in the selected route doesn't violate the delay requirement:

$$\sum_{i,j,t} l_{i,j,t}^{(s,d)} \leq H_{s,d} \quad \forall (s,d) \in SD \quad (4)$$

Throughput Constraints. To ensure that throughput requirements are met:

$$\sum_{(s,d)} l_{i,j,t}^{(s,d)} \cdot Th_{(s,d)} \leq \sum_{(p,\theta_i,\theta_r)} B_{i,j,t}^{p,\theta_i,\theta_r} \cdot (1 - BER_{i,j,t}^{p,\theta_i,\theta_r}) \cdot g_{i,j,t}^{p,\theta_i,\theta_r} \quad \forall i,j \in V, t \in T_i \quad (5)$$

Power Constraints. To ensure that power indicator $x_{i,t}^p = 1$ when transceiver t at node i is transmitting using power p .

$$\sum_{(j,\theta_i,\theta_r)} g_{i,j,t}^{p,\theta_i,\theta_r} \leq N x_{i,t}^p \quad \forall p \in P, \forall t \in T_i, \forall i \in V \quad (6)$$

$$\sum_{(j,\theta_i,\theta_r)} g_{i,j,t}^{p,\theta_i,\theta_r} \geq x_{i,t}^p \quad \forall p \in P, \forall t \in T_i, \forall i \in V \quad (7)$$

where N is the number of nodes.

Selector Constraints. To ensure that a single value has been selected for power and beam width at each node:

$$\sum_{(j,p,\theta_i,\theta_r)} g_{i,j,t}^{p,\theta_i,\theta_r} \leq 1 \quad \forall i \in V, \forall t \in T_i$$

where t is an FSO transceiver (8)

$$\sum_{(j,p,\theta_i,\theta_r)} g_{j,i,t}^{p,\theta_i,\theta_r} \leq 1 \quad \forall i \in V, \forall t \in T_i$$

where t is an FSO transceiver (9)

$$\sum_{(p,\theta_i,\theta_r)} g_{i,j,t}^p \leq 1 \quad \forall i, j \in V, \forall t \in T_i$$

where t is an RF transceiver (10)

Beam-opening Constraints. To ensure that transceiver t at node i is using the same beam opening during transmission and reception.

$$\sum_{(p,\theta_i,\theta_r)} g_{i,j,t}^{p,\theta_i,\theta_r} + \sum_{(p,\theta_m,\theta_n)} g_{j,i,t}^{p,\theta_m,\theta_n} \leq 1$$

When $\theta_n \neq \theta_i \quad \forall i, j \in V, t \in T_i, \forall \theta_i \in \Phi$ (11)

Alignment Constraints. To ensure that sender and receiver are in each other's transmission cones:

$$\sum_{(p,\theta_r)} g_{i,j,t}^{p,\theta_r,\theta_r} + \sum_{(p,\theta_m,\theta_n,t)} g_{i,k,t}^{p,\theta_m,\theta_n} \leq 1 \quad \text{When } \theta_n \neq \theta_i,$$

$\forall t \in T_i, \forall \theta_i \in \Phi, \forall i, j, k = i \text{ or } k \text{ in line of sight of } i \text{ and } j.$ (12)

	(Link , Transceiver)	Transmitted Power (mW)	Beam Opening (mrad)	Total Consumed Power
ILP	(1→2,0)	5	-	35
	(1→5,0)	5	-	
	(2→5,2)	5	80,80	
	(3→1,1)	10	80,240	
	(3→2,3)	5	80,80	
	(4→3,0)	5	-	
	(5→4,1)	5	80,160	

TABLE 1: Traffic used to design the topology

The complexity of any ILP problem depends on the number of variables and constraints in that problem. In the proposed formulation, the factors that determine the number of variables and constraints are the number of nodes (N), the number of transceivers (T), the number of source destination pairs (SD),

	Src	Dst	Route	Selected channels
ILP	1	2	1→2	0
	1	5	1→5	0
	2	4	2→5→4	2→1
	2	5	2→5	2
	3	1	3→1	1
	4	3	4→3	0
	4	2	4→3→2	0→3
	5	4	5→4	1

TABLE 2: Transmitted power and beam opening optimal solution using the proposed ILP formulation.

the transmission power granularity (P), and the beam width granularity (θ). The following two equations provide the number of variables and constraints involved in the ILP problem.

$$W = N[(N-1)(SDT + P(1 + T\theta^2)) + TP],$$

where W is the number of variables (13)

$$Z = N(N-1)[TSD + 2T + (T-1)\theta] + N[SD + 2PT + 2(T-1)] + SD,$$

where Z is the number of constraints. (14)

Request	S	D	Throughput (Mbps)	Delay
1	1	2	5	1
2	1	5	5	1
3	2	4	100	2
4	2	5	100	1
5	3	1	250	1
6	4	3	5	1
7	4	2	5	2
8	5	4	100	1

TABLE 3: Routing and channel selection for each requested connection using the proposed ILP.

4.1 ILP Results

We provide an experimental assessment for the ILP formulation of the topology control problem in hybrid RF/FSO mesh networks. In these experiments, we assumed that the capacity of FSO channel is 500 Mbps, the capacity of RF channel is 50 Mbps, the FSO receiver sensitivity is -43dBm, the RF receiver sensitivity is -84dBm, and the maximum beam opening is 240 mrad.

Table 1 presents the matrix of source-destination connection pairs and the desired QoS for each. Tables 2 and 3 (and Figure 13) provide the route found by the ILP solver. The solution uses low power (by adjusting the beam openings and transmission powers of carefully selected channels) and yet meets the joint throughput and end-to-end delay requirements. The solution obtained had a total power consumption of 35mW.

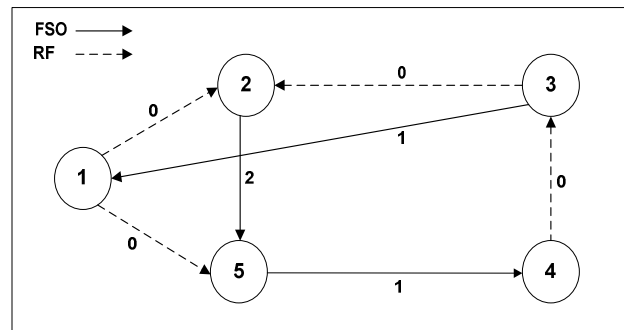


FIGURE 13:Topology generated by ILP solution

5. CONCLUSION AND FUTURE WORK

We presented the strengths and shortcomings of RF and FSO, and the challenges (and opportunities) facing researchers who seek to make hybrid RF/FSO communications a viable technology. We described a prototype implementation for a transceiver node in a hybrid RF/FSO system, illustrating the feasibility and accessibility of the approach. Our hardware was composed of readily available parts: a PIC microcontroller, RF transceivers, IR LEDs and photodiodes, and a few other IC components. The system was extended to provide selectable multichannel multi-frequency capabilities, and was successfully tested and found to be capable of sending and receive English text in 8-bit ASCII encoding. Based on the characteristics suggested by this prototype, we developed a mathematical model of the topology control problem for hybrid RF/FSO networks, as an integer linear program (ILP). We showed that this formulation, when presented to an ILP solver, is capable of adaptive adjustment of transmission power levels and beam openings, in a manner that minimizes total power consumption while satisfying end-to-end quality of service constraints.

Hybrid RF/FSO has the potential to be a “green solution” since it provides a secure, rapidly deployable communication infrastructure with reduced total energy consumption. The hardware prototype and the topology control solution we presented for hybrid RF/FSO can be readily extended and are an important first step in this direction. The hybrid RF/FSO approach could prove especially beneficial in challenging areas such as battlefield environment, intelligent transportation systems, and telesurgery.

6. REFERENCES

- [1] J. Juarez, A. Dwivedi, A. Hammons, S. Jones, V. Weerackody, R. Nichols, “Free Space Optical Communications for Next-Generation Military Networks,” IEEE Communication Magazine, November, 2006.
- [2] R. Raghavan, A. Kam, R. Mannepalli, “Modeling & Simulation to Study the Performance of Hybrid Free Space Optical/RF Military Communication Network”, Military Communications Conference, 2008. MILCOM 2008. IEEE
- [3] M.S Gupta, “What is RF?”, Microwave Magazine, IEEE, 12: 16, August 6, 2002
- [4] I. Djordjevic, H. Batshon, “Generalized Hybrid Subcarrier/Amplitude/Phase/Polarization LDPC-Coded Modulation Based FSO Networking”, ICTON 2010
- [5] O. Bouchet, H. Sizun, C. Boisrobert, “Free-space optics: propagation and communication”, John Wiley & Sons, 2006.
- [6] S. Sarkar, S. Dixit, B. Mukherjee, “Hybrid Wireless-Optical Broadband-Access Network (WOBAN): A Review of Relevant Challenge”, JOURNAL OF LIGHTWAVE TECHNOLOGY, VOL. 25, NO. 11, NOVEMBER 2007
- [7] I. Kim and E. Korevaar, “Availability of Free Space Optics (FSO) and hybrid FSO/RF systems,” Proc. SPIE Vol. 4530 No. 84 2001.

- [8] S. Khan, S. Arsha, "QoS Provisioning Using Hybrid FSO-RF Based Hierarchical Model for Wireless Multimedia Sensor Networks", (IJCSIS) International Journal of Computer Science and Information Security Vol. 4, No. 1 & 2, 2009
- [9] NetworkWorld, "FSO (Free Space Optics)", <http://www.networkworld.com/details/568.html>, July 22, 2002.
- [10] A. Dimakis, L. He, J. Musacchio, H. So, T. Tung, J. Walrand, "Adaptive Quality of Service for a Mobile Ad Hoc Network", MWCN 2003.
- [11] M. Gebhart, P. Schrotter, U. Birnbacher, E. Leitgeb, "Satellite communications, free space optics and wireless LAN combined: worldwide broadband wireless access independent of terrestrial infrastructure", Electrotechnical Conference, 2004. MELECON 2004. Proceedings of the 12th IEEE Mediterranean, November 1, 2004
- [12] H. Moradi, M. Falahpour, H. Refai, "Availability Modeling of FSO/RF Mesh Networks through Turbulence-induced Fading Channels", IEEE INFOCOM 2010, San Diego, CA, March 2010.
- [13] H. Willebrand and B. S. Ghuman, "Free Space Optics," Sams Pubs, 2001, 1st Edition.
- [14] J. Derenick, C. Thorne and J. Spletzer, "Hybrid Free-Space Optics/Radio Frequency (FSO/RF) Networks for Mobile Robots Teams," Springer Netherlands, 2005.
- [15] J. Derenick, C. Thorne, J. Spletzer, "On the deployment of a hybrid free space optic/radio frequency (FSO/RF) mobile ad-hoc network," IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS), 2005.
- [16] K. Hansen, "Wireless RF Design Challenges," IEEE [Radio Frequency Integrated Circuits \(RFIC\) Symposium, 2003.](#)
- [17] A. Kashyap, M. Shayman, "Routing and traffic engineering in hybrid RF/FSO networks, in: IEEE International Conference on Communications," (ICC), 2005.
- [18] CeTIC, "Wireless infrared and UWB impulse radio systems", <http://newyork.ing.uniroma1.it/neuwb2/DOCS/invited-03.pdf>, July 6, 2005.
- [19] MRV Communications Inc. "Data Security in Free Space Optics"
- [20] fSONA Optical Wireless, "FSO Security".
- [21] R. Nichols, "Protocol Adaptation in Hybrid RF/Optical Wireless Networks," IEEE MILCOM, 2005.
- [22] LXE Inc., "RF Wireless Basics: An Intro to the Latest in Wireless Data Collection Networks, Products, Standards and Solutions", Jan 2006.
- [23] H. Hakki, J. Sluss, Jr.1, H. Refai1, and M. Atiquzzaman, "Transporting RF Signals over Free-Space Optical Links", SPIE 2005.
- [24] J. Akella, Chang. Liu, D. Partyka, M. Yuksel, S. Kalyanaraman, P. Dutta. "Building blocks for mobile free-space-optical networks," Second IFIP International Conference on Wireless and Optical Communications Networks (WOCN), 2005.
- [25] R. Guerin, A. Orda, D. Williams, "QoS Routing Mechanisms and OSPF Extensions," IETF Internet Draft, 6 November 1996.
- [26] S. Ramamurthy and B. Mukherjee, "Survivable WDM mesh networks: Part I - Protection," IEEE INFOCOM, pp. 744-751, 1999.
- [27] A. Kashyap, A. Rawat, M. Shayman, "Integrated Backup Topology Control and Routing of Obscured Traffic in Hybrid RF/FSO Networks," IEEE Globecom, 2006.
- [28] Ramanathan and R. Hain, "Topology Control of Multihop Wireless Networks Using Transmit Power Adjustment," IEEE INFOCOM, 2000.
- [29] M. CaYang. Topology, "Control in Ad Hoc Wireless Networks Using Cooperative Communication," IEEE Transactions on Mobile Computing, Vol. 5 No. 6, June 2006.
- [30] X. Jia, D. Li and D.-Z. Du, "QoS Topology Control in Ad Hoc Wireless Networks," IEEE Infocom, 2004.
- [31] E., J. Llorca, S. Milner, C. Davis, "Topology Reconfiguration with Successive Approximations," IEEE Military Communications Conference (MILCOM), 2007
- [32] MAXIMSemiconductors. "+5V-Powered, Multichannel RS-232 Drivers/Receivers ", 2006.
- [33] Microchip. "MCP2120 infrared encoder/decoder", datasheet, 2001
- [34] LinCMOS Timers, "TLC555", datasheet, 1997.
- [35] MAXIMSemiconductors. "MAX3120 Low-power, 3V, 120 mA IrDA Infrared Transceiver", 1998.
- [36] Laipac Technology, Inc. <http://www.laipac.com/>
- [37] Laipac Technology, Inc. "TLP434A & RLP434A: RF ASK Hybrid Modules for Radio Control (New Version)", datasheet.
- [38] J. Buysse, M. Leenheer, C. Develder, B. Dhoedt, and P. Demeester, "Cost-effective Burst-Over-Circuit-Switching in a hybrid optical network," International Conference on Networking and Services (ICNS), April 2009.
- [39] C. Gauger, P. Kuhn, E. V. Breusegem, M. Pickavet, and P. D. "Hybrid optical network architectures: bringing packets and circuits together," Journal of High Speed Networks, Vol. 44 pp.36-42, August 2006.

- [40] S.D. Milner, C.C. Davis, "Hybrid free space optical/RF networks for tactical operations," Military Communications Conference (MILCOM), 2004.
- [41] V. Vishnevskiy, O. Semenova, "Queueing System with Alternating Service Rates for Free Space Optics-Radio Hybrid Channel.", MACOM 2010: 79-90
- [42] W. K. Seah, Z. A. Eu, and H.-P. Tan, "Wireless sensor networks powered by ambient energy harvesting (WSN-HEAP) – survey and challenges," in Proc. VITAE '09, May 2009.
- [43] H. Willebrand, B. S. Ghuman, "Free space optics: enabling optical connectivity in today's networks", Sams Publishing, 2002.
- [44] R. Ramirez-Iniguez, S. M. Idrus, Z. Sun, "Optical wireless communications: IR for wireless connectivity", CRC Press, 2008, 143:167
- [45] A. Umar, "Mobile Computing and Wireless Communications", nge solutions. inc, 2004
- [46] L. Dordova, O. Wilfert, "Free space optical link range determination on the basis of meteorological visibility", Radioelektronika, April 22, 2009.
- [47] I. Djordjevic, W. Ryan, B. Vasic, "Coding for Optical Channels", Springer, 2010, p. 84:121.
- [48] J. Hecht, "Understanding Fiber Optics", National Instruments, Sep 6, 2006.
- [49] B. Wang, G. Yan, C. Yan, "Fiber length and chromatic dispersion measurement technology using a novel optical frequency domain reflectometry", OSA-IEEE-COS, 2010.
- [50] A.M. Fraine, R. Egorov, D.S. Simon, O.V. Minaeva, A.V. Sergienko, "High resolution measurement of polarization mode dispersion with quantum interferometry", Photonics Society Summer Topical Meeting Series, IEEE, 2011.
- [51] I.B. Djordjevic, L.L. Minkov, X. Lei, T. Wang, "Suppression of Fiber Nonlinearities and PMD in Coded-Modulation Schemes With Coherent Detection by Using Turbo Equalization", Optical Communications and Networking, IEEE/OSA Journal, November 2009.
- [52] D.L. Polla, M.B. Wolfson, "RF MEMS integration present & future trends", Radio-Frequency Integration Technology, RFIT, 2009.

INSTRUCTIONS TO CONTRIBUTORS

The International Journal of Computer Networks (IJCN) is an archival, bimonthly journal committed to the timely publications of peer-reviewed and original papers that advance the state-of-the-art and practical applications of computer networks. It provides a publication vehicle for complete coverage of all topics of interest to network professionals and brings to its readers the latest and most important findings in computer networks.

To build its International reputation, we are disseminating the publication information through Google Books, Google Scholar, Directory of Open Access Journals (DOAJ), Open J Gate, ScientificCommons, Docstoc and many more. Our International Editors are working on establishing ISI listing and a good impact factor for IJCN.

The initial efforts helped to shape the editorial policy and to sharpen the focus of the journal. Starting with volume 5, 2013, IJCN aims to appear with more focused issues. Besides normal publications, IJCN intend to organized special issues on more focused topics. Each special issue will have a designated editor (editors) – either member of the editorial board or another recognized specialist in the respective field.

We are open to contributions, proposals for any topic as well as for editors and reviewers. We understand that it is through the effort of volunteers that CSC Journals continues to grow and flourish.

IJCN LIST OF TOPICS

The realm of International Journal of Computer Networks (IJCN) extends, but not limited, to the following:

- Algorithms, Systems and Applications
- ATM Networks
- Cellular Networks
- Congestion and Flow Control
- Delay Tolerant Networks
- Information Theory
- Metropolitan Area Networks
- Mobile Computing
- Multicast and Broadcast Networks
- Network Architectures and Protocols
- Network Modeling and Performance Analysis
- Network Security and Privacy
- Optical Networks
- Personal Area Networks
- Telecommunication Networks
- Ubiquitous Computing
- Wide Area Networks
- Wireless Mesh Networks
- Ad-hoc Wireless Networks
- Body Sensor Networks
- Cognitive Radio Networks
- Cooperative Networks
- Fault Tolerant Networks
- Local Area Networks
- MIMO Networks
- Mobile Satellite Networks
- Multimedia Networks
- Network Coding
- Network Operation and Management
- Network Services and Applications
- Peer-to-Peer Networks
- Switching and Routing
- Trust Worth Computing
- Web-based Services
- Wireless Local Area Networks
- Wireless Sensor Networks

CALL FOR PAPERS

Volume: 5 - Issue: 2

i. Submission Deadline : January 31, 2013 **ii. Author Notification:** February 15, 2013

iii. Issue Publication: April 2013

CONTACT INFORMATION

Computer Science Journals Sdn Bhd

B-5-8 Plaza Mont Kiara, Mont Kiara
50480, Kuala Lumpur, MALAYSIA

Phone: 006 03 6207 1607
006 03 2782 6991

Fax: 006 03 6207 1697

Email: cscpress@cscjournals.org

CSC PUBLISHERS © 2012
COMPUTER SCIENCE JOURNALS SDN BHD
M-3-19, PLAZA DAMAS
SRI HARTAMAS
50480, KUALA LUMPUR
MALAYSIA

PHONE: 006 03 6207 1607
006 03 2782 6991

FAX: 006 03 6207 1697
EMAIL: cscpress@cscjournals.org