

The availability and reliability of wireless multi-hop networks with stochastic link failures

Geir Egeland, *Member, IEEE*, Paal E. Engelstad *Senior Member, IEEE*

Abstract—The network reliability and availability in wireless multi-hop networks can be inadequate due to radio induced interference. It is therefore common to introduce redundant nodes. This paper provides a method to forecast how the introduction of redundant nodes increases the reliability and availability of such networks. For simplicity, it is assumed that link failures are stochastic and independent, and the network can be modelled as a random graph. First, the network reliability and availability of a static network with a planned topology is explored. This analysis is relevant to mesh networks for public access, but also provides insight into the reliability and availability behaviour of other categories of wireless multi-hop networks. Then, by extending the analysis to also consider random geometric graphs, networks with nodes that are randomly distributed in a metric space are also investigated. Unlike many other random graph analyses, our approach allows for advanced link models where the link failure probability is continuously decreasing with an increasing distance between the two nodes of the link. In addition to analysing the steady-state availability, the transient reliability behaviour of wireless multi-hop networks is also found. These results are supported by simulations.

Index Terms—Random graphs, multi-hop mesh, IEEE 802.11s, reliability, availability, transient behaviour.

I. INTRODUCTION

A wireless multi-hop network (typically with an *ad hoc* topology) is a network composed of a group of nodes interconnected via wireless links. The nodes in the network implements a routing protocol which enables them to communicate with each other over multiple link hops. The nodes in these networks are often self-configured and self-organised. Examples of such networks include wireless mesh networks [1], mobile ad hoc networks (MANETs) [2] and wireless sensor networks (WSNs) [3].

This paper analyses two important classes of static wireless multi-hop networks. They will be referred to as:

- Planned mesh networks
- Random mesh networks

Typically, in *planned mesh networks* the location of each node is carefully planned. These networks usually appear as a consequence of the high costs associated with interconnecting nodes in a network with wired links. For example, ad hoc technology can in a cost-efficient manner, extend the reach of a wired backbone through a wireless backhaul mesh network. Currently, there is a large number of commercial deployments

of such solutions for public access in urban areas, where the nodes are often located on roof-tops. These networks are referred to as *roof-top networks* or *backhaul mesh networks*.

Unlike a planned mesh, a *random mesh network* consists of nodes that are randomly distributed (i.e. geographically deployed in a random manner). They typically appear in situations where it is required that the network is set up in an ad hoc fashion. Such networks can be used to provide a self-formed temporary communication infrastructure for disaster recovery operations, for communication at conferences or conventions and in tactical military operations.

Common for these two types of wireless networks, is that radio-links are vulnerable to failures caused by radio induced interference, which in most cases has a negative impact on the network reliability and availability. This represents a considerable barrier for wide deployment of planned and random mesh networks.

In a planned mesh network it is therefore common to improve the network reliability and availability by introducing redundant nodes. Typically, the possible location of a redundant node in a planned mesh network is known, e.g. the operator of a roof-top network is in contact with a property owner that is willing to offer a specific roof-top as a site for the redundant node.

The site-acquisition and site-operation cost associated with each redundant node, is usually high. Therefore, the network planning phase should include a cost-benefit analysis, where the additional network reliability and availability of adding this redundant node is weighted against the additional costs. While the cost of adding a redundant node usually is easy to determine, it is more difficult to forecast the additional reliability and availability gained by adding the node. The main reason is that little have been published about network reliability and availability for mesh networks. By using the random graph analysis presented in this paper, however, it might be possible to estimate what will be the improved reliability and availability by adding one or more redundant nodes.

Indeed, the main contribution of this paper is a method to forecast how the introduction of redundant nodes increases the network reliability and availability of *planned mesh networks*. By using the random graph analysis presented in this paper, the network reliability and availability of a planned mesh network can be found as a function of the number of redundant nodes. By applying the same method, it is possible to forecast how the introduction of redundant nodes will increase the network reliability for any given topology. (The possible topologies are given, since the location of the redundant node normally

G. Egeland is with the Department of Electrical and Computer Engineering, University of Stavanger, 4036 Stavanger, Norway e-mail: (geir.egeland@gmail.com).

P. E. Engelstad is with UniK/UiO, Simula and Telenor, 1331 Fornebu, Norway e-mail: (paal.engelstad@telenor.com)

Manuscript received August 15, 2008; revised January 31, 2009.

is not arbitrary, but restricted to one or a limited number of alternatives.)

A similar method is also presented for *random mesh networks*, and it is based on a random geometric graph analysis. This method is useful in scenarios where the concept of redundant nodes makes sense. In a mobile ad hoc network, for example, the concept of redundant nodes translates into deploying more nodes in an area than strictly necessary in order to increase the network reliability and availability. Deploying these networks, it is useful to know the required node density in order to obtain the desired availability and reliability. Likewise, in a sensor dust network, where the nodes are static, it is also valuable to find the necessary sensor density that satisfies the required probability that the network is connected. Unlike many other random graph analyses, our approach allows for advanced link models where the link failure probability is continuously decreasing with an increasing distance between the two nodes of the link.

The rest of the paper is organised as follows: Section II presents the network models for planned and random mesh networks, corresponding to a random graph model and a random geometric graph model, respectively. It also presents link failure probability models used in the analyses. Reliability and availability metrics for a mesh network are presented in Section III using a random graph approach. The section also evaluates the effect of redundant MPs in a backhaul mesh network for two example topologies. The random graph approach is then extended to a random geometric graph analysis in Section IV, where networks of independently and uniformly distributed MPs are considered. While the analyses up until this point mainly focus on network availability and steady-state conditions, the time-dependency of the network reliability is addressed in Section V. This section also describes how to obtain the mean time to failure and the mean time between failures for an arbitrary mesh topology. Finally, related work is described in Section VI, before conclusions are drawn in Section VII.

II. NETWORK MODEL

A. Network Terminology

A network terminology is needed in order to describe the architecture of planned and random mesh networks, and in what way redundant nodes can be introduced in their architecture. This paper reuses the terminology of wireless mesh networks, more specifically of the IEEE 802.11s specification [1] of mesh networks, which is based on the IEEE 802.11 standard [4]. In this terminology a node in a mesh network is referred to as a *mesh point* (MP). Furthermore, an MP is referred to as a *mesh access point* (MAP) if it includes the functionality of an 802.11 access point, allowing regular 802.11 stations (STAs) access to the mesh infrastructure. Finally, an MP is referred to as a *mesh portal* (MPP) if it has additional functionality for connecting the mesh network to other network infrastructures.

Figure 1 illustrates the introduction of a redundant node in a backhaul mesh network. Under regular operation, each STA is connected to a MAP, and the STA's traffic is forwarded along

the shortest path between the MAP and the MPP. However, if a link in the shortest path becomes unavailable, the routing protocol ensures that a new path through the redundant node is constructed.

Given that the routing protocol is based on shortest-path, there is no load-sharing in the network. This means that the introduction of the redundant node is not intended to increase the overall throughput of the network. Its main purpose is to improve the network reliability and availability. Thus, this paper is solely concerned with the connectivity measures of a wireless mesh network, while other network performance metrics, such as throughput and delay, are not included in the analysis.

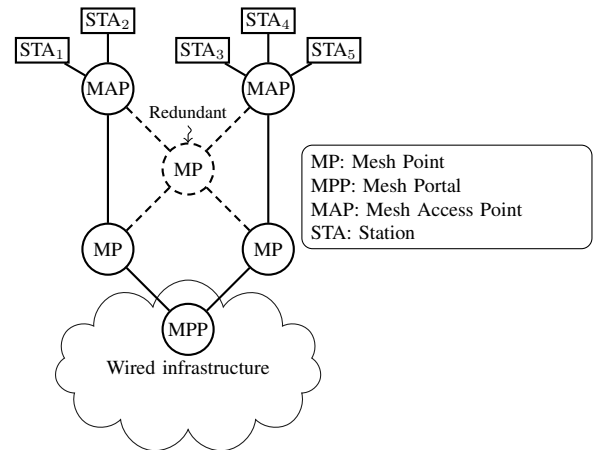


Fig. 1. A backhaul mesh network with a redundant node.

B. Graph description

A wireless mesh network can be described as a undirected¹ graph $G = (V, E)$, where the nodes in the network serve as the vertices $v_j \in V(G)$. Any two distinct nodes v_j and v_i create an edge $e_{i,j} \in E(G)$ if there is a link between them. For simplicity, we let ϵ denote the size of the graph, i.e. the number of edges, $\epsilon = |E(G)|$.

A minimal set of edges in the graph whose removal disconnects the graph is an *edge cutset*. The minimum cardinality of an edge cutset is the *edge connectivity* or *cohesion* $\beta(G)$. A (minimal) set of nodes that has the same property is a *node cutset*, and the minimum cardinality of this is the *node connectivity* $\chi(G)$. An edge *tieset* is a set of edges which forms a connection between a set of vertices. The minimum cardinality of a tieset is the shortest path between a set of vertices.

In order to provide an adequate measure of network reliability, the use of probabilistic reliability metrics and a *probabilistic graph* is necessary. This is an undirectional graph where each node has an associated probability of being in an operational state, and likewise for each edge.

In wireless multi-hop networks, link failures may be caused by radio fading, signal attenuation, radio interference, background noise and other inherent characteristics of the wireless

¹Most routing protocols have a symmetric neighbour discovery mechanism or other mechanisms that ensure that the links are undirected.

medium [5] [6]. As a result, the link failure frequency is in general much higher than the node failure frequency. For that reason, it is natural to model the nodes $v_j \in V(G)$ in the topology as invulnerable to failure, and concentrate on the link failures in the analysis. Because of these characteristics, this paper argues that a mesh network can be analysed as a random graph.

C. Random graph model for a planned mesh network

A mesh network of n nodes can be modelled as a random graph $G(V, E, p)$, where p is the link failure probability, i.e. the probability that a link between two nodes is *not* included in the graph, and n is the order of the graph, $|V(G)|$ (Generally, a random graph is denoted as $G(V, E, \hat{p})$, where $\hat{p}$ is the link existence probability (i.e. $\hat{p} = 1 - p$). In this paper, however, we consequently use the link failure probability p instead of the link existence probability $\hat{p}$.)

An underlying assumption in our analysis is that the frequent link failures seen in planned mesh networks can be modelled by the link failure probability p and that p is determined independently for each link. The latter means that a link $\epsilon_{s,d}$ connecting two nodes v_s and v_d may fail independently of $\epsilon_{i,j} \in E(G) \setminus \{v_s, v_d\}$.

D. Random geometric graph model for a random mesh network

A *random geometric graph* $G(V, E, r)$ is a geometric graph in which the $n = |V(G)|$ nodes are independently and uniformly distributed in a metric space [7]. In other words, it is a random graph for which a link between two nodes v_s and v_d exists if, and only if, their Euclidean distance is such that $\|v_s - v_d\| \leq r_0$.

However, this paper considers a generalised random geometric graph, $G(V, E, r, p)$, where the existence of a link is determined not only by the geographic distance between the uniformly distributed nodes, but also by the link failure probability p , which results from the salient characteristics of wireless communication.

Furthermore, we define:

- The *visibility graph*, as a graph where there is an edge between two vertices, if the two nodes are within radio transmission range r_0 of each other.
- The *connectivity graph*, as a graph where there is an edge between two vertices, if the two nodes are within radio transmission range of each other and if the communication link between the two nodes has not failed.

E. Link failure probability models

A basic link model, which is used in our analyses, is to assume that the link failure probability is equal for all links and independent of the Euclidean distance, i.e. a *distance-independent link model*. However, the fact that a link $\epsilon_{s,d}$ is defined, naturally means that the two nodes v_s and v_d are within each other's radio transmission range.

Although most of our analyses use a distance-independent link failure model, a *distance-dependent* model is also investigated. Especially for a random geometric graph it makes sense

to assume that the link failure probability of a link $\epsilon_{i,j}$ depends on the Euclidean distance between the two nodes v_i and v_j . In the following we present a more elaborate, but well-known link failure model that is distance-dependent and that is being used in research on ad hoc and mesh networks (e.g. cf. [8] [9]).

In radio communications, the average value of the received signal power on a link $\epsilon_{s,d}$ decreases with an increasing geometric distance between the two nodes v_s and v_d . This phenomenon is often referred to as pathloss, and the behaviour of the signal power is referred to as the *area mean power*, P_a . The area mean power is often modelled by a power law $P_a \propto r^{-\eta}$. Here, the pathloss exponent, η , depends on the terrain and the environment. It may vary between 2 in free space up to 6 in for example urban areas [10]. In ad hoc and mesh network research it is common to use a pathloss model and assume that a signal transmitted from v_s is received correctly at v_d if the received signal power exceeds a minimum required threshold. The result is a circular coverage area of radius r_0 around the transmitting node (Figure 2(a)).

In reality, however, the received power levels might vary significantly in space and time around the area mean power value of the pathloss model. Thus, to make the pathloss model more realistic, one might assume a log-normal shadowing radio propagation model [11] where the logarithmic value of the mean power at different locations is normally distributed with standard deviation σ around the logarithmic value of the area mean power (Figure 2(b)). The link *failure* probability is then given as a function of the distance r by [9]:

$$p_d(r) = \frac{1}{2} \left[1 + \operatorname{erf} \left(\frac{10 \log(r/r_0)}{\sqrt{2} \log(10) \psi} \right) \right], \psi \triangleq \frac{\sigma}{\eta} \quad (1)$$

where r_0 is given by the same threshold as for the area mean power model. Empirically ψ may vary in the range 0 to 6 [9]. In the analyses presented later, ψ will be varied and set to either of the values $\{0.25, 0.5, 1.0, 1.5, 2.5\}$.

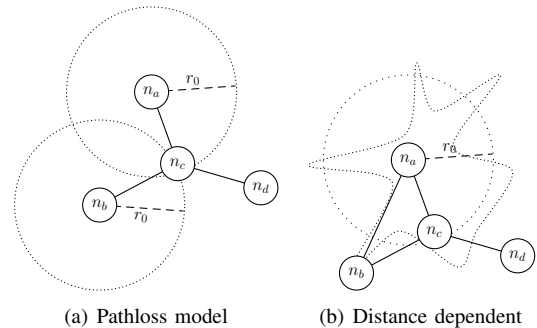


Fig. 2. Different connectivity for distance-dependent and distance-independent link models

The link failure probability $p_d(r)$ in Eq. (1) is implicitly also an expression for the characteristic of the two radio receivers on a link, and their ability to correctly receive data. As the bitrate on a link increase, the probability of correctly receiving data approaches a step function [12], similar to Eq. (1) with ψ approaching zero.

III. RANDOM GRAPH ANALYSIS OF PLANNED MESH NETWORKS

A. k -terminal reliability

We start the analysis of network reliability and availability by applying the k -terminal reliability, which is defined as the probability that a path exists and connects k nodes in a network. The k -terminal availability for the k nodes $\{n_1, \dots, n_k\} \subseteq V(G)$ can be found as:

$$R_c^{n_1, \dots, n_k}(G, p) = \sum_{i=w_{n_1, \dots, n_k}(G)}^{\epsilon} T_i^{n_1, \dots, n_k}(G) p^{\epsilon-i} (1-p)^i \quad (2)$$

$$= 1 - \sum_{i=\beta(G)}^{\epsilon} C_i^{n_1, \dots, n_k}(G) p^i (1-p)^{\epsilon-i} \quad (3)$$

where $\epsilon = |E(G)|$ is the size of the graph. In Eq. (2), $T_i^{n_1, \dots, n_k}(G)$ denotes the tieset with cardinality i , i.e. the number of subgraphs connecting the nodes $n_1, \dots, n_k$ with i edges. Furthermore, $w_{n_1, \dots, n_k}(G)$ is the size of the minimum tieset connecting the nodes $n_1, \dots, n_k$. In Eq. (3), $C_i^{n_1, \dots, n_k}(G)$ denotes the number of edge cutsets of cardinality i and $\beta(G)$ is the cohesion.

B. Network availability

The network reliability is defined as the probability that a network operates successfully, i.e. is connected, for a given period of time under environmental conditions [13]. More specifically, if the network operates successfully at the time t_0 , the network reliability yields the probability that *there were no failures* in the interval 0 to t .

The analysis of network reliability assumes for simplicity that there are no link repairs in the network. This is not exactly true for mesh networks, since a neighbour discovery mechanism or some other functionality will ensure that a failed link is restored when the radio conditions improve. The metric used to describe repairable networks is *availability*. The availability is defined as the probability that at any instant of time t , the network is up and available, i.e. not disconnected [13]. The network availability is a reliability measure of great significance, because it indicates the portion of the time the network is operational.

This paper, however, focus on the availability at steady-state, found as $t \rightarrow \infty$, i.e. when the transient effects from the initial conditions are no longer affecting the network. The steady-state availability is equal to $\text{MTBF}/(\text{MTBF} + \text{MTBR})$, where MTBF is the *mean-time-between failure* and MTBR is the *mean time between repair* [13]. It can also be expressed in terms of the *mean time to failure* (MTTF) and the *mean time to repair* (MTTR) [14].

When a link in the network is repairable, we can model it as a two-state Markov diagram, where one state represents the link being up, while the other represents the link being down. The link failures of an operational link $\epsilon_{i,j}$ are modelled according to an exponential distribution with a failure rate

parameter λ . In addition, we assume that the repair rate of a failed link is exponentially distributed with a rate parameter μ . (Figure 3)

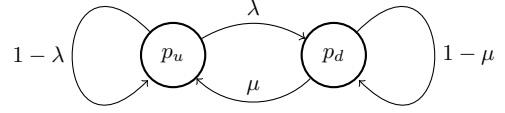


Fig. 3. A simple two state Markov model of the reliability of a link

At steady-state, the state probabilities for the link being up (p_u) or down (p_d) are easily found as:

$$p_u = \frac{\mu}{\mu + \lambda}, \quad p_d = \frac{\lambda}{\mu + \lambda} \quad (4)$$

C. Reliability and availability of a planned backhaul mesh network

Consider the mesh backhaul network illustrated in Figure 4. This network can be described as a graph G , that includes $k-1$ different distribution nodes $d_i \in D$, where $D = \{d_1, d_2, \dots, d_{k-1}\}$, and one root node r . According to the previous terminology of IEEE 802.11s, a distribution node corresponds to a MAP in an IEEE 802.11s network, while the root node corresponds to an MPP.

Under regular network operation, the transit traffic in the network is directed along the shortest path between the root node r and each distribution node, d_i in D . If any distribution node is disconnected from the root node, the network has failed, as it is not operating as intended. Thus, the network is fully operational only if there is an operational path between the root node and each of the distribution nodes. This is true if, and only if, the root node r and the $k-1$ distribution nodes are all connected. Thus, the reliability of the network may be analysed using the k -terminal reliability.

When considering reliability, the network availability is of highest interest. Using the expressions above, the k -terminal availability can now be found as:

$$A_c^{r, d_1, \dots, d_{k-1}}(G, p_d) = 1 - \sum_{i=\beta}^{\epsilon} C_i^{r, d_1, \dots, d_{k-1}} p_d^i (1-p_d)^{\epsilon-i} \quad (5)$$

A IEEE 802.11s network can be configured to allow the STAs to access the MAPs at one frequency band (e.g. using 802.11b or 802.11g) and use another frequency band for the communication between the MPs in the backhaul mesh network. Because the extra equipment cost of such a configuration often is minimal compared costs associated with site-acquisition, it is anticipated that many commercial mesh networks will implement a MAP at each MP in the network. For such a configuration, the *all-terminal* availability of the

network, i.e. $A_c(G, p_d) = A_c^{r, d_1, \dots, d_{k-1}}(G, p_d)$ for $k = |V(G)|$ in Eq. (5) has to be considered.

Alternatively, if there is only one root node r (i.e. MPP) and one distribution node d (i.e. MAP), the *two-terminal* availability is the metric of interest. The two-terminal availability is found as $A_c^{r, d}(G, p_d)$ in Eq. (5).

D. Evaluation by example

This section considers the topology in Figure 4, where two MAPs are distributed and connected to a MPP over a mesh network. The availability of the service provided by the MAPs is in general the quantity of most interest, and can be calculated using the three-terminal availability $A_c^{r, d_5, d_6}(G, p_d)$:

$$\begin{aligned} A_c^{r, d_5, d_6}(G, p_d) &= \sum_{i=w_{r, d_5, d_6}(G)}^{\epsilon} T_i^{r, d_5, d_6}(G) p_d^{\epsilon-i} (1-p_d)^i \\ &= 1 - \sum_{i=\beta_{r, d_5, d_6}(G)}^{\epsilon} C_i^{r, d_5, d_6}(G) p_d^i (1-p_d)^{\epsilon-i} \quad (6) \end{aligned}$$

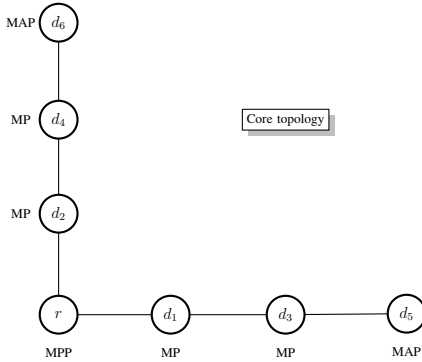


Fig. 4. The initial topology without redundant MPs

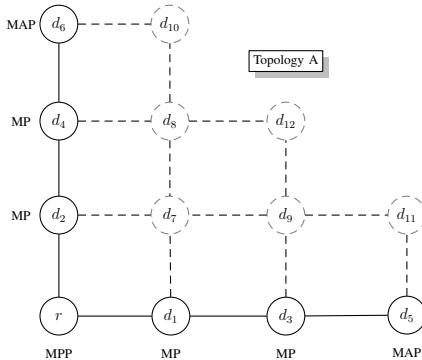


Fig. 5. The initial topology with redundant MPs (Topology A)

Figure 6 and Figure 7 show the calculated availability and the effect when an extra MP is added to the network. The redundant MPs are added in a particular order, consecutively $d_7, \dots, d_{12}$ (Figure 5). The results in Figure 6 indicates that the effect of adding redundant MPs is greatest when $p_d < 0.6$.

Figure 7 gives a different perspective of the same results, as it illustrates how the availability increases as redundant MPs are added *when the value of p_d is given*. Four different values of $p_d \in \{0.01, 0.1, 0.3, 0.5\}$ are considered. A link failure rate $p_d = 0.5$ is quite pessimistic, and will in most cases render the network as useless for any practical purposes. Although we concentrate our analysis on lower values of p_d , we have also considered setting $p_d = 0.5$, since this is easy to analyse analytically, as is shown below.

In telecommunications network, it is a requirement that a service must be available 99.999% of the time. Although it is very difficult to satisfy this requirement in multi-hop mesh networks, it is still interesting to analyse mesh networks with a relatively high availability, e.g. when the network operator has taken measures to ensure that the underlying link failure probability is low. In Figure 8 an enlarged version of Figure 6 is shown, for which the link failure probability is $p_d < 0.01$. It is observed that the curve for the redundant node d_{11} represents a leap in the availability, compared to a topology that only includes the nodes $d_7, \dots, d_{10}$. The reason is that when d_{11} is added, the entire network becomes *2-connected*. Thus, when d_{12} is then added, it results in only a negligible improvement of the availability. The improvement is most negligible at low values of p_d , since then with a very high probability one failover path (i.e. redundant path) will be sufficient to ensure connectedness of the network when a link fails. However, as p_d increases, the probability of multiple simultaneous link failures increases, leading to an increase in the probability of simultaneous failure of both the core path and the failover path. Thus, the availability increase resulting from the addition of d_{12} becomes less negligible as p_d increases.

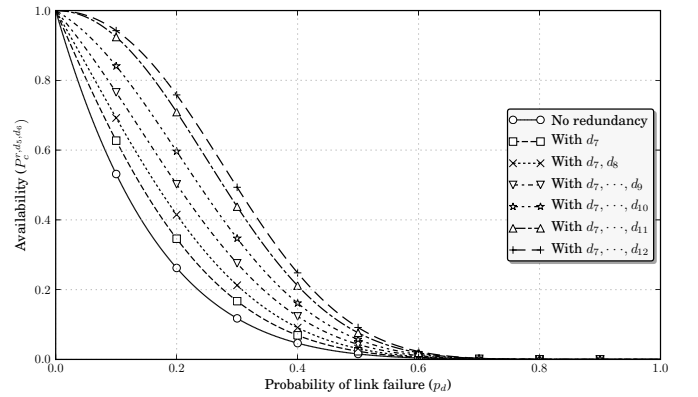


Fig. 6. Network availability when adding redundant MPs to Topology A

E. Assessing the importance of each redundant node: Their order of priority

When deploying several redundant nodes and there is a number of locations they can be situated, it is beneficial to find the order of priority between the possible locations. If we assume that the costs are equal for all locations, the order of priority will be determined by the increased reliability and

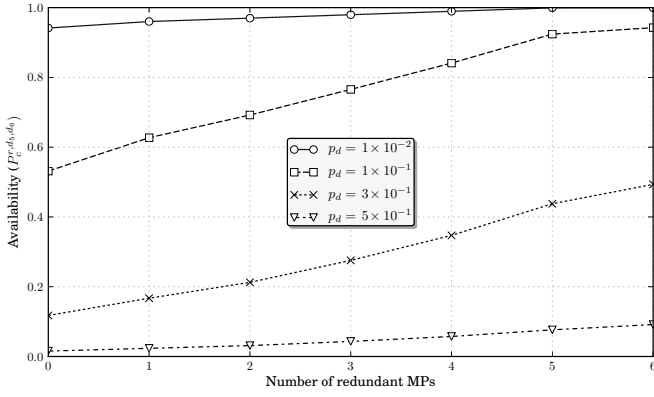


Fig. 7. The effect of redundant MPs on Topology A

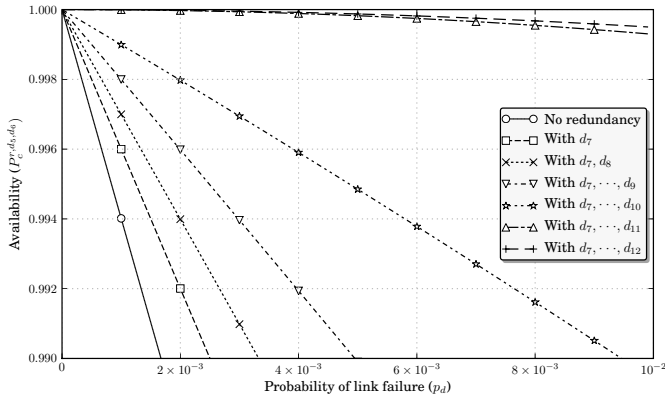


Fig. 8. Network availability at low link failure probabilities (an enlarged version of Figure 6)

availability alone. In this section, an order of priority is found based on the network availability.

Algorithm 1 describes a method for finding an order of priority. It starts with a connected graph, G , of nodes in all possible node locations (e.g. corresponding to $r, d_1, \dots, d_{12}$ in Topology A) and a connected subgraph of G , G_1 , of core nodes (e.g. $r, d_1, \dots, d_6$ in Topology A). For each non-core node in G that is a one-hop neighbour node of any of the core nodes, the availability is found for G_1 with a single one-hop neighbour added. The one-hop neighbour node that yields the highest availability is selected as the highest priority redundant node, and the location of this node is the highest priority location of a redundant node. The selected node is then added to the connected graph of core nodes, and the procedure is repeated to find the second highest priority node (or location). This is repeated until all redundant nodes are part of the graph of core nodes. (Step 8 and 14 of the algorithm can be omitted, as the value a_{2nd} is only used in the analysis of the algorithm below.)

The result of applying Algorithm 1 to Topology A is the ordered list $(\{d_7\}, \{d_8, d_9\}, \{d_{10}, d_{11}\}, \{d_{12}\})$. All actions within the WHILE-loop starting in step 3 comprise a *stage* of the algorithm. Due to the network symmetry, the two nodes $\{d_8, d_9\}$ are derived in the second stage of the algorithm, while

Algorithm 1 $\vec{L}(G, C)$

Require: A connected graph G , a set of connected core nodes $C \in V(G)$.

```

1:  $\vec{L} \leftarrow ()$ 
2:  $G_1 \leftarrow$  the connected subgraph of  $G$  containing only the nodes in  $C$ 
3: while  $G_1 \neq G$  do
4:   for  $c \in C$  do
5:      $N \leftarrow \{\text{one hop neighbours of } c \in V(G)\} \setminus C$ 
6:   end for
7:    $a_{max} \leftarrow 0$ 
8:    $a_{2nd} \leftarrow 0$ 
9:    $C_1 \leftarrow \emptyset$ 
10:  for  $n \in N$  do
11:     $G_1 \leftarrow$  the connected subgraph of  $G$  containing only the nodes in  $C \cup \{n\}$ 
12:     $a \leftarrow A_c^{r, d_5, d_6}(G_1, p_d)$ 
13:    if  $a > a_{max}$  then
14:       $a_{2nd} \leftarrow a_{max}$ 
15:       $a_{max} \leftarrow a$ 
16:       $C_1 \leftarrow \{n\}$ 
17:    end if
18:    if  $a = a_{max}$  then
19:       $C_1 \leftarrow C_1 \cup \{n\}$ 
20:    end if
21:  end for
22:   $\vec{L} \leftarrow \text{Append}(\vec{L}, C_1)$ 
23:   $C \leftarrow C \cup C_1$ 
24:   $G_1 \leftarrow$  the connected subgraph of  $G$  containing only the nodes in  $C$ 
25: end while
26: return  $\vec{L}$ 

```

the nodes $\{d_{10}, d_{11}\}$ are derived in the third stage.

The derived list $(\{d_7\}, \{d_8, d_9\}, \{d_{10}, d_{11}\}, \{d_{12}\})$ indicates that the redundant node d_7 should be added first. Because the topology is symmetric, it does not matter whether d_8 or d_9 are added second or third. Similarly, it does not matter whether d_{10} or d_{11} are added as number four or five. In any case, d_{12} should be the last node to be added.

TABLE I
ANALYSIS OF HOW REDUNDANT NODES IN TOPOLOGY A AFFECT THE AVAILABILITY A_c^{r, d_5, d_6}

	d_7	d_8	d_9	d_{10}	d_{11}	d_{12}
$A_c(G_1, p_d = 0.1)$	0.627	0.692	0.765	0.841	0.924	0.942
$A_c(G_1, p_d = 0.5)$	0.023	0.031	0.043	0.057	0.076	0.091
$ E(G_1) $	8	10	12	14	16	18
$\sum_{i=w(G_1)}^{ E(G_1) } T_i(G_1)$	6	32	177	940	5008	24000
ΔH	7	7.7	7.99	8.5	8.9	9.75
$\Delta A_c(G_1, p_d = 0.1)$	0.095	0.065	0.005	0.062	0.067	NA
$\Delta A_c(G_1, p_d = 0.5)$	0.007	0.007	0.002	0.008	0.009	NA
$\Delta_{G-G_2} A_c(p_d = 0.1)$	0.083	0.182	0.182	0.085	0.085	0.018

In the following, we will examine how the addition of each redundant node is improving the availability. First we change step 19 of Algorithm 1 to $C_1 \leftarrow \{n\}$ (which is the same action as performed in step 16). This change ensures that the algorithm produces a list where only one

node is added at a time, because only one node is selected in each stage of the algorithm. This will produce the list $(\{d_7\}, \{d_8\}, \{d_9\}, \{d_{10}\}, \{d_{11}\}, \{d_{12}\})$. The reason that d_8 here came out with a higher priority than d_9 , is because it was selected arbitrarily later than d_9 in step 10. For the same reason, d_{10} came out coincidentally with a higher priority than d_{11} .

In order to analyse how the addition of each redundant node improves the availability, we calculate the value $a_{max} = A_c^{r,d_5,d_6}(G_1, p_d)$ at the end of the stage (i.e. after step 24 in the algorithm). The results for a_{max} is shown in the first two rows of Table I for the link failure probability, $p_d \in \{0.1, 0.5\}$. For $p_d = 0.1$ it is observed that improvement in availability for adding node d_{12} is insignificant, as pointed out earlier.

At a considerably high link failure probability of $p_d = 0.5$ this is no longer the case. This setting of p_d is interesting to study, since the availability can be written as:

$$\left(\frac{1}{2}\right)^{|E(G_1)|} \times \sum_{i=w_{r,d_5,d_6}(G_1)}^{|E(G_1)|} T_i^{r,d_5,d_6}(G_1) \quad (7)$$

The first term in Eq. (7) reflects the availability of a path in the network, while the second term reflects the number of possible paths in the network. The first term is easy to find, since $|E(G_1)|$ of the core backbone in Topology A is equal to 6. Furthermore, each redundant node of the topology increases $|E(G_1)|$ by two, as shown in the third row of the table. Exploring the effect of the second term, we do the following: Instead of finding the availability in step 12 of Algorithm 1, we find the corresponding tiesets, i.e. $a \leftarrow \sum_{i=w_{r,d_5,d_6}(G_1)}^{|E(G_1)|} T_i^{r,d_5,d_6}(G_1)$ of G_1 . The results are listed in the fourth row of Table I. While the first term $(1/2)^{|E(G_1)|}$ decreases with a factor 1/4 each time a redundant node is added, it is observed that the second term $\sum_{i=w_{r,d_5,d_6}(G_1)}^{|E(G_1)|} T_i^{r,d_5,d_6}(G_1)$ increases with a factor around 5.

Having a clearer understanding of the case in which $p_d = 0.5$, it is easier to analyse the behaviour for when p_d is low. For sufficiently low values of p_d we can approximate $p_d^i(1 - p_d)^{\epsilon-i} \rightarrow p_d^i$ in the expression for the availability in Eq. (6). The implication of this is that when summing up all possible redundant paths of the tieset, the long paths contribute less to the overall sum compared to the shorter paths. The average hop length, $\Delta \bar{H}$, of the paths generated when adding a redundant node, is shown in the fifth row of Table I. Since the redundant nodes $d_7, \dots, d_{11}$ have locations close to the original core backbone and close to the center of the network, the redundant paths they generate do not increase the average hop length of the paths considerably. Node d_{12} on the contrary is located in the outskirts of the network, and the redundant paths it generates are considerably longer. The result is that the addition of d_{12} contributes less to the total availability, compared to the addition of the other nodes.

With the purpose of acquiring an in-depth understanding of the selection process of Algorithm 1, we also do the following: In each stage of the algorithm we calculate the value $\Delta A_c = a_{max} - a_{2nd}$ at the end of the stage. The value ΔA_c is an indicator of how much the availability improves by adding the most favourable redundant node, compared to adding the

second best redundant node in each stage. (A node that due to network symmetry yields exactly the same availability, is not counted as a second best node.) The calculated values for $\Delta A_c = a_{max} - a_{2nd}$ are shown in the sixth and seventh rows of Table I for $p_d = 0.1$ and $p_d = 0.5$, respectively. These results show that when node d_7 and d_8 have been added, it does not make any noticeable difference for the tree-terminal reliability whether d_9 or d_{10} are added first, even though they are not a symmetric pair of nodes. In light of the discussion above and by examining the topology, it is easy to accept that the reason is that both nodes will generate an approximately equal amount of redundant paths, and that the average length of these two sets of paths are quite similar.

F. Assessing the importance of each redundant node: Their importance in a given topology

When the number of redundant nodes and their location is decided, the next step in an availability analysis is to examine the weighted importance of each node in the topology. For example, the weighted importance of each node might be a guideline for real deployment on how to allocate the investments for each site, in terms of network equipment or physical measures to secure the site against threats of operation.

Assessing the weighted importance of each node in a *given graph* G , can be done as follows: A single node is temporarily removed from the graph G , resulting in a subgraph G_2 . By comparing the availability of the subgraph $A_c(G_2, p_d)$ with the availability of the original graph, $A_c(G, p_d)$, the impact on the availability caused by the node removal can be determined. The redundant node for which removal affects the availability the most (i.e. the node with the highest $\Delta_{G-G_2} A_c(p_d) = A_c(G, p_d) - A_c(G_2, p_d)$) is considered to be of highest importance. Note that when one redundant node d_i is removed, all other nodes $d_1, \dots, d_{i-1}, d_{i+1}, \dots, d_{12}$ are assumed to be part of the subgraph G_2 .

The final row of Table I shows the effect on the availability caused by removing a redundant node. The table shows the results for $p_d = 0.3$.

The weighted importance of a node might be higher than that of another node at a low p_d , while lower at a high p_d . This is illustrated in Figure 9. We can observe that for small link failure probabilities ($p_d < 0.1$), the failure for node d_7 has little impact, since the connectivity for r, d_5, d_6 is provided by longer paths. However, the failure of node d_9 has a considerably impact, since this will also affect paths through node d_{11} and d_{12} . As the link failure probability increases ($p_d > 0.1$), node d_7 will have a greater impact on the availability, since the longer paths connecting r, d_5, d_6 then will have a higher probability of failing.

Finally, it is observed that the importance of the core nodes $d_1, \dots, d_4$ can also be easily found and analysed in a similar manner. However, this aspect is considered out of scope here.

G. Alternative topologies

Although the set of node locations in a planned mesh is often predetermined in real deployments, it is of interest to

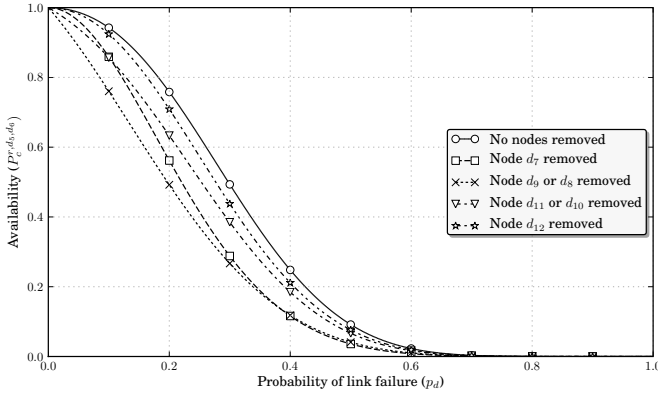


Fig. 9. The impact availability when a redundant node in Topology A fails

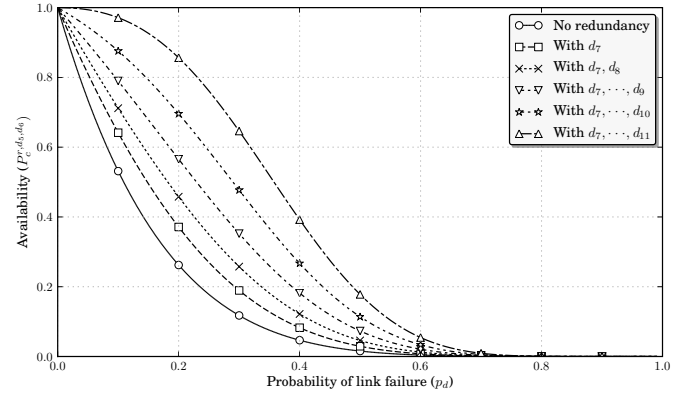


Fig. 11. Network availability when adding redundant MPs to Topology B

examine how the underlying topology affects the reliability and availability performance.

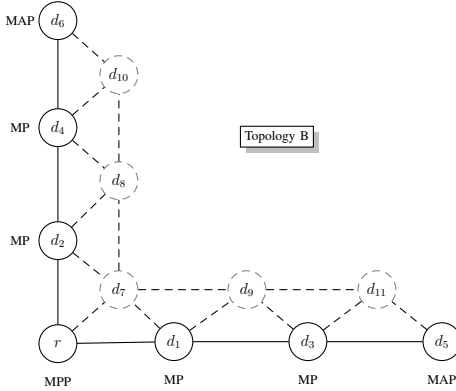


Fig. 10. An alternative deployment of redundant MPs (topology B)

Figure 10 (Topology B) illustrates an alternative configuration that deploys the redundant MPs geographically closer to the backhaul mesh. Assuming that the redundant MPs have a transmission range equivalent to the range in Figure 4, positioning the MPs closer to the backhaul mesh nodes will increase the number of links from the redundant MPs to the graph of connected nodes. Here, $|E(G_1)|$ increases with 3 (in contrast to 2 for Topology A) for each redundant node that is added to the network. The result is a higher number of available redundant paths in the network.

Results for Topology B are presented in Figure 11, Figure 12 and Table II. When comparing these with the corresponding results for Topology A, a higher level of availability is observed, which is expected in the light of the discussions above in Section III-E.

Furthermore, Figure 11 shows a leap in the availability when node d_{11} is added, which is most notable at relatively low values of p_d . The main reason is that d_{11} transforms the network into a two-connected topology. The same effect was observed for Topology A.

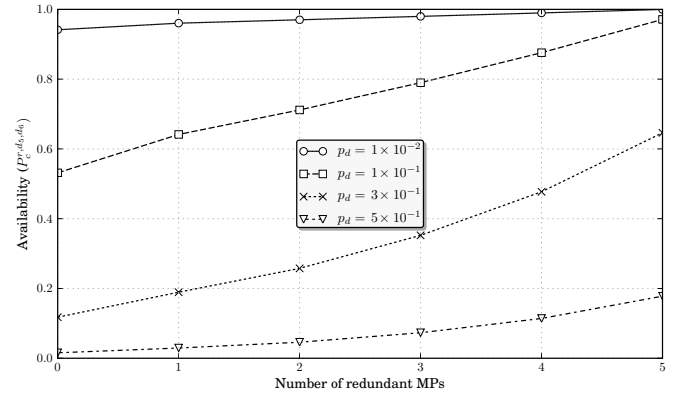


Fig. 12. The effect of redundant MPs on Topology B

IV. RANDOM GEOMETRIC GRAPH ANALYSIS OF RANDOM MESH NETWORKS

A. Ad hoc networks where some or all nodes are randomly distributed

The random graph analysis in the previous section considered network topologies with a pre-defined and planned *visibility graph* where the links in the visibility graph were included in the *connectivity graph* by a link failure probability p_d . The location of each node was planned in a way to ensure that the visibility graph was always connected.

This section, on the contrary, considers wireless multi-hop networks where the visibility graph is a *random geometric graph*. As before, the links in the visibility graph are included in the connectivity graph by the link failure probability p_d . However, as the visibility graph is a random geometric graph, the reliability and availability of the networks considered in this section, depend on both the topology of the nodes (i.e. whether the visibility graph is connected or not at a given point in time) and the failure probability p_d of the links in the visibility graph.

The connectivity of the visibility graph can be estimated separately. In [7], for example, the author shows that the probability for k -connectivity of a homogenous ad hoc network is given by:

TABLE II
ANALYSIS OF HOW REDUNDANT NODES IN TOPOLOGY B AFFECT THE
AVAILABILITY A_c^{r,d_5,d_6}

	d_7	d_8	d_9	d_{10}	d_{11}
$A_c(G_1, p_d = 0.1)$	0.641	0.711	0.789	0.875	0.971
$A_c(G_1, p_d = 0.5)$	0.029	0.045	0.073	0.114	0.177
$ E(G_1) $	9	12	15	18	21
$\sum_{i=w(G_1)} E(G_1) T_i(G_1)$	15	188	2395	29889	373106
$\Delta \bar{H}$	7	7.7	8.3	8.9	9.42
$\Delta_{G-G_2} A_c(p_d = 0.1)$	0.196	0.110	0.110	0.095	0.095

$$P(G \text{ is } k\text{-connected}) \cong P(d_{\min} \geq k) = \left(1 - \sum_{i=0}^{k-1} \frac{(\rho \pi r_0^2)^i}{i!} \cdot e^{-\rho \pi r_0^2}\right)^n \quad (8)$$

for high probabilities $P(d_{\min} \geq k)$. In Eq. (8), $d_{\min}$ is the minimum node degree, r_0 is the transmission range and $\rho = |V(G)|/A$ is the node density, where A is the size of the area. The equation can be used to determine the threshold value for the transmission range r_0 required to achieve an almost surely k -connected network.

In order to determine the network availability of a wireless multi-hop network with randomly distributed static nodes, one first needs to ensure a sufficiently high probability of the visibility graph being connected. In a network where the randomly distributed nodes are mobile, on the other hand, the network availability is determined both by the probability that the visibility graph is connected at a given instant of time, and by the link probabilities of the links in the visibility graph.

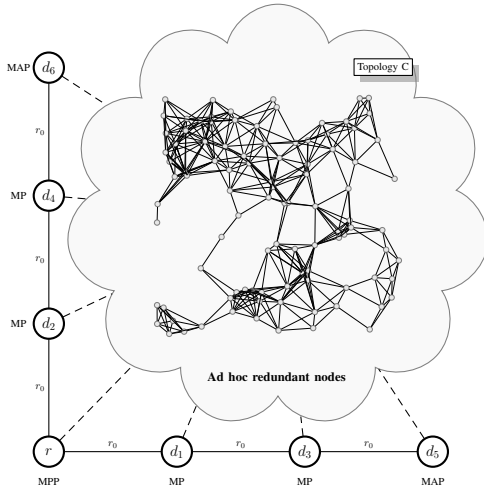


Fig. 13. Deployment of redundant MPs in an ad hoc fashion (topology C)

B. Example where the redundant MPs are randomly distributed

Consider, as an example, Topology C which is shown in Figure 13. This network architecture is a variation of the topologies analysed in the previous section. Here, and unlike Topology A and B, the redundant MPs are independently and uniformly distributed in the area, while the locations of the

core infrastructure nodes ($r, d_1, \dots, d_6$) are still planned (i.e. the visibility graph of the core infrastructure nodes is always connected). This network architecture might be relevant in scenarios where the self-configuration capability of a random ad hoc networks, enables the construction of failover paths for a wireless backhaul network.

In the figure, r_0 is the transmission range of the nodes, for both MPs and the ad hoc nodes which act as redundancies. This is not very realistic, but chosen for simplicity. The redundant nodes are added uniformly in an area of the size $A = 9r_0^2$.

Figure 14 illustrates the three-terminal network availability A_c^{r,d_5,d_6} as a function of a varying number of redundant MPs. Each curve in the figure consists of the mean values obtained from 100 randomly generated topologies of the redundant nodes. As expected, the availability improves gradually as the number of redundant nodes increases.

Figure 15 shows the effect of adding redundant nodes for a link failure probability of $p_d \in \{0.01, 0.1, 0.3, 0.5\}$. A comparison of Figure 15 with the results for Topology A and Topology B in Figure 6 and Figure 7, illustrates the advantage of carefully planning the locations of the redundant nodes, especially when the link failure probability $p_d > 0.1$.

In many networking scenarios, however, planning the location of each node is not always an option. Figure 15 establishes that a random distribution of the redundant nodes is a valid alternative at the cost of a higher number of redundant nodes.

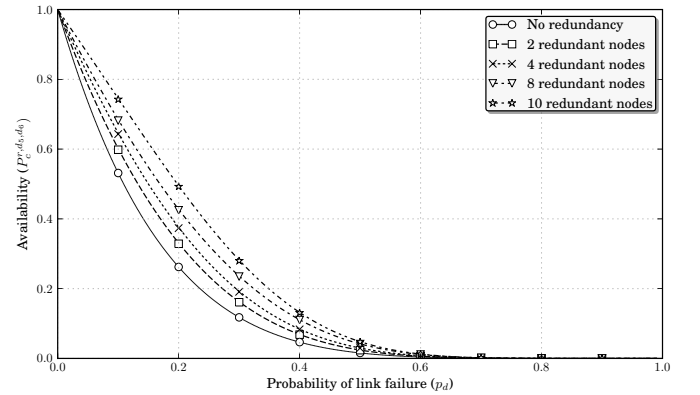


Fig. 14. Network availability when adding redundant MPs for Topology C

C. Example where all core MPs are randomly distributed

As another example, consider a network architecture where all MPs in the network (except the MPP and the MAPs) are randomly distributed. This is referred to as *Topology D* and is illustrated in Figure 16. Topology D is a variation of Topology C, except that now it is only the locations of nodes $\{r, d_5, d_6\}$ that are planned.

This scenario is relevant for mobile ad hoc networks targeted at providing connectivity between some strategic nodes (i.e. the MAPs) and an internet gateway (i.e. the MPP). In a military scenario, for example, ensuring connectivity between some strategic units (MAPs) and the commando head quarter

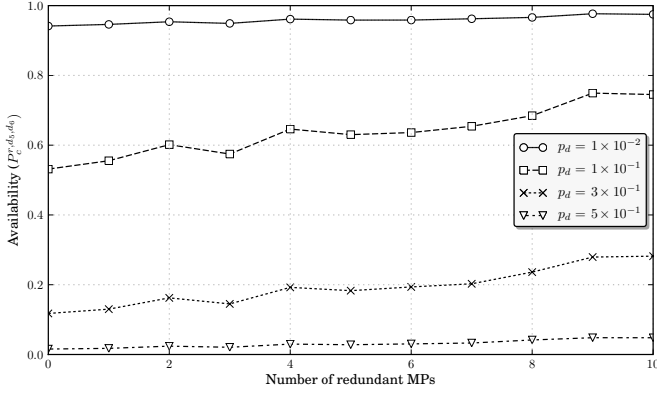


Fig. 15. The effect of redundant MPs on Topology C

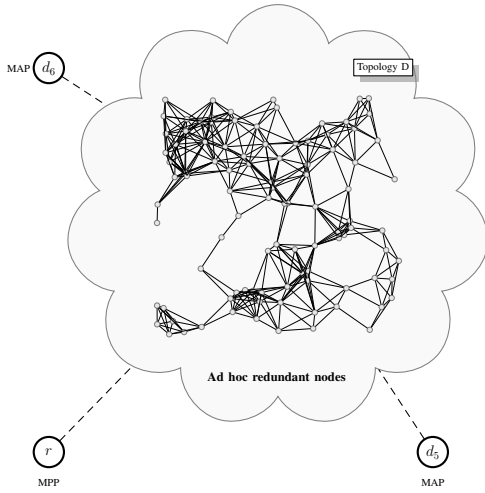


Fig. 16. Deployment of redundant MPs in an ad hoc fashion (Topology D)

(via the MPP) might be the primary objective of the network deployment.

Note that the visibility graph of the three nodes $\{r, d_5, d_6\}$ is not connected. Thus, it does not make sense to calculate the three-terminal availability before there is a certain density of the randomly distributed nodes. Otherwise, the visibility graph will never be connected. This is demonstrated in Figure 17, where the probability of Topology D being 1-connected is shown. As Figure 17 illustrates, the analytical and the simulated results differ. This is caused by using Euclidian distance metric for $P(k\text{-connected})$, where the redundant nodes at the edge of the area of the topology have fewer neighbour nodes than the ones in the middle. This was also the case in [7], where using a torodial distance metric gave similar results for analytical and simulated $P(1\text{-connected})$.

Assuming that the randomly distributed nodes in Topology D (Figure 16) are mobile, the network availability is determined both by the topology of the visibility graph and the link failure probabilities of the links in the visibility graph. Figure 18 shows the effect of adding redundant nodes for a link failure probability of $p_d \in \{0.01, 0.1, 0.3, 0.5\}$, as well as for

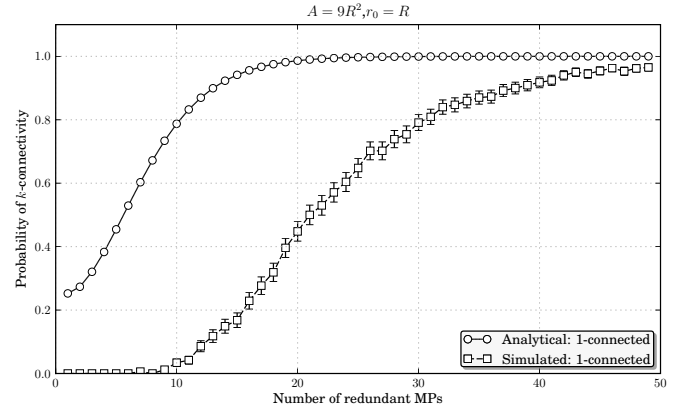


Fig. 17. The probability of Topology D being 1-connected

$p_d = 0$. Each curve in the figure is found as the mean of 10000 randomly generated topologies. Each of these topologies might represent a snapshot of the locations of the mobile nodes.

The curve for $p_d = 0$ (Figure 18) illustrates the probability that the visibility graph of the mobile nodes is connected at a given instant of time. Thus, the discrepancy between this curve and the four curves for $p_d \in \{0.01, 0.1, 0.3, 0.5\}$ illustrates how the link failure probability contributes to a reduction in the network availability.

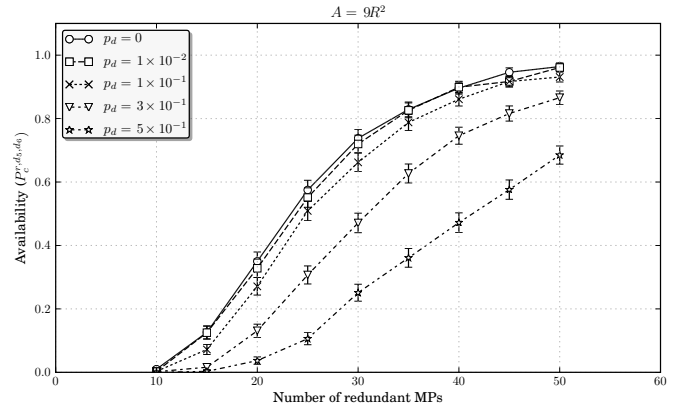


Fig. 18. The availability of Topology D with randomly distributed mobile nodes

D. Comparison of link probability models

In the analysis of topology A,B,C and D, we have assumed a fixed link failure probability. This can be justified for topology A and B, since the distance between the nodes is fixed. In topology C and D, however, the distance between the nodes varies.

In order to understand the effect of the distance-dependent link model, we simulated $M = 10000$ random generated variations of topology C and D. For every variation of the topologies, we generated $N = 100$ graph model representation (G_i), where the creation of the edges of each graph, was based on the link failure probability in Eq. (1). Each graph was then

checked for the connectivity between nodes $\{r, d_5, d_6\}$, and an estimate E^{r, d_5, d_6} for the availability was generated by:

$$E^{r, d_5, d_6} = \frac{1}{N} \times \left[\begin{array}{l} \text{Number of graphs where} \\ r, d_5, d_6 \text{ are connected} \end{array} \right] \quad (9)$$

Now, $\bar{P}_c^{r, d_5, d_6}$ can be estimated as:

$$\bar{P}_c^{r, d_5, d_6} = \frac{1}{M} \sum_{i=0}^{M-1} E_i^{r, d_5, d_6} \quad (10)$$

During the simulation, an estimate for the link failure probability $\bar{p}_d$ was calculated in a similar manner:

$$\hat{p}_d^{r, d_5, d_6} = \frac{1}{N} \sum_{\forall \epsilon_{i,j} \in G} \frac{1}{2} \left[1 + \operatorname{erf} \left(\frac{10 \log(r_{i,j}/r_0)}{\sqrt{2} \log(10) \psi} \right) \right] \quad (11)$$

Now, $\bar{p}_d^{r, d_5, d_6}$ can be estimated as:

$$\bar{p}_d^{r, d_5, d_6} = \frac{1}{M} \sum_{i=0}^{M-1} \hat{p}_{d,i}^{r, d_5, d_6} \quad (12)$$

The results from Eq. (10) for topology C and D is plotted in Figure 19 and Figure 20 respectively. The values of r_0 and ψ were varied to emulate different link failure probabilities. Using a heuristic approach, we observed the $\bar{p}_d$ generated by the simulation model and varied the value for r_0 and ψ until a link failure probability approximately to $\{0.01, 0.3, 0.5\}$ was achieved. These results are compared with simulating topology C and D using a fixed p_d equal to $\bar{p}_d$. The simulation parameters are shown in Table III.

TABLE III
SIMULATION PARAMETERS

	$\bar{p}_d$	r_0	ψ
Topology C	0.01	3.5	1
	0.27	8	1
	0.45	9	1
Topology D	0.05	9.5	0.25
	0.33	7.5	1.5
	0.52	4.5	2.5

Some inaccuracies are observed when comparing with a well-known distance-dependent link probability model. The difference in availability for $\bar{p}_d = 0.03$ in Figure 19 can be explained by the fact that when simulating the distance-dependent link model, the link failure probability for the links connecting the core nodes was set to a fixed value of $p = 0.01$, which was the mean link failure probability we wanted to achieve. However, the resulting $\bar{p}_d$ from the simulation was 0.03, because according to Eq. (1), the redundant nodes can only connect to the core nodes with a link failure probability of $p_d = 0.01$ in approximately 30% of the simulation area. Hence, the link failure probability for the distance-independent link model was calculated using $p_d = 0.03$ for all links, resulting in an availability that is approximately 10-20% lower.

For the other two curves in Figure 19, the simulation of the distance-dependent link model resulted in a mean link failure probability of $\bar{p}_d = 0.29$ and $\bar{p}_d = 0.48$. In these two cases

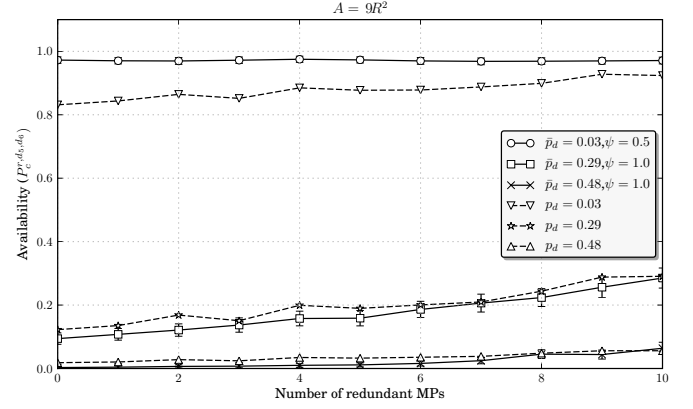


Fig. 19. The effect of redundant MPs on Topology C, with a distance-dependent link probability model

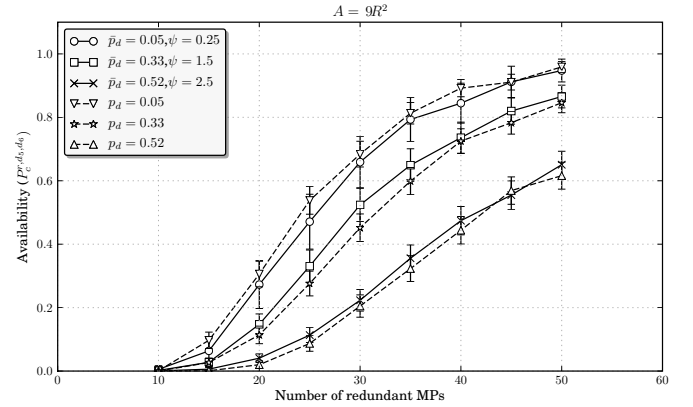


Fig. 20. The availability of Topology D with randomly distributed mobile nodes with distance-dependent link probability.

the link failure probability for the links connecting the core nodes, was set to $p = 0.3$ and $p = 0.5$ respectively. Thus, the resulting availabilities are approximately equal.

These comparisons suggest that a distance-independent link model might yield useful information about the network behaviour. The results in Figure 20 shows that for any given pair of r_0, ψ , the estimate for a mean link failure probability in Eq. (12) will result in the same availability by the use of a simple distance-independent link model.

In the analysis presented in the next section, we will revert back to using the distance-independent model.

V. TRANSIENT ANALYSES

A. Transient reliability analysis

While the network availability is a useful reliability measure, it does not provide every aspect of the network reliability, since it is a steady-state measure. For example, it gives the share of time the network is up, but it does not provide insight into the mean time to failure (MTTF) given that the network is connected at time $t = 0$, or into the mean time between failures (MTBF). To estimate such time-dependent measures, it is necessary to study the transient behaviour of the k -terminal reliability.

Analyses of transient behaviour normally become very complex when considering link failures *with* repair. Thus, when studying transient reliability, it is not uncommon to simplify the analysis by first considering link failures *without* repair. This will be done in the next section. Then, in the subsequent section, link failures *with* repairs will be considered.

B. The transient k -terminal reliability without link repair

It is assumed that networks start out with all links between two nodes within radio range of each other as fully operational, with the result that the *connectivity graph* is equal to the *visibility graph*. Furthermore, it is assumed that link failures are exponentially distributed with parameter λ and that when a link $\epsilon_{i,j}$ fails, it can not be repaired.

As illustrated in Appendix A, it can be shown that the expected time to when k -nodes will be disconnected, is given by:

$$E(t_d^{v_i, \dots, v_k}) = \frac{1}{\lambda} \times \sum_{n=1}^{\epsilon} \left[\binom{\epsilon}{n} \frac{(-1)^{n+1}}{n} - \sum_{i=\beta}^{\epsilon-1} C_i^{v_i, \dots, v_k} B(i+1, \epsilon-i) \right] \quad (13)$$

where $B(x, y)$ is the complete *Beta function*.

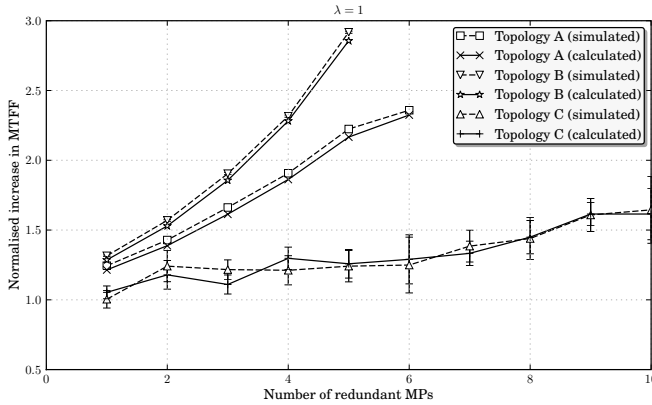


Fig. 21. Normalised (with respect to no redundancy) MTFF as the number of redundant MPs is increased for topologies A, B and C

Using Eq. (13), the mean time to first failure (MTFF) for topologies A, B and C can be found. These results are shown in Figure 21, illustrating that the time to the first network failure, i.e. node r, d_5, d_6 is disconnected, increases with an increasing number of redundant MPs. The results from a Monte Carlo simulation for the MTFF are also included in Figure 21 with a 95% confidence interval, which only reason for inclusion is to verify Eq. (13).

C. The transient k -terminal reliability with link repair

The assumption of no link repairs (i.e. that links that fail have infinite repair time) does not match well with the common operation of real networks. It is more realistic to

model the repair time according to an exponential distribution, which is done in this section.

In order to find the MTBF, the state of the network can be modelled using a Markov model with a link failure rate of λ and a repair rate μ . This is illustrated in Figure 23. The state S of the network is either connected or disconnected, $S_{i,j}, i \in [0, 1, \dots, \epsilon], j \in \{c, d\}$, and the transition probabilities are given by Eqs. (14)-(17).

$$W_i^{v_i, \dots, v_k} = \lambda(\epsilon - i) \cdot \frac{C_{i+1}^{v_i, \dots, v_k}}{\binom{\epsilon}{i+1}} \quad (14)$$

$$Q_i^{v_i, \dots, v_k} = \lambda(\epsilon - i) \cdot \frac{\binom{\epsilon}{i+1} - C_{i+1}^{v_i, \dots, v_k}}{\binom{\epsilon}{i+1}} \quad (15)$$

$$\overline{W}_i^{v_i, \dots, v_k} = \mu \cdot i \cdot \frac{C_{i-1}^{v_i, \dots, v_k}}{\binom{\epsilon}{i-1}} \quad (16)$$

$$\overline{Q}_i^{v_i, \dots, v_k} = \mu \cdot i \cdot \frac{\binom{\epsilon}{i-1} - C_{i-1}^{v_i, \dots, v_k}}{\binom{\epsilon}{i-1}} \quad (17)$$

Using Korolyuks theorem, the MTBF can be expressed as $1/\Lambda$, where:

$$\Lambda^{v_i, \dots, v_k} = \sum_{i=\beta}^{\alpha} p_{i,d} \cdot \overline{Q}_i^{v_i, \dots, v_k} \quad (18)$$

where β is the minimum cutset and α is the cutset where any link removal disconnects k nodes.

In Figure 22 the calculated MTBF for topologies A, B and C is plotted. The figure also shows simulation results for the same topologies. The simulation results were obtained by using a discrete event simulator with topologies A, B and C as input. The links of a topology fail and are repaired at events decided by $\lambda = 1$ and $\mu = 30$. The simulation results in Figure 22 are shown with 95% confidence intervals. As the figure shows, the calculations for the MTBF match well with the results from the simulation model. According to the results, Topology B shows superior performance as the redundant nodes are added to the topology.

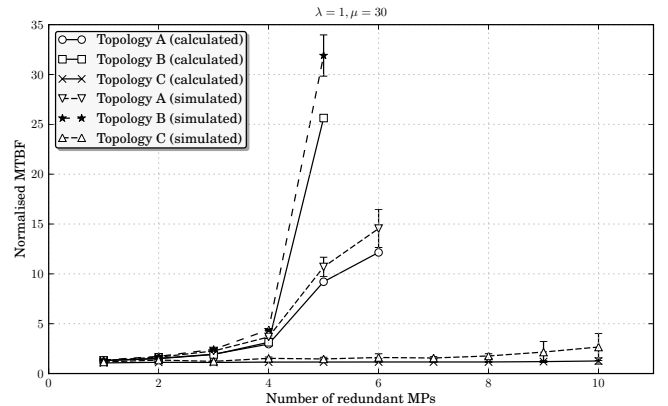


Fig. 22. Normalised (with respect to no redundancy) MTBF as the number of redundant MPs is increased for the topologies A, B and C

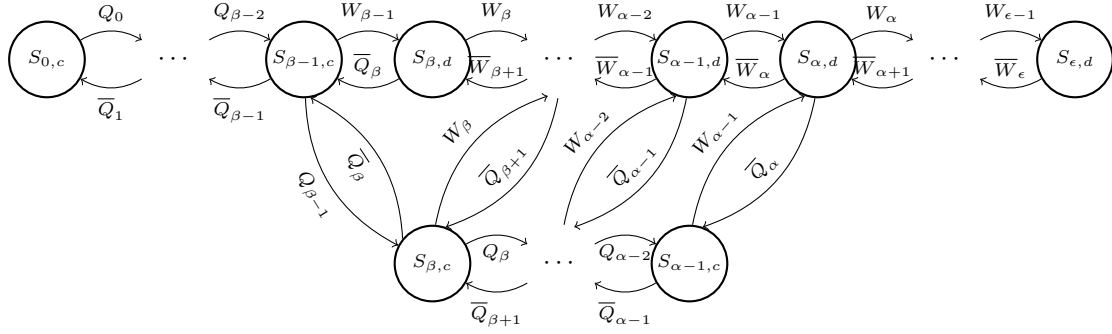


Fig. 23. A Markov model of the reliability of an undirectional graph where links fail with the rate λ and are repaired with the rate μ . The minimum cutset is given by β , and α is the cutset where any link removal disconnects the network.

VI. RELATED WORK

There have been several studies on k -terminal reliability for wired networks [15] [16] [17]. In [17], for example, link failures in fixed ring or double ring network structures are analysed. The work focuses on the design of a physical network topology that achieves a high level of reliability using unreliable network elements. It is shown that for independent link failures, network design should be optimised with respect to reliability under high stress, as the reliability under low stress is less sensitive to the network topology.

The results from most work on network reliability of fixed networks are not generally applicable to wireless networks. The main reason is that in fixed networks the probability of a link failure is so low compared to the probability of a node failure, that in the analysis it is common to consider the link as invulnerable to failure. In wireless networks, on the other hand, link failures may occur more frequently.

There are several other reasons why fixed network analyses are not applicable to wireless networks. For example, in fixed networks a link monitoring mechanism, or a *Link Integrity Test* operation, is often used to identify if a link has failed, while there is no equivalent to this in wireless networks. Even with a similar link monitoring mechanism available, a wireless link may also be affected by the hidden node problem [18]

There are, however, only a few studies on the reliability of wireless networks. The early work in [19] analyses radio broadcast networks, showing that computing the two-terminal problem for these networks is computationally difficult. The work in [20] deals with the problems of computing a measure for the reliability of distributed sensor networks and for the expected and the maximum message delay between data sources. The two-terminal reliability of ad hoc networks is computed in [21]. Their work focuses on the reliability of nodes and on the effects of node mobility, while the effects of link reliability in static topologies - which is investigated in this paper - are not considered.

The idea of modelling an ad hoc network as a random graph was already discussed in [22]. The work in [7] investigates the connectivity of a homogeneous ad hoc network where the nodes are independently and uniformly randomly distributed in a metric space. The network is modelled as a geometric random graph, and the work derives an analytical expression

that allows the determination of the required transmission range that creates, for a given node density, an almost surely k -connected network.

A problem with random geometric graphs is the fact that the analyses take place in regions with no obstructions, which means that a node can communicate with all other nodes within transmission range. This is at odds with the underlying constraints in many applications of distributed wireless networks, where there can generally be a large number of obstructions, limiting communication between nearby nodes. Also, connectivity is not governed by the radio transmission range alone, as it is also governed by higher layer protocols, which normally have mechanism for concealing link errors to a certain degree. In general, there is a lack of good models for random geometric graphs where the connectivity is decided by higher layer protocols.

VII. CONCLUSIONS AND FUTURE WORK

Fixed node topologies with static nodes are first investigated. This analysis might be pertinent to planned network structures, such as backhaul mesh networks. The network is modelled as a random graph, where the inherent characteristics of the wireless medium are assumed to result in stochastic link failures with a link failure probability p . The results show how redundant nodes improve network reliability and availability. Although analyses and results depend on the actual topology, the same analyses as are presented here, can be applied to any specific backhaul mesh topology of interest.

Topologies with a random distribution of redundant nodes are also analysed by the use of concepts from random geometric graphs. Such topologies are relevant to sensor networks, to mesh network scenarios where an ad hoc network is used as a failover network, and to mobile ad hoc networks (MANETs) targeted at providing connectivity between some strategically located nodes (i.e. the MAPs and the MPP). The analyses can act as a guideline for the required node density in these networks in order to meet a specific reliability and availability.

In addition to finding the steady-state availability, the time-dependent behaviour of the reliability is also analysed. Unlike our steady-state analysis, the transient analysis assumes that all nodes have a fixed position, as the time dependency of the location is not taken into account. Further work is needed to

analyse the time dependency derived from the node mobility. The presented transient analysis might be extended to also take node mobility into account, e.g. using the approach in [21] as a starting point for the model extension.

The paper demonstrates that if the network includes m MAPs and $k - m$ MPPs, the k -terminal reliability is a key performance metrics for determining the network reliability and availability. However, the k -terminal reliability problem is known to belong to the class of #P-complete problems, which was introduced in [23]. This is a definitive limitation of using the k -terminal reliability metric, since these problems have exponential time complexity.

In the analyses on both the planned and the random mesh networks it is an underlying assumption that the core topology is given. Thus, in both these cases the paper focuses on showing how the number of redundant nodes affect the availability. For the analysis of planned mesh networks, it is assumed that the redundant nodes can be placed in a limited set of locations, and the order of priority between these locations is found. For the random mesh network analysis it is assumed that the random distribution of redundant node is given, and the analysis focuses only on the number of nodes. A natural expansion of the work presented in this paper is to study scenarios with a higher degree of freedom.

First, in some scenarios of planned mesh networks, it is certainly possible that the locations of the redundant nodes are arbitrarily, and not in a limited set of destinations. Obtaining intuition about how to determine the optimal location for the redundant node, is an interesting issue for further work. The work should aim to further enhancing intuition about how to distribute a number of redundant nodes (e.g. organising them in different types of grid structures around the backbone/core) to best improve the reliability and availability of the resulting network.

Second, our assumption about a backbone with core nodes situated in fixed locations could also be relaxed. In planned mesh networks, one might not need to make a distinction between core nodes and redundant nodes, and instead analyse the types of node structures that provide the highest availability between a given set of source and destination nodes. It would also be an interesting issue for future work to see how a set of core nodes should be located around a random network with a given distribution of the randomly located nodes.

Another interesting topic for future work is to define a link model that is based on concepts like capture threshold. Then, the probability of correctly receiving a packet on a link varies with distance and interference caused by on-going transmission from other nodes. This could be used to define the link failure probability. In [24] the capture function is a unit step function, while other studies have shown that the capture function is more probabilistic [25] [26]. With such a model, it might be possible to incorporate protocol characteristic, and thus provide a more realistic approach to network reliability and availability analysis.

APPENDIX A MEAN TIME TO FAILURE CALCULATION

In order to find the expected value $E(t_d)$, the following can be used:

$$E(X^j) = \int_0^{\infty} jx^{j-1}(1 - F(x))dx, j = 1, 2, \dots \quad (19)$$

Then,

$$\begin{aligned} E(t_d) &= \int_0^{\infty} [1 - R_d(G, t)] dt \\ &= \int_0^{\infty} \left[1 - \sum_{i=\beta}^{\epsilon} C_i (1 - e^{-\lambda t})^i (e^{-\lambda t})^{\epsilon-i} \right] dt \\ &= \underbrace{\int_0^{\infty} 1 - C_{\epsilon} (1 - e^{-\lambda t})^{\epsilon} dt}_{E_1(t_d)} - \underbrace{\int_0^{\infty} \sum_{i=\beta}^{\epsilon-1} C_i (1 - e^{-\lambda t})^i (e^{-\lambda t})^{\epsilon-i} dt}_{E_2(t_d)} \end{aligned} \quad (20)$$

Using the binomial theorem, $E_1(t_d)$ can be written as:

$$E_1(t_d) = \int_0^{\infty} 1 - C_{\epsilon} \sum_{k=0}^{\epsilon} \binom{\epsilon}{k} (-1)^k e^{-\lambda t k} dt \quad (21)$$

For all topologies G , one has that $C_{\epsilon} = 1$, and $E_1(t_d)$ can now be expressed as:

$$E_1(t_d) = \frac{1}{\lambda} \sum_{k=1}^{\epsilon} \binom{\epsilon}{k} \frac{(-1)^{k+1}}{k} \quad (22)$$

The integral for $E_2(t_d)$ is easier to solve. Putting $u = e^{-\lambda t}$, $E_2(t_d)$ can be written as:

$$E_2(t_d) = \sum_{i=\beta}^{\epsilon-1} C_i \int_1^0 -\frac{1}{\lambda} (1 - u)^i u^{\epsilon-i-1} du \quad (23)$$

$$= \frac{1}{\lambda} \sum_{i=\beta}^{\epsilon-1} C_i B(i+1, \epsilon-i) \quad (24)$$

where $B(x, y)$ is the *Complete Beta function*.

The expected value $E(t_d)$ is then:

$$E(t_d) = \frac{1}{\lambda} \sum_{k=1}^{\epsilon} \binom{\epsilon}{k} \frac{(-1)^{k+1}}{k} - \frac{1}{\lambda} \sum_{i=\beta}^{\epsilon-1} C_i B(i+1, \epsilon-i)$$

ACKNOWLEDGMENT

This work has been financed by the Research Council of Norway through the SWACOM project.

REFERENCES

- [1] *LAN/MAN Specific Requirements - Part 11: Wireless Medium Access Control (MAC) and physical layer (PHY) specifications: Amendment: ESS Mesh Networking*, IEEE Std. 802.11s Draft 2.0, 2008.
- [2] (2008, August) Ietf mobile ad-hoc networks working group. [Online]. Available: <http://www.ietf.org/html.charters/manet-charter.html>
- [3] H. Gharavi and S. Kumar, "Special issue on sensor networks and applications," *Proceedings of the IEEE*, vol. 91, no. 8, pp. 1151–1153, Aug. 2003.
- [4] *Wireless LAN Medium Access Control (MAC) and Physical Layer (PHY) Specification*, IEEE Std. 802.11, 1997.
- [5] R. Hekmat and P. V. Miegheem, "Interference in wireless multihop ad-hoc network," in *in Med-hoc-Net*, 2002, pp. 389–399.
- [6] M. Gerharz, C. D. Waal, M. Frank, and P. Martini, "Link stability in mobile wireless ad hoc networks," in *In Proceedings of the 27th Annual IEEE Conference on Local Computer Networks (LCN'02*, 2002, pp. 30–39.
- [7] C. Bettstetter, "On the minimum node degree and connectivity of a wireless multihop network," in *MobiHoc '02: Proceedings of the 3rd ACM international symposium on Mobile ad hoc networking & computing*. New York, NY, USA: ACM, 2002, pp. 80–91.
- [8] R. Hekmat and P. Van Miegheem, "Degree distribution and hopcount in wireless ad-hoc networks," *Networks, 2003. ICON2003. The 11th IEEE International Conference on*, pp. 603–609, Sept.-1 Oct. 2003.
- [9] R. Hekmat and P. V. Miegheem, "Study of connectivity in wireless ad-hoc networks with an improved radio model," in *in Proc. of WiOpt*, 2004.
- [10] T. S. Rappaport and T. Rappaport, *Wireless Communications: Principles and Practice (2nd Edition)*. Prentice Hall PTR, December 2001.
- [11] H. L. Bertoni, *Radio Propagation for Modern Wireless Systems*. Prentice Hall Professional Technical Reference, 1999.
- [12] H. Holma and A. Toskala, Eds., *WCDMA for UMTS: Radio Access for Third Generation Mobile Communications*. New York, NY, USA: John Wiley & Sons, Inc., 2001.
- [13] M. L. Shooman, *Reliability of Computer Systems and Networks: Fault Tolerance, Analysis, and Design*. John Wiley and Sons, Inc., 2002.
- [14] A. Kershenbaum, *Telecommunications network design algorithms*. New York, NY, USA: McGraw-Hill, Inc., 1993.
- [15] C. J. Colbourn, "Reliability issues in telecommunications network planning," in *Telecommunications network planning*, B. S. P. Soriano, Ed. Kluwer Academic Publishers, 1999, ch. 9, pp. 135–146.
- [16] —, *The Combinatorics of Network Reliability*. Oxford Univ. Press, 1987.
- [17] G. Weichenberg, V. Chan, and M. Medard, "High-reliability topological architectures for networks under stress," *Selected Areas in Communications, IEEE Journal on*, vol. 22, no. 9, pp. 1830–1845, Nov. 2004.
- [18] F. Tobagi and L. Kleinrock, "Packet switching in radio channels: Part ii—the hidden terminal problem in carrier sense multiple-access and the busy-tone solution," *Communications, IEEE Transactions on [legacy, pre - 1988]*, vol. 23, no. 12, pp. 1417–1433, 1975.
- [19] H. AboElFotouh and C. J. Colbourn, "Computing 2-terminal reliability for radio-broadcast networks," *Reliability, IEEE Transactions on*, vol. 38, no. 5, pp. 538–555, Dec 1989.
- [20] H. AboElFotouh, S. Lyengar, and K. Chakrabarty, "Computing reliability and message delay for cooperative wireless distributed sensor networks subject to random failures," *Reliability, IEEE Transactions on*, vol. 54, no. 1, pp. 145–155, March 2005.
- [21] S. Kharbash and W. Wang, "Computing two-terminal reliability in mobile ad hoc networks," *Wireless Communications and Networking Conference, 2007.WCNC 2007. IEEE*, 11–15 March 2007.
- [22] I. Chlamtac and A. Faragó, "A new approach to the design and analysis of peer-to-peer mobile networks," *Wirel. Netw.*, vol. 5, no. 3, pp. 149–156, 1999.
- [23] L. G. Valiant, "The complexity of computing the permanent," *Theor. Comput. Sci.*, vol. 8, pp. 189–201, 1979.
- [24] The ns2 network simulator. <http://www.isi.edu/nsnam/ns/>.
- [25] M. Soroushnejad and E. Geraniotis, "Probability of capture and rejection of primary multiple-access interference in spread-spectrum networks," *Communications, IEEE Transactions on*, vol. 39, no. 6, pp. 986–994, Jun 1991.
- [26] A. Kochut, A. Vasan, A. U. Shankar, and A. Agrawala, "Sniffing out the correct physical layer capture model in 802.11b," in *ICNP '04: Proceedings of the 12th IEEE International Conference on Network Protocols*. Washington, DC, USA: IEEE Computer Society, 2004, pp. 252–261.



protocols for mobile ad hoc network. The focus for his PhD work is reliability in wireless ad hoc networks.



Dr. Paal E. Engelstad is associate Professor at the Unik, University of Oslo (UiO), a Senior Research Scientist at Telenor R&I, and working as a project manager at Simula Research Laboratory. He holds a number of patents and has been publishing a number of papers over the past years. Paal E. Engelstad has a PhD in Computer Science (UiO - 2005), a bachelor in Computer Science (UiO - 2001), a master in Physics (NTNU/Kyoto University, Japan - 1994 - Honors degree with distinction) and a bachelor in Physics (NTNU - 1993) The topic of his PhD was Middleware and Autoconfiguration functionality in Ad Hoc and Personal Area Networks. His research interests include mobility, IP technology, ad hoc and wireless networks. Nowadays, his main focus is on network reliability, on performance analysis of IEEE 802.11e and on wireless sensor networks.