

Analysis of the Self-Organization in Wireless Multi-Hops Networks

Nathalie Mitton - Anthony Busson - Eric Fleury

N° 5328

Octobre 2004

THÈME 1



***apport
de recherche***

Analysis of the Self-Organization in Wireless Multi-Hops Networks

Nathalie Mitton - Anthony Busson - Eric Fleury

Thème 1 — Réseaux et systèmes

Projet ARES

Rapport de recherche n° 5328 — Octobre 2004 — 42 pages

Abstract: Flat wireless multi-hops architectures are not scalable. In order to overcome this major drawback, hierarchical routing is introduced since it is found to be more effective. The main challenge in hierarchical routing is to group nodes into clusters. Each cluster is represented by its cluster head. Most of conventional methods use either the connectivity (degree) or the node Id to perform the cluster-head election. Such parameters are not really robust in terms of side effects. In this paper we introduce a novel measure that both forms clusters and performs the cluster-head election. Analytical models and simulation results show that this new measure proposed for cluster-head election induces less cluster-head changes as compared to classical methods.

Key-words: ad hoc, wireless, self-organization, stochastic geometry, scalability

Unité de recherche INRIA Rhône-Alpes

655, avenue de l'Europe, 38334 Montbonnot Saint Ismier (France)

Téléphone : +33 4 76 61 52 00 — Télécopie +33 4 76 61 52 52

Analyse de l'auto-organisation dans les réseaux sans fils

multi-sauts

Résumé : Les architectures à plat n'offrent que des possibilités d'utilisation sur de petites échelles. Dans le but de pourvoir à cet inconvénient, le routage hiérarchique s'est montré efficace. Le principal défi est donc de regrouper les nœuds en groupes dits "clusters". Chaque cluster est représenté par son chef de cluster. Les méthodes conventionnelles utilisent, entre autres, pour son élection le degré ou encore l'identifiant. De tels paramètres n'offrent pas toujours les bonnes propriétés. Nous proposons ici une nouvelle métrique permettant d'organiser le réseau de façon à pouvoir utiliser un réseau ad hoc sur des échelles plus larges. Les résultats analytiques et de simulation montrent que cette métrique induit moins de changements de topologie que les méthodes classiques.

Mots-clés : ad hoc, auto-organisation, géométrie stochastique, passage à l'échelle

1 Introduction

Wireless Multi-hops Networks (WMN *e.g.* ad-hoc, sensors...) consist of a set of mobile wireless nodes without the support of a pre-existing fixed infrastructure. Ad hoc networks find applications in battlefields coordination or on-site disaster relief management, sensor net in zone monitors... Each host/node acts as a router and is able to arbitrary move. This feature is a challenging issue for protocol design since the protocol must adapt to frequent changes of network topologies. More recently, researchers have applied ad hoc paradigms in sensor networks which induce to be able to set up a very large number of nodes.

In order to be able to use wireless multi-hops networks on very large scales, flat routing protocols (reactive or proactive) are not really suitable. Indeed, such routing protocols become ineffective for large scale wireless multi-hops networks, because of bandwidth (flooding of control messages) and processing overhead (routing table computation). The most common solution used to solve this scalability problem is to introduce a hierarchical routing by grouping geographically close nodes into clusters and by using an "hybrid" routing scheme: classically proactive approach inside each cluster and reactive approach between clusters ([12, 16, 9]). Such an organization also presents numerous advantages as to synchronize stations in a group or to attribute new service zones more easily.

In this report, we propose a novel metric suitable for organizing a WM network into clusters and we present a new distributed cluster-head election heuristic for a wireless multi-hops network. Our new metric does not rely on "static" parameters and thus our novel heuristic extends the notion of clusters formation. The proposed heuristic allows load balancing to insure a fair distribution

of load among cluster-heads. Moreover, we implement a mechanism for the cluster-head election that tries to favor their re-election in future rounds, thereby reducing transition overheads when old cluster-heads give way to new ones. We expect from network organization to be robust towards node mobility. If we want to keep overhead as low as possible, our organization must change as less as possible when nodes move and topology evolves. Moreover, we would like to be able to apply some localization process and inter-groups routing above our organization.

The remainder of this paper is organized as follows. Section 2 defines the system model and introduces some notations. Section 3 reviews several techniques proposed for cluster-head selection. Sections 4 and 5 present our main contribution, detail the distributed selection algorithm and give some formal analysis. Simulation experiments presented in section 6 demonstrate that the proposed heuristic is better than earlier heuristics in terms of stability. Finally, we conclude in section 7 by discussing possible future areas of investigation.

2 System model

In a WM network, all nodes are alike and may be mobile. There is no base station to coordinate the activities of subsets of nodes. Therefore, all the nodes have to collectively make decisions and the use of distributed algorithms is mandatory. Moreover, all communications are performed over wireless links. We classically model an ad hoc network by a graph $G = (V, E)$ where V is the set of mobile nodes ($|V| = n$) and $e = (u, v) \in E$ represents a wireless link between a pair of nodes u and v if and only if they are within communication range of each other.

For the sake of simplicity, let's first introduce some notations. Let's call $d(u, v)$ the distance between nodes u and v in the graph G (i.e. the number of hops between nodes u and v). We note $\mathcal{C}(u)$ the cluster owning the node u and $\mathcal{H}(u)$ the cluster-head of this cluster. We will also note $\Gamma_k(u)$ the k -neighborhood of a node u , i.e., $\Gamma_k(u) = \{v \in V | v \neq u, d(u, v) \leq k\}$ and will note $\delta_k(u) = |\Gamma_k(u)|$. Thus we have $\delta_1(u) = \delta(u)$ being the degree of node u . Note that node u does not belong to $\Gamma_k(u)$.

We will note $e(u/\mathcal{C}) = \max_{v \in \mathcal{C}(u)}(d(u, v))$ the *eccentricity* of a node u inside its cluster. Thus the *diameter* of a cluster will be $D(\mathcal{C}(u)) = \max_{v \in \mathcal{C}(u)}(e(v/\mathcal{C}))$.

3 Related work

Researchers have proposed several techniques for cluster formation and cluster-head selection. All solutions aim to identify a subset of nodes within the network and bind a leader to it. Each cluster-head is responsible for managing communication between nodes into their cluster as well as routing information to other cluster-heads in other clusters. Typically, backbones are constructed to connect neighborhoods in the network.

Some past solutions try to gather nodes into homogeneous clusters by using either an identity criteria (e.g., the lowest Id [10, 3]) or a fixed connectivity criteria (maximum degree [4], 1-hop clusters [2, 6, 11], k -hop clusters [7]) or based on both connectivity and identity criteria (Max-Min d-cluster [1]).

Such solutions based on a fixed cluster diameter [1, 4, 8], fixed cluster radius [13] or a constant number of nodes by cluster [17] are not adapted to large WM networks since they may generate a

large number of cluster-heads and changes when nodes move. Therefore, it is suitable to control the cluster-heads' density in the network. The underlying idea in solutions based on the degree [8, 15] is that nodes with a high degree are good candidates to be eligible as cluster-heads since the resulting clusters will be larger. However, even small changes in the network topology may result in large changes in the degree of the nodes. This means that the cluster-heads are not likely to remain as cluster-heads for a long time and the clustering structure becomes unstable. Using the lowest-ID algorithm, the nodes with a low ID remain cluster-heads most of the time but the clusters' shape may not be very suitable and this is an unfair distribution which could lead some nodes to loose power prematurely. Some previous clustering solutions also rely on synchronous clocks for exchange of data between nodes as for example in the Linked Cluster Algorithm (LCA) [2], but such an heuristic is developed for relatively small number of nodes (less than 100) and cannot be envisaged for a greater amount. Solutions were also envisaged to base the election on a pure mobility criteria [4] but even if mobility should be taken into account, electing only non mobile nodes may result in isolated cluster-heads, which may be useless.

In all previous works, the design of clusters selection appears to be similar with few variants. Each node locally computes its own value of a given criteria (degree, mobility...) and locally broadcasts this value in order to compete with its neighbors. All nodes are thus able to decide by their own if they win the tournament and can be declared cluster-head. In case of multiple winners, other criterion (*e.g.*, Id) are used until unambiguity. Every nodes which had joined the same node belong to the same cluster that can be identified by its cluster-head.

4 Main objectives

The main goal is to design a heuristic that selects some nodes as cluster-heads and computes clusters in a large WN network. As we mentioned in the previous section, the definition of a cluster should not be defined a priori by some fixed criteria but must reflect the *density* of the network. Indeed, if we need to change clusters structure every time we add or remove nodes, we will not be able to be scalable since changes will be too much frequent with a great amount of mobile nodes and will generate too much traffic. Wireless links will be quickly saturated. In order to be scalable, the heuristic should be completely distributed and asynchronous (avoiding any clock synchronization). The number of messages exchanges should be minimized. In fact, we use only local broadcast messages like HELLO PACKET ([5]) in order to discover the 2-neighborhood of a node. Finally, in order to ensure stability, it would be better to maintain cluster-heads as such whenever it is possible and nodes that are "too" mobile to initiate any communication will not participate in the ballot phase and won't belong to any cluster. Otherwise, they would break the clusters structure uselessly.

The criteria metric should gather and aggregate nodes into clusters not on an absolute criteria (like degree or diameter) and thus should be adaptive in order to reflect the features of the network. To elect a cluster-head, we need to promote node stability by limiting traffic overhead when building and maintaining the network organization. Secondly, the criteria should be robust, *i.e.* not be disturbed by a slightly topology change. Finally, the criteria should be computed locally by using only local traffic (intra-cluster routing) since it is cheaper than inter-clusters traffic. The main is to reconstruct the clusters topology as less as possible in spite of great nodes mobility.

Based on these general requirements we propose a novel heuristic based on a metric criteria which gathers the density of the neighborhood of a node. This density criteria reveals to be stable when the topology evolves slightly. As the network topology changes slightly the node's degree is much more likely to change than its density that smooths the relative topology changes down inside its own neighborhood.

5 Our contributions

5.1 The density metric criteria

In this section, we introduce our criteria called *density*. The notion of density should characterize the "relative" importance of a node in the WN network and in its k -neighborhood. As mentioned earlier, the node degree is not adequate. The density notion should absorb small topology changes. The underlying idea is that if some nodes move in $\Gamma_1(u)$ (i.e., a small evolution in the topology), changes will affect the microscopic view of node u (its degree $\delta_1(u)$ will change) but its macroscopic view will in fact not change a lot since globally the network does not drastically change and its $\Gamma_1(u)$ globally remains the same. The density is directly related to both the number of nodes and links in a k -neighborhood. Indeed, the density will smooth local changes down in $\Gamma_k(u)$ by considering the ratio between the number of links and the number of nodes in $\Gamma_k(u)$.

Definition 1 (density) *The k -density of a node $u \in V$ is*

$$\rho_k(u) = \frac{|e = (v, w) \in E \mid w \in \{u, \Gamma_k(u)\} \text{ and } v \in \Gamma_k(u)|}{\delta_k(u)} \quad (1)$$

The 1-density (also noted $\rho(u)$) is thus the ratio between the number of edges between u and its 1-neighbors (by definition the degree of u), the number of edges between u 's 1-neighbors and the number of nodes inside u 's 1-neighborhood.

To illustrate this definition, let's take the following example on Figure 1. Let's consider the node p and let's compute its density value $\rho(p)$. It is the ratio between the number of edges $L(p)$ and the number of nodes $\Gamma(p)$ in its 1-neighborhood. Nodes in the 1-neighborhood of node p are the dark gray ones ($\Gamma(p) = \{a, b, c, d, e, f\}$). We thus have $L(p)$ equal to the number of links between these nodes (dashed links) and also the number of links from node p toward these dark gray nodes (dotted links). So, $L(p) = 4 + 6 = 10$ and $\delta(p) = 6$ thus $\rho(p) = 10/6 = 5/3$. Note that to compute $\rho(p)$, node p needs to know $\Gamma_2(p)$ since it must be able to compute the number of edges that exist between all its 1-neighbors.

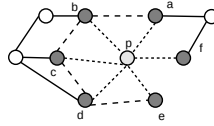


Figure 1: Density example.

In [14], we have shown that the most robust metric among the different k -density metrics is actually the 1-density ($k = 1$), which is also the cheapest in terms of message exchanges and time complexity. Indeed, note that to compute $\rho_k(u)$, the node u must know $\Gamma_{k+1}(u)$ since it must be able to compute the number of edges that exist between all its k -neighbors.

So, in the following, we mainly consider the 1-density.

5.2 Cluster-head selection and cluster formation

5.2.1 Basic idea

Each node locally computes its density value and broadcasts it to all its 1-neighbors (*e.g.* by piggybacking in `HELLO` packets). Each node is thus able to decide by itself whether it wins in its 1-neighborhood (the smallest Id will be used to decide between joint winners) or it joins one of its neighbors. This neighbor may have joined a other node and so one. The cluster-head will be the node which has joined itself. If node u has joined node w , we will say that w is node u 's parent and we will note it $w = \mathcal{F}(u)$. The cluster can then extend itself until it reaches a cluster frontier. The only two constraints that we introduce here to define a cluster is that, first two neighbors can not be both cluster-heads and then that, if node u is cluster-head, all of its 1-neighbors belong to its cluster. This ensures that a cluster-head is not too off-center in its own cluster, that a cluster has at least a diameter of two and that two cluster-heads will be distant of at least three hops. This will minimize packets collisions.

5.2.2 Heuristic

The heuristic process is quite simple. On a regular basis (frequency of `HELLO` packets for instance), each node computes its k -density based on its view of its $(k + 1)$ -neighborhood. To simplify the notation we describe here the 1-density heuristic in algorithm 1. The k -density is similar since the only modification to make is to gather the $(k + 1)$ -neighborhood which is given by sending `HELLO` packets within k -hops.

Algorithm 1 Cluster-head selection

For all node $u \in V$

▷ *Checking the neighborhood*

Gather $\Gamma_2(u)$

if ($\text{mobility}(u, \Gamma_1(u)) > \text{threshold}$) **then**

▷ *Checking the 1-neighborhood consistency. If this one changes too much, node u will not participate to the ballot phase since it is "relatively" too mobile.*

break

end

Compute $\rho(u)$

Locally broadcast $\rho(u)$

▷ *This local broadcast can be done by piggybacking $\rho(u)$ in HELLO packets.*

▷ *At this point, node u is aware of all of its 1-neighbors' density value and knows whether they are eligible.*

if ($\rho(u) = \max_{v \in \Gamma_1(u)}(\rho(v))$) **then**

$\mathcal{H}(u) = u$

▷ *u is promoted cluster-head.*

▷ *Note that if several nodes are joint winners, the winner will be the previous cluster-head whether it exists, otherwise, the less mobile node, otherwise, the smallest Id.*

$\forall v \in \Gamma_1(u), v \in \mathcal{C}(u), \mathcal{F}(v) = u, \mathcal{H}(v) = u$

▷ *All neighbors of u will join the cluster created by u as well as all nodes which had joined u 's neighbors. u become their cluster-head as well as their parent.*

else

▷ $\exists w \in \Gamma_1(u) | \rho(w) = \max_{v \in \Gamma_1(u)}(\rho(v))$

$\mathcal{F}(u) = w$

$$\mathcal{H}(u) = \mathcal{H}(w)$$

▷ Either $\mathcal{F}(w) = \mathcal{H}(w) = w$ and u is directly linked to its cluster-head, either $\exists x \in \Gamma_1(w) | \mathcal{F}(w) = x$ (w has joined another node x) and $\mathcal{H}(u) = \mathcal{H}(w) = \mathcal{H}(x)$.

▷ If there exist k ($k > 1$) nodes w_i such that $\rho(w_i) = \max_{v \in \Gamma_1(u)} (\rho(v))$ and such that $w_i \notin \Gamma_1(w_j)$ ($i \neq j$) then u will join the node w_i which Id is the lowest and all $\mathcal{C}(w_i)$ (for $i=1$ to k) will merge consequently.

end

Locally broadcast $\mathcal{F}(u)$ and $\mathcal{H}(u)$

5.3 Example

To illustrate this heuristic, let's take the following example (Fig 2). Let's suppose that the node E is too mobile to be eligible.

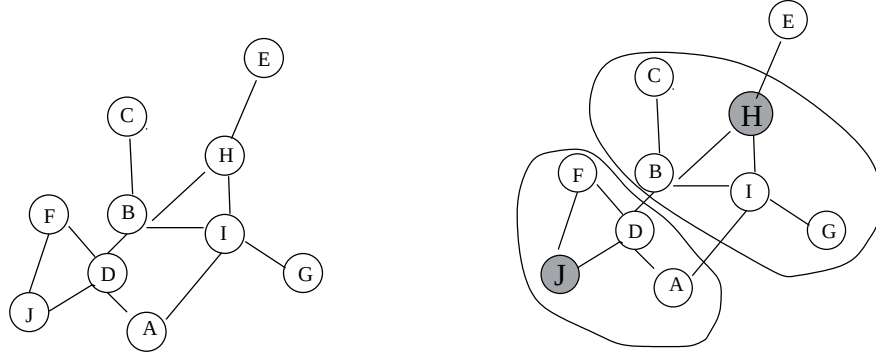


Figure 2: Clustering example.

In its 1-neighborhood topology, node A has two 1-neighbors ($\Gamma_1(A) = \{D, I\}$) and two links ($\{(A, D), (A, I)\}$); node B has four 1-neighbors ($\Gamma_1(B) = \{C, D, H, I\}$) and five links ($\{(B, C), (B, D), (B, H), (B, I), (H, I)\}$). Table 1 shows the final results.

Nodes	<i>A</i>	<i>B</i>	<i>C</i>	<i>D</i>	<i>E</i>	<i>F</i>	<i>G</i>	<i>H</i>	<i>I</i>	<i>J</i>
Neighbors	2	4	1	4		2	1	2	4	2
Links	2	5	1	5		3	1	3	5	3
1-density	1	1.25	1	1.25		1.5	1	1.5	1.25	1.5

Table 1: Results of our heuristics on the illustrative example.

In the illustrative example, node C joins its 1-neighbor which density is the highest: node B ($\mathcal{F}(C) = B$). Yet, the node with the highest density in node B 's neighborhood is H . Thus, $\mathcal{F}(B) = H$ and so $\mathcal{H}(C) = \mathcal{H}(H)$. As node H has the highest density in its own neighborhood, it becomes its own cluster-head: $\mathcal{H}(H) = H$. To sum up, C joins B which joins H and all three of them belong to the cluster which cluster-head is H : $\mathcal{H}(C) = \mathcal{H}(B) = \mathcal{H}(H) = H$. Moreover, we have $\rho_1(J) = \rho_1(F)$. As it is the first construction, none of J and F was cluster-head before. If we suppose that J has the smallest Id between both nodes $\mathcal{H}(F) = \mathcal{H}(J) = J$. At last, we obtain two clusters organized around two cluster-heads: H and J . (See figure on the right side on figure 2)

5.4 Maintenance

Given that every node is mobile and subject to move at any time, our cluster organization must adapt to topology changes. For this, our nodes have to periodically check their environment and so check their mobility. If they become too mobile, they will not join any cluster, if at the opposite, they were too mobile and now are able to communicate, they will join the cluster of their neighbor which has the highest density. Each node periodically checks its density and its neighbors' one. They continue joining their neighbor which has the highest density. If this last changes, the reconstruction will

be automatic without generating much additional traffic overhead. In our example (5.3), if node E slows down and becomes eligible, it will compute its density value. It has 1 neighbor (node H) and 1 link ($H - E$) in its 1-neighborhood, thus its density value is 1. Node H will thus be aware of it and compute its new density value in consequent: it then has 3 neighbors (the previous plus node E) and 4 links (the previous plus link $H - E$) in its 1-neighborhood. So we have $\rho(H) = 4/3 = 1,33$ and $\rho(E) = 1$. Following the heuristic, we have $\mathcal{F}(E) = H$ and still $\mathcal{H}(H) = H$. Node E will simply join the former cluster which cluster-head is H without any changes.

5.4.1 Departure and arrival

These processes are very simple. Indeed, when a node arrives or slows down in a way that it becomes enough stable to join a cluster, it checks its neighborhood listening and transmitting HELLO packets. It can then compute the heuristic algorithm and join a cluster. Then, it can join its neighbor which density value is the highest or create its own cluster according to the selection criterion.

When a node goes away, its former neighbors won't receive its HELLO packets anymore and then will change their density value consequently. For its own, either it becomes too mobile to joins any cluster, either it checks its new neighborhood and follows the heuristic.

Actually, as every node checks periodically its mobility and its k -density values over its k -neighborhood, every node movement is registered automatically without any supplementary message exchanges.

5.5 Analysis of the average density

In this section we compute two important characteristic factors of our cluster heuristic. We first compute the mean density of nodes and then we compute an upper bound on the expected number of cluster-heads.

We first analyze the average 1-density $\tilde{\rho}(u)$ of a node u . We consider a wireless multiple-hops network where nodes are distributed according to a Poisson point process of constant spatial intensity λ . Each node has a transmission range equal to R depending on its transmitting power $\mathcal{P}_u$.

We compute the mean density under Palm probability. We thus compute the density of a node located at the origin point (under Palm probability, there exists almost certainly a point in 0). Since we consider a stationary point process, this node's density is valid for every point. Let $\rho(0)$ be the density value of node 0. Φ is used to design the point process. $\Phi(B(u, R))$ thus represents the number of the process points in $B(u, R)$ where $B(u, R)$ is the ball centered in u with radius R . $\mathbb{E}^o$ and $\mathbb{P}^o$ design respectively the expectation and the probability under Palm distribution. We compute:

$$\tilde{\rho}(u) = \mathbb{E}^o [\rho(0)]$$

Lemma 1 *The mean 1-density of any node u is $\tilde{\rho}(u) = \mathbb{E}^o [\rho(0)]$ where:*

$$\mathbb{E}^o [\rho(0)] = 1 + \frac{1}{2} \left(\pi - \frac{3\sqrt{3}}{4} \right) \times \left(\lambda R^2 - \frac{1 - \exp\{-\lambda\pi R^2\}}{\pi} \right)$$

Proof 1 Let B'_u be the ball centered in $u \in \mathbb{R}^2$, with radius R minus the node 0, that is, $B'_u = B(u, R) \setminus u$. Let's note $(Y_i)_{i=1, \dots, \Phi(B'_0)}$ Φ 's nodes being in B'_0 . From the density definition, we have:

$$\mathbb{E}^o [\rho(0)] = 1 + \frac{1}{2} \mathbb{E}^o \left[\sum_{i=1}^{\Phi(B'_0)} \frac{\Phi(B'_0 \cap B'_{Y_i})}{\Phi(B'_0)} \right]$$

Moreover, we suppose that $\rho(0) = 1$ if $\Phi(B'_0) = 0$. Actually, we are going to consider the conditional expected value that we are looking for. More precisely, we consider the expected value conditioned on the number of nodes in B'_0 . Thus we have:

$$\begin{aligned} \mathbb{E}^o [\rho(0)] &= 1 + \frac{1}{2} \sum_{k=1}^{+\infty} \mathbb{E}^o \left[\sum_{i=1}^{\Phi(B'_0)} \frac{\Phi(B'_0 \cap B'_{Y_i})}{\Phi(B'_0)} \middle| \Phi(B'_0) = k \right] \mathbb{P}^o (\Phi(B'_0) = k) \\ &= 1 + \frac{1}{2} \sum_{k=1}^{+\infty} \sum_{i=1}^k \frac{1}{k} \mathbb{E}^o \left[\Phi(B'_0 \cap B'_{Y_i}) \middle| \Phi(B'_0) = k \right] \times \mathbb{P}^o (\Phi(B'_0) = k) \end{aligned} \quad (2)$$

Moreover, we know that $\Phi(B'_0) = k$, and that nodes $(Y_i)_{i=1, \dots, k}$ are independent one from each other and uniformly distributed in B'_0 . Thus, $\mathbb{E}^o \left[\Phi(B'_0 \cap B'_{Y_i}) \middle| \Phi(B'_0) = k \right]$ is the same for all $i, i = 1, \dots, k$. Knowing $\nu(B'_0 \cap B'_{Y_i})$ (ν is the Lebesgue measure in $\mathbb{R}^2$) and that $\Phi(B'_0) = k$, the amount of nodes in $B'_0 \cap B'_{Y_i}$ follows a binomial law with parameter $\left(k - 1, \frac{\nu(B'_0 \cap B'_{Y_i})}{\nu(B'_0)}\right)$ and the average number of points is $(k - 1) \frac{\nu(B'_0 \cap B'_{Y_i})}{\nu(B'_0)}$.

Thus we have, for all $i = 1, \dots, k$:

$$\begin{aligned} \mathbb{E}^o \left[\Phi(B'_0 \cap B'_{Y_i}) \middle| \Phi(B'_0) = k \right] &= \frac{(k - 1)}{\pi R^2} \mathbb{E}^o \left[\nu(B'_0 \cap B'_{Y_i}) \middle| \Phi(B'_0) = k \right] \\ &= \frac{(k - 1)}{\pi R^2} \mathbb{E}^o [\nu(B'_0 \cap B'_{Y_i})] \end{aligned} \quad (3)$$

This last equality comes from the fact that the area $\nu(B'_0 \cap B'_{Y_i})$ does not depend of the number of nodes in B'_0 , since all Y_i are independent.

Knowing that Y_i is at a distant r from the origin point, we can compute the area of the intersection

$$\nu(B'_0 \cap B'_{Y_i}) = A(r) = 2R^2 \arccos \frac{r}{2R} - r \sqrt{R^2 - \frac{r^2}{4}}.$$

and thus, since Y_i is uniformly distributed in B'_0 , we have

$$\mathbb{E}^o [\nu(\Phi(B'_0 \cap B'_{Y_i}))] = \mathbb{E}^o [A(r)] \quad (4)$$

$$= \int_0^{2\pi} \int_0^R \frac{A(r)}{\pi R^2} r \, dr \, d\theta \quad (5)$$

$$= R^2 \left(\pi - \frac{3\sqrt{3}}{4} \right) \quad (6)$$

The last quantity divided by πR^2 (i.e. $1 - \frac{3\sqrt{3}}{4\pi}$) will be denoted p in the rest of the report. p corresponds to the probability that two neighbors of the origin are themselves neighbors. p is approximately equal to 0.58.

Combined with equation 2, we obtain

$$\begin{aligned} \mathbb{E}^o [\rho(0)] &= 1 + \frac{1}{2} \sum_{k=1}^{+\infty} \sum_{i=1}^k \frac{1}{k} \frac{k-1}{\pi} \left(\pi - \frac{3\sqrt{3}}{4} \right) \mathbb{P}^o (\Phi(B'_0) = k) \\ &= 1 + \frac{1}{2} \sum_{k=1}^{+\infty} \frac{k-1}{\pi} \left(\pi - \frac{3\sqrt{3}}{4} \right) \mathbb{P}^o (\Phi(B'_0) = k) \\ &= 1 + \frac{1}{2\pi} \left(\pi - \frac{3\sqrt{3}}{4} \right) \times \left(\sum_{k=1}^{+\infty} k \mathbb{P}^o (\Phi(B'_0) = k) - \sum_{k=1}^{+\infty} \mathbb{P}^o (\Phi(B'_0) = k) \right) \\ &= 1 + \frac{1}{2\pi} \left(\pi - \frac{3\sqrt{3}}{4} \right) (\lambda \pi R^2 - (1 - \exp\{\lambda \pi R^2\})) \end{aligned} \quad (7)$$

This last equality comes from Slyvniack's theorem which says in this case that $\Phi(B'_0)$, under Palm probability, follows a discrete Poisson law with intensity $\lambda \pi R^2$. ■

Lemma 2 The variance of the number of edges between n ($n \geq 0$) points independently and uniformly distributed in $B(0, R)$ defined as $\mathbb{E}_n^o \left[\left(\sum_{i=1}^n \Phi(B'_0 \cap B'_i) \right)^2 \right] - \mathbb{E}_n^o \left[\sum_{i=1}^n \Phi(B'_0 \cap B'_i) \right]^2$ is:

$$\mathbb{E}_n^o \left[\left(\sum_{i=1}^n \Phi(B'_0 \cap B'_i) \right)^2 \right] - \mathbb{E}_n^o \left[\sum_{i=1}^n \Phi(B'_0 \cap B'_i) \right]^2 = 2n(n-1) [p - p^2 + 2(n-2)(p_2 - p^2)]$$

This quantity conditioning by a discrete Poisson law allows us to deduce the variance for the density:

$$\begin{aligned} \mathbb{E}^o [\rho(0)^2] - \mathbb{E}^o [\rho(0)]^2 &= \frac{1}{2} \left[(1 - \exp \{-\lambda \pi R^2\}) (2 + (p - p^2) - 3(p^2 - p_2)) + \lambda \pi R^2 (p_2 - p^2) \right. \\ &\quad \left. + (2(p_2 - p^2) - (p - p^2)) \sum_{i=1}^n \frac{1}{n} \frac{(\lambda \pi R^2)^n}{n!} \exp \{-\lambda \pi R^2\} \right] \end{aligned} \quad (8)$$

Proof 2 We note $\mathbb{E}_n^o$ the expectation under Palm and under the condition that $\Phi(B'_0) = n$. In a first time, we compute for $n > 1$

$$\mathbb{E}_n^o \left[\left(\sum_{i=1}^n \Phi(B'_0 \cap B'_i) \right)^2 \right]$$

We have,

$$\begin{aligned} \mathbb{E}_n^o \left[\left(\sum_{i=1}^n \Phi(B'_0 \cap B'_i) \right)^2 \right] &= \sum_{i=1}^n \mathbb{E}_n^o \left[\Phi(B'_0 \cap B'_i)^2 \right] + \sum_{i,j=1, i \neq j}^n \mathbb{E}_n^o \left[\Phi(B'_0 \cap B'_i) \Phi(B'_0 \cap B'_j) \right] \quad (9) \\ &= n \mathbb{E}_n^o \left[\Phi(B'_0 \cap B'_1)^2 \right] + n(n-1) \mathbb{E}_n^o \left[\Phi(B'_0 \cap B'_1) \Phi(B'_0 \cap B'_2) \right] \quad (10) \end{aligned} \quad (11)$$

with

$$\mathbb{E}_n^o \left[\Phi \left(B'_0 \cap B'_1 \right)^2 \right] = \mathbb{E}_n^o \left[\left(\sum_{i=2}^n \mathbf{1}_{|y_i - y_1| < R} \right)^2 \right] \quad (12)$$

$$= (n-1)\mathbb{E}_n^o [\mathbf{1}_{|y_i - y_1| < R}] + (n-1)(n-2)\mathbb{E}_n^o [\mathbf{1}_{|y_2 - y_1| < R} \mathbf{1}_{|y_3 - y_1| < R}] \quad (13)$$

$$= (n-1)p + (n-1)(n-2)p_2 \quad (14)$$

$$(15)$$

and

$$\mathbb{E}_n^o \left[\Phi \left(B'_0 \cap B'_1 \right) \Phi \left(B'_0 \cap B'_2 \right) \right] = \mathbb{E}_n^o \left[\left(\sum_{i=2}^n \mathbf{1}_{|y_i - y_1| < R} \right) \left(\sum_{i=1, i \neq 2}^n \mathbf{1}_{|y_i - y_2| < R} \right) \right] \quad (16)$$

$$= p + 3(n-2)p_2 + (n-2)(n-3)p^2 \quad (17)$$

$$(18)$$

p is the probability that two points independently and uniformly distributed in $B(0, R)$ are neighbors (as defined in the previous lemma). p_2 is the probability that for three points independently and uniformly distributed point in $B(0, R)$, one of them is neighbor of the two others.

Formally, let be Y_1, Y_2 and Y_3 three uniformly distributed points in B_0 , we have

$$p = \mathbb{P}(y_2 \in B_0 \cap B_{y_1}) = \frac{2}{\pi} \int_{u=0}^1 \left(2 \arccos \frac{u}{2} - u \sqrt{1 - \frac{u^2}{4}} \right) u du \approx 0.5865$$

$$p_2 = \mathbb{P}(y_2 \in B_0 \cap B_{y_1}, y_3 \in B_0 \cap B_{y_1}) = \frac{2}{\pi^2} \int_{u=0}^1 \left(2 \arccos \frac{u}{2} - u \sqrt{1 - \frac{u^2}{4}} \right)^2 u du \approx 0.3642$$

We observe that these two probabilities do not depend on R . Moreover, the probability that two points be neighbors of a same third point (probability defined as p_2) is different of p^2 ($p^2 \approx 0.344$ and $p_2 \approx 0.3642$).

We obtain,

$$\mathbb{E}_n^o \left[\left(\sum_{i=1}^n \Phi(B'_0 \cap B'_i) \right)^2 \right] = n(n-1) [2p + 4(n-2)p_2 + (n-2)(n-3)p^2] \quad (19)$$

$$(20)$$

Since $\mathbb{E}_n^o \left[\sum_{i=1}^n \Phi(B'_0 \cap B'_i) \right] = n(n-1)p$, we have for $n > 1$,

$$\mathbb{E}_n^o \left[\left(\sum_{i=1}^n \Phi(B'_0 \cap B'_i) - \mathbb{E}_n^o \left[\sum_{i=1}^n \Phi(B'_0 \cap B'_i) \right] \right)^2 \right] = 2n(n-1) [p - p^2 + 2(n-2)(p_2 - p^2)] \quad (21)$$

The results are true for $n = 1$ since in this case $\Phi(B'_0 \cap B'_1) = 0$. Moreover we assumed that $\sum_{i=1}^n \Phi(B'_0 \cap B'_i) = 0$ when $n = 0$ (equation 21 holds also for $n \geq 0$).

In section 6, we shall use this lemma to show that the clustering coefficient (defined later) is not appropriate for the selection of cluster head. ■

5.6 Analysis of the number of cluster-heads

We first need some technical lemmas.

Lemma 3 *The average number of cluster-heads that belong to a given Borel subset C is given by:*

$$\mathbb{E} [\text{Number of heads in a Borel subset } C] = \lambda \nu(C) \mathbb{P}_\Phi^o(0 \text{ is head}) \quad (22)$$

Lemma 4 *The probability that the origin is a cluster-head under Palm probability is given by:*

$$\mathbb{P}_\Phi^o(0 \text{ is head}) = \mathbb{P}_\Phi^o \left(\rho(0) > \max_{k=1, \dots, \Phi(B_0)} \rho(Y_k) \right) \quad (23)$$

where the sequence Y_k represents the points of Φ in B_0 , the ball centered at the origin with a radius R . We fix $\max_{k=u,\dots,v} \rho(Y_k) = 0$ if $v < u$.

We can now bound the quantity defined in Lemma 4.

Theorem 1 *An upper bound on the number of cluster-heads is given by:*

$$\mathbb{P}_\Phi^o \left(\rho(0) > \max_{k=1,\dots,\Phi(B_0)} \rho(Y_k) \right) \leq \left(1 + \sum_{n=1}^{+\infty} \frac{1}{n} \frac{(\lambda\pi R^2)^n}{n!} \right) \exp \{-\lambda\pi R^2\} \quad (24)$$

Proof 3 Let B'_0 be the ball centered in 0 with a radius of R minus the singleton 0. In the case of the node at the origin point is the only one in B_0 , it is obviously a cluster-head. Indeed, it has the highest density value among nodes in B_0 . We have,

$$\begin{aligned} \mathbb{P}^o \left(\rho(0) > \max_{k=1,\dots,\Phi(B_0)} \rho(Y_k) \right) &= \mathbb{P}^o \left(\rho(0) > \max_{k=1,\dots,\Phi(B_0)} \rho(Y_k) \middle| \Phi(B'_0) > 0 \right) \times \mathbb{P}^o \left(\Phi(B'_0) > 0 \right) \\ &\quad + \mathbb{P}^o \left(\Phi(B'_0) = 0 \right) \end{aligned} \quad (25)$$

Thus, we compute:

$$p_0 = \mathbb{P}^o \left(\rho(0) > \max_{k=1,\dots,\Phi(B_0)} \rho(Y_k) \middle| \Phi(B'_0) > 0 \right) \times \mathbb{P}^o \left(\Phi(B'_0) > 0 \right) \quad (26)$$

$\rho(Y_1)$ under Palm distribution knowing that $\Phi(B'_0) > 0$ where Y_1 is one of B'_0 's nodes, is the same that the one of the density value of the node at the origin point, knowing that $\Phi(B'_0) > 0$.

However, we have:

$$p_0 < \mathbb{P}^o \left(\rho(Y_1) > \max(\rho(0), \max_{k=2,\dots,\Phi(B_0)} \rho(Y_k)) \middle| \Phi(B'_0) > 0 \right)$$

The proof of this inequality is omitted here and will be presented in future paper.

Moreover, the event

$$\{\rho(Y_1) > \max(\rho(0), \max_{k=2, \dots, \Phi(B_0)} \rho(Y_k))\}$$

is included in the event

$$\{\rho(Y_1) > \max_{k=2, \dots, \Phi(B_0)} \rho(Y_k)\}$$

thus

$$\begin{aligned} p_0 &\leq \mathbb{P}^o \left(\rho(Y_1) > \max_{k=2, \dots, \Phi(B'_0)} \rho(Y_k) \middle| \Phi(B'_0) > 0 \right) \times \mathbb{P}^o \left(\Phi(B'_0) > 0 \right) \\ &= \sum_{n=1}^{+\infty} \mathbb{P}^o \left(\rho(Y_1) > \max_{k=2, \dots, \Phi(B'_0)} \rho(Y_k) \middle| \Phi(B'_0) = n \right) \times \mathbb{P}^o \left(\Phi(B'_0) = n \right) \\ &= \sum_{n=1}^{+\infty} \mathbb{P}^o \left(\rho(Y_1) > \max_{k=2, \dots, \Phi(B'_0)} \rho(Y_k) \middle| \Phi(B'_0) = n \right) \times \mathbb{P} \left(\Phi(B_0) = n \right) \\ &\leq \sum_{n=1}^{+\infty} \frac{1}{n} \frac{(\lambda \pi R^2)^n}{n!} \exp \{-\lambda \pi R^2\} \end{aligned}$$

The last equality is obtained thanks to the fact that the density of the points standing in B_0 is equi-distributed since the locations of these points are uniformly and independently distributed in B_0 . More precisely,

$$\mathbb{P}^o \left(\rho(Y_i) > \max_{k=1, \dots, n; k \neq i} \rho(Y_k) \middle| \Phi(B'_0) = n \right) \leq \frac{1}{n} \quad (27)$$

Moreover, the number of nodes under Palm distribution in a $\mathbb{R}^2$ Borel set which does not contain the origin point, follows a discrete Poisson law (Slivnyak's theorem [18] page 121).

We finally have:

$$\mathbb{P}^o \left(\rho(0) > \max_{k=1, \dots, \Phi(B_0)} \rho(Y_k) \right) \leq \exp \{-\lambda \pi R^2\} + \sum_{n=1}^{+\infty} \frac{1}{n} \frac{(\lambda \pi R^2)^n}{n!} \exp \{-\lambda \pi R^2\} \quad (28)$$



6 Simulation and results

We performed simulations in order to evaluate the performance of the proposed heuristic and compare it with some existing ones which have revealed to obtain pretty good results. The geometric approach used in the analysis allows to model the spatial organization of networks. As in Section 5.5, nodes are randomly deployed using a Poisson process in a 1×1 square with various levels of intensities λ (and thus various number of nodes) varying from 500 to 1000 which gives on average from 500 to 1000 nodes above our simulation square environment. Two nodes are said to have a wireless link between them if they are within communication range of each other. The communication range R is set to 0.1 in all tests. Some of the more noteworthy simulation statistics measured are : number of cluster-heads per surface unit, cluster diameter, node eccentricity in its cluster and cluster stability. These statistics provide a basis for evaluating the performance of the proposed heuristic. In each case, each statistic is the average over 1000 simulations. Note that as opposed to [1], for a given number of nodes, we fix a minimum radius such that the network is connected.

Results in Table 2 compare both theoretical analysis and simulated results of our heuristic for the average degree and node density. They match pretty well.

6.1 Clusters characteristics

Major characteristics of clusters and cluster-heads are presented in Table 3. Note that our heuristic based on the 1-density is scalable: when the number of nodes significantly increase (from 500 to

	500 nodes		600 nodes		700 nodes	
	Theory	Simulation	Theory	Simulation	Theory	Simulation
mean degree	14.7	14.3	17.8	17.3	21.0	20.2
mean 1-density	4.7	5.0	5.6	5.9	6.5	6.8
	800 nodes		900 nodes		1000 nodes	
	Theory	Simulation	Theory	Simulation	Theory	Simulation
mean degree	24.1	23.1	27.3	25.9	30.0	29.0
mean 1-density	7.5	7.1	8.4	8.6	9.3	9.4

Table 2: Average degree and density of nodes.

	500 nodes	600 nodes	700 nodes	800 nodes	900 nodes	1000nodes
# clusters	15	14.49	14.23	15.5	13.02	14
Nodes by cluster	31.2	41.4	49.2	51.5	69.1	72.7
$D(\mathcal{C})$	4.99	5.52	5.5	5.65	6.34	6.1
$\bar{e}(u/\mathcal{C})$	2.1	2.3	2.3	2.4	2.4	2.6

Table 3: Cluster characteristics for 1-density.

1000) and the node eccentricity remains the same, the number of clusters is stable. Figure 3 compares simulation results and analytic upper bound of the number of clusters for an observation area 1×1 and $R = 0.1$.

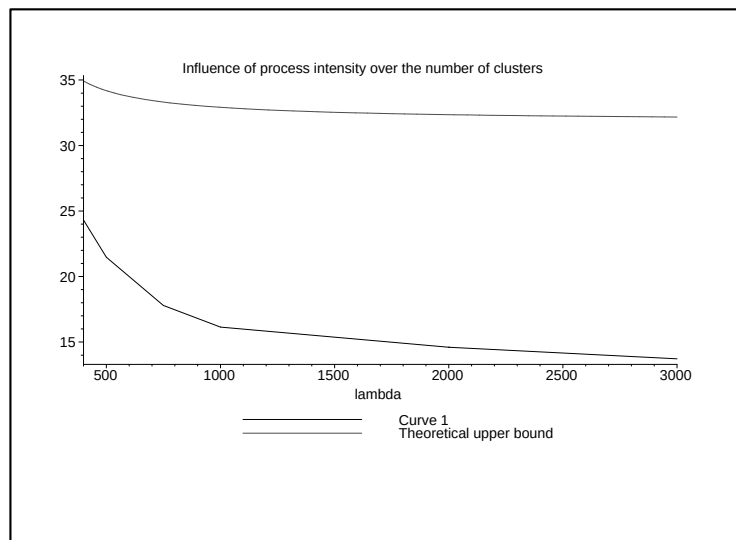


Figure 3: Number of clusters in function of Poisson process intensity.

Figures 4 and 5 show how density values are distributed over the network. We can note that in a cluster, the strongest values are situated around the cluster-head and the values decrease as the distance between a node and its cluster-head increases.

6.1.1 Clusters shape

As we can note in Figure 6, the clusters built by our metric match pretty well with the Voronoi diagram drawn from the cluster-heads whatever the process intensity. This means that once cluster-heads elected, most of nodes belong to the cluster which cluster-head is the closer to them in euclidean distance among every cluster-heads.

In order to check and quantify this characteristic, we performed simulations which give the percentage of nodes which belong to the "right Voronoi cell", *i.e.* which is closer to its cluster-

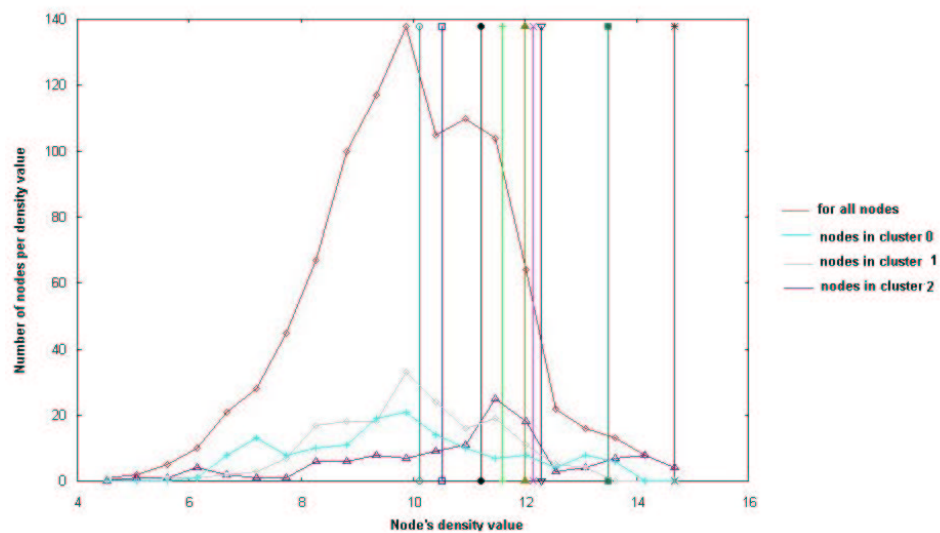


Figure 4: Density values distribution over nodes

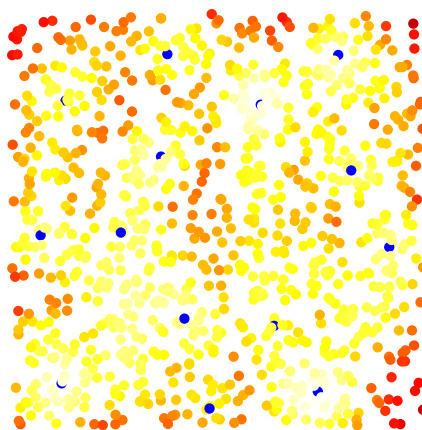


Figure 5: Density values distribution over nodes. Redder the color is, stronger the node's density value is. cluster-heads appear in blue.

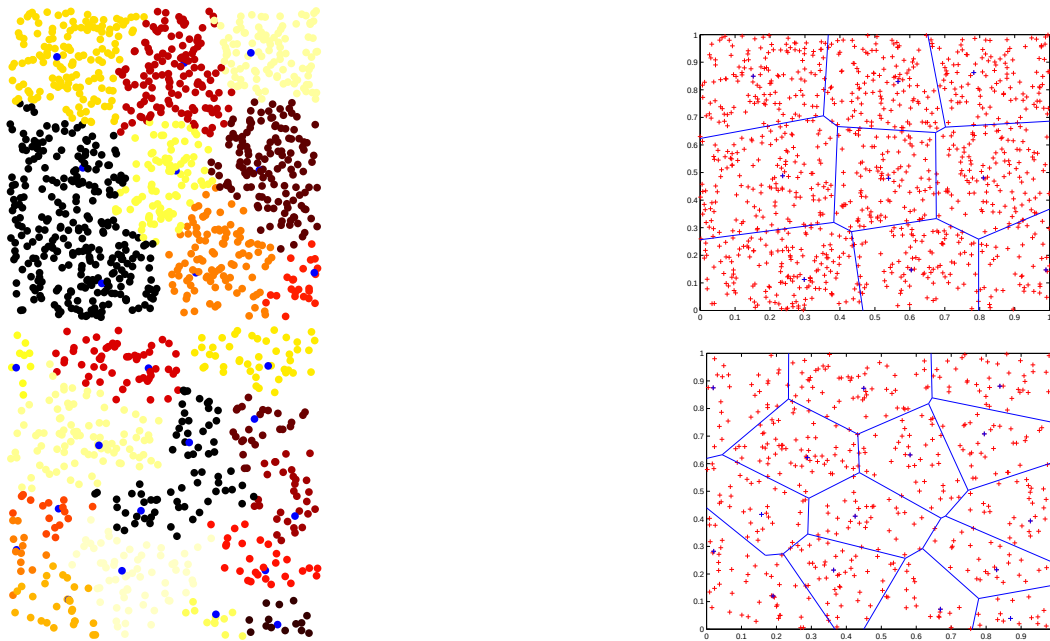


Figure 6: Density-based cluster organization (left schemes) and Voronoi diagram of cluster-heads (right schemes) for process of intensity 1000 (above) and 500 (below)

head than any other one in euclidean distance. As in ad hoc networks we mainly use the number of hops as distance value, we also performed similar simulations for the number of nodes, *i.e.* we computed the percentage of nodes closer in number of hops to their cluster-head than any other one. In Figure 7, cluster-heads appear in blue, nodes in the "right" Voronoi cell in black, other one in red. Results, presented in Table 4 and Figure 7, show that a great part of nodes lays in the Voronoi cell of their cluster-head whatever the process intensity. This characteristic will be very useful in terms of broadcast efficiency as if cluster-heads need to spread information over their own cluster, if most of nodes are closer than the one which sends the information, we save bandwidth, energy and latency.

	500nodes	600nodes	700nodes	800nodes	900nodes	1000nodes
Euclidean distance	84.17%	84.52%	84.00%	83.97%	83.82%	83.70%
Number of hops	85.43%	84.55%	84.15%	83.80%	83.75%	83.34%

Table 4: Percentage of nodes closer to their cluster-head than any other one in euclidean distance and in number of hops

6.2 Comparison with the degree heuristic

In the aim to evaluate our metric, we compare it over the organization obtained with the degree heuristic for a same distribution of nodes. Indeed, the degree heuristic is pretty close to ours and it is cheaper. Nevertheless, we haven't computed the degree heuristic exactly like in [4] as in this paper, clusters' diameter can not exceed 2 hops as a cluster-head is directly linked to every node in

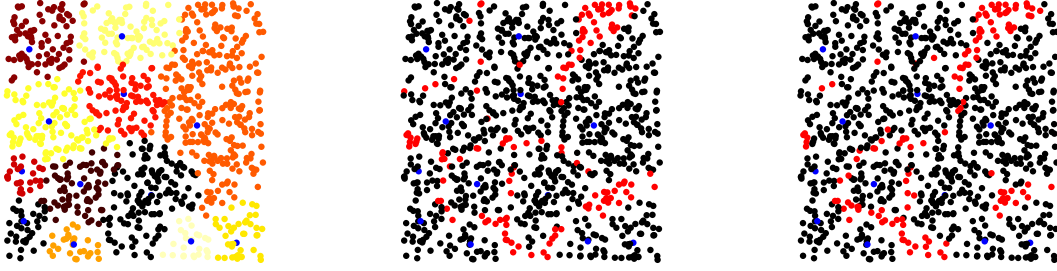


Figure 7: Nodes in the "right" Voronoi cell for a cluster organization on (a) in Euclidean distance (b) and in number of hops (c).

its cluster and it won't be comparable to our metric in such a way. Therefore, we performed the simulation using our heuristic but switching the node's density function by the degree function.

If we just have a look over the clusters organization formed by each metric, we can note that they are very similar as the Figure 8 and the Table 7 show. Therefore, we can wonder why using the density value rather than the degree value as this one is cheaper.

	# clusters	Nodes by cluster	$\tilde{D}(\mathcal{C})$	$\tilde{e}(u/\mathcal{C})$
degree	10.0	100	5.8	4.2
1-density	11.1	90.9	5.0	3.8

Table 5: Cluster characteristics for 1-density and degree heuristics.

Over nodes mobility, we expect the organization to change as less as possible, *i.e.*, that the cluster-heads remain cluster-heads as long as possible. Indeed a cluster is defined by its cluster-head, other



Figure 8: Example of clusters organization for 1000 nodes with a radius $R = 0.1$ with the degree heuristic (a) and with 1-density heuristic (b).

nodes can migrate from one cluster to another one, this will not break the cluster. Then, the most noteworthy factor is the percentage of former cluster-head re-elected after moves.

Therefore, we performed simulations in which nodes can move in a random way at a random speed from 0 to $10m/s$ (for cars) and from 0 to $1.6m/s$ (for pedestrians). We observe each 2 seconds for 15 minutes. Results presented in Table 6 show that in average, our metric reconstructs clusters less often than the degree heuristic, it's thus better since more robust towards node mobility.

	500 nodes		600 nodes		800 nodes		1000 nodes	
	Density	Degree	Density	Degree	Density	Degree	Density	Degree
from 0 to 1.6	68.7%	65%	67.2%	63.5%	64.5%	62.4%	62.2%	56.8%
from 0 to 10	30.1%	27.5%	27%	25.3%	26.2%	23.1%	24.8%	20.35%

Table 6: % of cluster-head reelection for two different speeds.

In order to understand why the cluster organization is more stable with the density rather than the degree, we analyse how degree and density of neighbor nodes are perturbed when a mobile node arrive. More formally, let's take a Poisson Point process of intensity λ distributed in $B(0, 2R)$. We assume that there is a point at the origin. We consider the degree and the density of two nodes, the node 0 at the origin and one node y among its neighbors (arbitrary chosen among all the neighbors). Then, we add a mobile node u . The added node u is a node uniformly distributed in $B(0, R)$ and is therefore supposed to be a neighbor of 0. We are then able to compute the probability that the order of the degrees of 0 and y has changed due to the presence of the mobile node. That occurs only in two cases:

- when $d(0) = d(y)$ and if the mobile node is not a neighbor of y ($d(0)$ and $d(y)$ are the degrees of 0 and y)
- when $d(0) = d(y) - 1$ and if the mobile node is not a neighbor of y .

We obtain

$$P_p = \mathbb{E} \left[\frac{\nu(C)}{\pi R^2} \sum_{n=0}^{+\infty} \frac{(\lambda \nu(C))^{2n}}{n!n!} \left(1 + \frac{\lambda \nu(C)}{n+1} \right) \right] (1 - \exp \{-\lambda \pi R^2\})$$

where $\nu(C)$ is the random variable which describes the area of $B(0, R) \setminus B(y, R)$ and where P_p is the probability that the order between $d(0)$ and $d(y)$ has changed.

For our metric, the density, we are not able to derive this probability analytically. However, we obtain an accurate approximation via simulation. In the Figure 9, we show the probabilities obtained by simulation for both metrics as well as the probability analytically computed i.e. P_p .

These probabilities match the results shown in Table 6. When a node is mobile, the probability that the order changes between points is more important with the degree as metric than the density. Thus, the degree metric is less stable than the density towards nodes mobility.

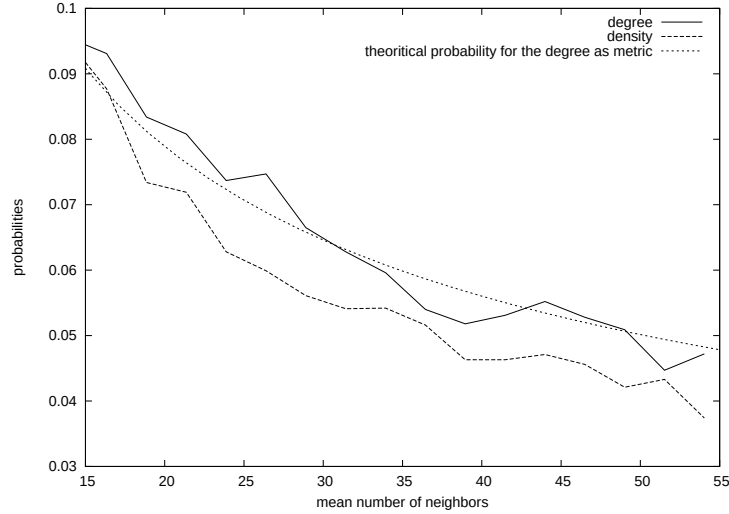


Figure 9: Probability that the order change between two neighbors for the two metrics: probabilities obtained by simulation with 10000 samples and by the analytical formula of P_p .

6.3 Comparison with the Max-Min d -cluster heuristic

We have wished to compare our heuristic with the Max-Min one because the Max-Min d -cluster presents very good results, is not based only on lowest/highest node's ID and can adapt the cluster diameter. This allows it to achieve a balance in the clusters size instead of having the clusters with the largest ID be much larger than the other.

6.3.1 Clusters organization

	# clusters	Nodes by cluster	$\tilde{D}(\mathcal{C})$	$\tilde{e}(u/\mathcal{C})$
1-density	11.1	90.9	5.0	3.8
Max-Min 2-cluster	28.6	34.9	3.6	3.1
Max-Min 3-cluster	13.3	75.2	4.9	3.4
Max-Min 4-cluster	8.2	122.0	6.5	4.9

Table 7: Clusters characteristics for 1-density and Max-Min d -clusters heuristics.

In Figure 10, we compare the amount of clusters produced by our metric and by the Max-Min d -clusters heuristic for $d = 3$ (the one the closest of ours as we have a mean cluster diameter around 6 hops) over a 1000 nodes topology for different values of R .

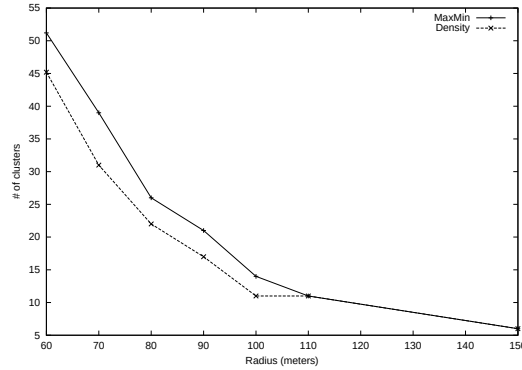


Figure 10: Number of clusters produced for 1000 nodes in function of radius R with our density metric ($- \times -$) and with the Max-Min 3-cluster metric ($- + -$).

We can see that the number of clusters computed by both metrics is similar when the radius is pretty high but that Max-Min d -Cluster computes more small clusters when the network is sparse. Thus, our metric has a better behavior towards sparse networks (less connected) as it produces less clusters and then generates less control traffic. Moreover, at the opposite of Max-Min, our heuristic does not allow clusters with only one node (the cluster-head). Nevertheless, in both cases, we can notice that the cluster-head is pretty centered in its cluster, which is useful for limiting services traffic.

Figures 11 (a) and (b) plot one example of clusters organization results obtained during a simulation, both from the same nodes distribution. We can notice that clusters are homogeneous and correspond to what we expected: cluster-heads are well distributed over the environment in a homogeneous way. Clusters gather close nodes with high connectivity in order to favor intra-cluster traffic.

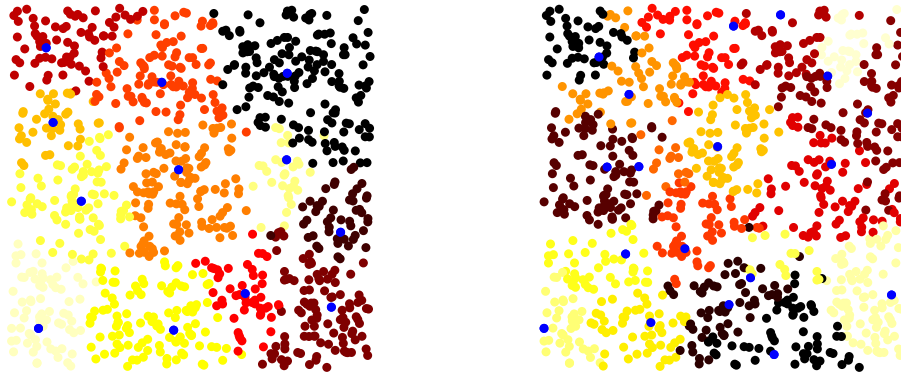


Figure 11: Example of clusters organization for 1000 nodes with a radius $R = 0.1$ with the 1-density heuristic (a) and with the Max Min 3-cluster (b).

6.3.2 Clusters organization over nodes mobility

As we did to compare our metric to the degree heuristic, we performed simulations where nodes can move in a random way at a random speed from 0 to $10m/s$ (for cars) and from 0 to $1.6m/s$ (for pedestrians). (Our metric's results are presented in Table 6.) We could notice that Max-Min d -cluster heuristic is more robust than ours in this case as cluster-heads are reelected in every case at a rate comprise between 90% and 98%. This was predictable as even when nodes move, their ID does not change unlike the density value. In this way, we can say that Max-Min d -cluster heuristic presents a better behavior as ours.

6.3.3 Clusters organization over nodes arrival

In this section, we compare both heuristics over arrival tests, that is, as opposed to classical scenarios where nodes only move, we start the scenario with an initial configuration (1000 nodes) and nodes arrives randomly in the network by groups of 100 with randomly distributed IDs. Results are presented in Table 8).

	% cluster-head reelection	# clusters first step	# clusters last step
1-density	94.3%	14	14
Max Min 3-cluster	100%	15	22

Table 8: Comparison Max Min 3-cluster and density heuristics over massive nodes arrival

We can note that the density heuristic presents a pretty good behavior as it has a good reelection percentage and that the number of clusters do not explode. New nodes are absorbed in existing

clusters. Max Min heuristic has an excellent reelection percentage but we can note that the number of clusters increase. This is due to the fact that if new nodes with a higher ID than the former ones arrive, it will create its own cluster but the existing clusters remain all the same. It may be the only one node in its cluster.

Indeed, the cluster-head election is based on a purely static data, the ID of a node, if one node vanishes or appears, it is enough to trigger a new election which it is not the case in our heuristic since the density measure is able to “absorb” local modification.

6.3.4 Non uniform distribution and arrival

The last test that we present is when nodes are not uniformly distributed but rather concentrated around few points, for example cities. Figures 12 (a) and (b) illustrate such scenarii. As we can see our heuristic generates less clusters and cluster-heads are much more centered inside their cluster. The Max Min heuristic generates sometimes several useless neighbor cluster-heads. For 1000 nodes, on average, our heuristics generates 8.7 clusters whereas the Max Min heuristics generates 15.25 clusters.

6.3.5 Complexity

Max Min d -cluster is composed of 3 phases that flood messages up to d hops in order to converge: one to compute the max, one to compute the min, and one to announce the winner. Our heuristic is purely local and we can implement it by doing piggy packing in Hello packets. So Max Min d -cluster algorithm is more costly in term of messages overhead and latency.



Figure 12: Non uniform distribution of nodes: Example of clusters organization for 1000 nodes with a radius $R = 0.1km$ with the 1-density heuristic (a) and with the Max Min 3-cluster (b).

6.4 Comparison with the clustering coefficient

As we model our ad hoc network by a graph, we wondered why not use a value used in graphs theory, the clustering coefficient. This coefficient noted $c(x)$ for a point at x is the ratio between the number of links between two node's 1-neighbors and the maximum number of links which can exist between them. For instance, the clustering coefficient for the point at the origin when it has n neighbors ($n > 0$) is

$$c(0) = \frac{1}{n(n-1)} \sum_{i=1}^n \Phi(B'_0 \cap B'_i).$$

It is interesting to observe the behavior of the variance of the clustering coefficient for the point located at the origin and for a given number of neighbors $n > 0$.

From the lemma 2, we have:

$$\begin{aligned}
\mathbb{E}_n^o \left[(c(0) - \mathbb{E}_n^o [c(0)])^2 \right] &= \frac{1}{n^2(n-1)^2} \mathbb{E}_n^o \left[\left(\sum_{i=1}^n \Phi(B'_0 \cap B'_i) - \mathbb{E}_n^o \left[\sum_{i=1}^n \Phi(B'_0 \cap B'_i) \right] \right)^2 \right] \\
&= \frac{2}{n(n-1)} [p - p^2 + 2(n-2)(p_2 - p^2)]
\end{aligned} \tag{30}$$

(31)

We can see that the variance of the clustering coefficient converges to zero as n goes to infinity. We may show with the Bienayme-Chebyshev inequality that this coefficient converges in probability to its mean p (as defined in the proof of the lemma 1). It is easy to show, that when the process intensity increases, the clustering coefficient converges in probability to the constant p . Therefore, when the process intensity is high, the clustering coefficients of the process points are quantitatively very close to p and do not traduce the degree, density or the region of concentration of the process. None node will emerge as cluster-head without competition as every node has the same value. We note that the result is different for the density since its mean and variance tend to infinity when n goes to infinity. In this way, even if the intensity of nodes increases, a single node will be able to free itself from the other ones.

Table 9 and Figure 13 present the simulation results when the clustering coefficient is used. We can see that this metric is absolutely not appropriate for our needs. Indeed, it constructs more clusters than necessary and cluster-heads are very eccentered in their cluster. In Table 9, the "% of cluster-heads median centres" represents the percentage of cluster-heads which are the node in their cluster which minimize the distance with every other node in its cluster. The higher this amount is, the better is.

	# clusters	% of cluster-heads median centres	$\tilde{D}(\mathcal{C})$	$\tilde{e}(u/\mathcal{C})$
1-density	11.1	81.2%	5.0	3.8
clustering coefficient	22.4	37%	3.7	2.9

Table 9: Cluster characteristics for 1-density and clustering coefficient.

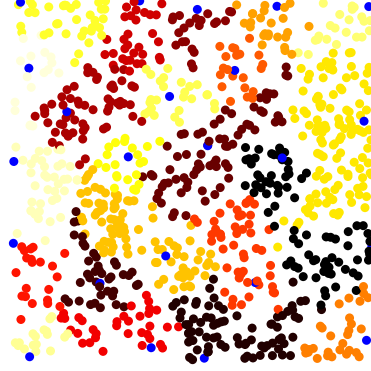


Figure 13: Clusters organization obtained with the clustering coefficient

7 Conclusion and perspectives

We have proposed a distributed algorithm for organizing ad hoc (or sensor) nodes into a flexible hierarchy of clusters with a strong objective of not using fixed and non adaptive criteria. Thanks to stochastic geometry and Palm distribution theory, we have performed formal analysis and we were able to compute the average density of nodes but also we can bound the number of clusters in a given area if nodes are randomly distributed. We have shown by simulation and analytic analysis that our metric based on the density gathers the dynamics of node neighborhood and outperforms classical static criteria used in past solutions (e.g., max degree).

In future, we intend to test deeper our metric, its behavior over different environments and mobility models and also its behavior from the point of view of the data link layer and collisions. We are currently investigating the use of purely distributed hash functions in order to solve the node localization problem once the clusterization is done. Once again, we should be able to apply stochastic geometry in order to derive formal bound on the number of hops (and not the euclidean distance) between nodes and their cluster head.

References

- [1] A. Amis, R. Prakash, T. Vuong, and D. Huynh. Max-min d-cluster formation in wireless ad hoc networks. In *Proceedings of the IEEE INFOCOM*, march 2000.
- [2] D. Baker and A. Ephremides. The architectural organization of a mobile radio network via a distributed algorithm. *IEEE Transactions on Communications*, 29(11):1694–1701, 1981.
- [3] P. Basu, N. Khan, and T. Little. A mobility based metric for clustering in mobile ad hoc networks. In *Proceedings of Distributed Computing Systems Workshop 2001*, 2001.
- [4] G. Chen and I. Stojmenovic. Clustering and routing in mobile wireless networks. Technical Report TR-99-05, SITE, June 1999.
- [5] T. Clausen and P. Jacquet. Optimized link state routing protocol (olsr). RFC 3626, October 2003.
- [6] B. Das and V. Bharghavan. Routing in ad-hoc networks using minimum connected dominating sets. In *ICC*, 1997.

-
- [7] Y. Fernandess and D. Malkhi. K-clustering in wireless ad hoc networks. In *Proceedings of the second ACM international workshop on Principles of mobile computing*, 2002.
 - [8] M. Gerla and J. Tzu-Chieh Tsai. Multicluster, mobile, multimedia radio network. *Baltzer Journals*, July 1995.
 - [9] M. Hayashi and C. Bonnet. A new method for scalable and reliable multicast system for mobile networks. In *1st Networking - ICN 2001 : First International Conference*, Colmar, France, July, 9-13th 2001.
 - [10] M. Jiang and Y. Tay. Cluster based routing protocol(CBRP). DRAFT draft-ietf-manet-cbrp-spec-01.txt, IETF, July 1999.
 - [11] L. Kleinrock and J. Silvester. Spatial reuse in multihop packet radio networks. In *Proceedings of the IEEE*, volume 75, pages 156–167. IEEE, 1987.
 - [12] P. Krishna, N. H. Vaidya, M. Chatterjee, and D. K. Pradhan. A cluster based approach for routing in dynamic networks. In *ACM SIGCOMM*, pages 49–65. ACM, ACM, April 1997.
 - [13] H.-C. Lin and Y.-H. Chu. A clustering technique for large multihop mobile wireless networks. In *Proceedings of the IEEE Vehicular technology conference*, Tokyo, Japan, may 2000. IEEE.
 - [14] N. Mitton and E. Fleury. Self - organization in ad hoc networks. Research Report RR-5042, INRIA, 2004.

- [15] N. Nikaein, H. Labiod, and C. Bonnet. Ddr-distributed dynamic routing algorithm for mobile ad hoc networks. In *1st ACM international symposium on mobile ad hoc routing and computing*, Boston, MA, USA, November, 20th 2000. ACM.
- [16] M. Pearlman, Z. Haas, and S. Mir. Using routing zones to support route maintenance in ad hoc networks. In *Wireless Communications and Networking Conference (WCNC 2000)*, pages 1280–1285. IEEE, September 2000.
- [17] R. Ramanathan and M. Steenstrup. Hierarchically-organized, multihop mobile wireless networks for quality-of-service support. *Mobile networks and applications*, 3:101–119, June 1998.
- [18] D. Stoyan, S. Kendall, and J. Mecke. *Stochastic geometry and its applications, second edition*. John Wiley & Sons, 1995.



Unité de recherche INRIA Rhône-Alpes

655, avenue de l'Europe - 38334 Montbonnot Saint-Ismier (France)

Unité de recherche INRIA Futurs : Parc Club Orsay Université - ZAC des Vignes

4, rue Jacques Monod - 91893 ORSAY Cedex (France)

Unité de recherche INRIA Lorraine : LORIA, Technopôle de Nancy-Brabois - Campus scientifique

615, rue du Jardin Botanique - BP 101 - 54602 Villers-lès-Nancy Cedex (France)

Unité de recherche INRIA Rennes : IRISA, Campus universitaire de Beaulieu - 35042 Rennes Cedex (France)

Unité de recherche INRIA Rocquencourt : Domaine de Voluceau - Rocquencourt - BP 105 - 78153 Le Chesnay Cedex (France)

Unité de recherche INRIA Sophia Antipolis : 2004, route des Lucioles - BP 93 - 06902 Sophia Antipolis Cedex (France)

Éditeur

INRIA - Domaine de Voluceau - Rocquencourt, BP 105 - 78153 Le Chesnay Cedex (France)

<http://www.inria.fr>

ISSN 0249-6399