



Message Drop and Scheduling in DTNs: Theory and Practice

Amir Krifa, Chadi Barakat, Thrasyvoulos Spyropoulos

► To cite this version:

Amir Krifa, Chadi Barakat, Thrasyvoulos Spyropoulos. Message Drop and Scheduling in DTNs: Theory and Practice. IEEE Transactions on Mobile Computing, 2012, 11 (9), pp.14. 10.1109/TMC.2011.163 . inria-00542309

HAL Id: inria-00542309

<https://inria.hal.science/inria-00542309>

Submitted on 2 Dec 2010

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

Message Drop and Scheduling in DTNs: Theory and Practice

Amir Krifa, Chadi Barakat, *Senior Member, IEEE*, and Thrasyvoulos Spyropoulos, *Member, IEEE*,

Abstract—In order to achieve data delivery in Delay Tolerant Networks (DTN), researchers have proposed the use of *store-carry-and-forward* protocols: a node there may store a message in its buffer and carry it along for long periods of time, until an appropriate forwarding opportunity arises. This way, messages can traverse disconnected parts of the network. Multiple message replicas are often propagated to further increase delivery probability. This combination of long-term storage and message replication imposes a high storage and bandwidth overhead. Thus, efficient scheduling and drop policies are necessary to: (i) decide on the order by which messages should be replicated when contact durations are limited, and (ii) which messages should be discarded when node buffers operate close to their capacity.

In this paper, we propose a practical and efficient joint scheduling and drop policy that can optimize different performance metrics, such as average delay and delivery probability. We first use the theory of encounter-based message dissemination to derive the optimal policy based on global knowledge about the network. Then, we introduce a method that estimates all necessary parameters using locally collected statistics. Based on this, we derive a distributed scheduling and drop policy that can approximate the performance of the optimal policy in practice. Using simulations based on synthetic and real mobility traces, we show that our optimal policy and its distributed variant outperform existing resource allocation schemes for DTNs. Finally, we study how sampled statistics can reduce the *signaling* overhead of our algorithm and examine its behavior under different congestion regimes. Our results suggest that close to optimal performance can be achieved even when nodes sample a small percentage of the available statistics.

Index Terms—Delay Tolerant Network, Congestion, Drop Policy, Scheduling Policy



1 INTRODUCTION

MOBILE ad hoc networks (MANETs) had been treated, until recently, as a connected graph over which end-to-end paths need to be established. This legacy view might no longer be appropriate for modelling existing and emerging wireless networks [1], [2], [3], [4]. Wireless propagation phenomena, node mobility, power management, etc. often result in intermittent connectivity with end-to-end paths either lacking or rapidly changing. To allow some services to operate even under these challenging conditions, researchers have proposed a new networking paradigm, often referred to as Delay Tolerant Networking (DTN [5]), based on the *store-carry-and-forward* routing principle [1]. Nodes there, rather than dropping a session (and respective packets) when no forwarding opportunity is available, store and carry messages until new communication opportunities arise.

Despite a large amount of effort invested in the design of efficient routing algorithms for DTNs, there has not been a similar focus on queue management and message scheduling. Yet, the combination of long-term storage and the, often expensive, message replication performed by many DTN routing protocols [6], [7] impose a high bandwidth and storage overhead on wireless nodes [8].

Moreover, the data units disseminated in this context, called *bundles*, are self-contained, application-level data units, which can often be large [5]. As a result, it is expected that nodes' buffers, in this context, will often operate at full capacity. Similarly, the available bandwidth during a contact could be insufficient to communicate all intended messages. Consequently, *regardless of the specific routing algorithm used*, it is important to have: (i) efficient drop policies to decide which message(s) should be discarded when a node's buffer is full, and (ii) efficient scheduling policies to decide which message(s) should be chosen to exchange with another encountered node when bandwidth is limited.

In this paper, we try to solve this problem in its foundation. We develop a theoretical framework based on Epidemic message dissemination [9], [10], [11], and propose an optimal joint scheduling and drop policy, GBSD (Global knowledge Based Scheduling and Drop) that can maximize the average delivery rate or minimize the average delivery delay. GBSD derives a per-message utility by taking into account all information that are relevant for message delivery, and manages messages accordingly. Yet, to derive these utilities, it requires *global* network information, making its implementation difficult in practice, especially given the intermittently connected nature of the targeted networks. In order to amend this, we propose a second policy, HBSD (History Based Scheduling and Drop), a distributed (*local*) algorithm based on statistical learning. HBSD uses network history to estimate the current state of required (*global*) network parameters and uses these estimates, rather

- Amir Krifa and Chadi Barakat are with the Project-Team Planète, INRIA Sophia-Antipolis, France.
E-mail(s): Amir.Krifa@sophia.inria.fr, Chadi.Barakat@sophia.inria.fr
- Thrasyvoulos Spyropoulos is with the Swiss Federal Institute of Technology (ETH), Zurich, Switzerland.
E-mail: spyropoulos@tik.ee.ethz.ch

than actual values (as in GBSD), to calculate message utilities for each performance target metric.

To our best knowledge, the recently proposed RAPID protocol [12] is the only effort aiming at scheduling (and to a lesser extent message drop) using a similar theoretical framework. Yet, the utilities derived there are sub-optimal, as we will explain later, and require global knowledge (as in GBSD), raising the same implementation concerns. Simulations using both synthetic mobility models and real traces show that our HBSD policy not only outperforms existing buffer management and scheduling policies (including RAPID), but can also approximate the performance of the reference GBSD policy, in all scenarios considered.

Furthermore, we look deeper into our distributed statistics collection solution and attempt to identify the available tradeoffs between the collection overhead and the resulting performance. Aggressively collecting statistics and exchanging them with every node encountered allows estimates to converge faster (and thus achieves good performance), but it can potentially result in high energy and bandwidth consumption, and also interfere with data transmissions. Our results suggest that close to optimal performance can still be achieved even when the signaling overhead is forced (through sampling) to take only a small percentage of the contact bandwidth.

Finally, we examine how our algorithm behaves under different congestion regimes. Interestingly, we find that (i) at low to moderately congested regimes, the optimal policy is simply equivalent to dropping the message with the oldest age (similarly to the findings of [13]), while (ii) at highly congested regimes, the optimal policy is not linear on message age; some young messages have to be dropped, as a means of indirect admission control, to allow older messages to create enough replicas and have a chance to be delivered. Hence, our framework can also explain what popular heuristic policies are doing, in this context, relative to the optimal one.

The rest of this paper is organized as follows. Section 2 describes the current state-of-the art in terms of buffer management and scheduling in DTNs. In Section 3, we describe the “reference”, optimal joint scheduling and drop policy that uses global knowledge about the network. Then, we present in Section 4 a learning process that enables us to approximate the global network state required by the reference policy. Section 5 discusses our evaluation setup and presents performance results for both policies (GBSD and HBSD) using synthetic and real mobility traces. In Section 6, we examine in detail the mechanism to collect and maintain network history statistics, and evaluate the signaling-performance trade-off. Section 7 studies the behavior of our HBSD policy in different congestion regimes. Finally, we conclude this paper and discuss future work in Section 8.

2 STATE OF THE ART

A number of sophisticated solutions have been proposed to handle routing in DTNs. Yet, the impact of buffer man-

agement and scheduling policies on the performance of the system has been largely disregarded, in comparison, by the DTN community.

In [14], Zhang et al. present an analysis of buffer constrained *Epidemic* routing, and evaluate some simple drop policies like drop-front and drop-tail. The authors conclude that drop-front, and a variant of it giving priority to source messages, outperform drop-tail in the DTN context. A somewhat more extensive set of combinations of *heuristic* buffer management policies and routing protocols for DTNs is evaluated in [13], confirming the performance of drop-front. In [15], Dohyung et al. present a drop policy which discards a message with the largest expected number of copies first to minimize the impact of message drop. However, all these policies are heuristic, i.e. not explicitly designed for optimality in the DTN context. Also, these works do not address scheduling. In a different work [16], we address the problem of optimal drop policy only (i.e. no bandwidth or scheduling concerns) using a similar analytical framework, and have compared it extensively against the policies described in [14] and [13]. Due to space limitations, we do not repeat these results here. We rather focus on the more general *joint scheduling and drop problem*, for which we believe the RAPID protocol [12] represents the state-of-the-art.

RAPID is the first protocol to explicitly assume both bandwidth and (to a lesser extent) buffer constraints exist, and to handle the DTN routing problem as an optimal resource allocation problem, given some assumption regarding node mobility. As such, it is the most related to our proposal, and we will compare directly against it. Despite the elegance of the approach, and performance benefits demonstrated compared to well-known routing protocols, RAPID suffers from the following drawbacks: (i) its policy is based on suboptimal message utilities (more on this in Section 3); (ii) in order to derive these utilities, RAPID requires the flooding of information about all the replicas of a given message in the queues of all nodes in the network; yet, the information propagated across the network might arrive stale to nodes (a problem that the authors also note) due to change in the number of replicas, change in the number of messages and nodes, or if the message is delivered but acknowledgements have not yet propagated in the network; and (iii) RAPID does not address the issue of signalling overhead. Indeed, in [12], the authors showed that whenever the congested level of the network starts increasing, their meta-data channel consumes more bandwidth. This is rather undesirable, as meta-data exchange can start interfering with data transmissions amplifying the effects of congestion. In another work [17], Yong et al. present a buffer management schema similar to RAPID. However they do not address the scheduling issue nor the trade-off between the control channel overhead and system performance. In this paper, we successfully address all these three issues.

3 OPTIMAL JOINT SCHEDULING AND DROP POLICY

In this section, we first describe our problem setting and the assumptions for our theoretical framework. We then use this framework to identify the optimal policy, GBSD (Global Knowledge based Scheduling and Drop). This policy uses global knowledge about the state of each message in the network (number of replicas). Hence, it is difficult to implement it in a real world scenario, and will only serve as reference. In the next section, we will propose a distributed algorithm that can successfully approximate the performance of the optimal policy.

3.1 Assumptions and Problem Description

We assume there are L total nodes in the network. Each of these nodes has a buffer, in which it can store up to B messages in transit, either messages belonging to other nodes or messages generated by itself. Each message has a Time-To-Live (TTL) value, after which the message is no more useful to the application and should be dropped by its source and all intermediate nodes. The message can also be dropped when a notification of delivery is received, or if an “anti-packet” mechanism is implemented [14]¹.

Routing: Each message has a single destination (unicast) and is assumed to be routed using a replication-based (multiple copy) scheme [8]. Such schemes include epidemic routing [6], gossiping [8], spray and wait [19], coding-base schemes [20], [21], etc.² During a contact, the routing scheme used will create a list of messages to be replicated among the ones currently in the buffer. Thus, different routing schemes might choose different messages. For example, epidemic routing will replicate all messages not already present in the encountered node’s buffer [6].

For the purposes of this paper, we will use epidemic routing as a case study, for the following reasons. First, its simplicity allows us to concentrate on the problem of resource allocation, which is the focus of this paper. Second, it consumes the most resources per message compared to any other scheme. As a result, it can be easily driven to medium or high congestion regimes, where the efficient resource allocation problem is most critical (this will also be more clear in Section 7). Third, given the nature of random forwarding schemes, unless a buffer is found full or contact capacity is not enough to transfer all messages, epidemic forwarding is optimal

in terms of delay and delivery probability. Consequently, epidemic routing along with appropriate scheduling and message drop policies, can be viewed as a new routing scheme that optimally adapts to available resources [12].

Finally, we note that our framework could be used to treat other types of traffic (e.g. multicast), as well. However, we focus on unicast traffic to elucidate the basic ideas behind our approach, and defer the treatment of multi-point traffic to future work.

Mobility Model: Another important element in our analytical framework is the impact of mobility. In the DTN context, message transmissions occur only when nodes encounter each other. Thus, *the time elapsed between node meetings is the basic delay component*. The meeting time distribution is a basic property of the mobility model assumed [11], [10]³. To formulate the optimal policy problem, we will not assume any specific mobility model used, but rather consider a generic model that has the following two properties:

- 1) meeting times are exponentially distributed or have at least an *exponential tail*;
- 2) nodes move *independently* of each other.

Regarding the former requirement, it has been shown that many popular mobility models like Random Walk, Random Waypoint and Random Direction [11], [10] have such a property. It is also known in the theory of random walks on graphs that hitting times on subsets of vertices usually have an exponential tail [24]. Moreover, it has recently been argued that meeting and inter-meeting times observed in many traces may also exhibit an exponential tail [25]. Regarding the independence assumption, although it might not always hold in some scenarios, it turns out to be a useful approximation (as will be seen in the results section). In fact, one could use a mean-field analysis argument to show that independence is not required, in the limit of large number of nodes, for the analytical formulas derived to hold (see e.g. [26]).

Buffer Management and Scheduling: Let us consider a time instant when a new contact occurs between nodes i and j . The following resource allocation problem arises when nodes are confronted with limited resources (i.e. contact bandwidth and buffer space)⁴.

(Scheduling Problem) If i has X messages in its *local* buffer that it should forward to j (chosen by the routing algorithm), but does not know if the contact will last long enough to forward all messages, which ones should it send first, so as to maximize the *global* delivery probability for *all* messages currently in the network?

1. Once a node delivers a packet to the destination, it should delete the packet from its buffer to save storage space and prevent the node from infecting other nodes. Moreover, to avoid being re-infected by the packet, a node can keep track of packet delivery. We refer to this information stored at the node as “anti-packet”; various algorithms have been proposed to propagate anti-packets to other infected and susceptible nodes [14], [18].

2. We do not treat here utility-based forwarding schemes (e.g. [7], [22], [8]). Their forwarding decisions may partially conflict with our utility-based buffer management policy, as some of the work is repeated during both routing and buffer management (sub-optimally).

3. By *meeting time* we refer to the time until two nodes starting from the stationary distribution come within range (“first meeting-time”); If some of the nodes in the network are static, then one needs to use *hitting times* between mobile and static nodes. Our theory can be easily modified to account for static nodes by considering, for example, two classes of nodes with different meeting rates (see e.g. [23]).

4. We note that, by “limited resources”, we do not imply that our focus is only small, resource-limited nodes (e.g. wireless sensors), but rather that the offered forwarding or storage load exceeds the available capacity. In other words, we are interested in congestion regimes.

GBSD Optimization Scheme (Input F, Input K, Output U_i)
Inputs:

- **F:** The expectation of the global metric to optimize.
- **K:** The total number of different messages within the network.

Output: U_i: The utility of message i.

- **Objective:** At each contact, maximize the gain $\Delta(F)$ on the global metric F

$$F = \sum_{i=0}^{K(t)} f_i \quad \text{Where } f_i \text{ is the per-message metric value (ex: message delivery probability or delivery delay)}$$

• In the case of congestion (limited transfer opportunity) a DTN node should take the drop (replication) decision leading to the best gain for F. To find this decision, we differentiate F with respect to n_i (number of copies of the message i), and then discretize and replace dn_i by $\Delta(n_i)$ to obtain:

$$\Delta(F) = \sum_{i=0}^{K(t)} \frac{\partial f_i}{\partial n_i} \Delta(n_i) = \sum_{i=0}^{K(t)} U_i \Delta(n_i) \quad \text{and} \quad U_i = \frac{\partial f_i}{\partial n_i}$$

$$\begin{cases} \Delta(n_i) = -1 & \text{If we drop an already existing message } i \text{ from the buffer} \\ \Delta(n_i) = 0 & \text{If no action for message } i \text{ is taken} \\ \Delta(n_i) = +1 & \text{If we replicate message } i \text{ during a contact (scheduling), or if we store a newly received message } i \text{ (buffer management)} \end{cases}$$

- Based on this, a DTN node should (i) replicate messages in the order of decreasing U_i (ii) drop the message with the lowest U_i when buffer is full

Fig. 1. GBSD Global optimization policy

(Buffer Management Problem) If one (or more) of these messages arrive at j 's buffer and find it full, what is the best message j should drop among the ones already in its buffer (locally) and the newly arrived one, in order to maximize the average delivery rate among all messages in the network (globally)?

To address these two questions, we propose the following policy. Given a routing metric to optimize (e.g. delivery delay or delivery probability), our policy, GBSD, derives a per-message utility that captures the marginal value of a given message copy, with respect to the chosen optimization metric. Based on this utility, two main functions are performed:

- 1) *Scheduling*: at each contact, a node should replicate messages in decreasing order of their utilities.
- 2) *Drop*: when a new message arrives at a node with a full buffer, this node should drop the message with the smallest utility among the one just received and all buffered messages for which the current node is not the source.

We will derive next such a per-message utility for two popular metrics: maximizing the average delivery probability (rate), and minimizing the average delivery delay. Table 1 contains some useful notation that we will use throughout the paper. Finally, the GBSD optimization policy is summarized in Figure 1.

3.2 Maximizing the average delivery rate

We first look into a scenario where each message has a finite *TTL* value. The source of the message keeps a copy of it during the whole *TTL* duration, while intermediate nodes are not obliged to do so. To maximize the average

TABLE 1
Notation

Variable	Description
L	Number of nodes in the network
K(t)	Number of distinct messages in the network at time t
TTL _i	Initial Time To Live for message i
R _i	Remaining Time To Live for message i
T _i = TTL _i - R _i	Elapsed Time for message i. It measures the time since this message was generated by its source
n _i (T _i)	Number of copies of message i in the network after elapsed time T _i
m _i (T _i)	Number of nodes (excluding source) that have seen message i since its creation until elapsed time T _i
λ	Meeting rate between two nodes; λ = $\frac{1}{E[H]}$ where E[H] is the average meeting time

delivery probability among all messages in the network the optimal policy must use the per message utility derived in the following theorem, in order to perform scheduling and buffer management.

Theorem 3.1. Let us assume that there are K messages in the network, with elapsed time T_i for the message i. For each message i ∈ [1, K], let n_i(T_i) be the number of nodes who have a copy of the message at this time instant, and m_i(T_i) those that have “seen” the message (excluding the source) since its creation⁵ (n_i(T_i) ≤ m_i(T_i) + 1). To maximize the average delivery rate of all messages, a DTN node should apply the GBSD policy using the following utility per message i:

$$U_i(DR) = (1 - \frac{m_i(T_i)}{L-1}) \lambda R_i \exp(-\lambda n_i(T_i) R_i). \quad (1)$$

Proof: Let the meeting time between nodes be exponentially distributed with parameter λ. The probability that a copy of a message i will not be delivered by a node is then given by the probability that the next meeting time with the destination is greater than R_i, the remaining lifetime of a message (R_i = TTL - T_i). This is equal to exp(-λR_i).

Knowing that message i has n_i(T_i) copies in the network, and assuming that the message has not yet been delivered, we can derive the probability that the message itself will not be delivered (i.e. none of the n_i copies gets delivered):

$$P\{\text{message } i \text{ not delivered} \mid \text{not delivered yet}\} = \prod_{i=1}^{n_i(T_i)} \exp(-\lambda R_i) = \exp(-\lambda n_i(T_i) R_i). \quad (2)$$

Here, we have not taken into account that more copies of a given message i may be created in the future through new node encounters. We have also not taken into account that a copy of message i could be

5. We say that a node A has “seen” a message i, when A had received a copy of message i sometime in the past, regardless of whether it still has the copy or has already removed it from the buffer.

dropped within R_i (and thus this policy is to some extent “greedy” or “locally optimal”, with respect to the time dimension). Predicting future encounters and the effect of further replicas created complicates the problem significantly. Nevertheless, the same assumptions are applied for all messages equally and thus can justify the relative comparison between the delivery probabilities for different messages.

We need to also take into consideration what has happened in the network since the message generation, in the absence of an explicit delivery notification (this part is not considered in RAPID [12], making the utility function derived there suboptimal). Given that all nodes including the destination have the same chance to see the message, the probability that a message i has been already delivered is equal to:

$$P\{\text{message } i \text{ already delivered}\} = m_i(T_i)/(L-1). \quad (3)$$

Combining Eq.(2) and Eq.(3) the probability that a message i will get delivered before its TTL expires is:

$$\begin{aligned} P_i &= P\{\text{message } i \text{ not delivered yet}\} * (1 - \exp(-\lambda n_i(T_i)R_i)) \\ &\quad + P\{\text{message } i \text{ already delivered}\} \\ &= (1 - \frac{m_i(T_i)}{L-1}) * (1 - \exp(-\lambda n_i(T_i)R_i)) + \frac{m_i(T_i)}{L-1}. \end{aligned}$$

So, if we take at instant t a snapshot of the network, the global delivery rate for the whole network will be:

$$DR = \sum_{i=1}^{K(t)} \left[\left(1 - \frac{m_i(T_i)}{L-1}\right) * (1 - \exp(-\lambda n_i(T_i)R_i)) + \frac{m_i(T_i)}{L-1} \right]$$

In case of a full buffer or limited transfer opportunity, a DTN node should take respectively a drop or replication decision that leads to the best gain in the global delivery rate DR. To define this optimal decision, we differentiate DR with respect to $n_i(T_i)$, then we discretize and replace dn by Δn to obtain:

$$\begin{aligned} \Delta(DR) &= \sum_{i=1}^{K(t)} \frac{\partial P_i}{\partial n_i(T_i)} * \Delta n_i(T_i) \\ &= \sum_{i=1}^{K(t)} \left[\left(1 - \frac{m_i(T_i)}{L-1}\right) \lambda R_i \exp(-\lambda n_i(T_i)R_i) * \Delta n_i(T_i) \right] \end{aligned}$$

Our aim is to maximize $\Delta(DR)$. In the case of message drop, for example, we know that: $\Delta n_i(T_i) = -1$ if we drop an already existing message i from the buffer, $\Delta n_i(T_i) = 0$ if we don't drop an already existing message i from the buffer, and $\Delta n_i(T_i) = +1$ if we keep and store the newly-received message i . Based on this, GBSD ranks messages using the per message utility in Eq.(1), then schedules and drops them accordingly. This utility can be viewed as the *marginal utility value* for a copy of a message i with respect to the total delivery rate. The value of this utility is a function of the global state of the message (n_i and m_i) in the network. $\square$

3.3 Minimizing the average delivery delay

We next turn our attention to minimizing the average delivery delay. We now assume that all messages generated have infinite TTL or at least a TTL value large enough to ensure a delivery probability close to 1. The following Theorem derives the optimal per-message utility, for the same setting and assumptions as Theorem 3.1.

Theorem 3.2. *To minimize the average delivery delay of all messages, a DTN node should apply the GBSD policy using the following utility for each message i :*

$$U_i(DD) = \frac{1}{n_i(T_i)^2 \lambda} \left(1 - \frac{m_i(T_i)}{L-1}\right). \quad (4)$$

Proof: Let us denote the delivery delay for message i with random variable X_i . This delay is set to 0 (or any other constant value) if the message has been already delivered. Then, the total expected delivery delay (D) for all messages for which copies still exist in the network is given by,

$$D = \sum_{i=1}^{K(t)} \left[\frac{m_i(T_i)}{L-1} * 0 + \left(1 - \frac{m_i(T_i)}{L-1}\right) * E[X_i | X_i > T_i] \right]. \quad (5)$$

We know that the time until the first copy of the message i reaches the destination follows an exponential distribution with mean $1/(n_i(T_i)\lambda)$. It follows that,

$$E[X_i | X_i > T_i] = T_i + \frac{1}{n_i(T_i)\lambda}. \quad (6)$$

Substituting Eq.(6) in Eq.(5), we get,

$$D = \sum_{i=1}^{K(t)} \left(1 - \frac{m_i(T_i)}{L-1}\right) \left(T_i + \frac{1}{n_i(T_i)\lambda}\right).$$

Now, we differentiate D with respect to $n_i(T_i)$ to find the policy that maximizes the improvement in D ,

$$\Delta(D) = \sum_{i=1}^{K(t)} \frac{1}{n_i(T_i)^2 \lambda} \left(\frac{m_i(T_i)}{L-1} - 1\right) * \Delta n_i(T_i).$$

The best drop or forwarding decision will be the one that maximizes $|\Delta(D)|$ (or $-\Delta(D)$). This results to the per message utility of Eq.(4). $\square$

Note that, the per message utility with respect to delivery delay is different than the one for the delivery rate. This implies (naturally) that both metrics cannot be optimized concurrently.

4 USING NETWORK HISTORY TO APPROXIMATE GLOBAL KNOWLEDGE IN PRACTICE

It is clear from the above description that the optimal policy (GBSD) requires global information about the network and the “spread” of messages, in order to optimize a specific routing metric. In particular, for each

message present in a node's buffer, we need to know the values of $m_i(T_i)$ and $n_i(T_i)$. In related work [12], it has been suggested that this global view could be obtained through a secondary, "instantaneous" channel (e.g. cellular network), if available, or by flooding ("in-band") all necessary meta-data. Regarding the former option, cellular network connections are known to be low bandwidth (measurements suggest only few kbps even for 2.5-3G technologies [27]) and high cost in terms of power and actual monetary cost per bit. In networks of more than a few nodes, the amount of signalling data might make this option prohibitive. Concerning flooding, our experiments show that the impact of the flooding delay on the performance of the algorithm is not negligible. In practice, intermittent network connectivity and the long time it takes to flood buffer status information across DTN nodes, make this approach inefficient.

A different, more robust approach is to find estimators for the unknown quantities involved in the calculation of message utilities, namely m and n . We do this by designing and implementing a learning process that permits a DTN node to gather knowledge about the global network state at different times in the past, by making in-band exchanges with other nodes. Each node maintains a list of encountered nodes and the state of each message carried by them as a function of time (i.e. its buffer state history). Specifically, it logs whether a given message was present at a given time T in a node's buffer (counting towards n) or whether it was encountered earlier but is not anymore stored, e.g. it was dropped (counting towards m). In Section 6, we describe our statistics maintenance and collection method, in more detail, along with various optimizations to considerably reduce the signalling overhead.

Since global information gathered thus about a specific message might take a long time to propagate (as mentioned earlier) and hence might be obsolete when we calculate the utility of the message, we follow a different route. Rather than looking for the current value of $m_i(T)$ and $n_i(T)$ for a specific message i at an elapsed time T , we look at what happens, *on average*, for all messages after an elapsed time T . In other words, the $m_i(T)$ and $n_i(T)$ values for message i at elapsed time T are estimated using measurements of m and n for the same elapsed time T but *measured for (and averaged over) all other older messages*. These estimations are then used in the evaluation of the per-message utility.

Let's denote by $\hat{n}(T)$ and $\hat{m}(T)$ the estimators for $n_i(T)$ and $m_i(T)$ of message i . For the purpose of the analysis, we suppose that the variables $m_i(T)$ and $n_i(T)$ at elapsed time T are instances of the random variables $N(T)$ and $M(T)$. We develop our estimators $\hat{n}(T)$ and $\hat{m}(T)$ so that when plugged into the GBSD's delivery rate and delay per-message utilities calculated in Section 3, we get two new per-message utilities that can be used by a DTN node without any need for global information about messages. This results in a new scheduling and

drop policy, called HBSD (History Based Scheduling and Drop), a deployable variant of GBSD that uses the same algorithm, yet with per-message utility values calculated using estimates of m and n .

4.1 Estimators for the Delivery Rate Utility

When global information is unavailable, one can calculate the average delivery rate of a message over all possible values of $M(T)$ and $N(T)$, and then try to maximize it. In the framework of the GBSD policy, this is equivalent to choosing the estimators $\hat{n}(T)$ and $\hat{m}(T)$ so that the calculation of the average delivery rate is unbiased:

$$E[(1 - \frac{M(T)}{L-1}) * (1 - \exp(-\lambda N(T)R_i)) + \frac{M(T)}{L-1}] = (1 - \frac{\hat{m}(T)}{L-1}) * (1 - \exp(-\lambda \hat{n}(T)R_i)) + \frac{\hat{m}(T)}{L-1}$$

Plugging any values for $\hat{n}(T)$ and $\hat{m}(T)$ that verify this equality into the expression for the per-message utility of Eq.(1), one can make sure that the obtained policy maximizes the average delivery rate. This is exactly our purpose. Suppose now that the best estimator for $\hat{m}(T)$ is its average, i.e., $\hat{m}(T) = \bar{m}(T) = E[M(T)]$. This approximation is driven by the observation we made that the histogram of the random variable $M(T)$ can be approximated by a Gaussian distribution with good accuracy. To confirm this, we have applied the Lillie test [28], a robust version of the well known Kolmogorov-Smirnov goodness-of-fit test, to $M(T)$ for different elapsed times ($T = 25\%, 50\%$ and 75% of the TTL). This test led to acceptance for a 5% significance level. Consequently, the average of $M(T)$ is at the same time the unbiased estimator and the most frequent value among the vector $M(T)$. Then, solving for $\hat{n}(T)$ gives:

$$\hat{n}(T) = -\frac{1}{\lambda R_i} \ln\left(\frac{E[(1 - \frac{M(T)}{L-1}) \exp(-\lambda N(T)R_i)]}{(1 - \frac{\bar{m}(T)}{L-1})}\right) \quad (7)$$

Substituting this expression into Eq.(1) we obtain the following new per message utility for our approximating HBSD policy:

$$\lambda R_i E[(1 - \frac{M(T)}{L-1}) \exp(-\lambda R_i N(T))] \quad (8)$$

The expectation in this expression is calculated by summing over all known values of $N(T)$ and $M(T)$ for past messages at elapsed time T . Unlike Eq.(1), this new per-message utility is a function of past history of messages and can be calculated locally. It maximizes the average message delivery rate calculated over a large number of messages. When the number of messages is large enough for the law of large numbers to work, our

history based policy should give the same result as that of using the real global network information.

Finally, we note that L , the number of nodes in the network, could also be calculated from the statistics maintained by each node in the network. In this work, we assume it to be fixed and known, but one could estimate it similar to n and m , or using different estimation algorithms like the ones proposed in [29].

4.2 Estimators for the Delivery Delay Utility

Similar to the case of delivery rate, we calculate the estimators $\hat{n}(T)$ and $\hat{m}(T)$ in such a way that the average delay is not affected by the estimation. This gives the following per-message utility specific to HBSD,

$$\frac{E\left[\frac{L-1-M(T)}{N(T)}\right]^2}{\lambda(L-1)(L-1-\bar{m}(T))} \quad (9)$$

This new per-message utility is only a function of the locally available history of old messages and is thus independent of the actual global network state. For large number of messages, it should lead to the same average delay as when the exact values for m and n are used.

5 PERFORMANCE EVALUATION

5.1 Experimental Setup

To evaluate our policies, we have implemented a DTN framework into the Network Simulator NS-2 [30]. This implementation includes (i) the Epidemic routing protocol with *FIFO* for scheduling messages queued during a contact and *drop-tail* for message drop, (ii) the RAPID routing protocol based on flooding (i.e. no side-channel) as described, to our best understanding, in [12], (iii) a new version of Epidemic routing enhanced with our optimal joint scheduling and drop policy (GBSD), (iv) another version using our statistical learning based distributed algorithm (HBSD), and (v) the VACCINE anti-packet mechanism described in [14]⁶.

In our simulations, each node uses the 802.11b protocol to communicate, with rate 11Mbps/sec. The transmission range is 100 meters, to obtain network scenarios that are neither fully connected (e.g. MANET) nor extremely sparse. Our simulations are based on three mobility scenarios, a synthetic one, based on the Random Waypoint model and two real-world mobility traces: the first trace was collected as part of the ZebraNet wildlife tracking experiment in Kenya described in [31]. The second mobility trace tracks San Francisco's Yellow Cab taxis. Many cab companies outfit their cabs with *GPS* to aid in rapidly dispatching cabs to their costumers. The Cabspotting system [32] talks to the Yellow Cab server and stores the data in a database. We have used an API

6. We have also performed simulations without any anti-packet mechanism, from which similar conclusions can be drawn.

provided by the Cabspotting system in order to extract mobility traces⁷.

To each source node, we have associated a CBR (Constant Bit Rate) application, which chooses randomly from $[0, TTL]$ the time to start generating messages of 5KB for a randomly chosen destination. We have also considered other message sizes (see e.g. [16]), but found no significant differences in the qualitative and quantitative conclusions drawn regarding the relative performance of different schemes⁸. Unless otherwise stated, each node maintains a buffer with a capacity of 20 messages to be able to push the network towards a congested state without exceeding the processing and memory capabilities of our simulation cluster. We compare the performance of the various routing protocols using the following two metrics: the average delivery rate and average delivery delay of messages in the case of infinite TTL ⁹. Finally, the results presented here are averages from 20 simulation runs, which we found enough to ensure convergence.

5.2 Performance evaluation for delivery rate

First, we compare the delivery rate of all policies for the three scenarios shown in Table 2.

TABLE 2
Simulation parameters

Mobility pattern:	RWP	ZebraNet	Taxis
Simulation's Duration(h):	7	14	42
Simulation' Area (m^2):	3000*3000	3000*3000	-
Number of Nodes:	70	70	70
Average Speed (m/s):	2	-	-
$TTL(h)$:	1	2	6
CBR Interval(s):	360	720	2160

TABLE 3
Taxi Trace & Limited buffer and bandwidth

Policy:	GBSD	HBSD	RAPID	FIFO\DT
D. Probability:	0.72	0.66	0.44	0.34
D. Delay(s):	14244	15683	20915	36412

Figure 2 shows the delivery rate based on the Random Waypoint model. From this plot, it can be seen that: the GBSD policy plugged into Epidemic routing gives the best performance for all numbers of sources. When

7. Note that this trace describes taxi's positions according to the *GPS* cylindrical coordinates (*Longitude*, *Latitude*). In order to use these traces as input for the NS-2 simulator, we have implemented a tool [30] based on the Mercator cylindrical map projection which permit us to convert traces to plane coordinates.

8. In future work, we intend to evaluate the effect of variable message size and its implications for our optimization framework. In general, utility-based scheduling problems with variable sized messages can often be mapped to Knapsack problems (see e.g. [33]).

9. By infinite TTL , we mean any value large enough to ensure almost all messages get delivered to their destination before the TTL expires.

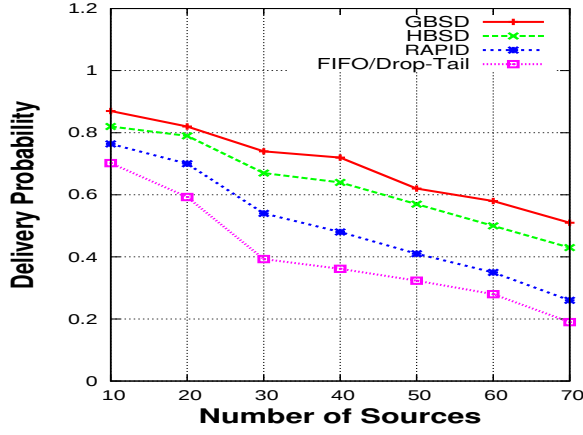


Fig. 2. Delivery Probability for Epidemic Routing with different scheduling and drop policies (both buffer and bandwidth constraints).

TABLE 4

ZebraNet Trace & Limited buffer and bandwidth

Policy:	GBSD	HBSD	RAPID	FIFO\DT
D. Probability:	0.68	0.59	0.41	0.29
D. Delay(s):	4306	4612	6705	8819

congestion-level decreases, so does the difference between GBSD and other protocols, as expected. Moreover, the HBSD policy also outperforms existing protocols (RAPID and Epidemic based on FIFO/drop-tail) and performs very close to the optimal GBSD. Specifically, for 70 sources, HBSD offers an almost 60% improvement in delivery rate compared to RAPID and is only 14% worse than GBSD. Similar conclusions can be also drawn for the case of the real Taxi traces or ZebraNet traces and 70 sources. Results for these cases are respectively summarized in Table 3 and Table 4.

5.3 Performance evaluation for delivery delay

To study delays, we increase messages' TTL (and simulation duration), to ensure almost every message gets delivered, as follows. Random Waypoint: (duration 10.5h, TTL = 1.5h). ZebraNet: (simulation duration = 28h, TTL = 4h). Taxi trace: (simulation duration = 84h, TTL = 12h). Traffic rates are as in Section 5.2.

For the random waypoint mobility scenario, Figure 3 depicts the average delivery delay for the case of both limited buffer and bandwidth. As in the case of delivery rate, GBSD gives the best performance for all considered scenarios. Moreover, the HBSD policy outperforms the two routing protocols (Epidemic based on FIFO/drop-tail, and RAPID) and performs close to GBSD. Specifically, for 70 sources and both limited buffer and bandwidth, HBSD average delivery delay is 48% better than RAPID and only 9% worse than GBSD.

Table 3 and Table 4 show that similar conclusions can be drawn for the delay under respectively the real Taxi(s) and ZebraNet traces.

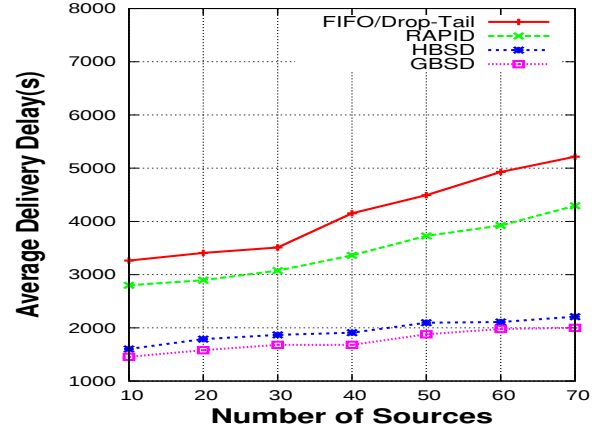


Fig. 3. Delivery Delay for Epidemic Routing with different scheduling and drop policies (both buffer and bandwidth constraints).

6 MAINTAINING NETWORK HISTORY

The results of the previous section clearly show that our distributed policy (HBSD) that uses estimators of global message state (rather than actual state) successfully approximates the performance of the optimal policy (GBSD). This is as an important step towards a practical implementation of efficient buffer management and scheduling algorithms on wireless devices. Nevertheless, in order to derive good estimators in a distributed manner, nodes need to exchange (a possibly large amount of) metadata during every node meeting. Potentially, each node needs to know the history of all messages having passed through a node's buffer, for every node in the network. In a small network, the amount of such "control" data might not be much, considering that large amounts of data transfers can be achieved between 802.11 transceivers during short contacts (data transfers of a few 10s of MBytes have been reported for experiments between vehicles moving at high speeds [34]). Nevertheless, in larger networks, this method can quickly become unsalable and interfere with (or starve) data transmissions, if statistics maintenance and collection is naively done.

In this section, we describe the type of statistics each node maintains towards calculating the HBSD utility for each message, and propose a number of mechanisms and optimizations to significantly reduce (and control) the amount of metadata exchanged during contacts. Finally, we explore the impact of reducing the amount of collected statistics on the performance of our buffer management and scheduling policy. Our results suggest that, with a carefully designed statistics collection and maintenance scheme, order(s) of magnitude less metadata can be exchanged (compared to maintaining a complete view about the network), without significantly affecting performance.

6.1 Maintaining Buffer State History

In order to keep track of the statistics about past messages necessary to assign appropriate utility values to messages considered for transmission or dropping, we propose that each node maintains the data structure depicted in Figure 4. Each node maintains a list of messages whose history in the network it keeps track of (we will see in the next section how a node chooses which messages to include in this list). For each message, it maintains its *ID* (a unique string resulting from the combination of some of its attributes), its *TTL* and the list of nodes that have *seen* it before (i.e. had stored the messages at some time in the past and should be accounted towards calculating m or n). Then, for each of the nodes in the list, it maintains a data structure with the following data: (i) the node's *ID*, (ii) a boolean array *Copies_Bin_Array*, and (iii) the version *Stat_Version* associated to this array.

The *Copies_Bin_Array* array (Fig. 5) enables nodes to maintain dynamically what each message experienced during its life time. For a given entry pair (message a and node b) in this list, the *Copies_Bin_Array*[k] indicates if the node a had already stored or not a copy of message b in its buffer during *Bin* k . In other words, time is quantized into “bins” of size *Bin_Size*, and bin k correspond to the period of time between $k * \text{Bin_Size}$ and $(k + 1) * \text{Bin_Size}$. As a result, the size of the *Copies_Bin_Array* is equal to $\text{TTL} / \text{Bin_Size}$.

How should one choose *Bin_Size*? Clearly, the larger it is, the fewer the amount of data a node needs to maintain and to exchange during each meeting; however, the smaller is also the granularity of values the utility function can take and thus the higher the probability of an incorrect (scheduling or buffer management) decision. As already described in Section 3, message transmissions can occur only when nodes encounter each other. This is also the time granularity at which buffer state changes occur. Hence, we believe that a good trade-off is to monitor the evolution of each message's state at a bin granularity in the order of meeting times¹⁰. This *already results in a big reduction of the size of statistics to maintain locally* (as opposed to tracking messages at seconds or milliseconds granularity), while still enabling us to infer the correct messages statistics.

Finally, the *Stat_Version* indicates the Bin at which the last update occurred. Let's assume that a message a is first stored at a node b during bin 3. It then creates a new entry in its list for pair (a, b) , inserts 0s in bins 0 – 2 of the new *Copies_Bin_Array* and 1s in the rest of the bins, and sets the *Stat_Version* to 3. If later, at in bin 5 node b decides to drop this message, then the list entry is maintained, but it sets all bins from 5 to $\text{TTL} / \text{Bin_Size}$ to 0, and updates the *Stat_Version* to 5.

10. According to the Nyquist-Shannon [35] sampling theorem, a good approximation of the size of a *Bin* would be equal to inter-meeting-time/2. A running average of the observed times between consecutive meetings could be maintained easily, in order to dynamically adjust the bin size [8].

Finally, when the *TTL* for message a elapses (regardless of whether a is still present in b 's buffer or not), b sets the *Stat_Version* to $\text{TTL} / \text{Bin_Size}$, which also indicates that all information about the history of *this* message in *this* buffer is now available. The combination of how the *Copies_Bin_Array* is maintained and the *Stat_Version* updated, ensures that only the minimum amount of necessary metadata for *this pair of (message, node)* is exchanged during a contact.

We note also that, in principle, a *Message_Seen_Bin_Array* could be maintained, indicating if a node a had *seen* (rather than *stored* a message b at time t , in order to estimate $m(T)$. However, it is easy to see that the *Message_Seen_Bin_Array* can be deduced directly from the *Copies_Bin_Array*, and thus no extra storage is required. Summarizing, based on this lists maintained by all nodes, any node can retrieve the vectors $N(T)$ and $M(T)$ and can calculate the HBSD per-message utilities described in Section 4 without a need for an *oracle*.

6.2 Collecting Network Statistics

We have seen so far what types of statistics each node maintains about each past (message ID, node ID) tuple it knows about. Each node is supposed to keep up-to-date the statistics related to the messages it stores locally (i.e. entries in the list of Fig. 4 corresponding to its own node ID). However, it can only update its knowledge (and the respective entry) about the state of a message a at a node b when it either meets b directly, or it meets a node that has more recent information about the (a, b) tuple (i.e. a higher *Stat_Version*). The goal of the statistics collection method is that, through such message exchanges, nodes converge to a unified view about the state of a given message at *any* buffer in the network, during its lifetime.

Sampling Messages to Keep Track of: We now look in more detail into what kind of metadata nodes should exchange. The first interesting question is the following: *should a node maintain global statistics for every message it has heard of or only a subset?* We argue that monitoring a dynamic subset of these messages is sufficient to quickly¹¹ converge to the correct expectations we need for our utility estimators. This dynamic subset is illustrated in Figure 6 as being the Messages Under Monitoring, which are stored in the MUM buffer; it is dynamic because its size is kept fixed while messages inside it change. When a node decides to store a message for the first time, if there is space in its MUM buffer, it also inserts it there and will track its global state. In other words, each node randomly chooses a few messages it will *sample*, for which it will attempt to collect global state, and does not keep track of all messages currently alive in the network. The actual *sampling rate* depends on the size of the MUM buffer and the offered traffic

11. While speed of convergence is not that important, due to our history-based approach, it becomes significant in non-stationary scenarios with traffic load fluctuations and node churn, as we shall see.

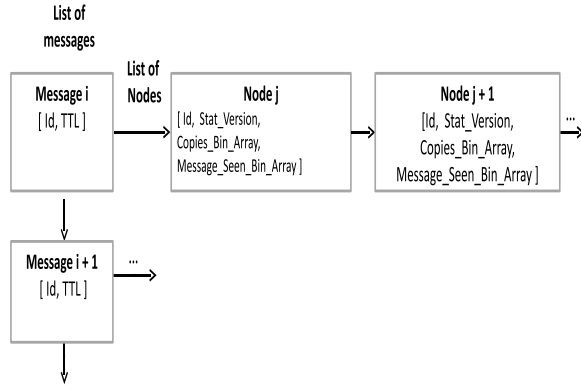


Fig. 4. Network History Data Structure

Example: TTL = 3600 (s), Bin_Size = 360 (s), Number of Bin(s) = 10

Copies_Bin_Array	0	0	1	1	1	0	0	0	0
Message_Seen_Bin_Array	0	0	1	1	1	1	1	1	1

Fig. 5. Example of Bin arrays

load, and results in significant further reduction in the amount of metadata exchanged. At the same time, a smaller MUM buffer might result to slower convergence (or even lack of). In Section 6.3 we study the impact of MUM buffer size on the performance of our algorithm.

Handling Converged Messages: Once the node collects an entire history of a given message, it removes it from the MUM buffer and pushes it to the buffer of Messages with a Complete History (MCH). A node considers that it has the complete history of a given message only when it gets the last version (i.e. $Stat_Version = TTL/Bin_Size$) of the statistics entries related to all the nodes the message goes through during its TTL ¹². Finally, note that, once a node decides to move a message to the MCH buffer, it only needs to maintain a short summary (i.e. number of nodes with a copy $n(T)$ and number of nodes having seen the message, $m(T)$, at time T) rather than the per node state as in Fig. 4.

Statistics Exchanged: Once a contact opportunity is present, both peers have to ask only for newer versions of the statistics entries (message ID, node ID) related to the set of messages buffered in their MUM buffer. This ensures that, even for the sampled set of messages, only new information is exchanged and no bandwidth is wasted. This optimization does not introduce any extra latency in the convergence of our approximation scheme.

12. Note that there is a chance that a node might “miss” some information about a message it pushes in its MCH. This occurs, for example, if it receives the last version for a subset of nodes which had the message, before it receives *any* version from another node that also had the message. This probability depends on the statistics of the meeting time (first and second moment) and the TTL value. Nevertheless, for many scenarios of interest, this probability is small and it may only lead to slightly underestimating the m and n values.

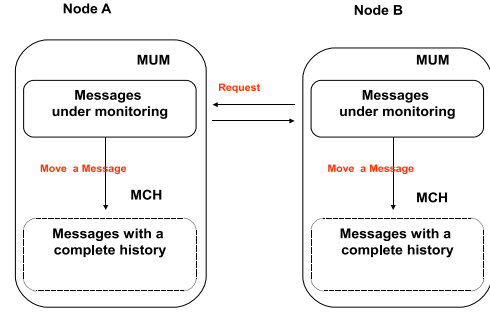


Fig. 6. Statistics Exchange and Maintenance.

6.3 Performance Tradeoffs of Statistics Collection

We have presented a number of optimizations to (considerably) reduce the amount of metadata stored and the amount of signalling overhead. Here, we explore the trade-off between the signalling overhead, its impact on performance, and the dynamicity of a given scenario. Our goal is to identify operation points where the amount of signalling overhead is such that it interferes minimally with data transmission, while at the same time it suffices to ensure timely convergence of the required utility metrics per message. We will consider throughout the random waypoint simulation scenario described in Section 5.2. We have observed similar behaviour for the trace-based scenarios.

Amount of Signalling Overhead per Contact: We start by studying the effect of varying the size of the MUM buffer (number of messages under monitoring) on the average size of exchanged statistics per-meeting. Figure 7 compares the average size of statistics exchanged during a meeting between two nodes for three different sizes of the MUM buffer (20, 50 and 80), as well as for the basic epidemic statistics exchange method (i.e. unlimited MUM). We vary the number of sources in order to cover different congestions regimes.

Our first observation is that increasing the traffic load (and thus the amount of congestion) results in decreasing the average amount of statistics exchanged per-meeting (except for the MUM size of 20 messages). This might be slightly counterintuitive, since a higher traffic load implies more messages to keep track of. However, note that a higher congestion level also implies that much fewer copies per message will co-exist at any time (and new versions are less frequently created). As a result, much less metadata per message is maintained and exchanged, resulting in a downward trend. In the case of a MUM size of 20, it seems that these two effects balance each other out. In any case, the key property here is that, in contrast with the flooding-based method of [12], *our distributed collection method scales well, not increasing the amount of signalling overhead during high congestion.*

A second observation is that, using our statistics collection method, a node can reduce the amount of signalling overhead per meeting up to an order of magnitude (e.g. for MUM = 20), compared to the unlimited

MUM case, even in this relatively small scenario of 70 nodes. (Note also that, the plot shown for the epidemic case, already implements the binning and versioning optimizations of Section 6.1).)

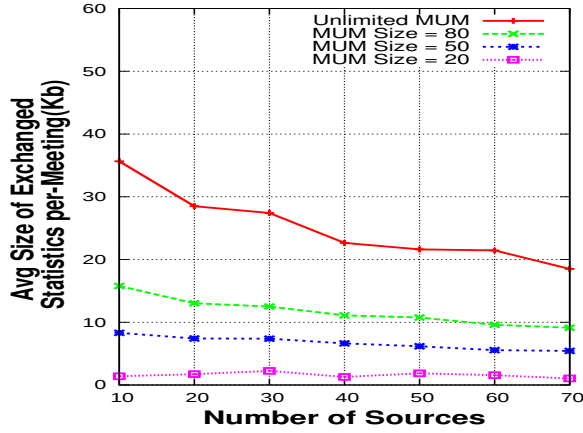


Fig. 7. Signalling overhead (per contact) resulting from HBSD statistics collection.

Finally, we plot in Figure 8 the average size of exchanged (non-signalling) data per-meeting. We can observe that increasing the size of the *MUM* buffer results in a slight decrease of the data exchanged. This is due to the priority we give to statistics exchange during a contact. We note also that this effect becomes less pronounced when congestion increases (in line with Fig. 7). Finally, in the scenario considered, we can observe that, for *MUM* sizes less than 50, signalling does not interfere with data transmissions (remember that packet size is 5KB). This suggests that, in this scenario, a *MUM* size of 50 messages represents a good choice with respect to the resulting signalling overhead. In practice, a node could find this value online, by dynamically adjusting its *MUM* size and comparing the resulting signalling overhead with average data transfer. It is beyond the scope of this paper to propose such an algorithm. Instead, we are interested in exposing the various tradeoffs and choices involved in efficient distributed estimation of statistics. Towards this goal, we explore next the effect of the *MUM* sizes considered on the performance of our HBSD algorithm.

Convergence of Utilities and Performance of the HBSD Policy : In this last part, we fix the number of sources to 50 and we look at the impact of the size of the *MUM* buffer on (i) the time it takes the HBSD delivery rate utility to converge, and (ii) its accuracy. We use the *mean relative square error* to measure the accuracy of the HBSD delivery rate utility, defined as follows:

$$\frac{1}{\#Bins} * \sum_{Bins} \frac{(A - B)^2}{B^2},$$

where, for each *bin*, *A* is the estimated utility value of Eq. (8) (calculated using the approximate values of *m* and *n*, collected with the method described previously)

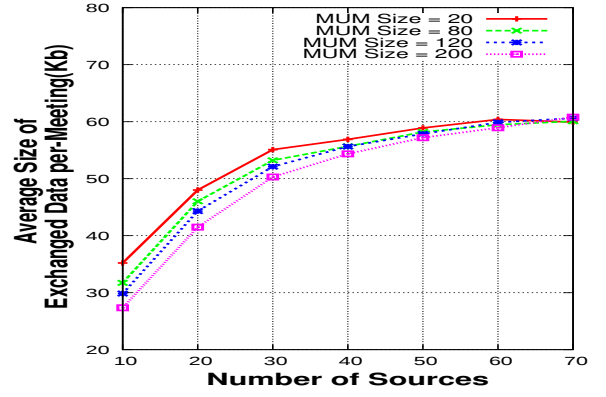


Fig. 8. Average size of exchanged (non-signalling) data per contact.

and *B* is the utility value calculated using the real values of *m* and *n*.

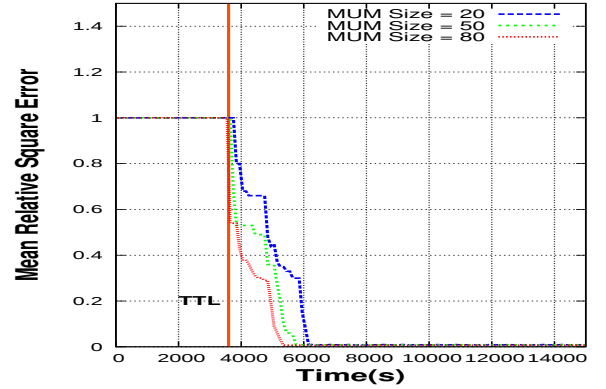


Fig. 9. Mean relative square errors for HBSD delivery rate utility.

Figure 9 plots the *mean relative square errors* for the HBSD delivery rate utility, as a function of time. We can observe that, increasing the size of the *MUM* buffer results in faster reduction of the *mean relative square error* function. With a *MUM* buffer of 80 messages, the delivery rate utility estimate converges 800 seconds faster than using an *MUM* buffer of 20 messages. Indeed, the more messages a node tracks in parallel, the faster it can collect a working history of past messages that it can use to calculate utilities for new messages considered for drop or transmission. We observe also that all plots converge to the same very small error value¹³. Note also that it is not the absolute value of the utility function (during different time bins) that we care about, but rather the *shape* of this function, whether it is increasing or decreasing, and the relative utility values. (We will look into the shape of this function at different congestion regimes in the next section.)

In fact, we are more interested in the end performance of our HBSD, as a function of how “aggressively” nodes

13. We speculate that this remaining error might be due to slightly underestimating *m* and *n*, as explained earlier.

collect message history. In Figures 10 and 11, we plot the delivery rate and delay of HBSD, respectively, for different MUM sizes. These results correspond to the scenario described in Section 5.2, where we have a fixed number of CBR sources. As is evident from these figures, regardless of the size of the MUM buffer sizes, nodes eventually gather enough past message history to ensure an accurate estimation of per message utilities, and a close-to-optimal performance. In such scenarios, where traffic intensity is relatively stable, even a rather small MUM size (i.e. very low sampling rate) suffices to achieve good performance. This is not necessarily the case when traffic load experiences significant fluctuations (e.g. due to new popular content appearing in the network).

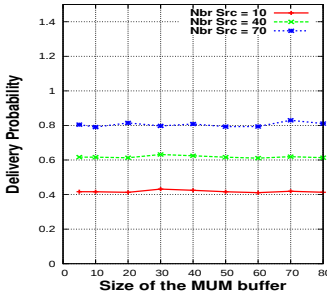


Fig. 10. Delivery Probability for HBSD with statistics collection (static traffic load).

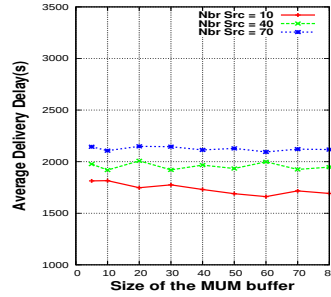


Fig. 11. Deliver Delay for HBSD with statistics collection (static traffic load).

When the offered traffic load changes frequently (or node churn is high, e.g. experiencing “flash crowds”), convergence speed becomes important. The bigger the MUM buffer the faster our HBSD policy react to changing congestion levels. We illustrate this with the following experiment. We maintain the same simulation scenario, but we vary the number of CBR sources among each two consecutive TTL(s), from 10 to 70 sources (i.e. the first and second TTL window we have 10 sources, the third and fourth window 70 sources, etc. — this is close to a *worst case* scenario, as there is a sevenfold increase in traffic intensity within a time window barely higher than a TTL, which is the minimum required interval to collect any statistics). Furthermore, to ensure nodes use non-obsolete statistics towards calculating utilities, we force nodes to apply a *sliding window* of one TTL to the messages with complete history stored in the MCH buffer, and to delete messages out of this *sliding window*¹⁴

Figures 12 and 13 again plot the HBSD policy delivery rate and delay, respectively, as a function of MUM buffer size. Unlike the constant load case, it is easy to see there that, increasing the size of the MUM buffer, results in considerable performance improvement. Nevertheless, even in this rather dynamic scenario, nodes manage to keep up and produce good utility estimates, with only a modest increase on the amount of signalling overhead

required.

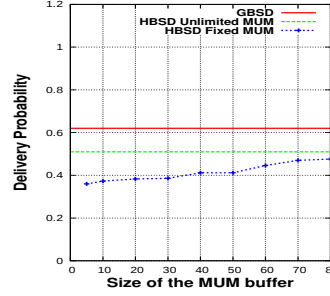


Fig. 12. Deliver Probability for HBSD with statistics collection (dynamic traffic load).

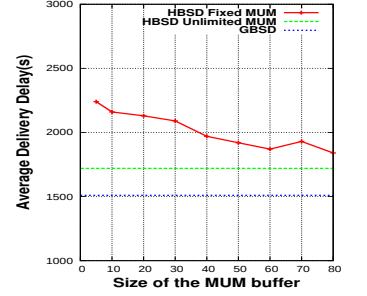


Fig. 13. Deliver Delay for HBSD with statistics collection (dynamic traffic load).

7 DISTRIBUTION OF HBSD UTILITIES

We have described how to efficiently collect the necessary statistics in practice, and derive good estimates for the HBSD utility distribution during the lifetime of a message. In this last section, we turn our attention to the utility distributions themselves. First, we are interested whether the resulting distributions for HBSD delivery rate and delivery delay utilities react differently to different congestion levels, that is, if the priority given to messages of different ages shifts based on the offered load. Furthermore, we are interested whether the resulting utility shape (and respective optimal policy) could be approximated by simple(r) policies, in some congestion regimes.

We consider again the simulation scenario used in Section 5.2 and Section 6.3. First, we fix the number of sources to 50, corresponding to a *high congestion regime*. In Figure 14 and Figure 15, we plot the distribution of the HBSD delivery rate and delivery delay utilities described in Sections 4.1 and 4.2. It is evident there that the optimal utility distribution has a non-trivial shape for both optimization metrics, resulting in a complex optimal scheduling and drop policy. This also helps explain why simple drop and scheduling policies (e.g. Drop Youngest or Oldest Message, DropTail, FIFO or LIFO scheduling, etc.), considered in earlier work [14], [16] lead to incorrect decisions during congestion and perform worse than the GBSB and HBSD policies [16].

Next, we consider a scenario with low congestion. We reduce the number of sources to 15, keep the buffer size of 20 messages, but we also decrease the CBR rate of sources from 10 to 2 messages/TTL. In Figures 16 and 17, we plot the distribution of the HBSD delivery rate and delivery delay utilities, respectively, for this low congestion scenario. Surprisingly, our HBSD policy behaves very differently now, with both utility functions decaying monotonically as a function of time (albeit not at constant rate). This suggests that the optimal policy in low congestion regimes could be approximated by the simpler “Drop Oldest Message” (or schedule

14. A running average could be used for smoother performance. We only care here to demonstrate the effect of dynamic traffic loads.

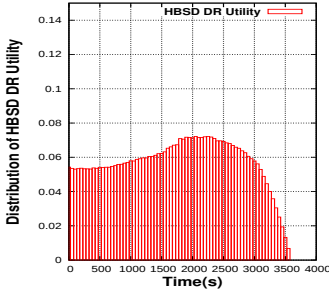


Fig. 14. Distribution of HBSD DR utility in a congested network.

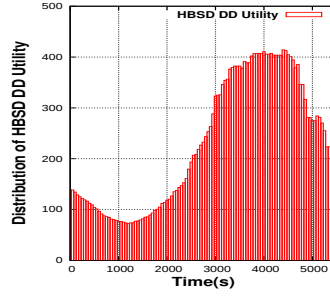


Fig. 15. Distribution of HBSD DD utility in a congested network.

younger messages first) policy, which does not require any signalling and statistics collection between nodes.

To test this, in Tables 5 and 6, we compare the performance of the HBSD policy against a simple combination of “Drop Oldest Message” (for Buffer Management) and “Transmit Youngest Message First” (for Scheduling during a contact). We observe, that in the low congestion regime (Tables 6) the two policies indeed have similar performance (4% and 5% difference in delivery rate and delivery delay, respectively). However, in the case of a congested network (Table 5), HBSD clearly outperforms the simple policy combination.

We can look more carefully at Figures 14 and 15, to understand what is happening in high congestion regimes. The number of copies per message created at steady state depends on the total number of messages co-existing at any time instant, and the aggregate buffer capacity. When too many messages exist in the network (for the provided buffer space per node), uniformly assigning the available messages to the existing buffers (which is what a random drop and scheduling policy would do), would imply that every message can have only a few copies created. Specifically, for congestion higher than some level, the average number of copies per message allowed is so low that most messages cannot reach their destination during their TTL (this depends only on the number of copies and mobility model). *Uniformly assigning resources between nodes is no more optimal.* Instead, to ensure that at least some messages can be delivered on time, the optimal policy gives higher priority to older messages that have managed to survive long enough (and have probably created enough copies), and “kills” some of the new ones being generated. This is evident by the values assigned at different bins (especially in the delivery delay case). In other words, when congestion is excessive *our policy performs an indirect admission control function.*

Contrary to this, when the offered load is low enough to ensure that all messages can on average create enough copies to ensure delivery, the optimal policy simply performs a fair (i.e. equal) distribution of resources (ensured by the utility functions of Figures 16 and 17).

The above findings suggest that it would be quite useful to find a generic way to signal the congestion level

TABLE 5
HBSD vs. “Schedule Younger First\Drop-Oldest” in a congested network.

Policies:	HBSD	“Schedule Younger First\Drop-Oldest”
D. Rate(%):	54	29
D. Delay(s):	1967	3443

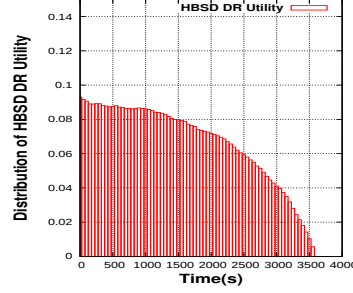


Fig. 16. Distribution of HBSD DR utility in a low congested network.

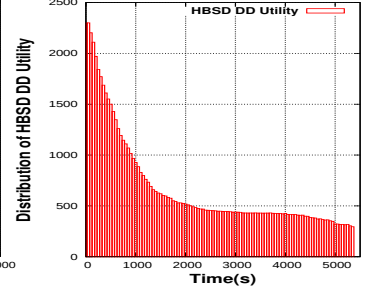


Fig. 17. Distribution of HBSD DD utility in a low congested network.

TABLE 6
HBSD vs “Schedule Younger First\Drop-Oldest” in a low congested network.

Policies:	HBSD	“Schedule Younger First\Drop-Oldest”
D. Rate(%):	87	83
D. Delay(s):	1530	1618

and identify the threshold based on which nodes can decide to either activate our HBSD scheme or just use a simple Drop/Scheduling policy. Suspending a complex Drop/Scheduling mechanism and its underlying statistics collection and maintenance methods, whenever not needed, can help nodes save an important amount of resources (e.g. energy), while maintaining the same end performance. Finally, we believe that the indirect signalling provided by the behaviour of the utility function during congestion, could provide the basis for an end-to-end flow control mechanism, a problem remaining largely not addressed in the DTN context.

8 CONCLUSION

In this work, we investigated both the problems of scheduling and buffer management in DTNs. First, we proposed an optimal joint scheduling and buffer management policy based on global knowledge about the network state. Then, we introduced an approximation scheme for the required global knowledge of the optimal algorithm. Using simulations based on a synthetic mobility model (Random Waypoint), and real mobility traces, we showed that our policy based on statistical learning successfully approximates the performance of the optimal algorithm. Both policies (GBSD and HBSD) plugged into the Epidemic routing protocol outperform

current state-of-the-art protocols like RAPID [12] with respect to both delivery rate and delivery delay, in all considered scenarios. Moreover, we discussed how to implement our HBSD policy in practice, by using a distributed statistics collection method, illustrating that our approach is realistic and effective. We showed also that unlike many works [12], [17] that also relied on the use of an in-band control channel to propagate metadata, our statistics collection method scales well, not increasing the amount of signalling overhead during high congestion.

Finally, we carried a study of the distributions of HBSD' utilities under different congestion levels and we showed that: when congestion is excessive, HBSD performs an indirect admission control function and has a non-trivial shape for both optimization metrics, resulting in a complex optimal scheduling and drop policy. However, when the offered load is low enough, HBSD can be approximated by a simple policy that does not require any signalling and statistics collection between nodes. The above findings suggest that it would be quite useful to find a generic way to signal the congestion level and identify the threshold based on which nodes can decide to either activate our HBSD scheme or just use a simple Drop/Scheduling policy. Suspending a complex Drop/Scheduling, whenever not needed, can help nodes save an important amount of resources, while maintaining the same end performance.

REFERENCES

- [1] S. Jain, K. Fall, and R. Patra, "Routing in a delay tolerant network," in *Proceedings of ACM SIGCOMM*, Aug. 2004.
- [2] S. Jain, M. Demmer, R. Patra, and K. Fall, "Using redundancy to cope with failures in a delay tolerant network," in *Proceedings of ACM SIGCOMM*, 2005.
- [3] N. Glance, D. Snowdon, and J.-L. Meunier, "Pollen: using people as a communication medium," *Computer Networks*, vol. 35, no. 4, Mar. 2001.
- [4] A. Pentland, R. Fletcher, and A. Hasson, "Daknet: Rethinking connectivity in developing nations," *IEEE Computer*, 2004.
- [5] "Delay tolerant networking research group," <http://www.dtnrg.org>.
- [6] A. Vahdat and D. Becker, "Epidemic routing for partially connected ad hoc networks," Duke University, Tech. Rep. CS-200006, 2000.
- [7] A. Lindgren, A. Doria, and O. Schelen, "Probabilistic routing in intermittently connected networks," *SIGMOBILE Mobile Computing and Communication Review*, vol. 7, no. 3, 2003.
- [8] T. Spyropoulos, K. Psounis, and C. S. Raghavendra, "Efficient routing in intermittently connected mobile networks: The multiple-copy case," *IEEE/ACM Transactions on Networking*, vol. 16, no. 1, pp. 77–90, 2008.
- [9] Z. J. Haas and T. Small, "A new networking model for biological applications of ad hoc sensor networks," *IEEE/ACM Transactions on Networking*, vol. 14, no. 1, pp. 27–40, 2006.
- [10] R. Groenevelt, G. Koole, and P. Nain, "Message delay in manet (extended abstract)," in *Proc. ACM Sigmetrics*, 2005.
- [11] T. Spyropoulos, K. Psounis, and C. S. Raghavendra, "Performance analysis of mobility-assisted routing," in *Proceedings of ACM/IEEE MOBIHOC*, 2006.
- [12] A. Balasubramanian, B. Levine, and A. Venkataramani, "Dtn routing as a resource allocation problem," in *Proceedings of ACM SIGCOMM*, 2007.
- [13] A. Lindgren and K. S. Phanse, "Evaluation of queuing policies and forwarding strategies for routing in intermittently connected networks," in *Proceedings of IEEE COMSWARE*, 2006.
- [14] X. Zhang, G. Neglia, J. Kurose, and D. Towsley, "Performance modeling of epidemic routing," in *Proceedings of IFIP Networking*, 2006.
- [15] H. P. Dohyung Kim and I. Yeom, "Minimizing the impact of buffer overflow in dtn," in *Proceedings International Conference on Future Internet Technologies (CFI)*, 2008.
- [16] A. Krifa, C. Barakat, and T. Spyropoulos, "Optimal buffer management policies for delay tolerant networks," in *IEEE SECON*, 2008.
- [17] D. J. L. S. L. Z. Yong L., Mengjiong Q., "Adaptive optimal buffer management policies for realistic dtn," in *IEEE GLOBECOM*, 2009.
- [18] T. Small and Z. Haas, "Resource and performance tradeoffs in delay-tolerant wireless networks," in *Proceedings of ACM SIGCOMM workshop on Delay Tolerant Networking (WDTN)*, 2005.
- [19] T. Spyropoulos, K. Psounis, and C. S. Raghavendra, "Spray and wait: Efficient routing in intermittently connected mobile networks," in *Proceedings of ACM SIGCOMM workshop on Delay Tolerant Networking (WDTN)*, 2005.
- [20] Y. Wang, S. Jain, M. Martonosi, and K. Fall, "Erasure coding based routing for opportunistic networks," in *Proceedings of ACM SIGCOMM workshop on Delay Tolerant Networking (WDTN)*, 2005.
- [21] Y. Lin, B. Li, and B. Liang, "Efficient network coded data transmissions in disruption tolerant networks," in *Proceedings of IEEE INFOCOM*, 2008.
- [22] J. Burgess, B. Gallagher, D. Jensen, and B. N. Levine, "Max-Prop: Routing for Vehicle-Based Disruption-Tolerant Networks," in *Proc. IEEE INFOCOM*, 2006.
- [23] T. Spyropoulos, T. Turetli, and K. Obrazcka, "Routing in delay tolerant networks comprising heterogeneous populations of nodes," *IEEE Transactions on Mobile Computing*, 2009.
- [24] D. Aldous and J. Fill, "Reversible markov chains and random walks on graphs. (monograph in preparation.)," <http://stat-www.berkeley.edu/users/aldous/RWG/book.html>.
- [25] T. Karagiannis, J.-Y. Le Boudec, and M. Vojnović, "Power law and exponential decay of inter contact times between mobile devices," in *Proceedings of ACM MobiCom '07*, 2007.
- [26] A. Chaintreau, J.-Y. Le Boudec, and N. Ristanovic, "The age of gossip: spatial mean field regime," in *Proceedings of ACM SIGMETRICS '09*, 2009.
- [27] S. K. S. B. A. Seth, M. Zaharia, "A policy-oriented architecture for opportunistic communication on multiple wireless networks (<http://blizzard.cs.uwaterloo.ca/keshav/home/papers/data/06/ocmp.pdf>)," University of Waterloo, Tech. Rep., 2006.
- [28] H. Lilliefors, "On the kolmogorov-smirnov test for normality with mean and variance unknown," *Journal of the American Statistical Association*, Vol. 62, pp. 399–402, 1967.
- [29] A. Guerrieri, A. Montresor, I. Carreras, F. D. Pellegrini, and D. Miorandi, "Distributed estimation of global parameters in delay-tolerant networks," in *in Proceedings of Autonomic and Opportunistic Communication (AOC) Workshop (colocated with WOW-MOM)*, 2009, pp. 1–7.
- [30] DTN Architecture for NS-2. [Online]. Available: <http://www-sop.inria.fr/members/Amir.Krifa/DTN>
- [31] Y. Wang, P. Zhang, T. Liu, C. Sadler, and M. Martonosi, "Movement data traces from princeton zebranet deployments," CRAWDAD Database. <http://crawdada.cs.dartmouth.edu/>, 2007.
- [32] Cabspotting Project. [Online]. Available: <http://cabspotting.org/>
- [33] C. Boldrini, M. Conti, and A. Passarella, "Contentplace: social-aware data dissemination in opportunistic networks," in *Proceedings of ACM MSWiM*, 2008.
- [34] D. Hadaller, S. Keshav, T. Brecht, and S. Agarwal, "Vehicular opportunistic communication under the microscope," in *Proceedings of ACM MobiSys'07*, New York, NY, USA, 2007.
- [35] Nyquist Shannon sampling theorem. [Online]. Available: <http://en.wikipedia.org/wiki/Nyquistrem>