

On the Delay Performance in a Large-Scale Wireless Sensor Network: Measurement, Analysis and Implications

Jiliang Wang, *Member, IEEE*, Wei Dong, *Member, IEEE*,
Zhichao Cao, *Member, IEEE*, and Yunhao Liu, *Senior Member, IEEE*

Abstract—We present a comprehensive delay performance measurement and analysis in a large-scale wireless sensor network. We build a light-weight delay measurement system and present a robust method to calculate the per-packet delay. We show that the method can identify incorrect delays and recover them with a bounded error. Through analysis of delay and other system metrics, we seek to answer the following fundamental questions: What are the spatial and temporal characteristics of delay performance in a real network? What are the most important impacting factors and is there any practical model to capture those factors? What are the implications to protocol designs? In this paper, we identify important factors from the data trace, and show that the important factors are not necessarily the same with those in Internet. Further, we propose a delay model to capture those factors. We revisit several prevalent protocol designs such as Collection Tree Protocol, opportunistic routing and Dynamic Switching based Forwarding, and show the our model and analysis are useful to practical protocol designs.

Index Terms—wireless sensor networks, delay measurement, large-scale, impacting factor.

I. INTRODUCTION

Recent advances in Wireless Sensor Networks (WSNs) have fostered a large number of applications, such as structural protection and health monitoring WSNs [1] [2], etc. Those WSN applications often require Quality of Service (QoS) guarantees to fulfill the system requirements, e.g., real-time data delivery. Of the major factors that affect system QoS, delay is an important one.

There are many research works on delay analysis and measurement. Kompella et al. [3] present a fine-grained latency measurement method in presence of packet losses for the Internet. Wilson et al. [4] present delay analysis results in data centers, providing guidelines to practical data center design. There are tremendous research efforts made to delay analysis and modeling in WSNs. For example, probabilistic delay bounds are presented in [5] [6] [7] [8] by extending network

This work is supported in part by NSFC under grant 61202359, NSFC Major Program 61190110, NSFC under grant 61373166, National High-Tech R&D Program of China (863) under grant 2011AA010100, NSFC under grant 61202402 and Research Fund for the Doctoral Program of Higher Education of China (20120101120179).

Jiliang Wang, Zhichao Cao and Yunhao Liu are with School of Software and TNLIST, Tsinghua University, China. E-mail: {jiliang, caozc, yunhao}@greenorbs.com

Wei Dong is with College of Computer Science, Zhejiang University, China. E-mail: dongw@zju.edu.cn

The preliminary version of this work was published at IEEE RTSS 2012.

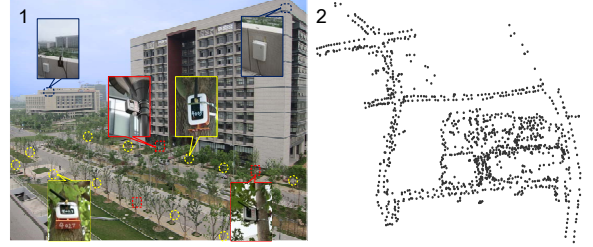


Fig. 1: Deployment and sensor nodes in CitySee: 1) overview of the deployment area, 2) node locations in the network.

calculus. Further, stochastic delay models are proposed by combining real-time theory and queuing theory [9] [10] [11] or applying Discrete Markov Process [12]. There are also some empirical network delay models such as [13] [14].

While there are excellent research works for WSNs, Internet and data centers, a practical end-to-end delay performance measurement and analysis in an operational large-scale WSN is still missing. On the other hand, considering the emerging demand of WSN applications, it is important to understand the delay performance in practical large-scale networks.

Delay performance measurement and analysis, in a large-scale WSN, face non-trivial challenges. First, different from the Internet and data centers [3] [15] [4], there are no effective methods in WSNs supporting per-packet delay measurement. Traditional delay measurement methods rely on network synchronization, which introduce additional overhead. Thus it is not always efficient to apply network synchronization to an operational network. Second, analyzing the collected information is challenging. Collecting all required information from the network incurs a high network overhead. Thus the information is usually incomplete due to resource constraints for sensor nodes and packet losses in the network. Meanwhile, there are various performance metrics in the network and a single delay change may be accompanied by variation of multiple metrics. Moreover, with low power listening, each node switches between wake up and sleep states to save energy. The delays exhibit intrinsic randomness in its distribution, introducing difficulty to efficient and automatical analysis.

In this work, we build an infrastructure for delay measurement in CitySee, a large-scale WSN consisting of 1200 nodes. The infrastructure does not rely on network synchronization and thus does not introduce additional over-

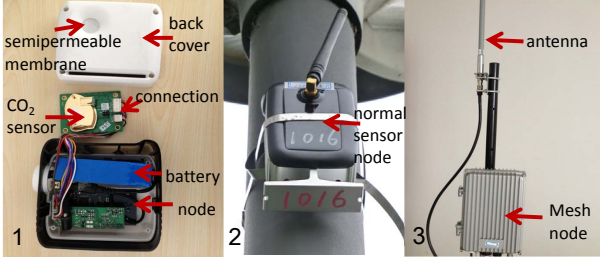


Fig. 2: Sensor nodes in CitySee: 1) CO_2 sensor node, 2) normal sensor node, and 3) mesh node.

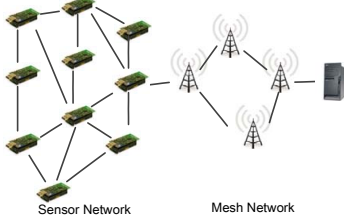


Fig. 3: Overview of network architecture.

head. We present basic statistical characteristics based on the collected data. To systematically and automatically identify important impacting factors from various parameters, we build a method based on Rulefit for the collected data trace. Further, we quantitatively calculate the correlation between different impacting factors and the delay performance. Based on those important factors, we build a practical delay model and validate the model using the collected data trace. Finally, we revisit three important protocols based on the measurement results and propose a practical delay model. In summary, the contributions of this paper are as follows.

- We build a measurement infrastructure in an operational large-scale WSN with little network overhead. Based on the collected data, we present the spatial and temporal characteristics of delay performance.
- We present an automatic method based on Rulefit [16] to identify important impacting factors to the delay performance.
- We propose a practical model and validate it with the collected data. We show the implications to protocol designs.

The rest of the paper is organized as follows. Section II presents the system overview. Section III shows the delay measurement method. Section IV presents the delay distribution overview. Section V introduces the method of identifying important factors from the data. Section VI builds a delay model based on the analysis. Section VII shows the implications of the analysis and the evaluation results. Section VIII introduces the related work. Finally, Section IX concludes this work.

II. SYSTEM OVERVIEW

A. The Network

The primary goal of CitySee is to precisely measure CO_2 emissions in a city-wide area. We started the project since July,

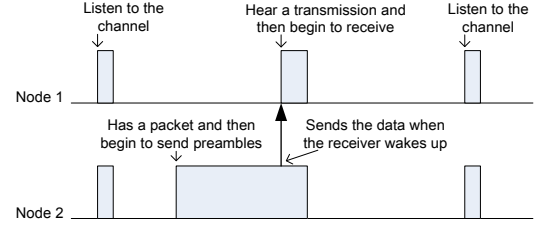


Fig. 4: Basic mechanism of LPL protocols.

2011. Figure 1 shows an overview of the network. Totally, we have deployed 1200 nodes. The network employs a tiered architecture with three kinds of nodes, i.e., normal telosB nodes, CO_2 nodes and mesh nodes.

In the network, each normal sensor node reads the sensing data and records system status. CO_2 nodes can also read the CO_2 concentration. The normal nodes and CO_2 nodes form a network and deliver their data to a sink node (normal node) in the network. Figure 2 shows the sensor nodes in our network.

The mesh nodes have a high bandwidth of several MB/s and a long transmission distance. They comprise the network backbone. Sink nodes of different subnets are connected to the network backbone in order to deliver packets to the base station. Figure 3 illustrates the overview of the system architecture.

B. Protocols

1) *Low Power Listening*: Low Power Listening (LPL) is widely adopted in WSNs to save energy. In LPL, each node switches between awake and sleep state to save energy. Most LPL protocols share the similar principle as shown in Figure 4. Each node samples the channel for a short duration in each cycle. If energy is detected, the node stays awake for another short duration to receive packets. Otherwise, the node turns off the radio and in the next cycle (e.g., 500ms later) re-samples the channel. To transmit a packet, the sender continues sending packets as preambles until the receiver wakes up. For broadcast, the preamble lasts for a cycle duration, in order to ensure that all neighboring nodes wake up once. For unicast with link layer ACK, the sender can stop the preambles until an ACK is received or the end of a cycle. Another type of LPL protocol is receiver-initiated low duty cycle protocol. Each node periodically wakes up and sends probe packets to see if there are transmissions intended for it. If a node has packets to send, it will keep awake and send the packets once receiving a probe packet from the receiver. Since a sender may begin to send packets at any time, the time the sender needs to wait is randomly distributed in the cycle. This introduces randomness to packet delay.

2) *Collection Tree Protocol*: Collection Tree Protocol (CTP) [17] is used to build routing tree in the network. CTP adopts the ETX metric [18], the expected transmission count, as the path quality metric. Each node selects a path with minimum ETX. The ETX of a link is calculated as $1/q$, where q is the packet reception ratio. The path ETX is calculated as the sum of all link ETXs along the path.

C. Measurement Infrastructure

Time synchronization can be used to measure delay in the network. However, time synchronization protocols in WSNs, such as FTSP [19], and etc [20] [21] [22], incurs additional traffic overhead into the network in order to maintain a global time stamp for all nodes.

We use a light-weight approach to measure the end-to-end packet delay without incurring synchronization traffic. For each packet, we define the delay as the time from the packet is generated at the source node to the time that the packet is received at the sink node. As shown in Figure 5, the delay for a packet on a path mainly consists of the following parts, 1) the packet transmission time for a packet, i.e., the time used to modulate the packet to signal, 2) the channel accessing time, i.e., the time used to contend for the channel, including the backoff time, 3) the queuing time and 4) the propagation time on each hop, i.e., the traveling time for the signal from the sender to the receiver which can be calculated as the distance between the sender and receiver divided by the light speed.

Our approach is based on the MAC layer time stamping technique (MLT) [19]. In MLT, a packet is time stamped just before the first byte is transmitted after backoff (with respect to sender's clock), and just before the first byte is received (with respect to the receiver's clock).

Our approach measures the end-to-end delay as follows. Figure 5 shows an example of 3 nodes. Assume a packet is generated at node 1's local time t_1 , we show how to measure the generation time to node 2's and node 3's clock. Suppose the packet is transmitted and time stamped at t_2 , with the time t_1 contained in the packet. Here, node 1 performs backoffs during the time $[t_1, t_2]$. Thus at time t_2 , the packet is modulated and emitted through the antenna. Then the packet is received at node 2 and time stamped at t_3 by MLT. Since the distance between two nodes is usually several hundreds of meters, we ignore the signal propagation time from node 1 to node 2. Node 2 calculates the packet generation time with respect to node 2's clock, by subtracting the time difference $t_2 - t_1$ from time t_3 , i.e., $t_3 - (t_2 - t_1)$. Intuitively, it seems that the packet is generated at time $t_3 - (t_2 - t_1)$ to node 2's clock.

Here, we denote the time $t_2 - t_1$ as the waiting time (t_w) at node 1. This is calculated as the time a packet is transmitted minus the time the packet is received at node 1. Similarly, the sink node, i.e. node 3, can calculate the packet generation time after receiving a packet from node 2 as $t_5 - (t_4 - (t_3 - (t_2 - t_1)))$, and the receiving time as t_5 , both are with respect to node 3's clock. Then the delay of the packet is calculated as $t_4 - (t_3 - (t_2 - t_1))$.

D. Collected Data

The network collects four types of packets, denoted as *C1*, *C2*, *C3* and *C4* respectively. Each node sends each type of packet in every 10 minutes. There is a common header with four fields, which records a common sequence number and the time stamps used for delay measurement. We will introduce in the next section how to use the time stamps for delay measurement.

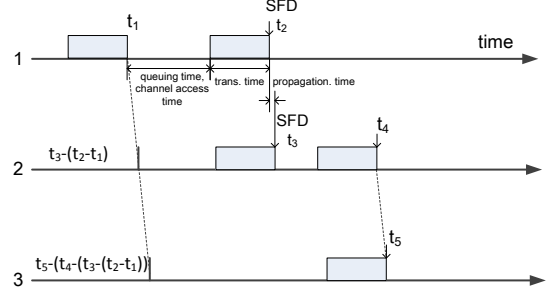


Fig. 5: Measuring delay in the network.

III. DELAY MEASUREMENT

In this section, we show how to derive the delay for each packet.

A. Delay Measurement Analysis

We assume the time t provided by a sensor node follows the linear clock model [19], i.e.,

$$t = (1 + \alpha)(\tau - t_o) \quad (1)$$

where τ is the time provided by a perfect clock, t_o is the offset and α is the clock drift bounded by $[-\hat{\alpha}, \hat{\alpha}]$. Each packet from a node has an unique increasing sequence number i . For packet i , we denote the receiving time of the packet i on node j with respect to the perfect clock as $t_r(i, j)$ and the transmitting time of the packet i on node j as $t_x(i, j)$. The transmitting/receiving time is the time just before the first byte of a packet is transmitted/received. We denote O as the source node and D as the sink node for a path. Thus, $t_r(i, O)$ is the generation time of packet i on the source node. Therefore, the end-to-end delay of packet i for a path can be calculated as

$$t_d(i) = t_r(i, D) - t_r(i, O). \quad (2)$$

Denote the waiting time for packet i at node j on the path as $t_w(i, j)$. We have $t_w(i, j) = t_x(i, j) - t_r(i, j)$. The waiting time $t_w(i, j)$ contains the backoff time on node i to contend for the channel. We ignore the signal propagation time in our analysis considering the distance between the sender and receiver is small. Therefore, the end-to-end delay is also calculated as

$$t_d(i) = \sum_{j=1}^{n-1} t_w(i, j). \quad (3)$$

In practical networks, we denote $\tilde{t}_X(i, j, k)$ the measured time of operation X for packet i on node j with respect to k 's clock. We can measure

- the SourceTime $\tilde{t}_r(i, O, O)$, i.e. the generation time of packet i to the source's clock,
- the SinkTime $\tilde{t}_r(i, D, D)$, i.e. the receiving time of packet i to the sink's clock,
- the waiting time $\tilde{t}_w(i, j, j)$ for packet i on node j on a path, and
- the SourceTimeAtSink as $\tilde{t}_d(i) = \tilde{t}_r(i, D, D) - \sum_{j=1}^{n-1} \tilde{t}_w(i, j, j)$.

We summarize the parameters in Table I.

TABLE I: Notations

parameter	description
$\hat{\alpha}$	The maximum clock drift.
O	The source node of a path.
D	The sink node of a path.
$t_X(i, j)$	The time of operation X for packet i on node j with respect to the perfect clock. $X=\{x, r\}$ denotes the time the packet is transmitted and the time the packet is received. $X=w$ denotes the waiting time. On the source node, we use $t_r(i, O)$ to denote the packet generation time.
$t_d(i)$	The end-to-end delay for packet i . Thus we have $t_d(i) = t_r(i, D) - t_r(i, O)$. We also have $t_d(i) = \sum_{j=1}^{n-1} t_w(i, j)$ for a path consisting of n nodes.
$\tilde{t}_X(i, j, k)$	The measured time of operation X for packet i on node j with respect to k 's clock.
$\tilde{t}_a(i)$	The measured packet generation time for packet i with respect to the sink's clock.

For two packets i_1 and i_2 , we denote the measured time difference on node j as $\Delta\tilde{t}_X(i_1, i_2, j) = \tilde{t}_X(i_1, j, j) - \tilde{t}_X(i_2, j, j)$ and the time difference with respect to the perfect clock as $\Delta t_X(i_1, i_2, j) = t_X(i_1, j) - t_X(i_2, j)$. Denote $\Delta\tilde{t}_a = \tilde{t}_a(i_1) - \tilde{t}_a(i_2)$, we have the following theorem.

Theorem 1. (Drift Constraint) *For two packets i_1 and i_2 ($i_1 > i_2$), the measured SourceTimeAtSink satisfies*

$$(1 - 2\hat{\alpha})\Delta\tilde{t}_r(i_1, i_2, O) \leq \Delta\tilde{t}_a \leq (1 + 2\hat{\alpha})\Delta\tilde{t}_r(i_1, i_2, O). \quad (4)$$

Proof. See Appendix A. $\square$

The intuition of the drift constraint theorem is as follows. Due to a maximum drift of $\hat{\alpha}$ for each node, the maximum relative drift between the SourceTimeAtSink and SourceTime should be less than $2\hat{\alpha}$. We can check the correctness of the measured time stamps in the received packets and thus filter the incorrect time stamps. Denote the calculated delay as

$$\tilde{t}_d(i) = \tilde{t}_r(i, D, D) - \tilde{t}_a(i) \quad (5)$$

we have the following theorem.

Theorem 2. (Delay Error Bound) *The calculated delay $\tilde{t}_d(i)$ satisfies*

$$(1 - \hat{\alpha})t_d(i) \leq \tilde{t}_d(i) \leq (1 + \hat{\alpha})t_d(i). \quad (6)$$

Proof. See appendix B. $\square$

This theorem shows that the calculated delay has a bounded error. Intuitively, the delay error is introduced by measuring the waiting time on each node. As the delay $t_d(i)$ is small, the error is also very small. This is different from time synchronization protocols which need to maintain a global clock all the time with periodical beacon packets.

Further, for two packets i_1 and i_2 ($i_1 > i_2$), we assume the time stamp of SourceTimeAtSink in packet i_1 is incorrect. We can use the correct time stamps in the packet i_2 to recover the i_1 . We calculate the delay of packet i_1 by

$$t'_d(i_1) = \tilde{t}_r(i_1, D, D) - (\tilde{t}_r(i_1, O, O) - (\tilde{t}_r(i_2, O, O) - \tilde{t}_a(i_2))). \quad (7)$$

It can also be shown that the error of the recovered delay is proportional to the delay $t_d(i_1)$ and the generation time difference between two packets.

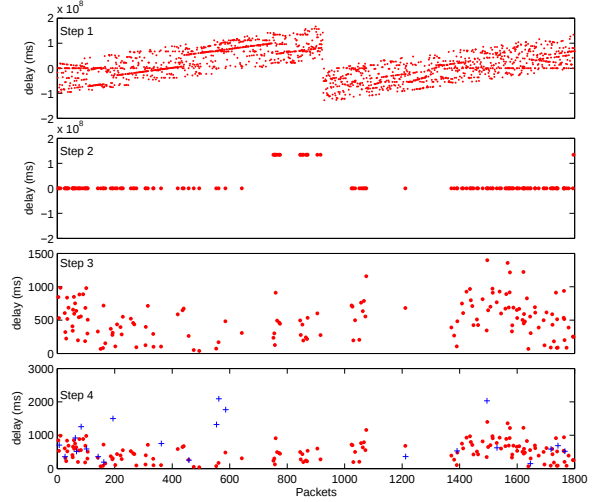


Fig. 6: Delay processing. Step 1: The original data. Step 2: Result after filtering the incorrect delays. Step 3: Result after recovering the overflow time stamps. Step 4: Result after recovering the incorrect delays based on the result in step 3. Blue crosses are the recovered delays.

B. Delay Processing

The data processing consists of the following steps.

- 1) We first calculate $\tilde{t}_d = \tilde{t}_r(i, D) - \tilde{t}_a(i)$. The calculated $\tilde{t}_d$ is shown in step 1 in Figure 6.
- 2) There exist various types of errors in the data. The first type of error comes from MLT. Due to packet overflow in the limited receiving buffer and packet losses, MLT cannot guarantee to provide correct stamps. To address this problem, we first validate the delay values and exclude the incorrect time stamps by Theorem 1. We group the delays satisfying the drift constraint into the same group. Since incorrect delays are randomly distributed, we omit those groups with much less elements than other groups. The result is shown in step 2.
- 3) After removing the incorrect delays, we find there are very large $\tilde{t}_d$. Those large values are due to time stamp overflow. In our measurement method, the SourceTimeAtSink provided by MLT is a 4-byte time stamp based on a 32KHz timer. Hence the maximum time is $t_{max} = 0xFFFFFFFF/32$ ms, i.e., about 1.5 days. Since normally $\tilde{t}_d$ is smaller than 1.5 days, we set $\tilde{t}_d = \tilde{t}_d \bmod t_{max}$. We show $\tilde{t}_d$ in step 3.
- 4) Till now, we have calculated the delays for the correct time stamps. At this step, we recover those incorrect delays according to Eq. (7) with previously received correct packets.

The final result is shown in Step 4 in Figure 6.

IV. DELAY OVERVIEW

A. Overall Distribution

Normally, the maximum single hop delay is L according to the LPL mechanism without other impacting factors, where L is the cycle length of LPL, e.g., $L = 500ms$ in our network. For a packet of k hops, the delay should be distributed between

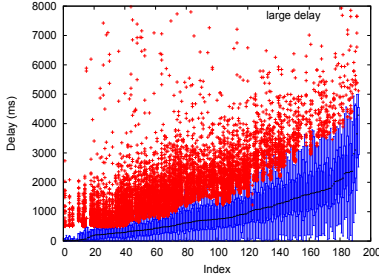


Fig. 7: Delay distribution for all nodes.

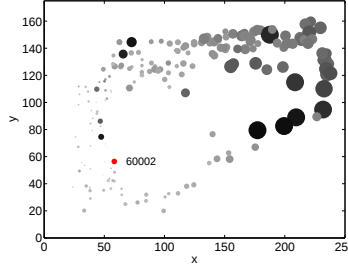


Fig. 8: Spatial distribution.

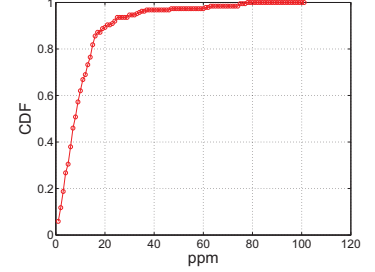


Fig. 9: CDF of drift for all nodes.

0 and kL without other impacting factors. The expected delay for such a path should be $kT/2$. Figure 7 shows the overall delay distribution for one subnet. The x-axis is the node ID and y-axis is the delay. For each node, We show the statistics of the delays with the median, 25th percentile, 75th percentile, $k \times 500$ and the lowest delay, where k is the average hop count for this node. We sort all the nodes with respect to the median delay. The red dots in the figure represent delays larger than $k \times 500$. We denote those delays as large delays. Overall, this figure presents several kinds of information. (1) The delay distribution exhibits randomness. (2) Though the delay of different nodes varies in a large range, the median delays are evenly distributed between 0 to 2 seconds. (3) There exist many large delays for most nodes. Later we will explain the reasons for those large delays.

B. Spatial Distribution

We further look at the spatial distribution of delays in the network. In Figure 8, each circle is plotted according to the physical location of different nodes. In the middle area, there is a building. The radius of each circle represents the average delay over the measurement period and the depth of the color represents the delay variation. A darker color indicates a larger variation. The red node (60002) is the sink node. We can see that nodes far away from the sink node have large average delays as well as large delay variations. Nodes in the right top area are farthest from the sink node and have the largest delays and delay variations.

C. Overall Clock Drift

Before examining the delay details, we look at the clock drift of sensor nodes. To calculate the relative clock drift, as in [19], we apply robust linear fitting for the collected time stamps and then calculate the slope as the relative clock drift for all nodes to sink node. Figure 9 shows the CDF of the relative clock drift for all nodes. The x-axis is the drift and y-axis is the CDF of nodes. More than 90% of nodes have a clock drift less than 40 ppm. This coincides with the result from [23]. The result shows that the clock drift of most nodes in the outdoor environment is relative stable.

V. ANALYSIS OF DIFFERENT FACTORS

In the data, we have various parameters from different aspects. We first collect features that may have an impact on

delay as reported in existing works, such as queue length, backoff time, parent change and etc. We also collect environment factors, e.g., temperature, humidity and light; routing parameters, e.g., routing loop; data transmission parameters, e.g., overflow, retransmission, duplicated counters; sensor node status, e.g., radio duty cycle, radio on time, task execution time and etc.

It is difficult to identify important impacting parameters to the delay performance. A single delay change may be accompanied by variations of different parameters. Moreover, the randomness introduced in LPL mechanism makes it even difficult to extract important factors. To address those issues, in this section, we leverage an automatic tool to identify important impacting factors and then investigate those important factors.

A. Important Factors

We use Rulefit [16] to find the important factors. Rulefit is a supervised learning approach to train predictors based on rule ensembles. Rulefit first trains a decision tree based on the input data. The decision tree can provide rules, each of which is a combination of one or more feature tests, i.e., combining one or more features into simple ‘and’ tests. Let x be a vector of n features and s_j be a subset of possible values for feature x_j . Then a rule takes the form of

$$r(x) = \prod_{j=1}^n I(x_j \in s_j) \quad (8)$$

where $I(\cdot)$ is an indicator function. A rule takes value one if all feature tests in the rule take value one. Rulefit provides the relative importance of different features based on rules. First, each rule is given an importance value according to its importance in the decision tree. Then the importance for a variable in a rule is calculated as the importance of the rule divided by the number of variables in the rule. The importance for a variable is the sum of the importance for the variable in all rules containing this variable.

Rulefit has two properties: (1) it can rank features by their relative importance to the prediction goal, and (2) it can provide easy-to-interpret rules (combinations of features) for user understanding. Rulefit has been adopted in recent works, e.g., [24], to understand different impacting factors. Here we present an overview of Rulefit approach. Interested readers can refer to [16] for more details.

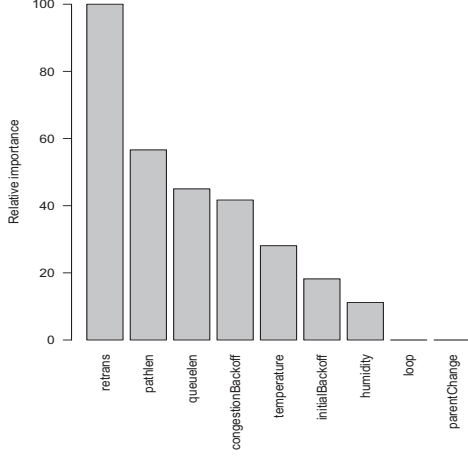


Fig. 10: Relative importance of the parameters.

We apply Rulefit to the collected parameters and delays. The ranking result of all parameters by Rulefit is shown in Figure 10. The top five most important factors are retransmission, hop count, queue length, congestion backoff and temperature. By applying Rulefit, we can filter those less important factors and then only focus on important factors for a real network.

B. Detailed Correlation

We now investigate those important factors and examine the relationship of those factors to the delay performance.

1) *Hop Count*: Figure 11 shows the delay distribution with respect to different hop counts. Overall, packets with a larger hop count to the sink node have a larger delay. More specifically, this figure shows two kinds of information. First, for hop count k , most delays (more than 80%) are less than $k \times 500$ ms, which coincides with the settings of LPL in our network. Second, this figure also shows that for each hop count, there exist many large delays, indicating other impacting factors to the delay performance.

2) *Retransmission*: Figure 12 (a) shows the average retransmission count at each hop. We can see that retransmissions per packet at different hops are almost equal. Though the traffic is high for nodes near the sink, the corresponding retransmission count for those nodes is similar to other nodes. We investigate the data and find the collisions near the sink are not severe. Note that the retransmission count for nodes at hop 1 is much lower than other nodes since the radio of the sink is always on.

Figure 12 (b) shows that the delay with different retransmission counts. It can be seen that the retransmission count and delay have a similar trend. The delay increases as the increasing of retransmission count. Meanwhile, we calculate the Pearson correlation for the retransmission and delay. The Pearson correlation between two variables X and Y is calculated as the covariance of the two variables divided by the product of their standard deviations, i.e.,

$$p_{X,Y} = \frac{\text{cov}(X,Y)}{\sigma_X \sigma_Y}. \quad (9)$$

Then we show the CDF of correlations for all nodes. As in Figure 12 (c), the x-axis is the correlation and y-axis is the

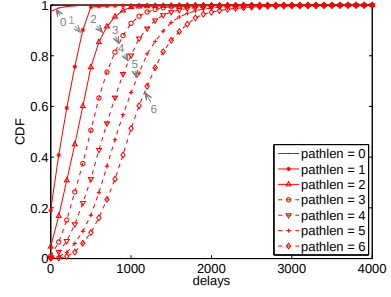


Fig. 11: Delay distributions to different hops.

CDF of nodes. We find that most nodes have high correlations between retransmission and delay.

Further, we look into those large delays. For hop k , we denote the delays larger than $k \times 500$ as large delays and other delays as normal delays. In order to correlate those large delays to retransmissions, we set a packet with retransmission count larger than 0 as a retransmission event. Then in the entire network, we calculate how retransmission events can be used to predict large delays. For a packet with a retransmission event, if such a packet has a large delay, we say the event is correlated to a large delay and call it a *true positive (TP)*. Otherwise, we call it a *false negative (FN)*. For a packet without retransmission event, if such a packet is correlated to a normal delay, we call it a *true negative (TN)*. Otherwise, we call it a *false positive (FP)*. Then we calculate the *accuracy* which gives the probability that a retransmission event can be used to predict a large delay, i.e.,

$$\text{accuracy} = \frac{TP + TN}{TP + FN + FP + TN}. \quad (10)$$

Figure 12 (d) shows the CDF of accuracy for all nodes. The x-axis is the accuracy and y-axis is the CDF of nodes. The result shows the accuracy for most nodes is high. More than 80% of nodes have an accuracy ratio higher than 60%.

Considering an extreme case when all packets have large delays, even randomly selecting a large enough subset of packets as retransmission events would lead to a high accuracy. To address such a problem, as in [25], we calculate the *balanced accuracy*, i.e.,

$$\text{balanced accuracy} = \frac{0.5 \times TP}{TP + FN} + \frac{0.5 \times TN}{TN + FP}. \quad (11)$$

Figure 12 (d) shows the CDF of balanced accuracy. We can see that for balanced accuracy, more than 70% of nodes still have a balanced accuracy higher than 0.6.

3) *Queuing*: We also examine the impact of queuing on the delay performance. There is only one queue on each node. For each packet along the path to sink node, we record the queue length when the packet arrives at each node. Figure 13 (a) shows the average queue length for different hops. Unlike the retransmission count, the average queue length varies for different hops. Nodes near the sink have larger average queue length. This indicates that though high traffic near the sink does not incur packet losses and retransmissions, it results in more congestions and thus a larger queue length. Since the sink is always on, nodes within 1 hop from sink node can quickly drain their packets and thus have a smaller queue length.

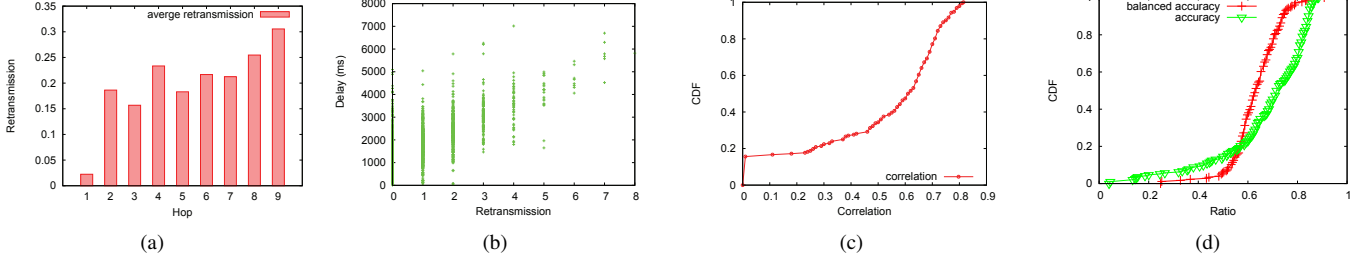


Fig. 12: Impact of retransmission to delay. (a) Hop count to average retransmissions. (b) Retransmission with respect to delay distribution. (c) CDF of correlation between retransmission and delay for all nodes. (d) Accuracy and balanced accuracy of using retransmission to predict large delay.

Figure 13 (b) shows the delay with respect to the total queue length along the path. We also find that many packets with a large queue length are correlated with large delays. We also calculate the Pearson correlation. Figure 13 (c) shows the CDF of correlation for all nodes. The x-axis is the correlation and y-axis is the CDF of nodes. We can see that for most nodes, delay has a positive correlation with the queue length. But the correlation is not as strong as that of retransmission. In LPL, the receiver is awake with a high probability after the first packet in the queue is transmitted. Thus consecutive packets may not need to wait until the receiver wakes up. Therefore, the impact of queue length on delay is relative small. We further explain the reason in the delay model in Section VI.

We further examine the predictability of queue length to large delays. Similar to retransmission, we set a packet with a queue length larger than 0 as a queuing event. Then we correlate those queuing events to large delays and calculate the accuracy and balanced accuracy. The CDF of accuracy and balanced accuracy are shown in Figure 13 (d). First, we can find that on many nodes, queuing events are correlated with large delays, e.g., about 80% of nodes have accuracy higher than 60%. For balanced accuracy, about 80% of nodes have balanced accuracy higher than 40%. In total, we can see that the correlation here is not as high as that for retransmission, which coincides with the result of Figure 13 (c).

4) *Other Factors*: From the result of Rulefit, we find that environment factors, MAC backoffs (including congestion backoff and initial backoff), and routing events (including parent change and loop event) are not as important as aforementioned factors.

It has been shown [23] [26] that environment factors such as temperature and humidity affect the clock drift and link quality. Figure 9 shows the drift of sensor nodes is relatively small and thus its impact to packet delay is also limited. The impact of environment on link quality has also been studied in different works such as [26]. Link quality is indeed related to delay performance. Such an impact is captured in retransmissions. Thus we do not consider environment as a direct impacting factor.

The impact of MAC backoffs on delay is also extensively studied in different works such as [27]. In our network, we find that the average backoff for each packet is relative small. Thus the impact of backoff is not as significant as other factors.

The impact of routing events, such as parent change and

loop events, has been studied [28] in Internet. It has been shown that those events may lead to a large end-to-end delay [28]. For different network deployments, the impacting factors may be different. The differences should be considered in protocol design. Different from Internet, our result shows that the impact of parent change is relatively small in WSNs. On the other hand, the impact of wireless link quality, retransmission and so on becomes high due to the following reasons. First, the number of those events is relative small. Second, nodes often switch among forwarders with similar hop counts. In practical protocol design, we should consider the scenario when forwarders with a similar hop count do not exist. Third, events such as loop events can be quickly recovered by the network protocol. For example, when a loop is detected in CTP, the beacon interval is decreased to minimum in order to propagate the information as quickly as possible, alleviating the impact of loop event on end-to-end delay.

C. Apply Rulefit to other Networks

To further show that our method can be applied in other networks, we implement our method with the recent proposed received-initiated low duty cycle protocol A-MAC [29]. We evaluate the protocol in an indoor network. Figure 14 shows the result of Rulefit. We can see that retransmission and queue length are still two of the most significant impacting factors. Different from the result in Figure 10, the impact of other factors is much lower. We investigate the data and find that there are less contentions for indoor environment with a small network size. Meanwhile, the temperature for the environment is more stable than the outdoor environment. The loop event and parent change event are also very rare. Nevertheless, we can still obtain impacting factors using Rulefit for A-MAC.

VI. DELAY MODEL

According to the analysis of different factors, in this section, we build a model for the end-to-end packet delay and validate it in our network.

A. Model

We first model the the single hop packet delay and then extend it to multi-hop end-to-end delay. To derive the model, we describe the following parameters.

- t_s : the sleep time in LPL.

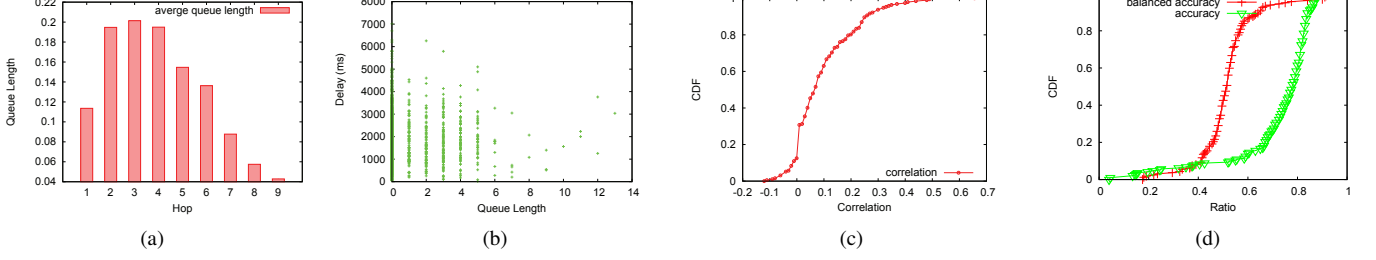


Fig. 13: Impact of queuing events to delay. (a) Queue length to average retransmissions. (b) Queue length with respect to delay distribution. (c) CDF of correlation between queue length and delay for all nodes. (d) Accuracy and balanced accuracy of using queuing events to predict large delay.

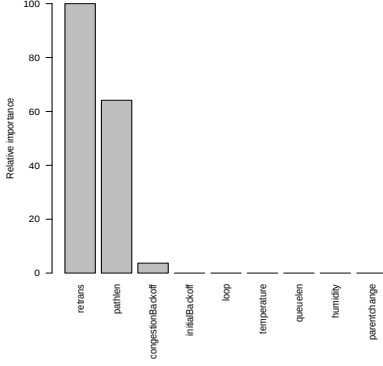


Fig. 14: Impacting factors of A-MAC.

- t_w : the awake time of the receiver in LPL.
- t_b : MAC layer backoff time, including the initial backoff time and congestion backoff time.
- u : the duty cycle ratio of the receiver, i.e., the awake time divided by the sleep time.
- t_x : the packet modulation/demodulation time. This is usually platform-dependent and is related to the packet size.
- r : the number of retransmissions for a packet.

We first calculate the delay for a single hop packet transmission. Assume the packet is retransmitted r times. According to the mechanism of LPL, each unsuccessful transmission takes of time of length $t_w + t_s$. Thus the time used for r retransmissions is $r(t_w + t_s)$. For the $(r+1)th$ transmission which is successful, there are two cases in LPL:

- **Case 1**: if the receiver is sleeping, the sender should wait until the receiver wakes up and then send the packet to the receiver.
- **Case 2**: otherwise, the packet can be sent directly.

For case 1, since the transmission can fall into any time during the sleeping period of the receiver, the delay is $U(0, t_s) + t_b + t_x$, where $U(0, t_s)$ is a random distribution between 0 and t_s . For case 2, the delay for a packet is $t_b + t_x$. According to the duty cycle ratio of each receiver, the probability for case 1 is $1 - u$ and for case 2 is u .

Consequently, the 1-hop delay is given by:

$$T(t_s, t_w, r) = \begin{cases} r(t_w + t_s) + t_b + t_x & \text{with prob. } u \\ r(t_w + t_s) + U(0, t_s) + t_b + t_x & \text{otherwise} \end{cases} \quad (12)$$

Based on the single hop delay, we derive the delay for a multi-hop path. A packet p is transmitted on a path consisting of n nodes from node 1 to n . At each node, the packet is first put into the transmission queue and then transmitted after prior packets in the queue are transmitted. To calculate the multi-hop delay, we first describe the following parameters.

- l_i : queue length at node i , i.e., the number of packets in the transmission queue, including the packet p . The packet needs to wait until prior $l_i - 1$ packets are transmitted.
- $r_{i,j}$: the number of retransmissions for the j -th packet in the queue on node i .

The time for the first packet in the queue can be calculated according to Eq. 12. We then calculate the time for following packets in the queue. In LPL, the receiver should stay awake after receiving a packet. When $r = 0$, the packet in the queue is directly sent to the receiver since the receiver is awake. Otherwise, the packet is retransmitted since the receiver is sleeping. Thus for packets except the first one in the queue, the delay is

$$T_q(t_s, t_w, r) = \begin{cases} t_b + t_x & r = 0 \\ T(t_s, t_w, r) & \text{otherwise} \end{cases} \quad (13)$$

Therefore, the delay for a path consisting of n nodes $(1, 2, \dots, n)$ is given by

$$D(n) = \sum_{i=1}^{n-1} T(t_s^{i+1}, t_w^{i+1}, r_{i,l_i}) + \sum_{i=1}^{n-1} \sum_{j=1}^{l_i-1} T_q(t_s^{i+1}, t_w^{i+1}, r_{i,j}) \quad (14)$$

where $T(t_s^{i+1}, t_w^{i+1}, r_{i,l_i})$ is the single hop delay for the first packet in the queue at node i , $T_q(t_s^{i+1}, t_w^{i+1}, r_{i,j})$ is the delay for the j th packet in the queue, t_s^{i+1} is the sleep time and t_w^{i+1} is the awake time of node $i+1$.

B. Model Validation

We validate our model with the collected data. In CitySee, we have recorded the queue length on each node of the path, i.e., l_i . We also recorded the ETX value from each node to its neighbors. This can be used to approximate $r_{i,j}$, for $j > 1$. For each packet, we also recorded the retransmissions on each hop which is $r_{i,j}$, for $j = 1$. On each node, we calculate the average duty cycle u and average backoff time t_b . Using those parameters as input, we calculate the delay with the model. Then we compare the result of the model with the delay calculated from the packets. The result is shown

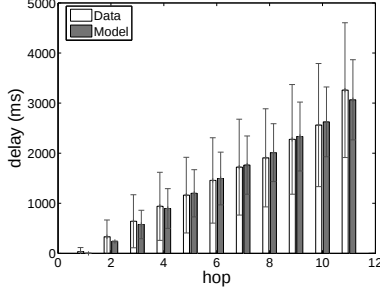


Fig. 15: Delay distribution of the model and collected data with respect to different hop counts.

in Figure 15. We show the average delay and variations for different hops. We can see that both results increase linearly with the hop count. The model and practical delay have very similar distribution, showing that our model is effective to capture those important factors.

VII. IMPLICATIONS

In this section, we revisit three prevalent protocols and discuss the implications of our analysis and delay model.

A. Routing Protocol

We first analyze the commonly used data collection protocol CTP in WSNs. Through analysis, we find that CTP protocol, with ETX as the routing metric, may not appropriately choose a good path. For brevity, we assume the queuing delay on each hop is 0, i.e., $l_i = 1$ on each hop. Then we have

$$D(n) = \sum_{i=1}^{n-1} T(i_s^{i+1}, t_w^{i+1}, r_{i,1})$$

According to Eq. 12, the expected delay is

$$E(D(n)) = E\left(\sum_{i=1}^{n-1} (r_{i,1}(t_w + t_s) + t_b + t_x)\right) + (1-u)E\left(\sum_{i=1}^{n-1} U(0, t_s)\right)$$

Denote ETX_i as the ETX for the link from node i to $i+1$. According to the definition of ETX, we have $ETX_i = E(r_{i,1}) + 1$. The expected delay on the path is

$$\begin{aligned} E(D(n)) &= \sum_{i=1}^{n-1} ((ETX_i - 1)(t_w + t_s) + t_b + t_x) + (1-u)(n-1)\frac{t_s}{2} \\ &= PathETX(t_w + t_s) - (n-1)(t_w + \frac{1+u}{2}t_s - t_b - t_x). \end{aligned} \quad (15)$$

We can see that even for paths with the same path ETX, the expected end-to-end delay can be different. In practical settings, the awake time t_w and sleep time t_s are usually larger than the backoff time t_b and data transmission time t_x . Thus we have $(n-1)(t_w + \frac{1+u}{2}t_s - t_b - t_x) > 0$. According to Eq. 15, for the same path ETX, $E(D(n))$ decreases as the hop count n increases. This indicates that longer paths, which are often prohibited previously, may even be better than shorter ones.

Considering a simple example for two paths with the same ETX of 2: the first path consists of one link with a link ETX 2; the second path consists of two links, each of which has a link ETX 1. CTP treats those two paths equally and thus randomly chooses one as the routing path (in current implementation, the one appears earlier in the routing table is

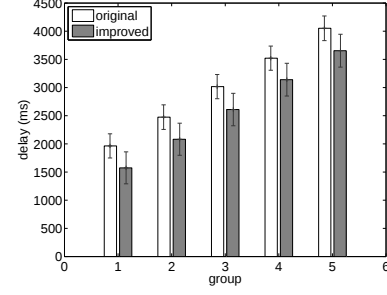


Fig. 16: Comparison of EDET with CTP protocol

selected). In fact, according to the delay model, the second path is better than the first one in terms of delay, even though the second one has a larger hop count. For the first path, there are 2 transmissions in expectation including 1 retransmission. For the second path, there are 2 transmissions in expectation without retransmissions. Intuitively, according to the mechanism of LPL, if there is no retransmission, the expected time is $L/2$, where L is the cycle in LPL. However, if a packet is retransmitted, the expected time is L . Hence, the second path is better. This also explains in Figure 10 why retransmission count is more important than hop count. In current design, the ETX metric only counts the expectation of transmissions and makes no difference for retransmissions and successful transmissions.

Accordingly, we propose a new path quality metric, i.e., expected delay based on ETX (EDET). According to equation 12, the $EDET_{i,j}$ for a single hop from node i to node j is calculated as $EDET_{i,j} = (ETX_{i,j} - 1)(t_w + t_s) + t_b + t_x + (1-u)t_s/2$, where $ETX_{i,j}$ is the ETX from node i to node j . Those parameters can be calculated by each node locally. Based on the single hop EDET, each node can calculate the path EDET to the sink node, which is the sum of link EDETs on the path. We improve the original CTP protocol with the EDET metric. In the improved protocol, to minimize the transmission cost, each node first selects a parent with minimal ETX. Then the node will select a parent with minimal EDET. Thus EDET will not incur additional transmission overhead. It is worth noting EDET can also be used without ETX to improve delay performance. Figure 16 shows the comparison result for the improved protocol and CTP. The x-axis is the group of packets with the same ETX. We can see that for different groups, the average delay with EDET is reduced.

B. Opportunistic Routing

Opportunistic routing can significantly improve system performance [30] [31]. We evaluate the practical performance of opportunistic routing with real trace from CitySee. For each node, we calculate the probability that the delay performance can be improved, i.e., the portion of neighbors with smaller delays than the node itself. We evaluate two different strategies as shown in Figure 17. First, we evaluate the strategy that a node can select the next forwarder in all neighbors. For more than 50% of nodes, the probability to select a better node is higher than 50%. However, opportunistic routing in

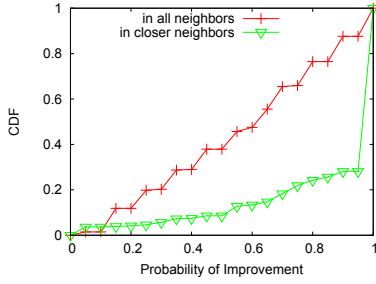


Fig. 17: Evaluation of opportunistic routing

all neighbors may still result in selection of a forwarder with a larger delay.

Further, we investigate the second strategy. Each node only selects neighbors closer to the sink node. We can see that it is better to select nodes from neighbors with smaller hop counts to the sink node. For more than 80% of nodes, the probability of improvement is larger than 95%.

C. Dynamic Switching based Forwarding

Dynamic switching based forwarding (DSF) [32] is proposed to optimize the expected end-to-end delay in WSNs by selecting a set of forwarders. In DSF, each node searches in its neighbors to find an optimal forwarding set with minimum expected end-to-end delays. More specifically, a node searches through its neighbors and calculates whether it is beneficial to add each neighbor into the forwarding set.

For each node, we calculate its link quality q_i to each neighbor i and the delay T_i of neighbor i . Assume the forwarding set is $F = \{v_1, v_2, \dots, v_{|F|}\}$ and $v_1, v_2, \dots, v_{|F|}$ are in ascending order of waking up time. For brevity, for a node which wakes up multiple times, we let it appear multiple times in the set. Then the expected delay is

$$D(F) = \left(\sum_{i=1}^{|F|} \prod_{k=0}^{i-1} (1 - q_k) q_i T_i \right), \quad (16)$$

where $q_0 = 0$.

DSF uses dynamic programming to find F such that $D(F)$ is minimized with a certain packet reception ratio being satisfied. We compare the delay of DSF using our trace with the practical measured packet delay in our network. Figure 18 shows the delay difference of measured delay in our network and the delay of applying DSF to our trace, i.e. practical delay—DSF delay. First, for most packets, the difference is larger than 0. We can see that for more than 75% of packets the practical delay is larger than the DSF delay, indicating that using DSF is beneficial in terms of delay performance. Meanwhile, note that the improvement for most nodes is less than 500 ms, indicating further improvement is possible. In LPL protocol, we can extend the delay model to the network without fixed scheduling for different nodes. In this case, as in Eq. 12, the probability for retransmission count r can be calculated [33] as

$$p(r=k) = \prod_{i \in F} (1 - q_i)^k \left(1 - \prod_{i \in F} (1 - q_i) \right) \quad (17)$$

Thus the expected number of retransmission can be calculated as

$$E(r) = \frac{\prod_{i \in F} (1 - q_i)}{1 - \prod_{i \in F} (1 - q_i)} \quad (18)$$

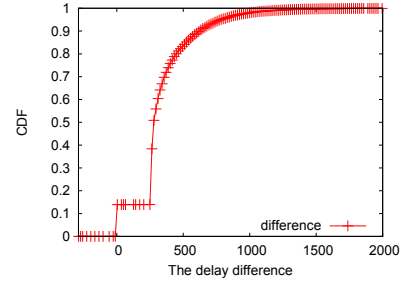


Fig. 18: Difference of the practical packet delay and DSF.

Then the single hop delay can be calculated as

$$D_{e,F} = E(r)(t_w + t_s) + t_b + t_x + \mu t_s \quad (19)$$

where e is the sender, μt_s is the expected delay for a successful transmission. We can further calculate the multiple hop delay and improve the performance according to techniques described in Section VI.

VIII. RELATED WORK

Delay Analysis. There are extensive works for delay performance analysis in WSNs. First, probabilistic delay bounds are proposed in [5] [6] [7] [8] by extending network calculus. In the second category, stochastic delay models are proposed. For example, in [9] [10] [11], different models are proposed by combining real-time theory and queuing theory. In those models, unreliable networks with heavy traffic are considered. There are some empirical network delay models [13] [14] proposed for end-to-end delay measurements. A delay model based on Discrete Markov Processing in the network is proposed in [12]. Besides those end-to-end delay models, the single hop channel access delay models are also analyzed in [27] [34] [35]. However, those works proposed in WSNs are often based on assumptions, e.g., traffic, routing path, and not evaluated in a real large-scale network. In this work, we are the first to propose a light-weight delay measurement and analysis in an operational LPL WSN.

Delay in Internet/Data Center. There are also a large number of research works in Internet and data centers. Pucha et al. [28] show the impact of routing events on end-to-end delay in Internet. Kompella et al. [3] present fine grain latency measurements in presence of packet losses for internet with a lossy difference aggregator. To measure the per-flow delay, Lee et al. [15] present a measurement method with reference delay interpolation. As the development of data center technologies, Wilson et al. [4] present delay analysis results in the data center and provide guidelines to the data center design.

Time Synchronization. There are also many global time synchronization methods in wireless sensor networks, e.g., [19] [36] [20]. Those methods can synchronize all nodes in the network and provide synchronized time stamps. Our measurement method is different from those methods in two aspects. First, our method does not require time synchronization among all nodes in the network. We recover the time stamp at the sink side. Thus this reduces the message exchange overhead among all nodes. Second, our method does not need to maintain a synchronized time stamp all the time. Therefore, our method does not require periodical message exchange.

Large Scale Sensor Network Deployment. There are many sensor network deployments in the world. Table II shows several network deployments in the world.

TABLE II: WSN deployments

System	Scale	Power supply	environment
Great Duck island [37]	100+	Battery	Outdoor
VigilNet [38]	200	Battery	Outdoor
MoteLab [39]	190	Tethered	Indoor
SensorScope [40]	97	Solar	Outdoor
Trio [41]	557	Solar/Battery	Outdoor
CitySee	1200	Battery	Outdoor

IX. CONCLUSION

We focus on delay measurement, analysis and implications in an operational large-scale LPL wireless sensor network. We propose a light-weight delay measurement method to efficiently calculate delay without time synchronization, which is applicable to operational networks. We analyze the spatial and temporal characteristics of delay distribution through carefully examine system metrics. Further, we extract different impacting parameters to delay performance with the incomplete data, propose a practical delay model to capture those factors and validate it in a large-scale network. Finally, we show the implications of measurement and analysis to protocol design.

REFERENCES

- [1] T. B. O. Chipara, C. Lu and G.-C. Roman, "Reliable clinical monitoring using wireless sensor networks: Experience in a step-down hospital unit," in *Proceedings of ACM SenSys*, 2010.
- [2] S. N. Pakzad, G. L. Fenves, S. Kim, and D. E. Culler, "Design and implementation of scalable wireless sensor network for structural monitoring," in *Journal of Infrastructure Systems*, 2008.
- [3] R. R. Kompella, K. Levchenko, A. C. Snoeren, and G. Varghese, "Every microsecond counts: Tracking fine-grain latencies with a lossy difference aggregator," in *Proceedings of ACM SIGCOMM*, 2009.
- [4] C. Wilson, H. Ballani, T. Karagiannis, and A. Rowstron, "Better never than late: Meeting deadlines in datacenter networks," in *Proceedings of ACM SIGCOMM*, 2011.
- [5] M. Fidler, "An end-to-end probabilistic network calculus with moment generating functions," in *Proceedings of IEEE IWQoS*, 2006.
- [6] A. Koubaa, M. Alves, and E. Tovar, "Modeling and worst-case dimensioning of cluster-tree wireless sensor networks," in *Proceedings of IEEE RTSS*, 2006.
- [7] A. Burchard, J. Liebeherr, and S. Patek, "A min-plus calculus for end-to-end statistical service guarantees," *IEEE Transactions on Information Theory*, vol. 52, no. 9, pp. 4105–4114, 2006.
- [8] J. Schmitt and U. Roedig, "Sensor network calculus - a framework for worst case analysis," in *Proceedings of IEEE DCOSS*, 2005.
- [9] J. Lehoczy, "Real-time queueing theory," in *Proceedings of IEEE RTSS*, 1996.
- [10] J. P. Lehoczy, "Real-time queueing network theory," in *Proceedings of IEEE RTSS*, 1997.
- [11] S.-N. Yeung and J. Lehoczy, "End-to-end delay analysis for real-time networks," in *Proceedings of IEEE RTSS*, 2001.
- [12] Y. Wang, M. C. Vuran, and S. Goddard, "Cross-layer analysis of the end-to-end delay distribution in wireless sensor networks," *IEEE/ACM Transactions on Networking*, vol. 20, no. 1, pp. 305–318, 2012.
- [13] E. Felemban, C.-G. Lee, E. Ekici, R. Boder, and S. Vural, "Probabilistic qos guarantee in reliability and timeliness domains in wireless sensor networks," in *Proceedings of IEEE INFOCOM*, 2005.
- [14] K. Gopalan, T.-c. Chiueh, and Y.-J. Lin, "Probabilistic delay guarantees using delay distribution measurement," in *Proceedings of ACM MULTIMEDIA*, 2004.
- [15] M. Lee, N. G. Duffield, and R. R. Kompella, "Not all microseconds are equal: Enabling per-flow measurements with reference latency interpolation," in *Proceedings of ACM SIGCOMM*, 2010.
- [16] J. Friedman and B. Popescu, "Predictive learning via rule ensembles," *Annals of Applied Statistics*, vol. 2, no. 3, pp. 916–954, 2008.
- [17] O. Gnawali, R. Fonseca, K. Jamieson, D. Moss, and P. Levis, "Collection tree protocol," in *Proceedings of ACM SenSys*, 2009.
- [18] D. Couto, D. Aguayo, J. Bicket, and R. Morris, "A high-throughput path metric for multi-hop wireless routing," in *Proceedings of ACM MobiCom*, 2003.
- [19] M. Maróti, B. Kusy, G. Simon, and Ákos Lédeczi, "FTSP: The Flooding Time Synchronization Protocol," in *Proceedings of ACM SenSys*, 2004.
- [20] J. Elson, L. Girod, and D. Estrin, "Fine-grained network time synchronization using reference broadcasts," in *Proceedings of USENIX OSDI*, 2002.
- [21] F. Ferrari, M. Zimmerling, L. Thiele, and O. Saukh, "Efficient Network Flooding and Time Synchronization with Glossy," in *Proceedings of ACM/IEEE IPSN*, 2011.
- [22] Y. Wang, Y. He, X. Mao, Y. Liu, Z. Huang, and X. Li, "Exploiting constructive interference for scalable flooding in wireless networks," in *Proceedings of IEEE INFOCOM*, 2012.
- [23] T. Schmid, Z. Charbiwala, J. Friedman, Y. H. Cho, and M. B. Srivastava, "Exploiting manufacturing variations for compensating environment-induced clock drift in time synchronization," in *Proceedings of the ACM SIGMETRICS*, 2008.
- [24] I. Cunha, R. Teixeira, D. Veitch, and C. Diot, "Predicting and tracking internet path changes," in *Proceedings of ACM SIGCOMM*, 2011.
- [25] I. Cohen, M. Goldszmidt, T. Kelly, J. Symons, and J. S. Chase, "Correlating instrumentation data to system states: a building block for automated diagnosis and control," in *Proceedings of USENIX OSDI*, 2004.
- [26] W. Xi, Y. He, Y. Liu, J. Zhao, L. Mo, Z. Yang, J. Wang, and X. Li, "Locating sensors in the wild: pursuit of ranging quality," in *Proceedings of ACM SenSys*, 2010.
- [27] G. Bianchi, "Performance analysis of the ieee 802.11 distributed coordination function," *IEEE Journal on Selected Areas in Communications*, vol. 18, no. 3, pp. 535–547, 2000.
- [28] H. Pucha, Y. Zhang, Z. M. Mao, and Y. C. Hu, "Understanding network delay changes caused by routing events," in *Proceedings of ACM SIGMETRICS*, 2007.
- [29] P. Dutta, S. Dawson-Haggerty, Y. Chen, C.-J. M. Liang, and A. Terzis, "Design and evaluation of a versatile and efficient receiver-initiated link layer for low-power wireless," in *Proceedings of ACM SenSys*, 2010.
- [30] S. Chachulski, M. Jennings, S. Katti, and D. Katabi, "Trading structure for randomness in wireless opportunistic routing," in *Proceedings of ACM SIGCOMM*, 2007.
- [31] S. Biswas and R. Morris, "Exor: Opportunistic multi-hop routing for wireless networks," in *Proceedings of ACM SIGCOMM*, 2005.
- [32] Y. Gu and T. He, "Dynamic switching-based data forwarding for low-duty-cycle wireless sensor networks," *IEEE Transactions on Mobile Computing*, vol. 10, pp. 1741–1754, 2011.
- [33] O. Landsiedel, E. Ghadimi, S. Duquennoy, and M. Johansson, "Low power, low delay: opportunistic routing meets duty cycling," in *Proceedings of ACM/IEEE IPSN*, 2012.
- [34] T. Sakurai and H. Vu, "Mac access delay of ieee 802.11 dcf," *IEEE Transactions on Wireless Communications*, vol. 6, no. 5, pp. 1702–1710, 2007.
- [35] O. Tickoo and B. Sikdar, "Modeling queueing and channel access delay in unsaturated ieee 802.11 random access mac based wireless networks," *ACM/IEEE Transactions on Networking*, vol. 16, no. 4, pp. 878–891, 2008.
- [36] P. Sommer and R. Wattenhofer, "Gradient clock synchronization in wireless sensor networks," in *Proceedings of ACM/IEEE IPSN*, 2009.
- [37] R. Szewczyk, A. Mainwaring, J. Polastre, J. Anderson, and D. Culler, "An analysis of a large scale habitat monitoring application," in *Proceedings of SenSys*, 2004.
- [38] T. He, S. Krishnamurthy, L. Luo, T. Yan, L. Gu, R. Stoleru, G. Zhou, Q. Cao, P. Vicaire, J. A. Stankovic, T. F. Abdelzaher, J. Hui, and B. Krogh, "Vigilnet: An integrated sensor network system for energy-efficient surveillance," *ACM Transactions on Sensor Networks*, vol. 2, no. 1, pp. 1–38, 2006.
- [39] G. Werner-Allen, P. Swieskowski, and M. Welsh, "Motelab: a wireless sensor network testbed," in *Proceedings of ACM/IEEE IPSN*, 2005.
- [40] G. Barrenetxea, F. Ingelrest, G. Schaefer, M. Vetterli, O. Couach, and M. Parlange, "Sensorscope: Out-of-the-box environmental monitoring," in *Proceedings of ACM/IEEE IPSN*, 2008.
- [41] P. Dutta, J. Hui, J. Jeong, S. Kim, C. Sharp, J. Taneja, G. Tolle, K. Whitehouse, and D. Culler, "Trio: enabling sustainable and scalable outdoor wireless sensor network deployments," in *Proceedings of ACM/IEEE IPSN*, 2006.

APPENDIX A

Proof. According to the linear clock model, we have

$$\Delta \tilde{t}_r(i_1, i_2, O) = (1 + \alpha_1) \Delta t_r(i_1, i_2, O) \quad (20)$$

$$\Delta \tilde{t}_r(i_1, i_2, D) = (1 + \alpha_2) \Delta t_r(i_1, i_2, D) \quad (21)$$

$$\sum_{j=1}^{n-1} \tilde{t}_w(i_1, j, j) = (1 + \alpha_3) \sum_{j=1}^{n-1} t_w(i_1, j), \quad (22)$$

$$\sum_{j=1}^{n-1} \tilde{t}_w(i_2, j, j) = (1 + \alpha_4) \sum_{j=1}^{n-1} t_w(i_2, j), \quad (23)$$

where α_1 , α_2 , α_3 and α_4 are clock drifts bounded by $[-\hat{\alpha}, \hat{\alpha}]$. Thus we have

$$\Delta \tilde{t}_a = (\tilde{t}_r(i_1, D, D) - \sum_{j=1}^{n-1} \tilde{t}_w(i_1, j, j)) - (\tilde{t}_r(i_2, D, D) - \sum_{j=1}^{n-1} \tilde{t}_w(i_2, j, j)) \quad (24)$$

According to Eq. (20)-(23), we have

$$\begin{aligned} \Delta \tilde{t}_a &= (\tilde{t}_r(i_1, D, D) - \sum_{j=1}^{n-1} \tilde{t}_w(i_1, j, j)) - (\tilde{t}_r(i_2, D, D) - \sum_{j=1}^{n-1} \tilde{t}_w(i_2, j, j)) \\ &= (1 + \alpha_2) \Delta t_r(i_1, i_2, D) \\ &\quad - (1 + \alpha_3) \sum_{j=1}^{n-1} t_w(i_1, j) + (1 + \alpha_4) \sum_{j=1}^{n-1} t_w(i_2, j) \end{aligned} \quad (25)$$

Since α_2 , α_3 and α_4 are bounded by $[\hat{\alpha}, \hat{\alpha}]$, we have

$$\Delta \tilde{t}_a \leq (1 + \hat{\alpha}) \Delta t_r(i_1, i_2, D) - (1 - \hat{\alpha}) \sum_{j=1}^{n-1} t_w(i_1, j) + (1 + \hat{\alpha}) \sum_{j=1}^{n-1} t_w(i_2, j) \quad (26)$$

According to Eq. (2) and (3), we have

$$\Delta \tilde{t}_a \leq (1 + \hat{\alpha}) \Delta t_r(i_1, i_2, O) + 2\hat{\alpha} \sum_{j=1}^{n-1} t_w(i_1, j) \quad (27)$$

Since $\hat{\alpha}$ is in the magnitude of 10^{-6} and $\sum_{j=1}^{n-1} t_w(i_1, j)$ is in the magnitude of 10^{-3} seconds, which are much smaller than $(1 + \hat{\alpha}) \Delta t_r(i_1, i_2, O)$, we ignore the second term and obtain

$$\Delta \tilde{t}_a \leq (1 + \hat{\alpha}) \Delta t_r(i_1, i_2, O) \quad (28)$$

Combining Eq. (20), for small $\hat{\alpha}$, we have

$$\begin{aligned} \Delta \tilde{t}_a &\leq \frac{(1 + \hat{\alpha}) \Delta \tilde{t}_r(i_1, i_2, O)}{1 - \hat{\alpha}} \\ &\approx (1 + 2\hat{\alpha}) \Delta \tilde{t}_r(i_1, i_2, O). \end{aligned} \quad (29)$$

Similarly, we have

$$\Delta \tilde{t}_a \geq (1 - 2\hat{\alpha}) \Delta \tilde{t}_r(i_1, i_2, O). \quad (30)$$

Combining Eq. (29) and (30), we prove the theorem. $\square$

APPENDIX B

Proof. According to Eq. (5), we have

$$\begin{aligned} \tilde{t}_d(i) &= (\tilde{t}_r(i, D, D) - \tilde{t}_a(i)) \\ &= \sum_{j=1}^{n-1} \tilde{t}_w(i, j, j) \end{aligned} \quad (31)$$

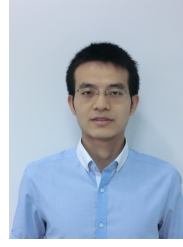
Based on Eq. (22), we have

$$\begin{aligned} \tilde{t}_d(i) &\leq (1 + \hat{\alpha}) \sum_{j=1}^{n-1} t_w(i, j) \\ &= (1 + \hat{\alpha}) t_d(i) \end{aligned} \quad (32)$$

Similarly, we have

$$\tilde{t}_d(i) \geq (1 - \hat{\alpha}) t_d(i) \quad (33)$$

Combining Eq. (32) and (33), we prove the theorem. $\square$



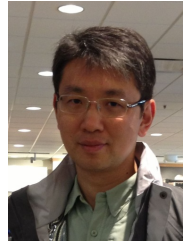
Jiliang Wang received his BE degree from University of Science and Technology of China and his PhD degree from Hong Kong University of Science and Technology. He is currently with School of Software and TNLIST, Tsinghua University. His research interests include wireless sensor networks, network measurement and pervasive computing. He is a member of the IEEE.



Wei Dong received his BS and PhD degrees from the College of Computer Science in Zhejiang University in 2005 and 2011, respectively. He is now an associate professor in the College of Computer Science in Zhejiang University. His research interests include networked embedded systems and wireless sensor networks. He is a member of the IEEE.



Zhichao Cao received the BE degree in Tsinghua University, and his PhD degree in Hong Kong University of Science and Technology. His current research interest is wireless sensor networks and mobile network. He is a member of the IEEE.



Yunhao Liu received the BS degree in automation from Tsinghua University, China, in 1995, the MS and PhD degrees in computer science and engineering from Michigan State University, in 2003 and 2004, respectively. He is now Cheung Kong Professor at Tsinghua University. His research interests include RFID and sensor network, the Internet, and pervasive computing.