

Estimating Cardinality of Arbitrary Expression of Multiple Tag Sets in a Distributed RFID System

Qingjun Xiao¹, Member, IEEE, ACM, Youlin Zhang², Member, IEEE, Shigang Chen³, Fellow, IEEE, Min Chen, Jia Liu, Member, IEEE, Guang Cheng, Senior Member, IEEE, and Junzhou Luo, Member, IEEE

Abstract—Radio-frequency identification (RFID) technology has been widely adopted in various industries and people's daily lives. This paper studies a fundamental function of spatial-temporal joint cardinality estimation in distributed RFID systems. It allows a user to make queries over multiple tag sets that are present at different locations and times in a distributed tagged system. It estimates the joint cardinalities of those tag sets with bounded error. This function has many potential applications for tracking product flows in large warehouses and distributed logistics networks. The prior art is either limited to jointly analyzing only two tag sets or is designed for a relative accuracy model, which may cause unbounded time cost. Addressing these limitations, we propose a novel design of the joint cardinality estimation function with two major components. The first component is to record snapshots of the tag sets in a system at different locations and periodically, in a time-efficient way. The second component is to develop accurate estimators that extract the joint cardinalities of chosen tag sets based on their snapshots, with a bounded error that can be set arbitrarily small. We formally analyze the bias and variance of the estimators, and we develop a method for setting their optimal system parameters. The simulation results show that, under predefined accuracy requirements, our new solution reduces time cost by multiple folds when compared with the existing work.

Index Terms—RFID, cardinality estimation, set expression.

Manuscript received August 31, 2017; revised April 8, 2018 and September 30, 2018; accepted January 18, 2019; approved by IEEE/ACM TRANSACTIONS ON NETWORKING Editor Y. Guan. Date of publication February 15, 2019; date of current version April 16, 2019. This work was supported in part by the National Key R&D Program of China under Grant 2018YFB0803400 and Grant 2017YFB1003000, in part by the National Natural Science Foundation of China under Grant 61872080 and Grant 61702257, in part by the National Science Foundation of the United States under Grant CNS-1718708, in part by a grant from Cyber Florida, and in part by the Key Lab of Computer Network and Information Integration of the Ministry of Education of China 93K-9. The preliminary version of this paper titled "Joint property estimation for multiple RFID tag sets using snapshots of variable lengths" was published in proceedings of the ACM MobiHoc (International Symposium on Mobile Ad Hoc Networking and Computing), pp. 151–160, 2016. (Corresponding author: Qingjun Xiao.)

Q. Xiao and G. Cheng are with the School of Cyber Science and Engineering, Jiangsu Key Laboratory of Computer Networking Technology, Southeast University, Nanjing 211189, China (e-mail: csqjxiao@seu.edu.cn; gcheng@njnet.edu.cn).

Y. Zhang and S. Chen are with the Department of Computer and Information Science and Engineering, University of Florida, Gainesville, FL 32611 USA, (e-mail: youlin@cise.ufl.edu; sgchen@cise.ufl.edu).

M. Chen was with the Department of Computer and Information Science and Engineering, University of Florida, Gainesville, FL 32611 USA. He is now with Google Inc., Mountain View, CA 94043 USA (e-mail: minchen@google.com).

J. Liu is with the State Key Laboratory of Novel Software Technology, Nanjing University, Nanjing 210023, China (e-mail: jialiu@nju.edu.cn).

J. Luo is with the School of Computer Science and Engineering, Key Laboratory of Computer Network and Information Integration, Ministry of Education, Southeast University, Nanjing 211189, China (e-mail: jluo@seu.edu.cn).

Digital Object Identifier 10.1109/TNET.2019.2894729

I. INTRODUCTION

OVER the past decade, radio-frequency identification (RFID) technology has been widely used by industries such as warehouse management, logistical control, and asset tracking in malls, hospitals or highways [1]. RFID systems can be conceptually divided into two parts: RFID tags (each carrying a unique ID) which are attached to physical objects, and RFID readers, which are deployed at places of interest to sense the existence of tags, quickly retrieve the tag IDs, or collect aggregate statistical information about a group of tags.

An important fundamental functionality of RFID system is called *cardinality estimation*, which is to count the number of tags in a physical region [1]–[9]. This function can be used to monitor the inventory level of a warehouse, the sales in a retail store, or the popularity of a theme park. Counting the number of tags can take much less time than a full system scan that collects all tag IDs. This is an important feature since RFID systems communicate via low-rate wireless channels and the execution time cost is the key performance metric in system design. In addition to its direct utility, tag estimation can work as a pre-processing step that improves the efficiency of tag identification process [10], [11]. More importantly, since it does not collect any tag IDs, the anonymity of tags can be preserved, particularly in scenarios where the party performing the operation (such as warehouse or port authority) does not own the tagged items.

Motivation

The previous works study relatively simple scenarios, i.e., estimate the cardinality of a tag set in the radio range of one reader [1]–[8], or estimate the union of tag sets near multiple readers [2], [8]. Our paper studies the cardinality estimation problem in a more generalized scenario: Multiple tag sets are captured by a distributed multi-reader system at different spatial or temporal domains. As requested by users, these tag sets may form any set expression using the operators of union ($\cup$), intersection ($\cap$) and relative complement ($\setminus$). We estimate the cardinality of any user-desired set expression, which is called *joint cardinality* of multiple sets.

For this joint cardinality estimation problem, we use two applications to better illustrate its usefulness. Just imagine that we are to manage a large logistics network, where tagged products are shipped from one location (factory, warehouse, port, or storage/retail facility) to another. Assume the reader, pre-deployed at each location, takes periodic snapshots of its local set of tags and keeps a series of such snapshots over time. The end users may want to know the quantity of goods flowing from one location to another. Clearly, we can address the query

by estimating the cardinality of *intersection between two snapshots* from different locations. However, a more complicated user query can be the quantity of goods traversing a routing path comprised of multiple locations. Then, to address the query, we may need to compute the cardinality of *intersection among more than two tag set snapshots*.

In the second application, imagine that we are monitoring a warehouse for its inventory dynamics over time. Users may want to know the amount of goods entering the warehouse (or called the number of new tags), and the amount of goods leaving (or the number of departed tags), between *two* user-defined time points. Suppose an RFID reader system has been deployed into the warehouse to take periodic snapshots about the existing tags. Then, a solution for this problem is to examine the difference between two snapshots taken at different time points, which contains the information about product inflow and outflow within the time interval. However, a key challenge is that a large warehouse inevitably needs more than one readers to achieve full coverage. Note that each reader can take a snapshot only about its local set of tags, and the snapshots taken by different readers may not be of an identical length. When a user queries for the dynamics of such a warehouse, he or she actually wants to know how the union of multiple tag sets (scanned by readers at different locations) fluctuates over time. Such a query will produce a complex set expression — to compute the cardinality of *the set difference between two union tag sets at different time points*.

Problem

From the above two applications, we can abstract a problem called *joint cardinality estimation*. It is to estimate the cardinality of an arbitrary set expression that involves multiple tag sets (whose number is denoted by k) existing in different temporal or spatial domains. The protocol designed to scan each tag set must be very time-efficient, and meanwhile, its absolute estimation error for the cardinality must be inside a preset bound with a good probability above a high threshold.

There are limited prior studies on this problem. The differential estimator (DiffEstm) [9] and the joint RFID estimation protocol (JREP) [12] can only handle set expression involving two tag sets. Although the recent work named composite counting framework (CCF) [13] can provide a generalized estimator for an arbitrary expression over multiple sets, it is designed based on a different, relative error model, resulting in large execution time, with unbounded worst-case time complexity. This is because, as multiple RFID readers are deployed at different places to scan their surroundings periodically, the tag sets they encounter may differ significantly in sizes. The biggest tag set can be many times larger than the average-sized tag sets. Both the prior solutions DiffEstm and CCF encode each tag set into a data structure whose length is determined by the size of the largest possible set (or even the union of several largest tag sets for union estimation), which causes unnecessarily long protocol execution time.

Our Contributions

Firstly, for the generalized joint cardinality estimation problem, we propose a solution with a novel design that supports versatile snapshot construction. It adopts a two-phase protocol design between a reader and its nearby tags to construct a snapshot of the tag set. The length of the snapshot can be (but not necessary) proportional to the size of the tag set, instead of being fixed to a large worst-case value. Given the snapshots of

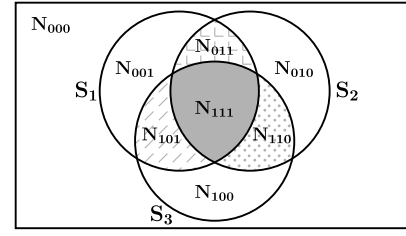


Fig. 1. Venn diagram of three tag sets S_1 , S_2 , S_3 , and an illustration of elementary subsets $N_{b_3 b_2 b_1}$ with $b_3 b_2 b_1$ falling between 000 and 111.

any k tag sets, although their lengths may be very different, we have derived closed-form formulas to estimate the joint cardinalities of the k sets.

Secondly, we analyze the means and variances of the estimated joint cardinalities computed from the formulas. We prove that the formulas produce asymptotically unbiased results and they estimate the joint cardinalities with an absolute error (probabilistic) bound that can be set arbitrarily. We also derive formulas for determining the optimal system parameters that minimize the execution time of taking snapshots, under a given accuracy constraint for joint cardinality estimation.

Thirdly, we perform extensive simulations, and the results show that, by allowing the snapshots to have variable lengths, the new solution significantly outperforms DiffEstm [9] and CCF [13], both of which assume a fixed length for snapshots. Under the same accuracy requirement, our new solution can reduce time cost by over 80% as compared with prior work.

The rest of this paper is organized as follows. Section II define the problem of joint cardinality estimation. Section III briefly describe the underlying ALOHA protocol for collecting raw data. Section IV describes a baseline solution, requiring all the frames to have an identical size, and section V presents our adaptive MJREP protocol. We analyze the estimation mean and variance of our protocol in section VI, and optimize its parameters in section VII. We provide simulation studies of our protocol in section X, and review related research in section XI. Section XII concludes the whole paper.

II. JOINT CARDINALITY ESTIMATION

In this section, we formally define the research problem of joint cardinality estimation for multiple RFID tag sets.

Definition of Joint Cardinalities

Suppose a distributed RFID system, where tagged objects are moved from one location to another. We use $S_1, S_2, \dots, S_k$ to denote the tag sets captured by RFID systems at different locations or time points. They can form an arbitrary set expression as connected by the union ($\cup$), intersection ($\cap$) and relative complement ($\setminus$) operations. The cardinality of such an expression is called a joint cardinality of the k tag sets.

A major difficulty is that the number of possible set expressions is really huge. In order to tame the high complexity, we start from a small group of expressions with special forms. Hence, we divide the union of all k sets $S_1 \cup S_2 \dots \cup S_k$ into subsets that are mutually disjoint. They are called *elementary subsets*, and the number of elementary subsets is only $2^k - 1$.

As an example, in Fig. 1, we illustrate the Venn diagram of three RFID tag sets S_1 , S_2 , S_3 , and we divide their union into $2^3 - 1 = 7$ elementary subsets. Each subset is denoted by $N_{b_3 b_2 b_1}$, where $b_3 b_2 b_1$ is a binary ranging from 001 to 111 to indicate whether this subset is included by S_3 , S_2 or S_1 . For instance, the subset N_{110} is included by S_3 and S_2 ,

but excluded by S_1 . In Fig. 1, N_{000} is a special case that corresponds to the tags not included by any sets S_1, S_2 or S_3 .

We formalize the concept of elementary subset $N_{b_k \dots b_2 b_1}$ as

$$N_{b_k \dots b_2 b_1} = \left(\bigcap_{b_i \neq 0}^{1 \leq i \leq k} S_i \right) \setminus \left(\bigcup_{b_i = 0}^{1 \leq i \leq k} S_i \right), \quad (1)$$

where the bit b_i indicates whether the elementary subset is included or excluded by the i th set S_i . It is equivalent to

$$N_{b_k \dots b_2 b_1} = \left(\bigcap_{b_i \neq 0}^{1 \leq i \leq k} S_i \right) \cap \left(\bigcap_{b_i = 0}^{1 \leq i \leq k} S_i^c \right), \quad (2)$$

if applying the rule of relative complement $A \setminus B = A \cap B^c$ to equation (1), where B^c is the absolute complement of B .

For a shorter notation, we replace $N_{b_k \dots b_2 b_1}$ by N_x , where x is a decimal that is equal to the binary number $b_k \dots b_2 b_1$. Hence, the definition of elementary subset in (2) becomes

$$N_x = \left(\bigcap_{2^{i-1} \wedge x \neq 0}^{1 \leq i \leq k} S_i \right) \cap \left(\bigcap_{2^{i-1} \wedge x = 0}^{1 \leq i \leq k} S_i^c \right), \quad (3)$$

where $2^{i-1} \wedge x$ extracts the i th bit from x by bitwise AND $\wedge$. There are two boundary cases: $N_{2^k-1} = S_1 \cap S_2 \dots \cap S_k$ is the intersection of all sets, and $N_0 = S_1^c \cap S_2^c \dots \cap S_k^c = (S_1 \cup S_2 \dots \cup S_k)^c$ is the complement of the union of all sets.

For an arbitrary elementary subset N_x with $1 \leq x < 2^k$, we denote its cardinality by n_x , and call it a *joint cardinality* of the k tag sets $S_1, S_2, \dots, S_k$. If the cardinalities n_x of all elementary subsets N_x ($1 \leq x < 2^k$) are known, we can compute the cardinality of an arbitrary set expression by summing up the cardinalities of elementary subsets it includes.

Any set expression can be rewritten as the union of several elementary subsets. As an example, let the queried tag set be $S_3 \cap (S_2 \cup S_1)$. It is equal to $S_3 \cap ((S_2^c \cap S_1) \cup (S_2 \cap S_1^c) \cup (S_2 \cap S_1))$. By applying the distributive law, it becomes $(S_3 \cap S_2^c \cap S_1) \cup (S_3 \cap S_2 \cap S_1^c) \cup (S_3 \cap S_2 \cap S_1)$. By the definition of N_x in (3), it equals $N_5 \cup N_6 \cup N_7$. Hence, the queried cardinality $|S_3 \cap (S_2 \cup S_1)|$ is equal to $n_5 + n_6 + n_7$.

The cardinality of such a set expression is called a *composite joint cardinality*. A special case of composite joint cardinality is the cardinality of N_0^c , which is equal to the union of all elementary subsets N_x , $1 \leq x < 2^k$. We denote such a union cardinality as n_0^c , which is equal to $\sum_{1 \leq x < 2^k} n_x$.

In a word, we put more focus on determining the $2^k - 1$ elementary joint cardinalities n_x . The composite joint cardinalities, whose number is huge, are left to a secondary position.

Probabilistic Estimation

In many practical applications, it often does not require to know the exact value of the joint cardinality n_x , and an approximated value $\hat{n}_x$ with desired accuracy is adequate. In the following problem definition, we require the absolute estimation error $\hat{n}_x - n_x$ must be bounded by a predefined range $\pm\theta$ at a probability of at least $1 - \delta$, which is called the (θ, δ) model. Besides n_x , we also consider to keep the absolute estimation error of n_0^c bounded by $\pm\theta$, which can act as a representative of composite joint cardinality.

Definition 1 (Joint Cardinality Estimation Problem): For joint cardinalities n_0^c or n_x ($1 \leq x < 2^k$), the joint estimation problem is to find an algorithm for generating estimations $\hat{n}_0^c$ and $\hat{n}_x$. They should satisfy the accuracy constraint:

$$\begin{aligned} \text{Prob}\{\hat{n}_0^c - \theta \leq n_0^c \leq \hat{n}_0^c + \theta\} &\geq 1 - \delta \\ \text{Prob}\{n_x - \theta \leq \hat{n}_x \leq n_x + \theta\} &\geq 1 - \delta, \end{aligned} \quad (4)$$

where $\hat{n}_0^c \pm \theta$ and $n_x \pm \theta$ are the confidence interval of estimations $\hat{n}_0^c$ and $\hat{n}_x$, respectively, and $1 - \delta$ is the confidence level.

An alternative way of specifying the estimation accuracy is based on a relative error bound $\epsilon \in (0, 1)$:

$$\begin{aligned} \text{Prob}\{n_0^c(1 - \epsilon) \leq \hat{n}_0^c \leq n_0^c(1 + \epsilon)\} &\geq 1 - \delta \\ \text{Prob}\{n_x(1 - \epsilon) \leq \hat{n}_x \leq n_x(1 + \epsilon)\} &\geq 1 - \delta. \end{aligned} \quad (5)$$

According to this model, the probabilities for the relative errors $\frac{\hat{n}_0^c - n_0^c}{n_0^c}$ and $\frac{\hat{n}_x - n_x}{n_x}$ to fall into the range $\pm\epsilon$ are at least $1 - \delta$.

This relative error model has been adopted by the previous work [9], [13], with a time complexity of $O(\frac{1}{\epsilon^2 J} \ln \frac{1}{\delta})$, where J is *Jaccard similarity* — the ratio of the intersection size of all tag sets to the union size of all sets. However, $\frac{1}{J}$ can be very large, since the intersection size can be small or even zero while the union size is very large, which may be the routine case in practice, instead of being rare. For the applications in the introduction, there can be two warehouses with few products moved between them, or there can be times when few products are moved in or out of a warehouse. In both cases, $\frac{1}{J}$ is very large or even tends to infinite.

In conclusion, the relative error model, as a remanent from the earlier literature on cardinality estimation of a single tag set, is no longer suitable for joint cardinality estimation of multiple sets. Therefore, this paper adopts the absolute error model in (4). The previous protocols named DiffEstm [9] and CCF [13] are not designed under this model.

III. ALOHA-BASED RFID PROTOCOL

In this section, we introduce a standardized RFID communication protocol based on slotted ALOHA, which can be used to take a snapshot of a tag set without collecting any tag IDs.

A. ALOHA Communication Protocol

A reader communicates with the tags in its radio range, using the following slotted ALOHA protocol, which is compliant with the de-facto RFID standard named EPC C1G2 [14].

Initially, the reader broadcasts a `Query` command to start an ALOHA frame with m time slots and using R as random seed. Upon receiving the command which carries the values of m and R , each tag selects a time slot pseudo-randomly through hashing $H(id \oplus R) \bmod m$, where m is the frame length, id is the tag ID, and $\oplus$ is bitwise XOR that mixes tag id and the random seed R . During the frame, the reader broadcasts a `QueryRep` command between any two adjacent slots, terminating the previous slot and starting the next slot. Each tag will transmit a response during its selected slot.

By executing the above slotted ALOHA protocol, from the perspective of the reader, the responses of all tags distribute uniformly at random across the m time slots in the frame. Furthermore, a sampling mechanism can be incorporated into the ALOHA frame as a field in the frame header. Due to the sampling, only p fraction ($0 < p \leq 1$) of tags respond in the frame and the rest of them just keep silent.

B. Empty/Busy Time Slots and Snapshot

Time slots can be classified into two types: a slot is called *empty* if it contains no tag response, or called *busy* if it has at least one tag response. Busy slots can be further classified into *singleton* slots (containing exactly one tag response) and *collision* slots (containing at least two tag responses). According to the EPCglobal RFID standard [14], each tag

response needs to contain a 16-bit random number in order to differentiate singleton slots from collision ones. In this paper, we only need to distinguish empty slots from busy slots. For this purpose, transmitting just one bit is sufficient to indicate the presence of a tag response in a slot, which significantly reduces the length of each time slot.

We use 0 to represent an empty slot and 1 a busy slot. By monitoring the empty/busy states of the slots, the reader can turn an ALOHA frame into a bitmap of m bits, a bit of 0 for each empty slot and a bit of 1 for each busy slot. This bitmap is a compact data structure, also called *snapshot*, which encodes the set of tags at the location of the reader and at the chosen time when the reader initiates the ALOHA frame. In the sequel, we use the terms “snapshot”, “bitmap” and “frame” interchangeably.

IV. A BASELINE SOLUTION

Suppose a distributed RFID system with multiple readers. The reader at each location can take a snapshot of the set of tags currently in its radio range using the ALOHA protocol. The reader will forward the periodically taken snapshots to a central server where queries on joint cardinalities over multiple tag sets are made. Consider an arbitrary query that considers k chosen snapshots, $B_1, B_2, \dots, B_k$. Let $S_1, S_2, \dots, S_k$ be the tag sets that are encoded in these bitmaps, respectively. Note that these tag sets may be present at different locations in the system or at the same location but different times.

In this section, we present a baseline solution that performs joint cardinality estimation over S_i , $1 \leq i \leq k$, based on their snapshots, B_i , $1 \leq i \leq k$. This baseline solution assumes that all the bitmaps must have the same length and use a common sampling probability. The key reason is that it retrieves useful information from these bitmaps by bitwise OR operation. Next, we describe its estimation formulas by details.

Union Cardinality Estimation

For ALOHA frames $B_1, B_2, \dots, B_k$, an important property is that, if a tag is sampled to respond, it will select the same time slot among all frames. This is because the tag uses the same hash function $H(id \oplus R) \bmod m$ for slot selection in different frames, where m is the common length of all frames and R is the random seed.

Out of the k bitmaps, we can arbitrarily select c bitmaps and combine them by bitwise OR, which is represented by $B_{i_1} \vee B_{i_2} \dots \vee B_{i_c}$ with $1 \leq i_1 < i_2 \dots < i_c \leq k$. Because a tag responds at the same slot in all these frames, when calculating the bitwise OR, its multiple responses in different frames will be automatically combined into one bit. Therefore, the OR of c bitmaps $B_{i_1} \vee B_{i_2} \dots \vee B_{i_c}$ is equivalent to the bitmap-based snapshot for the union of c tag sets $S_{i_1} \cup S_{i_2} \dots \cup S_{i_c}$ [1].

For estimating the cardinality of such a union set, a good method is to use the fraction of zero bits in the OR bitmap $B_{i_1} \vee B_{i_2} \dots \vee B_{i_c}$ [1]. Specifically, let z be the fraction of zero bits in the OR bitmap. Then, the cardinality of corresponding union set can be estimated as $-m \log(z) / p$, where m is the number of bits in the OR bitmap, and p is the sampling probability.

Joint Cardinality Estimation

Since the cardinality of any union set is known, we can easily derive $|S_1 \cap S_2 \dots \cap S_k|$, i.e., the cardinality of the intersection of all the k tag sets, which is denoted as the joint

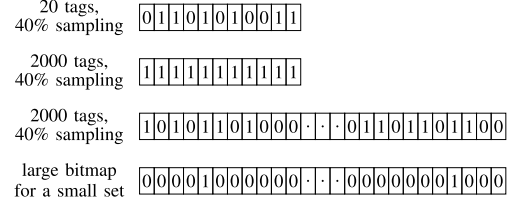


Fig. 2. Inefficiency problem of baseline protocol when handling small tag sets and large tag sets.

cardinality n_{2^k-1} . According to the *principle of inclusion and exclusion*, n_{2^k-1} is equal to

$$\begin{aligned} n_{2^k-1} &= |S_1 \cap S_2 \dots \cap S_k| = \sum_{1 \leq i_1 \leq k} |S_{i_1}| \\ &\quad - \sum_{1 \leq i_1 < i_2 \leq k} |S_{i_1} \cup S_{i_2}| + \sum_{1 \leq i_1 < i_2 < i_3 \leq k} |S_{i_1} \cup S_{i_2} \cup S_{i_3}| \\ &\quad + \dots + (-1)^{k-1} |S_1 \cup S_2 \dots \cup S_k|, \end{aligned} \quad (6)$$

where the union cardinalities $|S_{i_1}|$, $|S_{i_1} \cup S_{i_2}|$, $|S_{i_1} \cup S_{i_2} \cup S_{i_3}|$, ..., $|S_1 \cup S_2 \dots \cup S_k|$ all have unbiased estimators as stated.

Interestingly, Eq. (6) can be extended to estimating any joint cardinality n_x with $1 \leq x < 2^k$. By the definition of N_x in (3),

$$n_x = \left| \bigcap_{2^{i-1} \wedge x \neq 0}^{1 \leq i \leq k} S_i \right| - \left| \left(\bigcap_{2^{i-1} \wedge x \neq 0}^{1 \leq i \leq k} S_i \right) \cap \left(\bigcup_{2^{i-1} \wedge x = 0}^{1 \leq i \leq k} S_i \right) \right|.$$

The first term $\left| \bigcap_{2^{i-1} \wedge x \neq 0}^{1 \leq i \leq k} S_i \right|$ is the intersection of multiple tag sets and hence can be estimated similar to (6). The second term is the intersection of $\bigcap_{2^{i-1} \wedge x \neq 0}^{1 \leq i \leq k} S_i$ and $\bigcup_{2^{i-1} \wedge x = 0}^{1 \leq i \leq k} S_i$, where the latter can be treated as a single tag set encoded into the OR bitmap $\bigvee_{2^{i-1} \wedge x = 0}^{1 \leq i \leq k} B_i$. Hence, the second term can be regarded as the intersection of multiple sets with a union set, and can also be estimated similar to (6).

Because this baseline protocol depends on the INCLUSION-EXCLUSION principle in (6) for estimating the intersection of multiple sets from their unions, we call it INC-EXC for short.

We explain why this baseline protocol will have poor time efficiency. For a small tag set, if the sampling probability is very small, too few or even no tag will be sampled for the snapshot construction. Hence, the sampling probability has to be reasonably large, as depicted by the first bitmap in Fig. 2, where 20 tags are recorded with 40% sampling probability. However, for a large set, a significant sampling probability will cause all bits to be set as ones (as shown by the second bitmap in Fig. 2), unless the bitmap length is sufficiently large (see the third bitmap of Fig. 2). Now because the same large length has to be applied to all bitmaps, it becomes a great waste for small tag sets (the fourth bitmap of Fig. 2). Since each bit takes one time slot to determine, a large bitmap length implies a long time for taking a snapshot, even for a very small tag set. Hence, when the tag sets have dramatically different sizes, the baseline protocol will greatly waste protocol execution time.

V. ADAPTIVE ESTIMATION PROTOCOL

We present our MJREP, which is a Joint RFID Estimation Protocol to derive joint cardinalities for Multiple tag sets, even when they are encoded into bitmaps of different lengths.

Naturally, it is desirable to encode each tag set into a bitmap with a different length, depending on the cardinality of the tag set. This inspires us to develop a new algorithm to jointly analyze k bitmaps with variable lengths. The real difficulty is

not at how to combine k bitmaps; there are simple ways to combine them. The real difficulty comes after the combination — how to perform analysis on the information combined from non-uniformly sized snapshots, how to use that information for joint cardinality estimation, and most importantly, how to ensure the satisfaction of accuracy requirements in (4). These are the tasks that have not been fulfilled in the literature.

Our MJREP protocol can be divided into two components: an online encoding component for compressing each tag set into a bitmap, and an offline analysis component for estimating the joint cardinalities of multiple tag sets, using their bitmap-based snapshots. Whenever a bitmap that encodes a tag set is collected by RFID reader during the online phase, the reader offloads it immediately to a central server for long-term storage and for offline processing. Such an asymmetric design will push most complexity to the offline component at the server side, while keeping the online component at the side of reader and tags (for raw data collection) as efficient as possible. We will introduce the online component in the first subsection, and then the offline component in the subsequent subsection.

A. Online Encoding of a Tag Set

We use a two-phase protocol to produce a snapshot for tag set S_i whose length is proportional to the tag set size s_i . The first phase is to quickly generate an estimation with coarse accuracy for the number of tags s_i . In second phase, the RFID reader invokes the ALOHA protocol in Section III to encode the current tag set in a bitmap (snapshot), and the coarse estimation $\hat{s}_i$ can be used to configure the length of ALOHA frame proportional to $\hat{s}_i$. As described later in Section XI for related works, it becomes a common practice for RFID researchers to use such a two-phase protocol for estimating the cardinality of a single tag set [8]. In following, we will explain this two-phase protocol with more details.

Firstly, the RFID reader that covers the tag set S_i invokes an existing protocol to generate a coarse estimation of the set size s_i . Since this phase handles only a single tag set, its estimation accuracy can be specified by the (ϵ, δ) model in (5), in which the probability for the relative estimation error to fall within the range $\pm\epsilon$ is at least $1 - \delta$. Since the accuracy of this phase does not need to be high, we often configure $\epsilon = 20\%$ and $\delta = 5\%$. To implement this phase, many existing protocols can be used, such as LoF [2], GMLE [3], PET [5] and ZOE [7]. If LoF is used, then the needed number of time slots for scanning the tag set is $\mathcal{O}(\frac{1}{\epsilon^2} \log(s_{\max})) \cdot \log(\frac{1}{\delta})$ to attain the pre-defined cardinality estimation accuracy, where $s_{\max}$ is an upper bound for the size of any tag set. By a recent survey study [8], the time cost can be reduced to $\mathcal{O}(\frac{1}{\epsilon^2} \log \log(s_{\max})) \cdot \log(\frac{1}{\delta})$ if PET [5] is used. It is clear that the time expense is proportional to the *logarithm* or even the *log-logarithm* of the tag set size. Hence, the time cost of the first phase is negligibly small when its accuracy requirement is coarse. For instance, when the relative estimation error ϵ is 20% and δ is 5%, the time cost of the LoF algorithm [2] is about $32 \log(s_{\max})$.

Secondly, because the size of tag set has been coarsely estimated by the first phase as $\hat{s}_i$, the reader can scan the tag set using an ALOHA frame B_i whose length m_i is linearly proportional to $\hat{s}_i$. Then, the frame length m_i satisfies

$$m_i = \min_{p \in (0,1]} \{2^{\lceil \log_2(\frac{\hat{s}_i}{p}) \rceil}\}, \quad (7)$$

where ρ is the load factor of the frame. Here, the frame length m_i is configured to a power of two. Our purpose is to ensure, when jointly analyzing two frames, the length of the longer one is always integer multiple of the length of the shorter one.

Equation (7) can be regarded as the major proportion of time cost for the online component to encode a tag set. Later in section VII-A, we will prove an optimized value of the load factor ρ that minimizes (7) while keeping the pre-set accuracy constraint of joint cardinality estimation in (4) satisfied is

$$\rho = \frac{1}{2p k_{\max}} \left(-3 + \sqrt{3} \sqrt{8p \left(\frac{\theta^2}{k_{\max} s_{\max} Z_{\delta}^2} + 1 \right)} - 5 \right), \quad (8)$$

where θ is the bound of absolute estimation error, Z_{δ} is the $1 - \frac{\delta}{2}$ quantile of standard Gaussian distribution (e.g., $Z_{0.05} \approx 1.96$), $s_{\max}$ is the upper bound of the cardinality of a tag set, $k_{\max}$ is the largest number of tag sets that may involve in any user query, and p is the sampling probability. Since $s_{\max}$, $k_{\max}$, θ and Z_{δ} are all fixed values, (8) can be regarded a function of only one variable p . Let p^* be the optimal sampling probability that maximizes (8), which in turn will minimize the time cost of encoding a tag set in (7). It is clear that the value of p^* only depends on $s_{\max}$, $k_{\max}$, θ and δ . Hence, p^* is pre-determined for a system once these parameters are set, and the best ρ is also pre-determined.

B. Offline Estimation of Joint Cardinalities of Multiple Tag Sets

In this subsection, we present an offline analysis algorithm that derives the joint cardinalities n_0^c and n_x , $1 \leq x < 2^k$, of k tag sets, using the bitmaps $B_1, B_2, \dots, B_k$. Without loss of generality, we assume the bitmaps are sorted by their lengths in non-descending order that satisfies $m_1 \leq m_2 \leq \dots \leq m_k$.

Although all these bitmaps are assumed by equation (7) to have the same load factor ρ , our offline algorithm to describe later can in fact work well if each bitmap B_i has its own load factor ρ_i . We will prove in Section VI that, only when $\rho_1 = \rho_2 = \dots = \rho_k = \rho$, can we minimize the protocol time cost.

Expanded Bitwise OR: We introduce two bitwise operations which will be used later. In the binary representation of a value y , let $\text{lo}(y)$ be the location of the lowest-order 1-bit, and let $\text{hi}(y)$ be the location of the highest-order 1-bit. For example, if the binary representation of y is 1010001, then we have $\text{lo}(y) = 1$ and $\text{hi}(y) = 7$. A boundary case is that y is equal to zero. In this case, we define $\text{hi}(y) = 0$ and $\text{lo}(y) = 0$.

For the k bitmaps, we introduce an auxiliary bitmap called *expanded OR* to combine them, which is noted by $OR_{b_k \dots b_2 b_1}$.

$$OR_{b_k \dots b_2 b_1} = \bigvee_{b_i \neq 0}^{1 \leq i \leq k} \text{Expand}(B_i, m_{\text{hi}(b_k \dots b_2 b_1)})$$

The subscript $b_k \dots b_2 b_1$ indicates for each bitmap whether it is involved in the expanded OR: If b_i is one, the bitmap B_i is involved. Among the chosen bitmaps, the length of the longest bitmap is $m_{\text{hi}(b_k \dots b_2 b_1)}$, since $m_1 \leq m_2 \leq \dots \leq m_k$ and $\text{hi}(b_k \dots b_2 b_1)$ is the position of the highest 1-bit in $b_k \dots b_2 b_1$. The function $\text{Expand}(B_i, m_{\text{hi}(b_k \dots b_2 b_1)})$ increases the length of B_i to the largest bitmap length $m_{\text{hi}(b_k \dots b_2 b_1)}$ by self replication, such that all bitmaps after expansion have an equal length and can be combined by bitwise OR operator $\bigvee$.

Figure 3 illustrates an example of applying expanded OR to three bitmaps B_1 , B_2 and B_3 . Among them, B_3 is the longest. We replicate B_1 for one time and B_2 for three times,

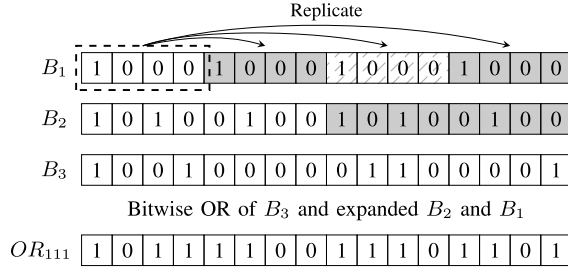


Fig. 3. Expanded OR of three bitmaps B_1 , B_2 , B_3 , whose sizes are 4, 8, 16.

such that after expansion, all bitmaps are of the same length and the bitwise OR operation can be used to combine them.

For simplicity, we replace $OR_{b_k \dots b_2 b_1}$ by a shorter notation OR_y , where y is a decimal equal to $b_k \dots b_2 b_1$. Therefore,

$$OR_y = \bigvee_{2^{i-1} \wedge y \neq 0}^{1 \leq i \leq k} \text{Expand}(B_i, m_{hi(y)}). \quad (9)$$

It is always feasible to expand the bitmap B_i to have the same length with the longest bitmap $B_{hi(y)}$, because both their lengths m_i and $m_{hi(y)}$ are the powers of two by equation (7), and the ratio $m_{hi(y)}/m_i$ is definitely an integer.

Expected Zero Fraction of OR_y : We analyze the expected fraction of zero bits in the bitmap OR_y . Let us focus on just one bit in OR_y , and we have the following property.

Property 1 (Probability of a Tag Setting a Bit in OR Bitmap): Considering an arbitrary bit in bitmap OR_y and an arbitrary tag from elementary subset N_x , when the bitwise AND of x and y is non-zero, the probability of the tag assigning the bit to one is $\frac{p}{m_{lo(x \wedge y)}}$; when $x \wedge y$ is zero, the probability is zero.

Proof: Please check out the Appendix. ■

According to Property 1, a tag has the chance to assign a bit of OR_y to one, only when it is in a subset N_x that satisfies $x \wedge y \neq 0$. Because in such a subset N_x , any tag assigns the j th bit of OR_y at a probability of $\frac{p}{m_{lo(x \wedge y)}}$, the probability that all tags in N_x do not assign this bit is $(1 - \frac{p}{m_{lo(x \wedge y)}})^{n_x}$.

Let $X_y^{(j)}$ be the event that the j th bit in OR_y remains zero. The occurrence of the event needs all tags in an arbitrary N_x with $x \wedge y \neq 0$ do not assign this bit. Thus, its probability is

$$\text{Prob}\{X_y^{(j)}\} = \prod_{1 \leq x < 2^k}^{x \wedge y \neq 0} (1 - \frac{p}{m_{lo(x \wedge y)}})^{n_x}. \quad (10)$$

Let z_y be the fraction of zero bits in OR_y . Then,

$$z_y = \frac{1}{m_{hi(y)}} \sum_{0 \leq j < m_{hi(y)}} 1_{X_y^{(j)}}, \quad (11)$$

where $m_{hi(y)}$ is the number of bits in bitmap OR_y (see (9)), and $1_{X_y^{(j)}}$ is the indicator function of the event $X_y^{(j)}$, whose value is one if the event occurs and is zero otherwise. Since the zero fraction z_y is the arithmetic mean of a large number of independent variables, according to the central limit theorem, z_y approximates a Gaussian distribution. Its expected value is

$$E(z_y) = E(\frac{1}{m_{hi(y)}} \sum 1_{X_y^{(j)}}) = \frac{1}{m_{hi(y)}} \sum_{0 \leq j < m_{hi(y)}} E(1_{X_y^{(j)}}).$$

Clearly, $E(1_{X_y^{(j)}}) = \text{Prob}\{X_y^{(j)}\}$. Hence, applying (10),

$$\begin{aligned} E(z_y) &= \frac{1}{m_{hi(y)}} \sum_{0 \leq j < m_{hi(y)}} \text{Prob}\{X_y^{(j)}\} \\ &= \text{Prob}\{X_y^{(j)}\} = \prod_{1 \leq x < 2^k}^{x \wedge y \neq 0} (1 - \frac{p}{m_{lo(x \wedge y)}})^{n_x}. \end{aligned}$$

Applying the approximation $(1 - \frac{p}{m})^n \approx (1 - \frac{p}{m'})^{\frac{m'}{m}n}$ which works when m and m' are both large, we have

$$\begin{aligned} E(z_y) &\approx \prod_{1 \leq x < 2^k}^{x \wedge y \neq 0} (1 - \frac{p}{m_{hi(y)}})^{\frac{m_{hi(y)}}{m_{lo(x \wedge y)}} \cdot n_x} \\ &= (1 - \frac{p}{m_{hi(y)}})^{\sum_{1 \leq x < 2^k}^{x \wedge y \neq 0} \frac{m_{hi(y)}}{m_{lo(x \wedge y)}} \cdot n_x}. \end{aligned}$$

Using the sign function sgn (which equals to 1, 0 or -1 when its input parameter is positive, zero or negative), we have

$$E(z_y) \approx (1 - \frac{p}{m_{hi(y)}})^{\sum_{1 \leq x < 2^k} \text{sgn}(x \wedge y) \cdot \frac{m_{hi(y)}}{m_{lo(x \wedge y)}} \cdot n_x}. \quad (12)$$

Estimator of Joint Cardinality n_x : Using the fraction of zero bits in OR_y with $1 \leq y < 2^k$, we are able to estimate each joint cardinality n_x with $1 \leq x < 2^k$. We will show later that this essentially is to solve a *fully determined linear system*, which puts $2^k - 1$ constraints over $2^k - 1$ unknown variables.

By (11), we know the variance of z_y is inversely proportional to the number of observations $m_{hi(y)}$. Hence, when the number of slots $m_{hi(y)}$ in the frame OR_y is sufficiently large, we can approximate $E(z_y)$ by z_y . Then, (12) becomes

$$z_y \approx (1 - \frac{p}{m_{hi(y)}})^{\sum_{1 \leq x < 2^k} \text{sgn}(x \wedge y) \cdot \frac{m_{hi(y)}}{m_{lo(x \wedge y)}} \cdot n_x}.$$

We will prove later that such an approximation indeed produces unbiased estimators. Taking the logarithm of both sides,

$$\log(z_y) / \log(1 - \frac{p}{m_{hi(y)}}) \approx \sum_{1 \leq x < 2^k} \text{sgn}(x \wedge y) \cdot \frac{m_{hi(y)}}{m_{lo(x \wedge y)}} \cdot n_x.$$

Applying the approximation $\log(1 - \frac{p}{m}) \approx -\frac{p}{m}$ for large m ,

$$-\frac{m_{hi(y)}}{p} \log(z_y) \approx \sum_{1 \leq x < 2^k} \text{sgn}(x \wedge y) \cdot \frac{m_{hi(y)}}{m_{lo(x \wedge y)}} \cdot n_x.$$

If we define the measurement of the number of tags in OR_y as

$$\hat{u}_y = -\frac{m_{hi(y)}}{p} \log(z_y), \quad (13)$$

then the above equation can be simplified as

$$\sum_{1 \leq x < 2^k} \text{sgn}(x \wedge y) \cdot \frac{m_{hi(y)}}{m_{lo(x \wedge y)}} \cdot n_x \approx \hat{u}_y. \quad (14)$$

Applying (13) to each OR_y bitmap, we can collect a vector of measurements $\hat{\mathbf{u}} = [\hat{u}_1, \hat{u}_2, \dots, \hat{u}_y, \dots, \hat{u}_{2^k-1}]^T$. If putting together the sizes of all elementary subsets, we can obtain the vector of unknowns: $\mathbf{n} = [n_1, n_2, \dots, n_x, \dots, n_{2^k-1}]^T$. With $\hat{\mathbf{u}}$ and $\mathbf{n}$ properly defined, equation (14) can be rewritten as

$$\mathbf{M} \mathbf{n} \approx \hat{\mathbf{u}}, \quad (15)$$

where $\mathbf{M}$ is the coefficient matrix defined in (16). Its element uses y as row index and x as column index, $1 \leq x, y < 2^k$.

$$\mathbf{M} = \left[\text{sgn}(x \wedge y) \cdot \frac{m_{hi(y)}}{m_{lo(x \wedge y)}} \right] \quad (16)$$

For elements of $\mathbf{M}$, if $x \wedge y$ is zero, then $\text{sgn}(x \wedge y)$ is zero, and the expression $\text{sgn}(x \wedge y) \cdot \frac{m_{hi(y)}}{m_{lo(x \wedge y)}}$ is treated as zero. By Property 2 to present later, we know that the coefficient matrix $\mathbf{M}$ is always non-singular. Hence, we can solve the equation system in (15) and obtain an estimator of joint cardinalities.

$$\hat{\mathbf{n}} = \mathbf{M}^{-1} \hat{\mathbf{u}} \quad (17)$$

Estimator of Composite Joint Cardinality n_0^c : Among the numerous composite joint cardinalities, the largest and the most important cardinality is n_0^c — the number of tags in the union of all the sets. We estimate it as $\hat{n}_0^c = \sum_{1 \leq x < 2^k} \hat{n}_x$, the sum of all elements in the vector $\hat{\mathbf{n}}$. After simplification,

$$\hat{n}_0^c = \widehat{u_{2^k-1}} - \frac{m_2-m_1}{m_1} \widehat{u_1} - \frac{m_3-m_2}{m_2} \widehat{u_3} \dots - \frac{m_{j+1}-m_j}{m_j} \widehat{u_{2^j-1}} \dots - \frac{m_k-m_{k-1}}{m_{k-1}} \widehat{u_{2^{k-1}-1}}. \quad (18)$$

Property 2: The coefficient matrix $\mathbf{M}$ defined in (16) is always non-singular, and its inverse matrix $\mathbf{M}^{-1}$ always exists.

Proof: Firstly, we consider a simple case that the sizes of all frames are equal, i.e., $m_1 = m_2 \dots = m_k$. In this case, the coefficient matrix $\mathbf{M}$ degrades to $\mathbf{G} = [\text{sgn}(x \wedge y)]$, $1 \leq x, y < 2^k$, by (16). We use $\mathbf{G}_k$ to denote the $\mathbf{G}$ matrix when there are k tag sets to perform joint estimation. We have

$$\mathbf{G}_2 = \begin{bmatrix} 1 & 0 & 1 \\ 0 & 1 & 1 \\ 1 & 1 & 1 \end{bmatrix}, \quad \mathbf{G}_{k+1} = \begin{bmatrix} \mathbf{G}_k & \mathbf{0} & \mathbf{G}_k \\ \mathbf{0} & 1 & 1 \\ \mathbf{G}_k & 1 & 1 \end{bmatrix}.$$

Clearly, $\mathbf{G}_{k+1}$ has a recursive form, whose left top, right top and left bottom blocks are $\mathbf{G}_k$. It is easy to verify its left top block is $\mathbf{G}_k$, since when the $(k+1)$ th bit of x and the $(k+1)$ th bit of y are both 0, other bits of x and y decide the value of $\text{sgn}(x \wedge y)$. The right bottom block of $\mathbf{G}_{k+1}$ is 1, since the bitwise AND of the $(k+1)$ th bit of x and y is 1.

Next, we use mathematical induction to prove that matrix $\mathbf{G}$ is non-singular. The base case is $\mathbf{G}_2$, which is non-singular since it can be transformed to a triangular matrix by Gaussian

elimination: $\mathbf{G}_2 \rightarrow \begin{bmatrix} 1 & 0 & 1 \\ 0 & 1 & 1 \\ 0 & 1 & 0 \end{bmatrix} \rightarrow \begin{bmatrix} 1 & 0 & 1 \\ 0 & 1 & 1 \\ 0 & 0 & -1 \end{bmatrix}$. For the inductive

step, we assume $\mathbf{G}_k$ is non-singular, and prove that $\mathbf{G}_{k+1}$ is also non-singular. By Gaussian elimination, $\mathbf{G}_{k+1} \rightarrow \begin{bmatrix} \mathbf{G}_k & \mathbf{0} & \mathbf{G}_k \\ \mathbf{0} & 1 & 1 \\ \mathbf{0} & 1 & 1 - \mathbf{G}_k \end{bmatrix} \rightarrow \begin{bmatrix} \mathbf{G}_k & \mathbf{0} & \mathbf{G}_k \\ \mathbf{0} & 1 & 1 \\ \mathbf{0} & 0 & -\mathbf{G}_k \end{bmatrix}$. Since $\mathbf{G}_k$ can be transformed to triangular form, so can $\mathbf{G}_{k+1}$. Thus, the non-singularity of $\mathbf{G}_{k+1}$ has been proved.

Secondly, we remove the assumption of identical frame size $m_1 = m_2 \dots = m_k$, and prove $\mathbf{M} = [\text{sgn}(x \wedge y) \cdot \frac{m_{\text{hi}(y)}}{m_{\text{lo}(x \wedge y)}}]$ is non-singular. By (16), the base case and inductive rule are

$$\mathbf{M}_2 = \begin{bmatrix} 1 & 0 & 1 \\ 0 & 1 & 1 \\ \frac{m_2}{m_1} & 1 & \frac{m_2}{m_1} \end{bmatrix}, \quad \mathbf{M}_{k+1} = \begin{bmatrix} \mathbf{M}_k & \mathbf{0} & \mathbf{M}_k \\ \mathbf{0} & 1 & 1 \\ \mathbf{D}_k \mathbf{M}_k & 1 & \mathbf{D}_k \mathbf{M}_k + 1 - \mathbf{G}_k \end{bmatrix},$$

where $\mathbf{D}_k$ is a diagonal matrix with $\mathbf{D}_k = \text{diag}(\{\frac{m_{k+1}}{m_1}\}^1, \{\frac{m_{k+1}}{m_2}\}^2, \dots, \{\frac{m_{k+1}}{m_k}\}^{2^{k-1}})$, where $\{\cdot\}^i$ means the component between two parentheses repeats for i times. Gaussian elimination can be applied to get a triangular matrix: $\mathbf{M}_{k+1} \rightarrow \begin{bmatrix} \mathbf{M}_k & \mathbf{0} & \mathbf{M}_k \\ \mathbf{0} & 1 & 1 \\ \mathbf{0} & 0 & -\mathbf{G}_k \end{bmatrix}$, which can prove $\mathbf{M}_{k+1}$ is non-singular. ■

VI. THEORETICAL ANALYSIS OF MJREP

It is easy to prove that the estimators $\hat{n}_x$ in (17) and $\hat{n}_0^c$ in (18) are asymptotically unbiased, due to the rigid process by

which they are derived. In this section, we focus on analyzing their variances, which determine their estimation errors.

A. Probabilistic Distribution of Measurements z_y

Since the zero fraction z_y of bitmap OR_y is the input of MJREP estimator, we need to analyze its mean and variance. By (11), the zero ratio z_y is the arithmetic mean of independent variables. When the number of variables $m_{\text{hi}(y)}$ is large, by the central limit theorem, z_y approximates a Gaussian distribution. The expected value of z_y is given in (12), and can be further simplified as $E(z_y) \approx e^{-p\omega_y}$, where ω_y is defined below and its physical meaning is the load factor of OR_y .

$$\omega_y = \sum_{1 \leq x < 2^k} \text{sgn}(x \wedge y) \cdot \frac{n_x}{m_{\text{lo}(x \wedge y)}} \quad (19)$$

Note that the symbol ω_y is different from ρ_i , which later will be used to denote the load factor of the frame B_i .

In the previous conference version, we have derived that the covariance of zero ratio z_{y_1} of OR_{y_1} and zero ratio z_{y_2} of OR_{y_2} is approximately

$$\text{Cov}(z_{y_1}, z_{y_2}) \approx \frac{e^{-p\omega_{y_1 \vee y_2}} - (1+p^2\omega_{y_1, y_2}^*) e^{-p(\omega_{y_1} + \omega_{y_2})}}{\min(m_{\text{hi}(y_1)}, m_{\text{hi}(y_2)})}, \quad (20)$$

where ω_{y_1, y_2}^* is density of common tags of OR_{y_1} and OR_{y_2} .

$$\omega_{y_1, y_2}^* = \min(m_{\text{hi}(y_1)}, m_{\text{hi}(y_2)}) \sum_{x \wedge y_1 \neq 0, x \wedge y_2 \neq 0} \frac{n_x}{m_{\text{lo}(x \wedge y_1)} m_{\text{lo}(x \wedge y_2)}} \quad (21)$$

B. Variance of Cardinality Measurements $\hat{u}_y$

Previously in (13), we have defined $\hat{u}_y$ as the measurement of the number of tags in the OR bitmap OR_y . We will analyze the covariance between cardinality measurement $\hat{u}_{y_1}$ of bitmap OR_{y_1} and cardinality measurement $\hat{u}_{y_2}$ of bitmap OR_{y_2} .

$$\begin{aligned} \text{Cov}(\hat{u}_{y_1}, \hat{u}_{y_2}) &= \text{Cov}\left(-\frac{m_{\text{hi}(y_1)}}{p} \log(z_{y_1}), -\frac{m_{\text{hi}(y_2)}}{p} \log(z_{y_2})\right) \\ &= \frac{m_{\text{hi}(y_1)} m_{\text{hi}(y_2)}}{p^2} \text{Cov}(\log(z_{y_1}), \log(z_{y_2})) = \frac{m_{\text{hi}(y_1)} m_{\text{hi}(y_2)}}{p^2} \\ &\quad \cdot [E(\log(z_{y_1}) \log(z_{y_2})) - E(\log(z_{y_1})) E(\log(z_{y_2}))] \end{aligned} \quad (22)$$

We use Taylor series to expand $\log(z_y)$ about the point $E(z_y) \approx e^{-p\omega_y}$. For simplicity, the point $e^{-p\omega_y}$ is denoted by q_y . Since z_y is an observation of ratio of zero slots from a large number of time slots, it will be very close to the point q_y . Then, $\log(z_y) = -p\omega_y + \frac{z_y - q_y}{q_y} + \mathcal{O}((z_y - q_y)^2)$. Since $E(\frac{z_y - q_y}{q_y}) \approx 0$, the expected value of $\log(z_y)$ is roughly

$$E(\log(z_y)) \approx -p\omega_y + E(\frac{z_y - q_y}{q_y}) \approx -p\omega_y.$$

After substituting $\log(z_y)$ by $-p\omega_y + \frac{z_y - q_y}{q_y}$ and $E(\log(z_y))$ by $-p\omega_y$, the equation (22) becomes

$$\begin{aligned} \text{Cov}(\hat{u}_{y_1}, \hat{u}_{y_2}) &\approx \frac{m_{\text{hi}(y_1)} m_{\text{hi}(y_2)}}{p^2} [E((-p\omega_{y_1} + \frac{z_{y_1} - q_{y_1}}{q_{y_1}}) \\ &\quad (-p\omega_{y_2} + \frac{z_{y_2} - q_{y_2}}{q_{y_2}})) - p\omega_{y_1} p\omega_{y_2}] \approx \frac{m_{\text{hi}(y_1)} m_{\text{hi}(y_2)}}{p^2} \\ &\quad \cdot E(\frac{z_{y_1} - q_{y_1}}{q_{y_1}} \frac{z_{y_2} - q_{y_2}}{q_{y_2}}) = \frac{m_{\text{hi}(y_1)} m_{\text{hi}(y_2)}}{p^2 q_{y_1} q_{y_2}} \text{Cov}(z_{y_1}, z_{y_2}). \end{aligned}$$

By substituting $Cov(z_{y_1}, z_{y_2})$ with its approximation in (20),

$$Cov(\hat{u}_{y_1}, \hat{u}_{y_2}) \approx \frac{1}{p^2} \max(m_{hi(y_1)}, m_{hi(y_2)}) \cdot \\ \times (e^{p(\omega_{y_1} + \omega_{y_2} - \omega_{y_1 \vee y_2})} - (1 + p^2 \omega_{y_1, y_2}^*)), \quad (23)$$

where ω_y can be found in (19), and ω_{y_1, y_2}^* is defined in (21). As a special case, when $y_1 = y_2 = y$, (23) becomes $Var(\hat{u}_y)$.

$$Var(\hat{u}_y) \approx \frac{1}{p^2} m_{hi(y)} (e^{p\omega_y} - (1 + p^2 \sum_{x \wedge y \neq 0} \frac{m_{hi(y)}}{m_{10(x \wedge y)}} n_x)) \quad (24)$$

C. Estimation Variance of Joint Cardinality

We analyze the variance of MJREP for estimating the cardinality n_x of elementary subset. By (17), we know each cardinality estimation $\hat{n}_x$ is a linear combination of OR bitmap measurements $\hat{\mathbf{u}} = [\hat{u}_1, \hat{u}_2, \dots, \hat{u}_y, \dots, \widehat{u_{2^k-1}}]^T$. Thus, the variance of each cardinality estimation $\hat{n}_x$ is a linear combination of covariance $Cov(\hat{u}_{y_1}, \hat{u}_{y_2})$ in (23). Let $Cov(\hat{\mathbf{u}}, \hat{\mathbf{u}})$ be the covariance matrix of measurement vector $\hat{\mathbf{u}}$. Then,

$$Var(\hat{n}_x) = \mathbf{M}^{-1}[x] Cov(\hat{\mathbf{u}}, \hat{\mathbf{u}}) (\mathbf{M}^{-1}[x])^T, \quad (25)$$

where $\mathbf{M}^{-1}[x]$ is the x th row of matrix $\mathbf{M}^{-1}$, and $(\mathbf{M}^{-1}[x])^T$ is its transpose and hence is a column vector.

Next, we analyze the estimation variance of the composite joint cardinality n_0^c . Equation (18) can be rewritten as

$$\hat{n}_0^c = \hat{u}_1 + (-\frac{m_2}{m_1} \hat{u}_1 + \hat{u}_3) + (-\frac{m_3}{m_2} \hat{u}_3 + \hat{u}_7) \dots + (-\frac{m_j}{m_{j-1}} \cdot \\ \times \widehat{u_{2^{j-1}-1}} + \widehat{u_{2^j-1}}) \dots + (-\frac{m_k}{m_{k-1}} \widehat{u_{2^{k-1}-1}} + \widehat{u_{2^k-1}}). \quad (26)$$

To simplify this equation, we define $\hat{d}_j$ as

$$\hat{d}_j = -\frac{m_j}{m_{j-1}} \widehat{u_{2^{j-1}-1}} + \widehat{u_{2^j-1}}, \quad \text{with } 2 \leq j \leq k. \quad (27)$$

We define d_j as the cardinality of the relative complement of tag set S_j with respect to union set $S_1 \cup S_2 \dots \cup S_{j-1}$.

$$d_j = |S_j \setminus (S_1 \cup S_2 \dots \cup S_{j-1})| \quad \text{with } 2 \leq j \leq k \quad (28)$$

It is easy to verify that $\hat{d}_j$ is an unbiased estimation of d_j . Applying (27) to (26), we have

$$\hat{n}_0^c = \hat{u}_1 + \sum_{2 \leq j \leq k} \hat{d}_j. \quad (29)$$

From (29), we have $Var(\hat{n}_0^c) = Var(\hat{u}_1 + \sum_{2 \leq j \leq k} \hat{d}_j)$.

We have proved in the previous conference version that the three kinds of estimations $\hat{u}_1$, $\hat{d}_i$ and $\hat{d}_j$, $2 \leq i < j$, are mutually uncorrelated. That is $Cov(\hat{u}_1, \hat{d}_i) \approx 0$, and $Cov(\hat{d}_i, \hat{d}_j) \approx 0$ for arbitrary i, j values with $2 \leq i < j$.

Applying the property of uncorrelatedness to (29), we have

$$Var(\hat{n}_0^c) \approx Var(\hat{u}_1) + \sum_{2 \leq j \leq k} Var(\hat{d}_j), \quad (30)$$

where $Var(\hat{u}_y)$ is given in (24). By definition of $\hat{d}_j$ in (27),

$$Var(\hat{d}_j) = \frac{m_j^2}{m_{j-1}^2} Var(\widehat{u_{2^{j-1}-1}}) + Var(\widehat{u_{2^j-1}}) \\ - 2 \frac{m_j}{m_{j-1}} \cdot Cov(\widehat{u_{2^{j-1}-1}}, \widehat{u_{2^j-1}}).$$

By (23), $Cov(\widehat{u_{2^{j-1}-1}}, \widehat{u_{2^j-1}}) \approx \frac{m_j}{m_{j-1}} Var(\widehat{u_{2^{j-1}-1}})$. Then,

$$Var(\hat{d}_j) \approx -\frac{m_j^2}{m_{j-1}^2} Var(\widehat{u_{2^{j-1}-1}}) + Var(\widehat{u_{2^j-1}}). \quad (31)$$

Applying the above equation of $Var(\hat{d}_j)$ to (30), we have

$$Var(\hat{n}_0^c) \\ \approx Var(\hat{u}_1) + \sum_{2 \leq j \leq k} (Var(\widehat{u_{2^j-1}}) - \frac{m_j^2}{m_{j-1}^2} Var(\widehat{u_{2^{j-1}-1}})) \\ = Var(\widehat{u_{2^k-1}}) - \sum_{1 \leq j < k} \frac{m_{j+1}^2 - m_j^2}{m_j^2} Var(\widehat{u_{2^j-1}}). \quad (32)$$

VII. PROTOCOL PARAMETERS

In this section, we optimize the parameters of MJREP, under the constraints of estimation accuracy in (4). Among various parameters, the size of largest tag sets $s_{\max}$ is a phenomenon of physical world and is out of the control of RFID system. The accuracy model (θ, δ) and the number of tag sets k totally depends on the demand of users. Hence, only two parameters are under the control of RFID readers, i.e., the load factor ρ_i of the frame B_i and the common sampling probability p .

In the first subsection, we investigate the appropriate configuration for the load factor ρ_i for frame B_i . In the second subsection, we study how to optimize the sampling probability p to minimize the execution time (or the length of frame B_i).

A. Configuration of Load Factors

Equation (4) requires that the probability for the absolute estimation errors of joint cardinalities n_0^c and n_x to fall within $\pm \theta$ is at least $1 - \delta$. We proved before that both $\hat{n}_x$ and $\hat{n}_0$ are asymptotically unbiased estimations and they approximate Gaussian distributions. Hence, Eq. (4) can be translated to

$$Var(\hat{n}_0^c) \leq (\theta / Z_\delta)^2 \quad \text{and} \quad Var(\hat{n}_x) \leq (\theta / Z_\delta)^2, \quad (33)$$

where Z_δ is $1 - \frac{\delta}{2}$ quantile of standard Gaussian distribution.

Property 3 (Variance Upper Bounds): The tight upper bound of $Var(\hat{n}_x)$, $1 \leq x < 2^k$, and $Var(\hat{n}_0^c)$ is $Var(\widehat{u_{2^k-1}})$.

$$Var(\hat{n}_x), Var(\hat{d}_j) \leq Var(\hat{n}_0^c) \leq Var(\widehat{u_{2^k-1}}) \quad (34)$$

Meanwhile, $Var(\widehat{u_{2^k-1}})$ is tightly upper bounded by

$$Var(\widehat{u_{2^k-1}}) \leq \frac{1}{p^2} \frac{s_k}{\rho_k} (e^{p \sum_{1 \leq i \leq k} \rho_i} - (1 + p^2 \sum_{1 \leq i \leq k} \rho_i)). \quad (35)$$

Proof: Refer to the previous conference version. ■

By property 3, $Var(\hat{n}_0^c)$ and $Var(\hat{n}_x)$ are tightly bounded by (35). Thus, the two constraints in (33) can be tightened as

$$\frac{1}{p^2} \frac{s_k}{\rho_k} (e^{p \sum_{1 \leq i \leq k} \rho_i} - (1 + p^2 \sum_{1 \leq i \leq k} \rho_i)) \leq (\theta / Z_\delta)^2, \quad (36)$$

which guarantees that (33) is satisfied even in the worst case.

In our system design, we shall configure $\rho = \rho_1 = \rho_2 \dots = \rho_k$ as a system-wide optimal load factor, which will be explained at the end of this subsection. Then, (36) becomes

$$\frac{1}{p^2} \frac{s_k}{\rho} (e^{pk\rho} - (1 + p^2 k\rho)) \leq (\theta / Z_\delta)^2. \quad (37)$$

This subsection focuses on the optimization of the load factor ρ , and keep the sampling probability p temporally fixed, whose optimization is postponed to the next subsection. Equation (37) has no explicit solution for ρ due to the existence of exponential term $e^{pk\rho}$. Hence, we apply the Taylor series $e^x \approx 1 + x + \frac{x^2}{2!} + \frac{x^3}{3!} + \mathcal{O}(x^4)$ to (37). Then,

$$\frac{1}{p^2} \frac{s_k}{\rho} (1 + pk\rho + \frac{(pk\rho)^2}{2} + \frac{(pk\rho)^3}{6} - (1 + p^2 k\rho)) \leq \frac{\theta^2}{Z_\delta^2} \\ \rho \leq \frac{1}{2pk} (-3 + \sqrt{3} \sqrt{8p(\frac{1}{ks_k} \frac{\theta^2}{Z_\delta^2} + 1) - 5}).$$

Assume the size of any tag set s_i is at most $s_{\max}$, and the number of tag sets k involved in any query is at most $k_{\max}$. In the worst case $s_k = s_{\max}$ and $k = k_{\max}$, we must ensure

$$\rho \leq \frac{1}{2pk_{\max}} \left(-3 + \sqrt{3} \sqrt{8p \left(\frac{1}{k_{\max}s_{\max}} \frac{\theta^2}{Z\delta^2} + 1 \right)} - 5 \right). \quad (38)$$

Since $\rho = \rho_i = \frac{s_i}{m_i}$, it is inversely proportional to the frame length m_i , which measures the protocol time cost for encoding the tag set S_i . Thus, we configure the target load factor as large as possible under the constraint, and obtain Eq. (8).

We justify our choice of setting $\rho = \rho_1 = \rho_2 \dots = \rho_k$. The left side of (36) is an increasing function in each ρ_i . If we allow these load factors to be unequal and still set their values to be as small as possible, then some of them will be greater than the right side of (38) and others will be small. Because $S_1, S_2, \dots, S_k$ are arbitrary tag sets under consideration, it means that some tag sets will be encoded with their load factors greater than the right side of (38) and some other will have smaller load factors. Let $S'_1, S'_2, \dots, S'_k$ be the k tag sets with load factors greater than the right side of (38). We should be able to perform joint estimation on any k encoded tag sets without violating the accuracy requirement. However, if we perform joint estimation on $S'_1, S'_2, \dots, S'_k$, because their load factors are larger than (38), the constraint of (36) will not hold.

B. Optimization of Sampling Probability p

Because $\rho = \rho_i = \frac{s_i}{m_i}$, we have $m_i = \frac{s_i}{\rho}$. Recall that the value of m_i must be a power of two to support the expanded OR of multiple ALOHA frames. Hence, $m_i = 2^{\lceil \log_2(\frac{s_i}{\rho}) \rceil}$. We want to choose the optimal sampling probability that minimizes the protocol execution time by keeping the frame length m_i as small as possible. Hence, we have the formula for the frame length defined by section V-A as (7) and quoted here: $m_i = \min_{p \in (0,1]} \{2^{\lceil \log_2(\frac{s_i}{p}) \rceil}\}$, where the load factor ρ is determined by (8), and the sampling probability p is hidden inside (8) and needs to be optimized. The optimal p^* that can minimize m_i depends on the pre-determined parameters $s_{\max}$, $k_{\max}$, θ and δ . We can numerically compute from (7) the optimal sampling probability p^* that minimizes m_i .

VIII. MJREP WITH UNRELIABLE CHANNELS

Our initial design of MJREP protocol makes an implicit assumption that the wireless transmission between a reader and a tag has no error, such that the reader can reliably detect the empty/busy state of each time slot. However, it is inevitable for real-world RFID systems to suffer from environmental noise or interference in wireless channels. Due to the unwanted radio energy of noise or interference, an empty slot which should be detected as a zero bit may be misinterpreted as a one bit. Besides, environmental noise may also affect a busy slot when a reader takes a snapshot of one tag set. But the noise signal will only deform the radio wave of tag's backscattered response in the slot. It is quite unlikely for the random noise to exactly cancel a tag response and turn the busy slot into an empty one. Therefore, we mainly focus on the impact of channel noise on empty slots. Let P_e be the probability for an empty slot to be corrupted by channel error, i.e., detected as a busy slot. Below we will analyze its impact.

We stress that our MJREP protocol is more resilient to random channel noise (which is uniformly distributed in all time slots) than other tag counting work, such as LoF [2]

and PET [5]. This is because LoF and PET use an *exponential distribution*, in which the tag sampling probability reduces exponentially in a sequence of time slots. Higher-order slots have greater impact on the tag counting result than lower-order slots. As a result, the channel noise in higher-order slots will greatly disturb the counting result, even though the mean error rate of the slots can be measured. By contrast, our protocol and ZOE [7] adopt a *uniform distribution* of tags in an ALOHA frame. All slots in the frame are equally important for the counting result. There does not exist any especially important slots whose error rates sensitively affect the counting result. But our protocol is different from ZOE [7], which counts the number of tags in a single tag set. Our aim is to perform joint cardinality estimation for k tag sets, which are encoded into bitmaps B_i , $1 \leq i \leq k$. The impact of unreliable channels must be characterized by a different set of formulas, so that their induced bias in the joint cardinality estimation can be compensated.

Suppose each frame corresponding to the bitmap B_i will experience environmental noise during wireless transmission, turning some of its empty slots into busy slots at a probability P_e . When performing joint estimation, as shown in (9), we need to calculate the bitmap OR_y , the expanded OR of bitmaps B_i satisfying $2^{i-1} \wedge y \neq 0$. Clearly, the bitmap OR_y may be influenced by channel error.

Recall that $X_y^{(j)}$ is the event that the j th bit in OR_y remains zero, whose original probability is given in (10). With the presence of channel error, this probability needs to be updated, as the occurrence of the event now requires two conditions.

- This bit is not assigned by the tags in an arbitrary subset N_x with $x \wedge y \neq 0$, whose probability is given in (10).
 - For any ALOHA frame B_i satisfying $2^{i-1} \wedge y \neq 0$, its $(j \bmod m_i)$ th bit does not experience channel error. Due to the independence of each slot in frame B_i , the probability of this condition is $(1 - P_e)^{\text{bc}(y)}$, where $\text{bc}(y)$ is the number of one-bits in the binary representation of y .
- Thus, the probability of the event $X_y^{(j)}$ is

$$\text{Prob}\{X_y^{(j)}\} = (1 - P_e)^{\text{bc}(y)} \cdot \prod_{1 \leq x < 2^k}^{x \wedge y \neq 0} \left(1 - \frac{p}{m_{10}(x \wedge y)}\right)^{n_x}.$$

Recall that z_y is the fraction of zero bits in the bitmap OR_y . Then, using $z_y \approx E(z_y) = \text{Prob}\{X_y^{(j)}\}$, we have

$$z_y / (1 - P_e)^{\text{bc}(y)} \approx \prod_{1 \leq x < 2^k}^{x \wedge y \neq 0} \left(1 - \frac{p}{m_{10}(x \wedge y)}\right)^{n_x}.$$

Applying logarithm operator to both sides of the above equation, and using $\log(1 + x) \approx x$ for small x value,

$$-\frac{m_{\text{hi}}(y)}{p} \log(z_y / (1 - P_e)^{\text{bc}(y)}) \approx \sum_{1 \leq x < 2^k} \text{sgn}(x \wedge y) \frac{m_{\text{hi}}(y)}{m_{10}(x \wedge y)} n_x,$$

where $\text{sgn}(x)$ is the sign function which equals 1, 0 or -1 when its parameter x is positive, zero or negative. For the above equation, its left-hand side is a constant value once OR_y is determined and channel error rate P_e is empirically known, while the right-hand side is a linear combination of the unknown variables n_x , $1 \leq x < 2^k$. Collocating $2^k - 1$ such equations established from the OR_y bitmaps with $1 \leq y < 2^k$, we can obtain a fully-determined linear equation system:

$$\mathbf{M} \mathbf{n} \approx \mathbf{\hat{u}}_e$$

where $\mathbf{M}$ is given in (16), and $\mathbf{\hat{u}}_e$ is defined as a column vector:

$$\mathbf{\hat{u}}_e = \left[-\frac{m_{\text{hi}}(y)}{p} \log(z_y / (1 - P_e)^{\text{bc}(y)}) \right]^T, \quad 1 \leq y < 2^k.$$

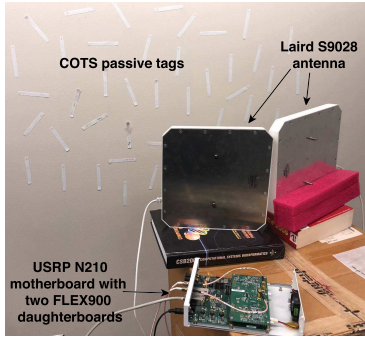


Fig. 4. The setup of our testbed based on USRP reader and COTS tags.

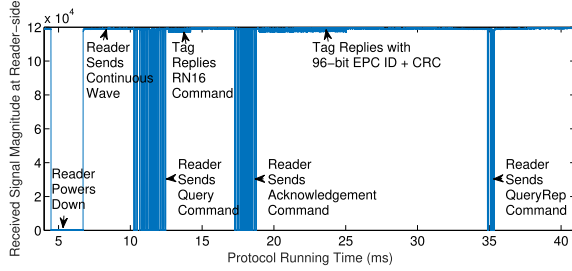


Fig. 5. A reader communicates with a single tag to collect its EPC ID.

Since Property 2 proves the non-singularity of coefficient matrix $\mathbf{M}$, the above linear equation system can always be solved.

IX. PROTOCOL IMPLEMENTATION ISSUES

In this section, we discuss several technical issues for implementing our MJREP protocol in real-world RFID systems.

A. USRP Reader to Query a Small Tag Set

We implement a UHF (ultra high frequency) RFID reader based on USRP (universal software radio peripheral) platform following [15]. Our testbed setup is illustrated in Fig. 4, consisting of an USRP reader and tens of commercial passive RFID tags. We control the USRP reader to interrogate its surrounding RFID tags by EPC C1G2 protocol [14].

In Fig. 5, we illustrate an ALOHA frame when a reader communicates with a single tag to collect its ID. The horizontal axis is time in ms, and the vertical axis is signal magnitude sensed by the reader. When the reader broadcasts a command, the signal magnitude fluctuates significantly. When the tag responds by backscattering, the continuous carrier wave is modulated slightly. From the reader's perspective, an ALOHA frame is composed of four periods: a power-down period (to force each tag to consume up energy reserve and hibernate), a continuous wave transmission period to energize surrounding tags, a *Query* command to notify powered-up tags that a frame with m time slots has started, and $m - 1$ *QueryRep* commands to mark the boundary of each two neighboring slots.

The ALOHA frame in Fig. 5 has only two time slots. The first time slot (roughly from 12.5ms to 35ms) is a singleton slot where only one tag responds. The second time slot is an empty slot where no tags respond. In the singleton slot, a tag reports its ID to the reader. When the tag receives *Query/QueryRep* command that starts its chosen slot, it sends a *RN16* command with a 16-bit random number to help

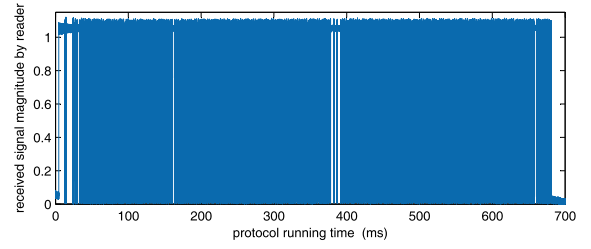


Fig. 6. A reader communicates with multiple tags for snapshot construction.

the reader detect tag collision. If the reader can successfully decode the *RN16*, it sends an *ACK* command to the tag containing the decoded random number. When the tag receives the acknowledgement and finds the number consistent with its own, it responds its 96-bit EPC ID plus CRC checksum code.

Clearly shown in Fig. 5, transmitting a 16-bit random number takes multi-fold less time than collecting a tag ID. To save time, after receiving *RN16* command, the reader can be configured to preempt the ID transmission by immediately sending *QueryRep* command (to start the next slot), without transmitting any *ACK* command which initiates ID transmission. The reader can make a determination of whether the slot is empty or busy based on the 16-bit number alone.

In Fig. 6, we depict a frame when a reader queries more than one tags. In this experiment, we configure the number of time slots in a frame to 256, and each tag only sends in its chosen slot a *RN16* command to reveal its presence. Figure 6 shows that an RFID reader can encode an arbitrary tag set into a bitmap, by checking whether each slot contains a *RN16* response. As long as each tag set has been pseudo-randomly encoded into a bitmap, our MJREP can accept multiple such bitmaps to perform joint cardinality estimation as in Fig. 3.

B. Multi-Reader System for a Large Tag Set

In the previous subsection, a single reader with one antenna is used to encode a small tag set within radio range. For a typical commercial RFID reader, the effective radio range is as small as 3~8 meters when querying battery-free passive tags. It is impossible to use just one reader to cover an entire warehouse or a cargo port. Below, we discuss how to query a large tag set in such a place by leveraging a distributed RFID system with multiple readers and multiple antennas.

In the scenario of a single reader with multiple antennas (for example, Impinj R420 reader can have at most 32 antennas), taking a snapshot of its tag set is not complicated. For most off-the-shelf readers, the antennas on a same reader are activated by a round robin fashion. Hence, for one antenna, when its activation time comes, it can broadcast RFID commands without interference, and all other antennas will keep listening and help the activated antenna to receive tag responses. If any antennas (including the activated one) senses a busy time slot, then the corresponding bit is one. Given multiple bitmap snapshots when different antennas are activated and having the same length, we can merge them by Bitwise OR to construct a snapshot of tags covered by all antennas of the reader.

In the scenario of multiple readers with multiple antennas, taking a snapshot of all tags needs more efforts. The key challenge is that multiple readers have the chance to be activated at the same time and have radio collisions. For example, if a tag locating within the intersected radio zone

of two readers receives commands from both readers at the same time, then the tag with very simple circuit may not correctly resolve reader commands. To address this problem, many reader scheduling algorithms have been proposed by literature [16], [17], which schedules conflicting readers with intersected radio zone to different time intervals. Then, each reader can take a snapshot of its own set of tags without interference by other readers. We can easily combine the snapshots from multiple readers by bitwise OR, as long as these snapshots are of the same length, which can be easily implemented inside a warehouse. From user perspective, these multiple readers can be regarded as one “virtual” reader that takes a bitmap snapshot for the tag set in the entire warehouse.

X. SIMULATION STUDIES

We evaluate the performance of our MJREP protocol by simulations. The most related work include CCF (Composite Counting Framework [13]), which is based on the relative error model, and DiffEstm (Differential Estimator [9]), which is also based on the relative error model and can only handle two tag sets. Please check Section II for discussion on the absolute error model and the relative error model. In this section, we will compare our MJREP protocol with CCF and INC-EXC. Note that INC-EXC protocol in Section IV degrades to DiffEstm [9], when it handles only two tag sets.

We will consider two performance metrics. First, given the same accuracy requirement defined in (4), we compare the execution times of all the three protocols. For MJREP, its execution time is measured as the number of time slots it takes the reader to encode a tag set into a bitmap, including the frame length m_i and other slots needed to give an initial rough estimation of tag set size s_i . We adopt GMLE [3] to generate an initial estimate with a 95% confidence interval of $\pm 20\%$ error. The time cost of GMLE hence is approximately $1.544 \cdot Z_{0.05}^2 / 0.2^2 \approx 148$ slots [3].

Second, when the three protocols are given the same execution time, we compare their probabilities of meeting a given error bound $\pm\theta$. The probability is measured as the number of joint estimations that meet the error bound divided by the total number of joint estimations performed in the simulation. When presenting simulation results, we only show the probability of successfully bounding the estimation error of union cardinality n_0^c , and omit the bounding probability of n_x , since by Property 3, $Var(\hat{n}_x) \leq Var(\hat{n}_0^c)$ and the bounding probability of n_x is always larger than n_0^c .

The system model is a distributed RFID system of multiple business places. Each place is deployed with a reader array to periodically take a snapshot of its set of tags, whose number ranges from 0 to 50000 with $s_{max} = 50000$. The average cardinality of a tag set is $s_{avg} = 10000$, which reflects that the normal business flow of tagged objects is smaller than the worst-case number that the system is designed to handle. The size of each tag set will be taken from a Gaussian distribution $\mathcal{N}(10000, 2000^2)$ truncated by the range $(0, s_{max}]$. For the accuracy requirement, we configure $\delta = 5\%$ and $\theta = 800$ by default. We will perform simulation studies with other values of system parameters δ , θ , s_{max} and s_{avg} as well.

A. Protocol Time Cost to Achieve Same Estimation Accuracy

We compare the average time cost of the three protocols, when they are forced to meet the same accuracy constraint. A critical parameter that affects protocol performance is k_{max} ,

TABLE I
PARAMETER SETTINGS FOR MJREP PROTOCOL

Number of Sets k_{max}	2	4	6	8
Optimal Sampling Probability p^*	1	1	1	1
Theoretical Value of Load factor ρ	0.86	0.28	0.14	0.09
Empirical Value of Load factor ρ	1.39	0.68	0.35	0.13

TABLE II
COMPARISON OF TIME COST AMONG PROTOCOLS

Number of Sets k_{max}	2	4	6	8
Time Cost of MJREP	10,274	21,072	40,234	77,044
Time Cost of INC-EXC	50,920	124,725	230,112	384,616
Time Cost of CCF for Union	42,244	168,976	380,196	675,904

the largest number of tag sets involved in a query. We perform simulations with k_{max} assigned to the value 2, 4, 6 or 8.

Before presenting simulation results, we explain how to theoretically configure the parameters of MJREP. With known error bound θ and δ , we can compute the value of load factor ρ from (8) and the optimal sampling probability p^* from (7). For different k_{max} values, we show the corresponding ρ and p^* in Table I. We find the optimal sampling probability p^* should be configured close to one in order to eliminate sampling error, in most circumstances with reasonably high accuracy.

The theoretical values of ρ are set conservatively (on the third row of Table I) to guarantee that the accuracy constraint is satisfied even in the worst case. Alternatively, their values can be set empirically through simulations for normal situations. In our simulation, we first compute the initial value of ρ from (8) and then perform bi-section search to increase it as large as possible such that the resulting value of m_i will still satisfy the accuracy requirement. Consequently, on the last row of the above table, when k_{max} equals 2, 4, 6 or 8, the load factor ρ is empirically configured to 1.39, 0.68, 0.35 or 0.13, respectively. It shows that, to support the queries that involve more tag sets, the load factor ρ must reduce, which is consistent with Eq. (8).

In the second row of Table II, we present the average execution time of MJREP in simulations (using the empirical parameters in Table I). When INC-EXC and CCF realize the same estimation accuracy, their execution times are shown in the third and fourth rows of Table II, respectively. Because they are not designed for absolute error bound, there is no formula to compute their frame length or the number of hash values stored. With $s_{max} = 50,000$, we use exhaustive search by simulation to find their minimum time cost that can meet the error bound. Table II shows that the frame length used by INC-EXC is at least 500% larger than that of MJREP, and the time cost of CCF is even higher. This is because INC-EXC (or CCF) has to adopt a large frame length (or store a large amount of tag hash values) to tolerate the worst case of estimating joint cardinalities for k_{max} tag sets whose cardinalities range between 45,000 and 50,000. This expensive time cost is fixed even when encoding small tag sets whose average size is only about 10,000. The time costs of CCF can get even worse than the results shown in the last row of Table II, if it is applied to another worst scenario of estimating the intersection of multiple sets, which is empty.

To give a straightforward impression on the time costs of MJREP, INC-EXC and CCF, we configure k_{max} to 4 and show

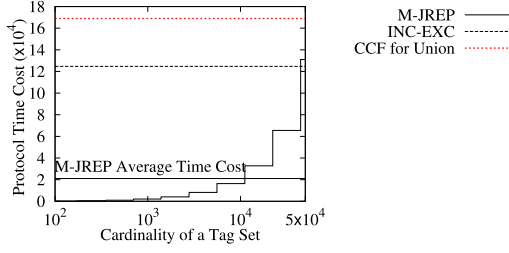


Fig. 7. Protocol execution time under the parameter settings: $k_{\max} = 4$, $s_{\max} = 50000$, $\theta = 800$, $\delta = 5\%$.

TABLE III
ACCURACY WHEN HANDLING $k_{\max}$ LARGE SETS

Number of Sets $k_{\max}$	2	4	6	8
Bounding Probability of MJREP	95%	95%	95%	95%
Bounding Probability of INC-EXC	7.8%	9.6%	18.4%	28.8%

TABLE IV
ACCURACY COMPARISON WHEN HANDLING TWO LARGE TAG SETS AND $k_{\max} - 2$ SMALL TAG SETS

Number of Sets $k_{\max}$	2	4	6	8
Bounding Probability of MJREP	95%	96%	96.6%	99.6%
Bounding Probability of INC-EXC	7.8%	50.4%	86.6%	99%

their comparison results in Fig. 7, where the horizontal axis is the size of a tag set, which varies from 100 to 50,000, and the vertical axis is the number of time slots needed (or hash values stored for CCF) to take a snapshot of the tag set. Due to the nature of their designs, INC-EXC uses a constant frame length of 124,725 slots, and CCF uses constant time cost of recording 168,976 hash values. The frame length of MJREP is variable. It is small when the tag set is small. For example, for a set of 10,000 tags, the number of time slots needed by MJREP is $2^{\lceil \log_2(10,000/0.68) \rceil} + 148 = 16,532$, only 13% of what is needed by INC-EXC. The average time cost of MJREP is 21,072 shown by the solid horizontal line.

B. Estimation Accuracy Under the Same Execution Time

We compare the estimation accuracy of INC-EXC and MJREP, when giving them the same execution time that is configured by the second row of Table II. Here, we omit the results of CCF for space, which are worse than INC-EXC.

When presenting simulation results, a difficulty is that the estimation accuracy is strongly affected by the sizes of tag sets involved and their ways of overlapping. It is impossible to present the simulation results of all the cases, and hence we focus on only two of them. The first is an extreme case that deals with $k_{\max}$ large sets (sized from 45,000 to 50,000) which are slightly overlapped. The second case can be regarded as a normal case that handles two large sets and $k_{\max} - 2$ small sets whose sizes randomly distribute between 0 and 5,000.

The simulation results of the extreme case are shown in Table III, and the results of the normal case is in Table IV. In both tables, our MJREP protocol performs well, because its probability of bounding absolute estimation error with θ is always above $1 - \delta = 95\%$. In contrast, the accuracy of INC-EXC is non-satisfactory for the normal case in Table IV, and severely degrades when handling the extreme case in Table III.

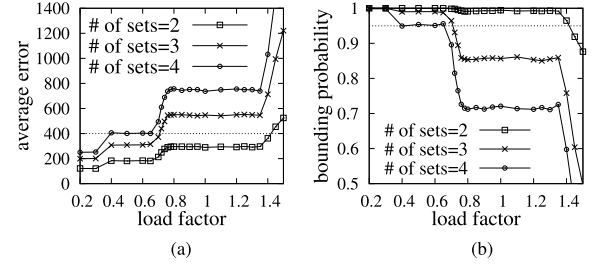


Fig. 8. Impact of load factor ρ on estimation accuracy of joint cardinalities. (a) Absolute estimation error. (b) Successful bounding rate.

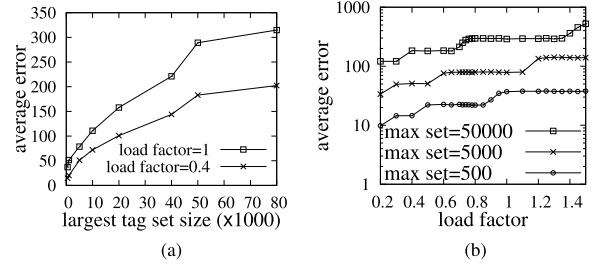


Fig. 9. Impact of largest tag set size $s_{\max}$ on average estimation error of MJREP, while the number of tag sets k is fixed to two. (a) Absolute error vs. tag set size. (b) Absolute error vs. load factor.

C. Impact of Load Factor on MJREP's Estimation Accuracy

In previous subsections, we have studied the impact of the number of tag sets $k_{\max}$ on estimation accuracy. Below we evaluate how the accuracy of MJREP is influenced by the load factor ρ . The length of each ALOHA frame m_i is inversely proportional to ρ as shown in equation (7).

We illustrate in Fig. 8(a) the relation between the load factor ρ and the average estimation error of union set n_0^c , which has the largest error among all joint cardinalities n_x by Property 3. Clearly, in the above plot, as the load factor ρ grows, the average estimation error increases, due to the smaller frame length m_i shown in (7). It is surprising that the error curves in the plot are staircase functions of ρ , instead of some smooth functions. This is because our MJREP protocol requires the length of an arbitrary frame to be a power of two, so that it is always an integer multiple of the length of another shorter frame when we perform the expanded OR operation in Fig. 3. Note that it is also a requirement by EPC C1G2 standard that the length of a frame should be a power of two [14].

We depict in Fig. 8(b) the relation between the load factor ρ and the probability of keeping estimation error within a threshold $\theta = 800$. In the plot, the largest load factor ρ that can ensure 95% bounding probability is about 1.4, 0.7 or 0.6, when the number of tag sets k is 2, 3 or 4, respectively. These ρ values are consistent with the crossing points in Fig. 8(a) between curves and a horizontal line $y = 400$, which keep the average estimation error under a threshold $\theta/Z_{0.95} = 800/Z_{0.95} = 400$, where Z_δ is defined in (33). This verifies that MJREP's estimation results follow Gaussian distributions.

D. Impact of Tag Set Size on MJREP's Estimation Error

We additionally evaluate the impact of the largest tag set size $s_{\max}$ on the average estimation error of MJREP. We illustrate the evaluation result in Fig. 9, where the y-axis is the average absolute estimation error of union set size. The plot (a) shows that MJREP's error increases as the upper bound of tag set size $s_{\max}$ grows, when the load factor is fixed to

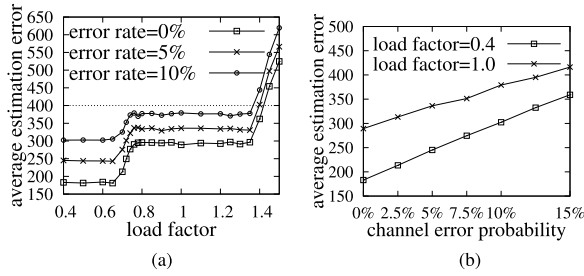


Fig. 10. Impact of channel error rate P_{err} on average estimation error of MJREP, while the number of tag sets k is fixed to two. (a) Absolute error vs. load factor. (b) Absolute error vs. error rate.

1 or 0.4. The plot (b) shows that MJREP's error increases as s_{max} grows, regardless of the value of load factor ρ . The reason, as shown by (32) and (24), is that the estimation variance for the union of all tag sets $Var(\hat{n}_0^c)$ is linear to the length of ALOHA frame m_i , which is roughly proportional to the tag set size s_i as in Eq. (7). Hence, the standard deviation of the union set estimation $\sqrt{Var(\hat{n}_0^c)}$ is roughly linear to $\sqrt{s_{max}}$, the square root of the upper bound of tag set size.

E. MJREP's Estimation Error Under Unreliable Channels

In this subsection, we evaluate the average estimation error of MJREP under unreliable wireless channels. To simulate the unreliable channel model whose expected error rate is P_{err} , we generate random numbers uniformly distributed in the range of $[0, 1]$. For each zero bit in a snapshot, if the random number generated is smaller than P_{err} , we flip that bit to one, which simulates one occurrence of time slot corruption.

In Fig. 10(a), we illustrate the average estimation error of MJREP, when the channel error rate P_{err} is configured to 0%, 5% and 10%. The plot shows that the error of MJREP increases as P_{err} grows. As a result, when P_{err} grows, the load factor ρ must be set with a smaller value, if we want to ensure the average estimation error is smaller than $\theta/Z_\delta = 400$.

In Fig. 10(b), We fix the load factor ρ to 0.4 or 1, and evaluate the increasing speed of estimation error relative to the channel error rate P_{err} . It shows that MJREP's estimation error linearly increases as P_{err} grows, when ρ is fixed.

XI. RELATED WORK

Much existing RFID work concentrates on how to efficiently collect the IDs of a group of tags, which is called *tag identification*. Since the tags communicate with a reader through wireless medium, inevitably collisions will happen when multiple tags respond to the same reader simultaneously. Collision arbitration protocols mainly fall into two categories, i.e., tree-based protocols [10], and slotted ALOHA protocols [11]. The de-facto RFID standard, named EPC C1G2, is a variant of the slotted ALOHA protocol [14].

Another branch of RFID research studies how to accurately estimate the cardinality of a tag set at low time cost without collecting any IDs. To minimize the time cost, a plethora of protocols have been developed, including unified probabilistic estimator [1], lottery frame protocol [2], generalized maximum likelihood estimation [3], first non-empty slot based estimation [4], probabilistic estimating tree [5], average run based tag estimation [6], and zero-one estimator [7]. An important recent

work proposed a two-phase protocol named SRC_s [8], which uses a first phase to quickly make a rough estimation of the tag cardinality, and a second phase based on ALOHA frame for achieving better accuracy. Our MJREP also adopts a two-phase protocol for efficiently encoding a tag set into a bitmap.

A recent research direction is to extend the tag counting problem from a single set to multiple sets. Some researchers focus on two tag sets scanned by a reader at different times, and estimate the cardinalities of their intersection/differences [9], [18]–[21]. Such information can help count the number of missing tags (which exist in the previous tag set, but no longer in the current), remaining tags (existing in both sets), and new tags (opposite to missing tags). Another work named CCF can estimate the cardinality of an arbitrary set expression [13]. It encodes each tag set into a sketch named *k-min hash values* [22], which can support the merging of multiple sketches. Its shortcoming is that it needs the sketches of all tag sets to be configured with the same k value.

However, for this tag counting problem with multiple sets, the previous studies are still limited from three perspectives. Firstly, most of them are not designed to estimate cardinality of a general set expression, except the work in [13]. Secondly, all of them specifies the accuracy requirement by the relative error model. Unfortunately, when the quantity to estimate approaches zero, their time cost to meet the accuracy requirement skyrockets to infinity (see Section II for detailed discussion). The correct choice is to adopt instead the absolute error model. Thirdly, previous work requires all tag sets must be compressed into data structures (called snapshots) with the same length, such that multiple snapshots can be merged easily to estimate the union of multiple sets [9], [13], [18], [19]. However, these snapshots may not have an equal length, especially when the tag sets they encode dramatically differ in sizes, which regrettably is commonly seen in real worlds.

A very recent work [12] addresses the above third problem by allowing the bitmap-based snapshots of tag sets to have adaptively different lengths, in order to reduce protocol time cost. But the paper is still inadequate in that it only deals with the joint cardinality estimation of two tag sets. In many applications, it often requires to estimate the cardinality of a general set expression that may involve an arbitrary number of tag sets. Our paper can solve this problem time-efficiently.

XII. CONCLUSION

In this paper, we have formulated a problem called *joint cardinality estimation*, in which the cardinality of an arbitrary set expression (involving multiple tag sets from different spatial or temporal domains) is estimated with bounded absolute error. We propose a protocol named MJREP with a novel design that allows multiple tag sets to be encoded into bitmaps with varied lengths. It provides a new method called *expanded OR* to combine the multiple bitmaps, and it designs formulas to exploit the combined information, estimate the cardinalities of all elementary subsets, and finally calculate the cardinality of the desired set expression. We have analyzed the bias and variance of MJREP, and also the optimal setting of its protocol parameters under predefined accuracy requirements. We have performed extensive simulation studies. The results show that our protocol can reduce the execution time by multiple folds as compared with INC-EXC and CCF protocols, which require all tag sets must be encoded into equal-length snapshots.

APPENDIX

PROBABILITY OF A TAG IN N_x SETTING A BIT IN OR_y

We firstly define a few notations. In the binary format of x , the series of one-bits from low end to high is at positions

$$\ell(x, 1), \ell(x, 2), \dots, \ell(x, \text{bc}(x)), \quad (39)$$

where $\ell(x, i)$ is the location of the i th one-bit in x , and $\text{bc}(x)$ is the number of one-bits in x (or called the bit count of x). For simplicity, we denote

- the location of the lowest-order 1-bit $\ell(x, 1)$ by $\text{lo}(x)$, and
- the location of the highest-order 1-bit $\ell(x, \text{bc}(x))$ by $\text{hi}(x)$.

By the definition of elementary subset N_x in (3), the one-bits in binary format of x decide which tag sets will include N_x , i.e., $S_{\ell(x,1)}, S_{\ell(x,2)}, \dots, S_{\ell(x,\text{bc}(x))}$. Since the tag set S_i are encoded by the ALOHA frames B_i , the one-bits in x also decide which frames may receive responses from tags in N_x .

$$B_{\ell(x,1)}, B_{\ell(x,2)}, \dots, B_{\ell(x,\text{bc}(x))} \quad (40)$$

Property 4: For an arbitrary tag from the elementary subset N_x , it may respond in (or be encoded by) the frames $B_{\ell(x,1)}, B_{\ell(x,2)}, \dots, B_{\ell(x,\text{bc}(x))}$. The tag will be sampled to respond either in all these frames or in none of them, since the sampling process is performed in a pseudorandom fashion.

Property 5: Assume a tag in N_x is sampled to respond, and its list of encoding frames in (40) includes $B_{\ell(x,i)}$ and $B_{\ell(x,i')}$, where the former frame is no longer than the latter $m_{\ell(x,i)} \leq m_{\ell(x,i')}$. For an arbitrary slot number j , if the tag does not pick the $(j \bmod m_{\ell(x,i)})$ th slot in frame $B_{\ell(x,i)}$, it will neither select the $(j \bmod m_{\ell(x,i')})$ th slot in $B_{\ell(x,i')}$.

Proof: Suppose a tag id is sampled and does not select the $(j \bmod m_{\ell(x,i)})$ th slot in frame $B_{\ell(x,i)}$, i.e., $H(id \oplus R) \neq j \bmod m_{\ell(x,i)}$. Since both $m_{\ell(x,i)}$ and $m_{\ell(x,i')}$ are the powers of two and $m_{\ell(x,i)} \leq m_{\ell(x,i')}$, the former is able to divide the latter. Thus, $H(id \oplus R) \neq j \bmod m_{\ell(x,i')}$, implying that the j th slot in $B_{\ell(x,i')}$ is not selected by the tag. ■

Property 6: Among the list of frames in (40), if in the first frame $B_{\ell(x,1)}$, a tag in subset N_x does not select the $(j \bmod m_{\ell(x,1)})$ th slot, then in any subsequent frame $B_{\ell(x,i)}$ with $i > 1$, the tag neither selects the $(j \bmod m_{\ell(x,i)})$ th slot.

Consider an arbitrary tag id in elementary subset N_x . As mentioned in (40), the frames that may receive responses of the tag id are $B_{\ell(x,1)}, B_{\ell(x,2)}, \dots, B_{\ell(x,\text{bc}(x))}$. Further consider the bitmap OR_y , which is the expanded OR of bitmaps $B_{\ell(y,1)}, B_{\ell(y,2)}, \dots, B_{\ell(y,\text{bc}(y))}$ as defined in equation (9). Among these selected bitmaps as marked by y , the bitmaps that may receive the response of tag id in the subset N_x are

$$B_{\ell(x \wedge y, 1)}, B_{\ell(x \wedge y, 2)}, \dots, B_{\ell(x \wedge y, i)}, \dots, B_{\ell(x \wedge y, \text{bc}(x \wedge y))}. \quad (41)$$

According to (9), the j th bit of bitmap OR_y is the OR of the $(j \bmod m_{\ell(y,i)})$ th bit in bitmap $B_{\ell(y,i)}$ with the index i ranging from 1 to $\text{bc}(y)$. But the tag in subset N_x only appears in the list of bitmaps in (41). Hence, we only need to analyze the probability for the tag to assign the $(j \bmod m_{\ell(x \wedge y, i)})$ th bit in bitmap $B_{\ell(x \wedge y, i)}$ with $i \in [1, \text{bc}(x \wedge y)]$.

As explained in Property 6, if the tag id does not select the $(j \bmod m_{\ell(x \wedge y, 1)})$ th slot in $B_{\ell(x \wedge y, 1)}$, then it neither selects the $(j \bmod m_{\ell(x \wedge y, i)})$ th slot in $B_{\ell(x \wedge y, i)}$ for any i value. Hence, the probability that the tag id picks the j th slot in OR_y equals the probability that the tag assigns the $(j \bmod m_{\ell(x \wedge y, 1)})$ th bit in $B_{\ell(x \wedge y, 1)}$, i.e., $\frac{p}{m_{\ell(x \wedge y, 1)}} = \frac{p}{m_{\text{lo}(x \wedge y)}}$, where $\text{lo}(x \wedge y)$ and $\ell(x \wedge y, 1)$ are the location of lowest 1-bit in $x \wedge y$.

REFERENCES

- [1] M. Kodialam and T. Nandagopal, "Fast and reliable estimation schemes in RFID systems," in *Proc. 12th Annu. Int. Conf. Mobile Comput. Netw.*, Sep. 2006, pp. 322–333.
- [2] C. Qian, H. Ngan, and Y. Liu, "Cardinality estimation for large-scale RFID systems," in *Proc. 6th Annu. IEEE Int. Conf. Pervasive Comput. Commun. (PerCom)*, Mar. 2008, pp. 30–39.
- [3] T. Li, S. Wu, S. Chen, and M. Yang, "Energy efficient algorithms for the RFID estimation problem," in *Proc. IEEE INFOCOM*, Mar. 2010, pp. 1–9.
- [4] H. Han *et al.*, "Counting RFID tags efficiently and anonymously," in *Proc. IEEE INFOCOM*, Mar. 2010, pp. 1–9.
- [5] Y. Zheng, M. Li, and C. Qian, "PET: Probabilistic estimating tree for large-scale RFID estimation," in *Proc. 31st Int. Conf. Distrib. Comput. Syst. (ICDCS)*, Jun. 2011, pp. 37–46.
- [6] M. Shahzad and A. X. Liu, "Every bit counts: Fast and scalable RFID estimation," in *Proc. 18th Annu. Int. Conf. Mobile Comput. Netw. (MOBICOM)*, Aug. 2012, pp. 365–376.
- [7] Y. Zheng and M. Li, "Towards more efficient cardinality estimation for large-scale RFID systems," *IEEE/ACM Trans. Netw.*, vol. 22, no. 6, pp. 1886–1896, Dec. 2014.
- [8] B. Chen, Z. Zhou, and H. Yu, "Understanding RFID counting protocols," in *Proc. 19th Annu. Int. Conf. Mobile Comput. Netw. (MOBICOM)*, Sep./Oct. 2013, pp. 291–302.
- [9] Q. Xiao, B. Xiao, and S. Chen, "Differential estimation in dynamic RFID systems," in *Proc. IEEE INFOCOM*, Apr. 2013, pp. 295–299.
- [10] J. Myung and W. Lee, "Adaptive splitting protocols for RFID tag collision arbitration," in *Proc. 7th ACM Int. Symp. Mobile Ad Hoc Netw. Comput. (MOBIHOC)*, May 2006, pp. 202–213.
- [11] S.-R. Lee, S.-D. Joo, and C.-W. Lee, "An enhanced dynamic framed slotted ALOHA algorithm for RFID tag identification," in *Proc. 2nd Annu. Int. Conf. Mobile Ubiquitous Syst., Netw. Services*, Jul. 2005, pp. 166–172.
- [12] Q. Xiao, M. Chen, S. Chen, and Y. Zhou, "Temporally or spatially dispersed joint RFID estimation using snapshots of variable lengths," in *Proc. 16th ACM Int. Symp. Mobile Ad Hoc Netw. Comput. (MOBIHOC)*, Jun. 2015, pp. 247–256.
- [13] H. Liu *et al.*, "Generic composite counting in RFID systems," in *Proc. IEEE 34th Int. Conf. Distrib. Comput. Syst. (ICDCS)*, Jun./Jul. 2014, pp. 597–606.
- [14] *EPC Radio-Frequency Identity Protocols Generation-2 UHF RFID Protocol for Communications at 860 MHz–960 MHz v2.0.0*, EPCglobal, Apr. 2014. [Online]. Available: <https://en.wikipedia.org/wiki/EPCglobal>
- [15] M. Buettner and D. Wetherall, "A software radio-based UHF RFID reader for PHY/MAC experimentation," in *Proc. IEEE Int. Conf. RFID*, Apr. 2011, pp. 134–141.
- [16] J. Waldrop, D. W. Engels, and S. E. Sarma, "Colorwave: An anticollision algorithm for the reader collision problem," in *Proc. IEEE Int. Conf. Commun. (ICC)*, May 2003, pp. 1206–1210.
- [17] Z. Zhou, H. Gupta, S. R. Das, and X. Zhu, "Slotted scheduled tag access in multi-reader RFID systems," in *Proc. IEEE Int. Conf. Netw. Protocols (ICNP)*, Oct. 2007, pp. 61–70.
- [18] C. C. Tan, B. Sheng, and Q. Li, "How to monitor for missing RFID tags," in *Proc. 28th Int. Conf. Distrib. Comput. Syst. (ICDCS)*, Jun. 2008, pp. 295–302.
- [19] H. Liu, W. Gong, X. Miao, K. Liu, and W. He, "Towards adaptive continuous scanning in large-scale RFID systems," in *Proc. IEEE INFOCOM*, Apr./May 2014, pp. 486–494.
- [20] X. Liu, B. Xiao, S. Zhang, and K. Bu, "Unknown tag identification in large RFID systems: An efficient and complete solution," *IEEE Trans. Parallel Distrib. Syst.*, vol. 26, no. 6, pp. 1775–1788, Jun. 2015.
- [21] X. Liu, S. Zhang, B. Xiao, and K. Bu, "Flexible and time-efficient tag scanning with handheld readers," *IEEE Trans. Mobile Comput.*, vol. 15, no. 4, pp. 840–852, Apr. 2016.
- [22] Z. Bar-Yossef, T. S. Jayram, R. Kumar, D. Sivakumar, and L. Trevisan, "Counting distinct elements in a data stream," in *Proc. 6th Int. Workshop Randomization Approximation Techn. (RANDOM)*, Sep. 2002, pp. 1–10.



Qingjun Xiao (M'10) received the B.Sc. degree from the Department of Computer Science, Nanjing University of Posts and Telecommunications, China, in 2003, the M.Sc. degree from the Department of Computer Science, Shanghai Jiao Tong University, China, in 2007, and the Ph.D. degree from the Department of Computer Science, The Hong Kong Polytechnic University, in 2011. After graduation, he joined Georgia State University and the University of Florida, where he was a Post-Doctoral Researcher for three years combined. He is currently

an Assistant Professor with Southeast University, China. His research interests include protocol and algorithm design in wireless sensor networks, RFID systems, and network traffic measurements. He is a member of the ACM.



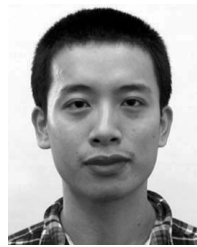
Youlin Zhang received the B.S. degree in electronic information engineering from the University of Science and Technology of China, Hefei, China, in 2014. He is currently pursuing the Ph.D. degree in computer and information science and engineering with the University of Florida, Gainesville, FL, USA, under the supervision of Prof. S. Chen. His research interests include the Internet of Things and big network data.



Shigang Chen (M'02–SM'12–F'16) received the B.S. degree in computer science from the University of Science and Technology of China, Hefei, China, in 1993, and the M.S. and Ph.D. degrees in computer science from the University of Illinois at Urbana-Champaign, USA, in 1996 and 1999, respectively.

After graduation, he was with Cisco Systems, Inc., San Jose, CA, USA, for three years before joining the University of Florida, Gainesville, FL, USA, in 2002, where he is currently a Professor with the Department of Computer and Information Science and Engineering. He has served on the Technical Advisory Board of Protego Networks, Sunnyvale, CA, USA, from 2002 to 2003. He has published over 100 peer-reviewed journals/conference papers. He holds 11 U.S. patents. His research interests include computer networks, Internet security, wireless communications, and distributed computing.

Dr. Chen is an Associate Editor of the IEEE/ACM TRANSACTIONS ON NETWORKING, *Computer Networks*, and the IEEE TRANSACTIONS ON VEHICULAR TECHNOLOGY. He has served on the Steering Committee of the IEEE IWQoS from 2010 to 2013. He has received the IEEE Communications Society Best Tutorial Paper Award in 1999 and the NSF CAREER Award in 2007.



Min Chen received the B.E. degree in information security from the University of Science and Technology of China in 2011, and the M.S. and Ph.D. degrees in computer science from the University of Florida in 2015 and 2016, respectively. His Ph.D. advisor was Prof. S. Chen. He is currently a Software Engineer with Google Inc., Mountain View, CA, USA. His research interests include the Internet of Things, big network data, next-generation RFID systems, and network security.



Jia Liu (M'12) received the B.E. degree in software engineering from Xidian University, Xi'an, China, in 2010, and the Ph.D. degree in computer science and technology from Nanjing University, Nanjing, China, in 2016. He is currently a Research Assistant Professor with the Department of Computer Science and Technology, Nanjing University. His research interest mainly focuses on RFID technology. He is a member of CCF.



Guang Cheng (SM'07) received the B.S. degree in traffic engineering from Southeast University, Nanjing, China, in 1994, the M.S. degree in computer application from the Hefei University of Technology in 2000, and the Ph.D. degree in computer network from Southeast University in 2003. He is currently a Full Professor with the School of Cyber Science and Engineering, Southeast University. His research interests are network security, network measurements, and traffic behavior analysis.



Junzhou Luo (M'07) received the B.S. degree in applied mathematics and the M.S. and Ph.D. degrees in computer network from Southeast University, Nanjing, China, in 1982, 1992, and 2000, respectively. He is currently a Full Professor with the School of Computer Science and Engineering, Southeast University. His research interests are next generation network architecture, network security, and cloud computing. He is a member of the IEEE Computer Society. He is also a member of the ACM.

He is the Chair of the ACM SIGCOMM in China and the Co-Chair of the IEEE SMC Technical Committee on Computer Supported Cooperative Work in Design.