

Chapter 2

Per-Flow Size Measurement

Per-flow size measurement, which is to count the number of packets for each active flow during a certain measurement period, has many applications in usage accounting, traffic engineering, service provision and anomaly detection. In order to maintain the high throughput of routers or switchers, the per-flow size measurement module should use high-bandwidth SRAM that allows fast memory accesses. Due to the limited SRAM space, exact counting, which requires to keep a counter for each flow, does not scale to measure big network data consisting of numerous flows. Some recent work takes a different design path to accurately estimate the flow sizes using counter architectures that can fit into tight SRAM. However, existing counter architectures have some limitations, either still requiring considerable SRAM space, or having a very small estimation range. This chapter presents a scalable counter architecture, called Counter Tree, which leverages a two-dimensional counter sharing technique to achieve far better memory efficiency and significantly extend estimation range. Furthermore, we improve the performance of Counter Tree by adding a status bit to each counter. The extensive experiments with real network trace demonstrate that the new architecture can produce accurate estimates for flows of all sizes even under a very tight memory space, e.g., 1 bit per flow.

2.1 Problem Statement

Per-flow size measurement is one of the fundamental problems in network traffic measurement [8, 10–13, 15–17]. It is to count the number of packets (or called flow size) for each active flow during a measurement period, e.g., 1 min or the time for processing 10 million packets. The flows under measurement can be per-source flows, per-destination flows, per-source/destination flows, TCP flows, http flows, or any user-defined logical flows. Each flow is uniquely identified by its flow label—for example, the flow labels for per-source flows are the source addresses.

Three metrics are used to evaluate the performance of different per-flow size measurement schemes:

1. **Memory requirement:** Due to the constraint of SRAM space, we want to use as small memory as possible for per-flow size measurement. Here we focus on the memory requirement for implementing the counter architectures, while the memory overhead for flow label collection is not our concern. Some memory-efficient approaches [12] for flow label collection can be found in literature.
2. **Processing time:** To keep up with the line speeds, the processing time for encoding a packet should be small, such that the implementation of the measurement module will not deteriorate throughput. In most counter architectures [8, 10, 13], the processing time for encoding a packet mainly results from the memory accesses and hash computations.
3. **Measurement accuracy:** Given a particular memory space, the measurement results of flow sizes should be as accurate as possible. Suppose the true size of a flow is s , and the measured size is $\hat{s}$. We use the relative bias $Bias(\frac{\hat{s}}{s})$ and relative standard error $StdErr(\frac{\hat{s}}{s})$ to evaluate the measurement accuracy, which are defined as follows:

$$Bias(\frac{\hat{s}}{s}) = E(\frac{\hat{s}}{s}) - 1, \quad (2.1)$$

$$StdErr(\frac{\hat{s}}{s}) = \sqrt{Var(\frac{\hat{s}}{s})} = \frac{\sqrt{Var(\hat{s})}}{s}. \quad (2.2)$$

Notations used in the chapter are given in Table 2.1 for quick reference.

Table 2.1 Notations

Symbols	Descriptions
s	True size of a flow
$\hat{s}$	Estimated size of a flow
M	Available memory space in bits
b	Number of bits in each physical counter
m	Number of leaf nodes in the Counter Tree
d	Degree of each non-leaf node in the Counter Tree
r	Number of virtual counters for each flow
h	Height of the Counter Tree
$C[i]$	The i th physical counter
$V[i]$	The i th virtual counter
k	Number of leaf nodes in a subtree counter
n	Total number of packets in the measurement period

2.2 Prior Art

With tremendous number of flows under measurement for big network data, it is impossible to keep an individual counter for each flow in SRAM. Therefore, exact counting generally adopts a hybrid SRAM–DRAM architecture [15–17], where small counters in SRAM are incremented at high speed, and occasionally written back to larger counters in DRAM. However, the hybrid architecture incurs very costly SRAM-to-DRAM updates. Furthermore, the flow-to-counter association requires considerable SRAM [13].

To fit the measurement module in tight SRAM, some schemes only give the distribution of flow sizes [3, 7], or measure the sizes of large flows [5, 6]. Some recent work takes a different design path to accurately estimate the flow sizes instead of counting their exact sizes, thereby reducing memory overhead. The state-of-the-art estimation approaches include bitmap-based MSCBF, and counter-based Counter Braids and randomized counter sharing scheme.

The *Multiresolution Space-Code Bloom Filter* (MSCBF) [8] employs multiple Bloom filters to encode packets with different sampling probabilities. Filters with high sampling probabilities can keep track of small flows, while filters with low sampling probabilities can track large flows. However, the bitmap nature of MRSCBF determines that it is not memory efficient for counting [10]. TinyTable [4] is a novel hash table-based data structure that represents multiset membership. It improves the query and update efficiency of Bloom filters. However, for per-flow traffic measurement, it still requires a high memory cost, which is tens of bits per flow [4].

The *Counter Braids* (CB) [12, 13] is a counter architecture for flow size measurement. It avoids the storage of flow-to-counter association by hashing flows to counters on the fly, and it reduces memory requirement by sharing counters among flows. A typical implementation of CB consists of two layers of counters, and employs three hash functions. To encode a packet, it is hashed to three counters based on its flow label, which are all incremented by one. If any of the first-layer counter overflows, another three second-layer counters will be used. Since each counter is shared by multiple flows, it counts all associated flows. Therefore, the counters essentially form a set of linear equations of the flow sizes. A message passing reconstruction algorithm was used to estimate the flow sizes in an iterative way. CB can recover the exact flow sizes when sufficient memory is available, e.g., 10 bits per flow. However, CB has three limitations. First, it performs 6 (occasionally 12) memory accesses to encode one packet. Second, it yields very biased or even meaningless estimates under a tight memory, e.g., less than 4 pits per flow. In fact, we find that the estimation results of CB do not converge even with 8 bits per flow, though it may occasionally produce very accurate results if we manually terminate the process after some iterations. Third, CB does not support instantaneous queries of flow sizes. All flow sizes must be decoded together at the end of a measurement period.

A new data encoding/decoding scheme, called *randomized counter sharing* [10], was designed to further reduce the memory requirement and processing time of per-flow size measurement. The idea is to split each flow among a number of counters (called the *storage vector* of the flow) that are randomly selected from a counter pool. When encoding a packet of a particular flow, it is randomly mapped to a counter of the flow’s storage vector, and the counter is then incremented by one. This scheme requires only 2 memory accesses for encoding one packet, achieving the optimal processing speed. Moreover, it can still yield reasonably accurate estimates under a tight memory space where CB no longer works. Two estimation methods CSM and MLM are used to estimate flow sizes. The most serious problem of this scheme is that its estimation range is limited, e.g., a few thousands in a typical implementation. For large flows with sizes beyond the estimation range, the scheme leads to very negatively biased estimates since overflowed counters lose information. In the journal version [11], some approaches were provided to extend the estimation range, which, however, cannot address the issue fundamentally. The first approach is to increase the length of each counter or the size of the storage vector. However, this approach degrades estimation accuracy since fewer counters are available or each counter is shared by more flows. Following a reasonable parameter setting, the estimation range is still very limited. The second approach employs a sampling module. Each arriving packet is sampled with a probability p before being encoded to a counter. Aggressive sampling not only introduces significant error [8], but also fails to measure some small size or even moderate-size flows. For example, if we let $p = 0.001$, flows with sizes less than 1000 are hardly be captured. The final approach resorts to the hybrid SRAM/DRAM design, which requires costly SRAM-to-DRAM updates.

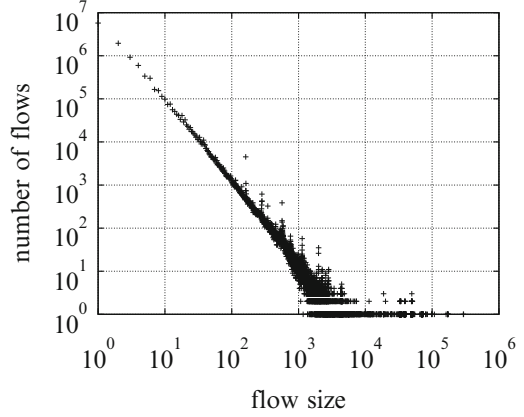
2.3 Design of Counter Tree Architecture

In this section, we provide a novel scalable counter architecture called Counter Tree [1].

2.3.1 Motivation

In spite of the large number of flows in networks, many studies reveal a common observation that a small percentage of large flows account for a high percentage of the traffic (also known as the heavy-tailed distribution). The study in [14] showed that the top 15 % of the destination prefixes (per-destination flows) account for over 95 % of the byte traffic. As an example, we use a network trace obtained from the main gateway of University of Florida, which contains about 68 million TCP flows and 750 million packets. The distribution of flow sizes is illustrated in Fig. 2.1, where each point represents the number (y coordinate) of flows that have a particular

Fig. 2.1 Distribution of flow sizes, where each point represents the number (y coordinate) of flows that have a particular size (x coordinate)



size (x coordinate). This log-scale figure demonstrates that the vast majority of flows have small sizes, while only a small number of flows have large sizes. Without knowing the flow sizes beforehand (which are in fact what we want to measure), the length of counters should be set according to the maximum flow size, which may need to be as large as 64 bits [15]. However, if a flow turns out to be small, e.g., with a size of 1, most of the bits in its counter will be wasted. This observation motivates us to save memory by utilizing the waste.

2.3.2 Two-Dimensional Counter Sharing

To reduce the memory waste caused by small flows, we should enable counter sharing. Therefore, we develop a novel counter sharing technique called two-dimensional counter sharing, which includes horizontal counter sharing and vertical counter sharing. The available memory space is divided into small physical counters, which are logically organized on different layers. Horizontal counter sharing means that any counter on any layer can be shared by different flows, and vertical counter sharing means that each higher-layer counter can be shared by multiple lower-layer counters, acting as their high-order bits.

In horizontal counter sharing, each counter is shared by multiple flows. The rationale is to let large flows borrow memory from small or medium flows that will not fully use their counters. But each counter will have to store the combined size of multiple flows particularly if the number of counters is much smaller than the number of flows. To alleviate this problem, the CountMin approach [2] maps each flow to multiple counters and use the smallest counter value as the flow size. The problems are that each packet of a flow will result in multiple counter updates (which reduces the processing throughput multiple folds) and that the smallest counter may still be the sum of multiple flow sizes, causing positive bias in estimation. Our new design in the next subsection will keep the benefit of counter sharing, while ensuring that approximately one counter is updated per packet and in the meantime avoiding

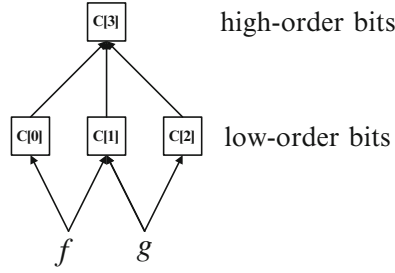


Fig. 2.2 An illustration of two-dimensional counter sharing, where $C[0]$, $C[1]$, $C[2]$, and $C[3]$ are four small physical counters, and f and g are two flows. Each low-order counter can be shared by multiple flows, and each high-order counter can be shared by multiple low-order counters as their high-order bits

the positive bias in estimation. The idea of our design is to split each flow to multiple counters through random mapping, and each counter can therefore be shared by different flows. In the example of Fig. 2.2, the counter $C[1]$ is shared by two flows f and g . If f is a small flow and g is a large flow, g can help make full use of $C[1]$.

Horizontal counter sharing allows some bits wasted by small or medium flows to be used by large flows. However, since small flows account for a dominant percentage of network flows, many bits in counters occupied only by small flows can still be wasted. This observation leads to the idea of *vertical counter sharing* below, which allows the more significant bits (i.e., higher-order bits) to be shared by more flows. As an example, the counter $C[3]$ in Fig. 2.2 serves as extra high-order bits for $C[0]$, $C[1]$, and $C[2]$, and any flows mapped to $C[0]$, $C[1]$, and $C[2]$ can use $C[3]$ when necessary.

We introduce the concept of virtual counters, each of which is the concatenation of multiple small physical counters, which are also called the *component counters* of the virtual counter. Physical counters do not share their bits, but virtual counters share bits by sharing their component counters, particularly those representing more significant bits in virtual counters. In essence, vertical counter sharing implements seamless dynamic memory allocation based on flow sizes, without having to allocate or deallocate physical counters of different sizes on the fly (which is too costly to implement on chip). Suppose we have three physical counters, each with 3 bits. The first counter is allocated to f , the second is for g , while the third is reserved for whoever needs it. As a result, g will use the third counter when the second counter overflows. These three component counters only require 9 bits to successfully record f and g .

The scheme of two-dimensional counter sharing not only contributes to significant memory saving, but it also introduces noise among virtual counters due to space sharing. Fortunately, we can employ statistical tools to remove such noise as we will show shortly. In the sequel, when we use the word “counter” without preceding it with “virtual,” we mean a physical (component) counter by default unless the context clearly suggests otherwise.

2.3.3 Counter Tree Architecture

We design a Counter Tree architecture to implement two-dimensional counter sharing. The architecture can be used to record flows of all sizes (small, medium, or large). Given a memory space of M bits, we divide it into small counters, each consisting of b bits. We organize those counters into a tree structure from the bottom up. Let the leaf nodes be the layer 0 which consists of m counters. The degree of each non-leaf node is d . Therefore, the number of counters on a upper layer is $\frac{1}{d}$ of it lower layer. Denote the height of the tree by h , with layers indexed from 0 to $h-1$. We have the following constraint:

$$\sum_{j=0}^{h-1} \frac{m}{d^j} \times b \leq M. \quad (2.3)$$

Therefore,

$$m \leq \frac{d^{h-1}M}{(d^h - 1)b}. \quad (2.4)$$

If the $(h-1)$ th layer contains more than one counter, namely $\frac{m}{d^{h-1}} > 1$, we put an extra node as the root of the tree, called *virtual root*. Figure 2.3 gives an example of organizing the 14 counters into a binary tree with three layers, where $m = 8$ and $d = 2$. Starting from a leaf node $C[i]$ ($0 \leq i < m$), the h counters along the path to the root form a virtual counter, denoted as $V[i]$. As a result, there will be m virtual counters in total, denoted as an array V .

Starting from $C[i]$ at layer 0, the counter at layer j ($0 \leq j < h$) that $V[i]$ will include, denoted by $V[i][j]$, is

$$V[i][j] = C[\lfloor \frac{i}{d^j} \rfloor] + \sum_{t=1}^j \frac{m}{d^{t-1}}. \quad (2.5)$$

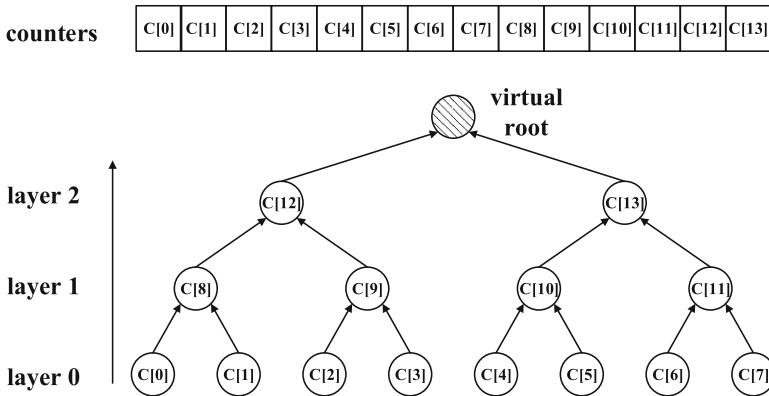


Fig. 2.3 An example of organizing counters into a binary tree

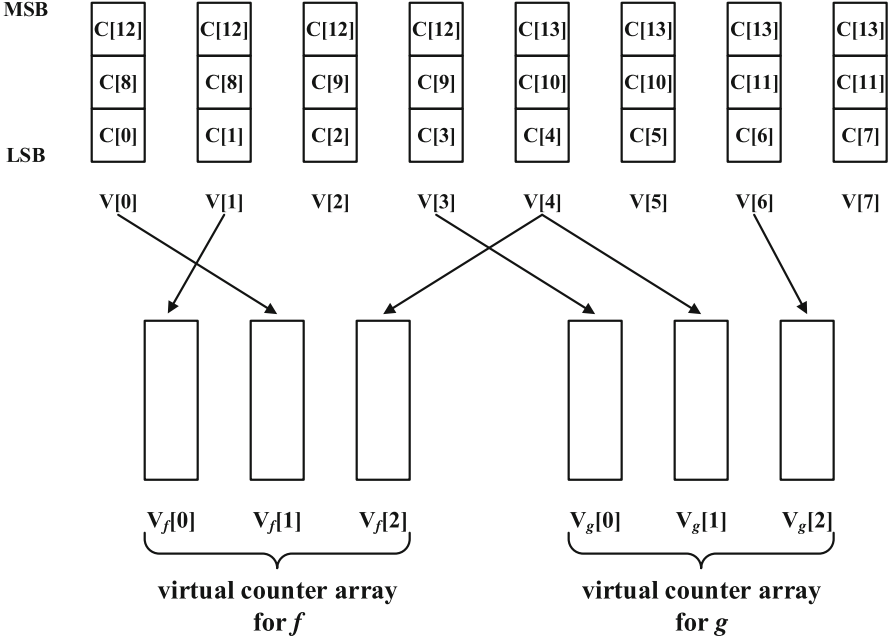


Fig. 2.4 Virtual counters and virtual counter arrays for flows. The virtual counter array for a particular flow consists of multiple counters pseudo-randomly chosen from the counters

Therefore, the Counter Tree in Fig. 2.3 can yield 8 virtual counters as shown in the upper half of Fig. 2.4. We can see that a counter located at a higher layer (which corresponds to more significant bits in a virtual counter) is shared by more virtual counters but will be used only by large flows. This embodies the idea of vertical counter sharing. Later we will show that each virtual counter can be shared by multiple flows, corresponding to the idea of horizontal counter sharing.

We note that multiple layers of counters are also used by Counter Braids [12, 13] under a different design. For Counter Braids, each flow is randomly mapped to u counters at the bottom level, where u may be 3. Any packet of a flow will cause all u counters of the flow to increase by one. Hence, each counter records the total number of packets for all flows that are mapped to it. That is, each counter represents a linear equation on the sizes of some flows (that are mapped to the counter). When there are enough counters, there will be enough linear equations, which we can solve (by the method of message passing in [13]) for the sizes of all flows. But this approach requires updating u counters per packet and it needs a sufficient number of counters. What about the memory space is too tight and thus the number of counters is insufficient? Our method can work under such tight space where Counter Braids no longer work. In Counter Tree, each flow is randomly mapped to r counters at the bottom level, where r should be large, e.g., in hundreds. Any packet of a flow will cause one of the r counters to increase by one. Therefore, each counter no longer

represents a linear equation on the sizes of some flows, which also means that the estimation method of Counter Braids cannot be applied here. Instead, we view the r counters of a flow f as a whole, whose sum carries both the size of flow f and the noise from other flows due to counter sharing. We do not know exactly who those other flows are, but nevertheless the noise can be statistically measured and removed.

2.3.4 Counting Range

The counting range of each virtual counter is 2^{bh} . To extend the counting range, one way is to increase b . However, larger b mean fewer counters are available, e.g., if we double b , the number of available counters will be reduced by half. Moreover, if b is set overly large, some bits in each counter will be wasted. Alternatively, we can increase h for the purpose of extending the counting range, which is more memory efficient as we will show shortly. Suppose M' bytes are allocated for layer-0 counters, which translates into $\frac{M'}{b}$ counters. Hence, the height of the Counter Tree can be up to $\log_d \frac{M'}{b} + 1$. Since the number of counters on the j th layers is reduced to $\frac{1}{d}$ of the $(j-1)$ th layer, the following memory constraint should hold:

$$\sum_{j=0}^{h-1} \frac{M'}{d^j} \leq M.$$

Since $\sum_{j=0}^{h-1} \frac{M'}{d^j} < \frac{d}{d-1} M'$, and we have

$$M' > \frac{d-1}{d} M, \quad (2.6)$$

which means that only $\frac{1}{d}$ of the memory needs to be reserved for non-leaf counters. Therefore, each virtual counter can have up to $b \times (\log_d \frac{M'}{b} + 1)$ bits, which translates to a counting range of $2^{b(\log_d \frac{M'}{b} + 1)}$. In contrast, the counting range of each counter is only $2^b - 1$. This means that as long as we spare $\frac{1}{d}M$ memory for upper layers and set $M' = \frac{d-1}{d}M$, we can extend the counting range from $2^b - 1$ to $2^{b(\log_d \frac{M'}{b} + 1)}$. The randomized counter sharing scheme [11] is a special case of Counter Tree with $h = 1$. Since it only records flows by one-layer physical counters, the estimation range of randomized counter sharing scheme is very limited, whereas the estimation range of Counter Tree with multiple layers is much larger. As an example, suppose $M = 1$ MB, $b = 4$, $d = 2$, and M' is therefore 0.5 MB. Hence, the counting range of each counter is $2^4 - 1 = 15$, while each virtual counter can count up to $2^{4(\log_2 \frac{0.5 \text{ MB}}{4 \text{ bit}} + 1)} = 2^{84}$ packets. Therefore, Counter Tree can significantly extend the estimation range to accommodate large flows.

2.3.5 Design Overview

Our traffic measurement function using Counter Tree consists of two modules. The online packet recording module stores the information of arriving packets in the Counter Tree. For each packet, it is mapped to a virtual counter by one hash computation and then the virtual counter will be updated, which needs approximately two memory accesses. At the end of each measurement period, the Counter Tree is stored to the disk and all counters are then reset to zeros. The offline data decoding module estimates the flow sizes. It is performed by a designated offline computer. Two methods are designed for separating the information about the size of a flow from the noise in the virtual counters. The first one is called Counter Tree base Estimation (CTE). The second one is based on the maximum likelihood estimation method (CTM). Both methods can yield accurate estimates for flow sizes.

2.4 Online Packet Recording

In this section, we show how to record a packet in the Counter Tree.

2.4.1 Recording

Consider an arbitrary flow f . We pseudo-randomly choose r ($r \ll m$) out of the m virtual counters to logically form a *virtual counter array* of f , denoted by V_f . The selection can be achieved by applying r independent hash functions to the flow label. The i th counter of V_f , denoted by $V_f[i]$, is chosen from V as follows:

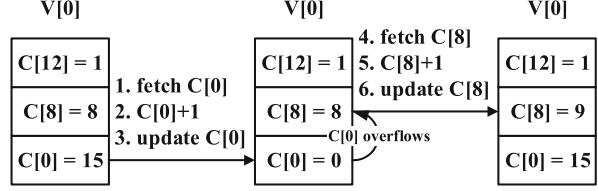
$$V_f[i] = V[h_i(f)], \quad (2.7)$$

where $0 \leq i < r$ and $h_i(\cdot)$ is a hash function $\in [0, m-1]$. To reduce the overhead of implementing r independent hash functions, we can use one master hash function H and a set S of random seeds, and let

$$h_i(f) = H(f \oplus S[i]), \quad (2.8)$$

where $\oplus$ is the XOR operator. Since $r \ll m$, the probability that r distinct virtual counters are selected by the hash functions to form the virtual counter array of f is $\frac{\binom{m}{r}}{m^r} \approx 1$. The bottom half of Fig. 2.4 illustrates the virtual counter arrays for f and g , where $r = 3$ and the virtual counter $V[4]$ is shared by both flows. We point out that our design does not limit the number of flows to be supported. There can be many more flows than the number of virtual vectors, m .

Fig. 2.5 The process for recording a packet to a virtual counter



At the beginning of each measurement period, all counters are initialized to 0s. When a packet of flow f arrives, the router extracts its flow label f , chooses a virtual counter from V_f uniformly at random, and increments that virtual counter by 1. More specifically, the router generates a random number $i \in [0, r - 1]$, computes the hash value $u = h_i(f)$, and sets $V[u] \leftarrow V[u] + 1$. Note that the update of $V[u]$ may involve the updates of multiple counters. Based on (2.5), the router first fetches counter $C[u]$ from memory and increases it by 1. If $C[u]$ does not overflow, the recording for this packet is done. Otherwise, the router further fetches $C[\lfloor \frac{u}{d} \rfloor + m]$, and adds the overflowed 1 to $C[\lfloor \frac{u}{d} \rfloor + m]$. The process continues until no overflow happens or the counter on the root has been reached. In the latter case, the virtual counter is overflowed beyond the upper bound that it is designed to handle. Figure 2.5 gives an example of the online recording process for a packet of f . Suppose $b = 4$ (i.e., the counting range of each counter is 15), $h = 3$, and $V[0]$ is chosen for recording that packet. The router first fetches $C[0]$ whose value is 15. After adding 1 to $C[0]$, $C[0]$ becomes 0 and leads to an overflow. Hence, the router writes back $C[0] = 0$, further fetches $C[8]$ with value 9, and assigns $C[8] \leftarrow C[8] + 1$. Since $C[8] = 10$ does overflow, the router writes it back and the recording process terminates.

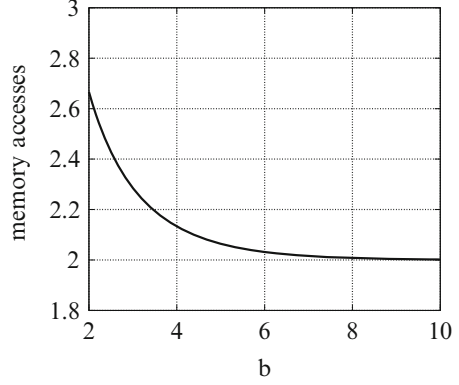
2.4.2 Number of Memory Accesses

To record a packet, the router at least needs to read and write 1 counter, which requires 2 memory accesses. Hence, the lower bound of the number of memory accesses for recording a packet is 2. In the worst case, the router needs to update h counters, which requires $2h$ memory accesses. The good thing is that the router needs to fetch another counter only when the current counter overflows, which happens after recording at least 2^b packets. Hence, the amortized number of memory accesses per packet is much smaller than the worst case. We have the following theorem:

Theorem 1. *A tight upper bound of the amortized number of memory accesses for recording a packet in the Counter Tree is $2 + \frac{2}{2^b - 1}$, where b is length of each counter.*

Proof. Suppose n packets are recorded. Each packet causes one counter at layer 0 to be updated, requiring 2 memory accesses. In total, there are $2n$ memory accesses at layer 0. The number of counter overflows at layer 0 is at most $\lfloor \frac{n}{2^b} \rfloor \leq \frac{n}{2^b}$, which means that counters at layer 1 will be updated for no more than $\frac{n}{2^b}$ times, requiring

Fig. 2.6 Amortized number of memory accesses per packet with respect to b



no more than $\frac{2n}{2^b}$ memory accesses. By simple induction, we know that the number of memory accesses at layer j is no more than $\frac{2n}{2^{jb}}$. Hence, the total number of memory accesses over all layers is no more than

$$\sum_{j=0}^{h-1} \frac{2n}{2^{jb}} = \frac{2n(1 - \frac{1}{2^{bh}})}{1 - \frac{1}{2^b}} < 2n(1 + \frac{1}{2^b - 1}).$$

Hence, the amortized number of memory accesses per packet is no more than

$$\frac{2n(1 + \frac{1}{2^b - 1})}{n} = 2 + \frac{2}{2^b - 1}. \quad (2.9)$$

The upper bound is tight when $n \equiv 0 \pmod{2^{(h-1)b}}$ and exactly $\frac{n}{2^{(h-1)b}}$ overflows happen at the j th layer, $0 \leq j < h - 1$.

Figure 2.6 shows the upper bound of the amortized number of memory accesses for recording a packet with respect to b . We find that it quickly converges to 2 (the lower bound) with the increase of b . In contrast, Counter Braids need to perform at least $2u$ memory accesses since each packet is recorded by u counters on each layer, where u can be 3.

The non-deterministic counter access time per packet may introduce stalls into the datapath pipeline, resulting in reduced datapath bandwidth. Counter Braids [12, 13] and any other counter architectures with variable access time [17] face the same problem. The key issue is how frequent the variable access time will occur. In case that access time for most packets is constant but only varies for a tiny fraction of packets, the impact on datapath bandwidth will be limited, particularly if we are able to reduce the chance of variable access time to an arbitrarily small level through a system parameter. This is the case for Counter Tree. More specifically, the overflow of a layer-0 counter occurs once every 2^b packets. On average, one out of 2^b packets has longer access time because a counter at the higher layer is involved, but the other $2^b - 1$ packets have constant access time because they do not cause counter

overflow. The fraction of all packets that have longer access time is $\frac{1}{2^b}$, which can be exponentially reduced by increasing the value of b , i.e., the number of bits in a counter. For example, when $b = 6$, only 1.6 % of all packets have variable access time and 98.4 % of the packets have constant access time. Suppose each counter overflow causes the pipeline to stall for $10\times$ of the normal access time. The overall access time will increase by approximately 14.4 %, and the throughput therefore decreases by approximately 12.6 %.

2.5 Counter Tree-Based Estimation

After the measurement period, offline estimation should be performed to recover flow sizes from the Counter Tree. Counter Tree estimates flow sizes through some statistical methods, and it supports efficient query on the size of an arbitrary flow, without having to compute the sizes of other flows as Counter Braids do (the computation overhead is high). In this section, we first present and analyze the Counter Tree-based Estimation (CTE) method.

2.5.1 CTE Method

Consider the i th virtual counter $V_f[i]$ in the virtual counter array of flow f . According to (2.7), $V_f[i] = V[u]$, where $u = h_i(f)$. We know $V[u]$ records some of f 's packets, as well as the noise introduced by other flows. There are two sources of noise: First, $V[u]$ is shared by other flows; Second, all component counters in $V[u]$ except $C[u]$ are shared by other virtual counters. To accurately recover the number of packets in f recorded by $V[u]$, we need to figure out how to remove such noise.

The component counter of $V[u]$ at the highest layer is $C[v]$, where $v = \lfloor \frac{u}{d^{h-1}} \rfloor + \sum_{t=1}^{h-1} \frac{m}{d^{t-1}}$ according to (2.5). Consider the subtree rooted at $C[v]$, denoted as T , consisting of d^{h-1} leaf nodes, which correspond to d^{h-1} virtual counters. Let $k = d^{h-1}$. Due to counter sharing, flows mapped to those k virtual counters may introduce noise to $V[u]$. We treat T as an aggregate, called a *subtree counter*, when dealing with noise. We denote the value of T by a random variable X_i . As an example, in Fig. 2.3, the value of the subtree counter rooted at $C[12]$ is $C[12] \times 2^{2b} + (C[8] + C[9]) \times 2^b + (C[0] + C[1] + C[2] + C[3])$. Note that the total number n of packets in all flows can be obtained from the entire Counter Tree in a similar way. Let random variable Y_i be the portion of X_i contributed by flow f , and Z_i be the portion of X_i contributed by all other flows. Hence, $X_i = Y_i + Z_i$.

Let s be the true flow size of f during the measurement period. Assume there is a large number of flows, n is large, the size of each flow is negligibly small when comparing with n , r is much larger than 1, and $r \ll m$.

Flow f has r virtual counters (including $V[u]$) in its virtual counter array. Each packet of f has a probability $\frac{1}{r}$ to be mapped to $V[u]$ and increment it by one. Therefore, Y_i follows a binomial distribution:

$$Y_i \sim B(s, \frac{1}{r}). \quad (2.10)$$

Consider an arbitrary packet belonging to a different flow g . The probability for the virtual counter array of g to include a particular virtual counter in T is approximately $1 - \frac{\binom{m-1}{r}}{\binom{m}{r}} = \frac{r}{m}$. When that happens, the conditional probability for this particular virtual counter to record the packet is approximately $\frac{1}{r}$. Combining the above analysis, the probability for this particular virtual counter to record the packet is approximately $\frac{r}{m} \times \frac{1}{r} = \frac{1}{m}$. Since there are k virtual counters in T , the probability that T records the packet is $(1 - \frac{1}{m})^k \approx \frac{k}{m}$ for $k \ll m$. There are $n - s$ packets outside of f . Since there are numerous flows under measurement, the total number n of packets can be much larger than the size s of any single flow, even for large flows, as we observe in our traffic traces. With $s \ll n$ and $n - s \approx n$, Z_i roughly follows a binomial distribution

$$Z_i \sim B(n, \frac{k}{m}). \quad (2.11)$$

Given the distributions of Y_i and Z_i , we know $E(Y_i) = \frac{s}{r}$, and $E(Z_i) = \frac{nk}{m}$. Therefore,

$$E(X_i) = E(Y_i + Z_i) = E(Y_i) + E(Z_i) = \frac{s}{r} + \frac{nk}{m}.$$

Hence, we have

$$s = rE(X_i) - \frac{nkr}{m}. \quad (2.12)$$

We do not know the exact value of $E(X_i)$, but we have the r instance values, also denoted as X_i , $0 \leq i < r$, that can be directly counted from the Counter Tree as subtree counters. Replacing $E(X_i)$ with the measured average $\frac{\sum_{i=0}^{r-1} X_i}{r}$, we obtain an estimate of s , denoted as $\hat{s}$, as follows:

$$\hat{s} = \sum_{i=0}^{r-1} X_i - \frac{nkr}{m}, \quad (2.13)$$

where the first term is the number of all packets recorded by the r subtree counters and the second term captures the average noise presented in r counters.

2.5.2 Analysis of $\hat{s}$

Since Y_i and Z_i follow binomial distributions, we have

$$\text{Var}(Y_i) = \frac{s}{r}(1 - \frac{1}{r}), \quad \text{Var}(Z_i) = \frac{nk}{m}(1 - \frac{k}{m}). \quad (2.14)$$

In addition, Y_i and Z_i are independent with each other. Hence, $\text{Cov}(Y_i, Z_i) = 0$. Hence, we have

$$\begin{aligned} \text{Var}(X_i) &= \text{Var}(Y_i) + \text{Var}(Z_i) + 2\text{Cov}(Y_i, Z_i) \\ &= \frac{s}{r}(1 - \frac{1}{r}) + \frac{nk}{m}(1 - \frac{k}{m}). \end{aligned} \quad (2.15)$$

Therefore,

$$E(\hat{s}) = E(\sum_{i=0}^{r-1} X_i) - \frac{ndr}{m} = r(\frac{nr}{m}) - \frac{nr}{m} = s, \quad (2.16)$$

which means the estimator $\hat{s}$ is unbiased. In addition, we have

$$\begin{aligned} \text{Var}(\hat{s}) &= \text{Var}(\sum_{i=0}^{r-1} X_i) = r^2 \text{Var}(X_i) \\ &= s(r-1) + \frac{nr^2}{m}(1 - \frac{k}{m}) \\ &= s(r-1) + \frac{nr^2 b(d^h-1)}{M}(1 - \frac{b(d^h-1)}{M}), \end{aligned} \quad (2.17)$$

where we have used $m = \frac{d^{h-1}M}{b(d^h-1)}$ and $k = d^{h-1}$. Hence, the standard error of the ratio $\frac{\hat{s}}{s}$ is

$$\text{StdErr}(\frac{\hat{s}}{s}) = \frac{\sqrt{s(r-1) + \frac{nr^2 b(d^h-1)}{M}(1 - \frac{b(d^h-1)}{M})}}{s}. \quad (2.18)$$

When n is sufficiently large, the binomial distribution, $Z_i \sim B(n, \frac{k}{m})$, can be approximated by a normal distribution, $N(\frac{nk}{m}, \frac{nk}{m}(1 - \frac{k}{m}))$. Similarly, $Y_i \overset{\text{approx}}{\sim} N(\frac{s}{r}, \frac{s}{r}(1 - \frac{1}{r}))$. Since the linear combination of independent normal distributions also follows normal distribution, $X_i \overset{\text{approx}}{\sim} N(\mu, \sigma^2)$, where $\mu = \frac{nk}{m} + \frac{s}{r}$, and $\sigma^2 = \frac{nk}{m}(1 - \frac{k}{m}) + \frac{s}{r}(1 - \frac{1}{r})$. According to (2.13), we have

$$\hat{s} \overset{\text{approx}}{\sim} N(s, s(r-1) + \frac{nr^2}{m}(1 - \frac{k}{m})). \quad (2.19)$$

Therefore, the α confidence interval for s is

$$\hat{s} \pm Z_\alpha \sqrt{s(r-1) + \frac{nr^2}{m} \left(1 - \frac{k}{m}\right)}, \quad (2.20)$$

where Z_α is the $\frac{1+\alpha}{2}$ percentile for the standard normal distribution. For example, when $\alpha = 95\%$, $Z_\alpha = 1.96$.

2.6 Counter Tree-Based Maximum Likelihood Estimation

In this section, we provide and analyze another estimator for flow sizes called Counter Tree-based Maximum likelihood Estimation (CTM).

2.6.1 CTM Method

From (2.11), the probability of $Z_i = z_i$ is

$$\text{Prob}\{Z_i = z_i\} = \binom{n}{z_i} \left(\frac{k}{m}\right)^{z_i} \left(1 - \frac{k}{m}\right)^{n-z_i}.$$

The value of n is known from the Counter Tree. The values of m and k are determined by prescribed system parameters M , b and d , h . Hence, $P\{Z_i = z_i\}$ can be written as a function of z_i , denoted as $p(z_i)$. Therefore, the probability for observing $X_i = x_i$ can be calculated by

$$\begin{aligned} \text{Prob}\{X_i = x_i\} &= \text{Prob}\{Y_i + Z_i = x_i\} \\ &= \sum_{z_i=0}^{x_i} p(z_i) \text{Prob}\{Y_i = x_i - z_i\} \\ &= \sum_{z_i=0}^{x_i} p(z_i) \binom{s}{y_i} \left(\frac{1}{r}\right)^{y_i} \left(1 - \frac{1}{r}\right)^{s-y_i} \\ &= \sum_{z_i=0}^{x_i} p(z_i) q(s, y_i), \end{aligned} \quad (2.21)$$

where $y_i = x_i - z_i$, $q(s, y_i) = \binom{s}{y_i} (\frac{1}{r})^{y_i} (1 - \frac{1}{r})^{s-y_i}$, and we have used $Y_i \sim B(s, \frac{1}{r})$. Hence, the likelihood function for observing $X_0 = x_0, X_1 = x_1, \dots, X_{r-1} = x_{r-1}$ is

$$L(s; x_0, x_1, \dots, x_{r-1}) = \prod_{i=0}^{r-1} \sum_{z_i=0}^{x_i} p(z_i) q(s, y_i). \quad (2.22)$$

Taking the logarithm for both sides of the likelihood function, we obtain the log-likelihood as follows:

$$\ln L = \sum_{i=0}^{r-1} \ln \left(\sum_{z_i=0}^{x_i} p(z_i) q(s, y_i) \right).$$

Using the logarithmic differentiation,¹ we can calculate

$$\frac{d \binom{s}{y_i}}{ds} = \binom{s}{y_i} \sum_{j=0}^{y_i-1} \frac{1}{s-j}.$$

Hence,

$$\begin{aligned} \frac{dq(s, y_i)}{ds} &= \binom{s}{y_i} \left(\frac{1}{r} \right)^{y_i} \left(1 - \frac{1}{r} \right)^{s-y_i} \left(\sum_{j=0}^{y_i-1} \frac{1}{s-j} + \ln \left(1 - \frac{1}{r} \right) \right) \\ &= q(s, y_i) \left(\sum_{j=0}^{y_i-1} \frac{1}{s-j} + \ln \left(1 - \frac{1}{r} \right) \right). \end{aligned}$$

Finally, we obtain

$$\frac{d \ln L}{ds} = \sum_{i=0}^{r-1} \frac{\sum_{z_i=0}^{x_i} p(z_i) q(s, y_i) \left(\sum_{j=0}^{y_i-1} \frac{1}{s-j} + \ln \left(1 - \frac{1}{r} \right) \right)}{\sum_{z_i=0}^{x_i} p(z_i) q(s, y_i)}. \quad (2.23)$$

The value of s that satisfies $\frac{d \ln L}{ds} = 0$ will maximize $\ln L$ and thereby the likelihood function L . In addition, we observe that $\frac{d \ln L}{ds}$ is monotonically decreasing with respect to s . Therefore, the bisection search method can be used to find the s value such that $\frac{d \ln L}{ds} = 0$. As a result, we obtain an estimator for s as follows:

$$\hat{s} = \arg \max_s \{\ln L\} = \{s \mid \frac{d \ln L}{ds} = 0\}. \quad (2.24)$$

¹Since $[\ln(f)]' = \frac{f'}{f}$, we know $f' = f[\ln(f)]'$.

2.6.2 Analysis of $\hat{s}$

Following the analysis in Sect. 2.5.2, $X_i \overset{\text{approx}}{\sim} N(\mu, \sigma^2)$. Hence, the probability density function of X_i is $f(x_i) = \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{(x_i-\mu)^2}{2\sigma^2}}$. Therefore, the likelihood function for observing $X_0 = x_0, X_1 = x_1, \dots, X_{r-1} = x_{r-1}$ can also be written as

$$L(s; x_0, x_1, \dots, x_{r-1}) = \prod_{i=0}^{r-1} \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{(x_i-\mu)^2}{2\sigma^2}}.$$

Taking the logarithm of the likelihood function, we have

$$\begin{aligned} \ln L &= \sum_{i=0}^{r-1} \ln\left(\frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{(x_i-\mu)^2}{2\sigma^2}}\right), \\ &= \sum_{i=0}^{r-1} -\ln \sqrt{2\pi} - \ln \sigma - \frac{(x_i - \mu)^2}{2\sigma^2}. \end{aligned}$$

The first and second order derivatives of $\ln L$ with respect to s are

$$\begin{aligned} \frac{d \ln L}{ds} &= \sum_{i=0}^{r-1} -\frac{r-1}{2r^2\sigma^2} + \frac{x_i - \mu}{r\sigma^2} + \frac{(x_i - \mu)^2(r-1)}{2r^2(\sigma^2)^2}, \\ \frac{d^2 \ln L}{ds^2} &= \sum_{i=0}^{r-1} \left(\frac{(r-1)^2}{2r^4(\sigma^2)^2} - \frac{1}{r^2\sigma^2} - \frac{2(r-1)(x_i - \mu)}{r^3(\sigma^2)^2} \right. \\ &\quad \left. - \frac{(r-1)^2(x_i - \mu)^2}{r^4(\sigma^2)^3} \right), \end{aligned}$$

where we have used $\frac{d\mu}{ds} = \frac{1}{r}$ and $\frac{d\sigma^2}{ds} = \frac{r-1}{r^2}$. Hence,

$$\begin{aligned} E\left(\frac{d^2 \ln L}{ds^2}\right) &= r\left(\frac{(r-1)^2}{2r^4(\sigma^2)^2} - \frac{1}{r^2\sigma^2} - \frac{(r-1)^2\sigma^2}{r^4(\sigma^2)^3}\right) \\ &= -\frac{2r\sigma^2 + (r-1)^2}{2r^3\sigma^4}, \end{aligned}$$

since $E(x_i - \mu) = 0$ and $E(x_i - \mu)^2 = \sigma^2$. Hence, the fisher information [9] is $I(s|x_0, x_2, \dots, x_{r-1}) = -E\left(\frac{d^2 \ln L}{ds^2}\right) = \frac{2r\sigma^2 + (r-1)^2}{2r^3\sigma^4}$. According to the asymptotic properties of maximum likelihood estimators [9], we have

$$\hat{s} \xrightarrow{d} N\left(s, \frac{1}{I(s|x_0, x_2, \dots, x_{r-1})}\right) = N\left(s, \frac{2r^3\sigma^4}{2r\sigma^2 + (r-1)^2}\right) \quad (2.25)$$

Therefore, the standard relative error is

$$StdErr(\frac{\hat{s}}{s}) = \frac{\sqrt{\frac{2r^3\sigma^4}{2r\sigma^2 + (r-1)^2}}}{s}, \quad (2.26)$$

and the α confidence interval for s is

$$\hat{s} \pm Z_\alpha \sqrt{\frac{2r^3\sigma^4}{2r\sigma^2 + (r-1)^2}}. \quad (2.27)$$

2.7 Enhanced Counter Tree Architecture

2.7.1 Motivation

Recall that the design of vertical counter sharing in the Counter Tree reserves counters at higher layers for large flows. However, the question that which flows have actually used the high-layer counters is left open since each high-layer counter can be shared by multiple virtual counters. If two flows share a high-layer counter and only one of them uses the counter for recording its packets, there is no way to figure out which of the two actually uses that counter from the values in the Counter Tree, which leads to ambiguity. Consider a simple example in which f is a small flow with size 1 and g is large flow with size 30,000. As shown in Fig. 2.4, we assume the packet of f is recorded in $V[1]$ (more specifically $C[1]$ since other counters in $V[1]$ are not used), and 10,000 of g 's packets are recorded in $V[3]$. As a result, $C[12]$ is used by g to accommodate such a large number of packets. We know that $V[1]$ and $V[3]$ share the component counter $C[12]$. Although $C[12]$ has never been used by f , it is also included as a component in f 's virtual counter $V[1]$. As a result, a great noise is introduced by g when we estimate the size of f since $C[12]$ is included in the estimation of f . To handle this problem, our previous design relies on mapping each flow to many virtual counters (which helps amortizing the noise) and using statistical means to remove the noise. However, we suspect that large noise introduced at higher layers can nevertheless degrade the estimation accuracy. In this section, we introduce additional mechanism to address this issue.

We define the *length* of each virtual counter as the number of component counters it truly uses. In the previous example, the length $V[1]$ is 1 since only $C[1]$ is used, while the length of $V[3]$ is 3 since $C[3]$, $C[9]$, and $C[12]$ have been used. We observe that the aforementioned ambiguity results from the uncertainty of the length of each virtual counter after recording. In the previous design of Counter Tree, we simply assume every virtual counter has the same length, which is equal to the height h of the tree. The experiment results in Sect. 2.8 will demonstrate that Counter Tree works well when the tree height h is set small, e.g., $h = 2$. However, the performance of Counter Tree seriously deteriorates when the tree height grows for the purpose of extending estimation range. Theoretically, this observation is

embodied by the fact that the variance of $\hat{s}$ increases exponentially with h according to (2.17). To accommodate large flows, reasonably large h should be used, leading to large estimation error. Therefore, we must figure out how to resolve the ambiguity.

2.7.2 Counters with Status Bits

To address this issue, we design an Enhanced Counter Tree (ECT) architecture which determines the length of each virtual counter and thereby resolves the ambiguity by adding a status bit to each counter. For each b -bit counter, its lower $(b - 1)$ bits are used for counting packets, and its most significant bit is used as the *status bit*. See Fig. 2.7 for illustration. Only if the counter overflows, will the status bit of that counter be set. With the status bits, we can calculate the true length for each individual virtual counter, thereby reducing the noise caused by vertical counter sharing among virtual counters. More specifically, to build a virtual counter according to its length, we traverse along the path from its layer-0 counter to the root, and include all counters until the first counter (included) whose status bits have not been set. Following the aforementioned example, Fig. 2.8 shows that the introduction of status bits can alleviate the ambiguity in sharing at higher layers. In the figure, $C[1]$ is a component counter of $V[1]$ that f is mapped to. Since $C[1]$ does not overflow (i.e., its status bit is not set), the length of $V[1]$ must be 1 and therefore

Fig. 2.7 A b -bit counter consists of $(b - 1)$ counting bits and a status bit

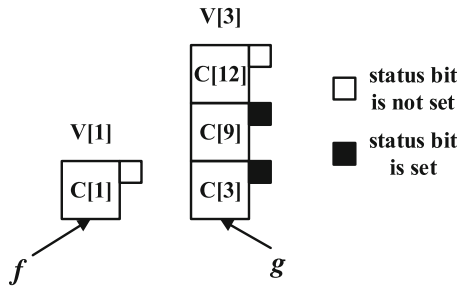
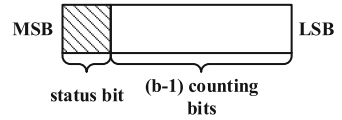


Fig. 2.8 After introducing status bits, we can calculate the lengths of virtual counters based on the status bits in their component counters. The length of $V[1]$ that f is mapped to is 1 since the status bit of $C[1]$ is not set, and the length of $V[3]$ that g is mapped to is 3 since the status bits of $C[3]$ and $C[9]$ are set while the status bit of $C[12]$ is not set

the component counter $C[12]$ should have been included. In contrast, the length of $V[3]$ (which g is mapped to) is 3 since the status bits of both $C[3]$ and $C[9]$ are set while the status bit of $C[12]$ is not. In conclusion, the extra status bits can provide better resolution for virtual counters and help resolving ambiguity about whether a high-layer component counter should be included in a virtual counter or not.

If a counter overflows due to a large flow, the small flows that share the counter will observe a large counter value. However, each flow is assigned to r counters at the bottom layer, where r is large, e.g., in hundreds. The number of large flows is much smaller than the number of small/medium flows in real traffic. For a small flow, if a few of its counters are overflowed due to sharing with large flows, the values of these counters will be large. But as long as most of the small flow's counters have small values, the effect of counter crosstalk will be dampened by the maximum likelihood estimation (2.24), which discounts the outliers (large counter values).

For the first estimator (2.13), when there are many, many flows, even without a large flow, the sheer number of small flows mapped to a counter may cause the counter to overflow. But this noise can be statistically measured and removed when r is sufficiently large. Among the r counters, statistically, some will carry larger-than-average noise and some will carry smaller-than-average noise. When r is large enough, such difference will be evened out due to the law of large numbers. The average noise term is captured by the second term in (2.13).

2.7.3 Recording and Estimation

After employing status bits, the process of online packet recording remains the same except that we need to set its status bit when a counter overflows. For the offline estimation process, we should slightly modify the estimators given in (2.13) and (2.24) to reflect the change. Consider the estimation of flow f . Let l_i be the length of the virtual counter $V_f[i]$ after the measurement period, where $0 \leq i < r$. Instead of using a unified height h for all subtree counters, the height of the subtree counter corresponding to $V_f[i]$ should be l_i , and the value k , meaning the number of leaf nodes in the subtree, is therefore d^{l_i-1} . For example, in Fig. 2.8, the height of the subtree counter corresponding to $V[1]$ is 1 and it contains only one leaf node, while the height of the subtree counter corresponding to $V[3]$ is 3 and it contains 2^{3-1} leaf nodes. Everything else remains the same for the offline estimation. Our experiments in Sect. 2.8 will show that the ECT with status bits remarkably improves the estimation accuracy of CT.

2.8 Experimental Evaluation

2.8.1 *Experiment Setup*

We have implemented the two estimators based on the Counter Tree, i.e., CTE and CTM, from Sects. 2.5 and 2.6, respectively. CTE and CTM share the same module for online packet recording, which performs the operations described in Sect. 2.4. Hence, when we evaluate the online operations of the Counter Tree, we use CT as the abbreviation. We have also implemented the Enhanced Counter Tree architectures. We compare them with the most related counter architectures: (1) randomized counter sharing (MLM) [10, 11] and (2) Counter Braids (CB) [12, 13]. The bitmap-based MSCBF is less memory efficient than the counter architectures [10]. Hence, we do not include it for comparison. Without losing generality, we use TCP flows for presentation, and we have obtained similar results when carrying out experiments with other types of flows. The network trace we use was captured by Cisco’s NetFlow at the main gateway of our university, and we are authorized to store the packets headers in disk for our experiments. The trace contains about 68 million TCP flows and 750 million packets. We implement those counter architectures in software and evaluate their performance by running experiments on the trace. Suppose each measurement period contains 10 million packets. We divide the trace into measurement periods, and perform different estimators in each period. We use all 750 million packets in our experiments, for 75 periods. We randomly pick the results from one period to report. In fact, the results from different periods are quite similar. During the chosen period, there are 1,070,632 TCP flows and 10,053,234 packets, and the minimum, average, and maximum flow size is 1, 9.39, and 10,972 packets, respectively.

We conduct four sets of experiments. The first set is used for comparing the estimation accuracy and range of CT, MLM, and CB. We vary the available memory space M from 0.125, 0.25, 0.5, to 1 MB, which translate to about 1bit/flow, 2bits/flow, 4bits/flow, and 8bits/flow, respectively. For Counter Braids, we use the same settings as [13]: A two-layer CB and three hash functions at both layers. The layer-1 counters are 8 bits deep and the layer-2 counters are 56 bits deep. For MLM, we set the counter length to 6 bits and the size of each storage vector to 100 as in [10]. For CTE and CTM, we implement a 2-layer tree (which is sufficient for our experiments) with $d = 2$ and $b = 4$ by default. For fair comparison with MLM, we set the size r of each virtual counter array to 100. The second set of experiments compare the processing overhead of online packet recording of these counter architectures. In the third set of experiments, we will vary the values of d , b , and r to study their respective impact on performance. In the fourth set of experiments, we compare the performance of Counter Tree and Enhanced Counter Tree under different tree heights.

Table 2.2 Comparison of average processing time for encoding a packet by CB, MLM, and CT

Memory size (MB)	Number of memory accesses			Number of hash computations		
	CB	MLM	CT	CB	MLM	CT
0.25	6.01	2	2.09	3.01	1	1
0.5	6.01	2	2.06	3.00	1	1
1	6.01	2	2.03	3.00	1	1
2	6.01	2	2.02	3.00	1	1

2.8.2 Processing Time for Recording a Packet

The processing time for encoding a packet mainly results from memory accesses to read and write counters and the computations of hash values. A typical implementation of Counter Braids requires 3 hash functions on each layer, mapping each flow to the corresponding counters. To encode a packet, the router needs to read the 3 associated counters on the first layer, increment them by 1, and then write them back to the memory. If any of the 3 counters overflows, the router has to read and write another 3 counters on the second layer, which requires another 3 hash computations. Hence, the lower bounds of the number of memory accesses and the number of hash computations by CB are 6 and 3, respectively. In contrast, MLM aligns all counters on the same layer, and each packet is hashed to only one counter, which requires 2 memory access and 1 hash computation. CT also only requires 1 hash computation to determine the virtual counter for a packet. Recall that (2.9) gives an upper bound of amortized number of memory accesses by CT. When $b = 4$, the amortized number of memory accesses is bounded by $2 + \frac{1}{2^4 - 1} \approx 2.13$. In the first set of experiments, we record the average number of memory accesses and average number of hash computations for encoding a packet by CB, MLM, and CT. The results are shown in Table 2.2. We can see that CT is almost as efficient as MLM, and they achieve approximately $3\times$ efficiency of CB. Moreover, the average number of memory accesses of CT decreases when more memory (counters) are available since each counter is shared by fewer flows, which reduces the overflows.

2.8.3 Estimation Accuracy and Range

Recall that our main objective is to design a counter architecture that can work in very tight space where existing counter architectures no long work well. So we first compare Counter Tree with CB and MLM in terms of estimation accuracy and range to see how they work under the same available memory. The estimation results of CB are shown in Fig. 2.9 which includes four plots for different values of M . Each point in the plots represents an $(s, \hat{s})$ pair for a particular flow, where the x coordinate is the true flow size s and the y coordinate is the estimated flow size $\hat{s}$. The equality line, $y = x$, is presented for reference: The closer a point is to the

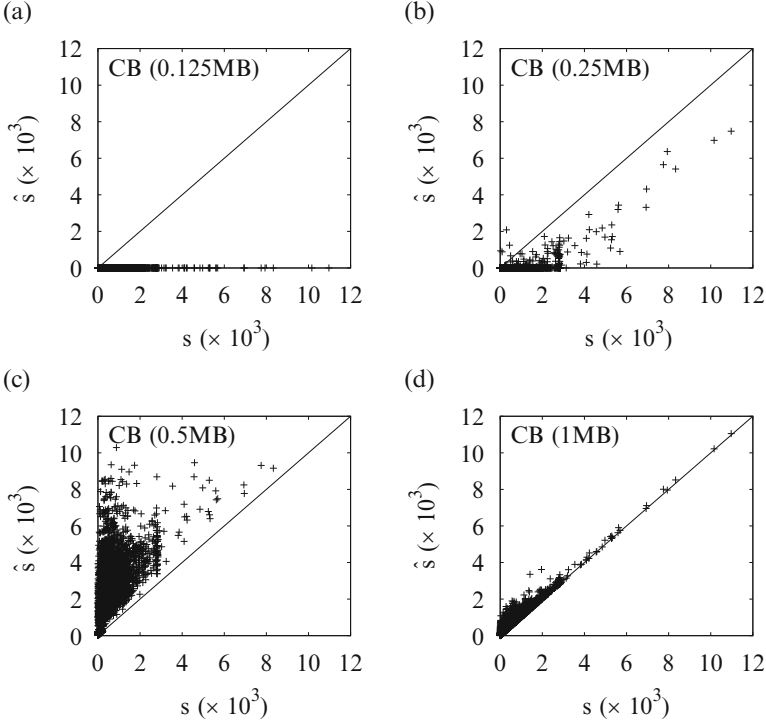


Fig. 2.9 (a) Shows estimation results by Counter Braids when $M = 0.125$ MB. Each point in the plot represents an $(s, \hat{s})$ pair for a particular flow, where the x coordinate is the true flow size s and the y coordinate is the estimated flow size $\hat{s}$. The equality line, $y = x$, is presented for reference: a point closer to the equality line is more accurate. (b) Shows estimation results by Counter Braids when $M = 0.25$ MB. (c) Shows estimation results by Counter Braids when $M = 0.5$ MB. (d) Shows estimation results by Counter Braids when $M = 1$ MB

equality line, the more accurate the estimate is. We can see that when a very tight memory $M = 0.125$ MB is used, CB cannot produce any meaningful results. When $M = 0.5$ MB, CB generates positively biased results that are all above the equality line. When the available memory space increases to 1 MB, the estimation results of CB do not converge. So we terminate the process after 1000 iterations. We find that when $M \geq 2$ MB, CB can yield very accurate estimates (which is not shown in the figure). Therefore, CB does not suite for traffic measurement under very tight memory.

Figure 2.10 presents the estimation results of MLM. MLM does not work when $M = 0.125$ MB. Although MLM can yield accurate estimates for small or moderate flows when more memory is available, it produces very negatively biased results for large flows. Because large flows may lead to counter overflows, some packets cannot be recorded when the counters are fully used. Although the increase of M can

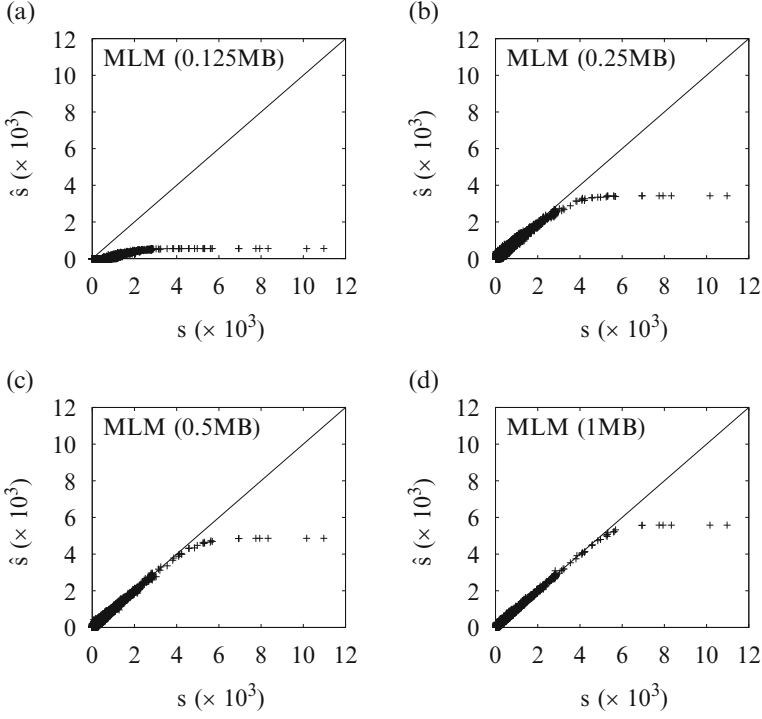


Fig. 2.10 Estimation results by MLM when $M = 0.125$ MB, 0.25 MB, 0.5 MB, and 1 MB (a)–(d), respectively

enlarge the estimation range of MLM to some extent, it does not address the problem fundamentally. For example, the estimation range is about 6000 when $M = 1$ MB.

The estimation results of CTE and CTM are given in Figs. 2.11 and 2.12, respectively. As expected, the employment of the Counter Tree architecture significantly extends the counting range than MTM. Both CTE and CTM can yield very accurate estimates for all flows, including flows with very large sizes, even under a tight memory. The estimates become more accurate when more memory space is available. The relative estimation bias $Bias(\frac{\hat{s}}{s})$ and the relative standard error $\frac{\sqrt{Var(\hat{s})}}{s}$ of CTE and CTM are presented in Fig. 2.13. We find that CTE and CTM in fact have comparable estimation accuracy. Figure 2.13 shows that the relative errors for large flows are small. Generally, $Bias(\frac{\hat{s}}{s})$ and $\frac{\sqrt{Var(\hat{s})}}{s}$ decrease with the increase of s . Although the relative errors for small flows can be large, the results in Figs. 2.11 and 2.12 demonstrate that no small flows will deviate significantly from the equality line for large absolute errors (which would cause mis-classification). It is expected and true that the relative errors for small flows are large for virtually all estimation methods. We use an extreme example to bring out the idea: There is a difference in interpreting the results for small flows and those for large flows. Consider a small

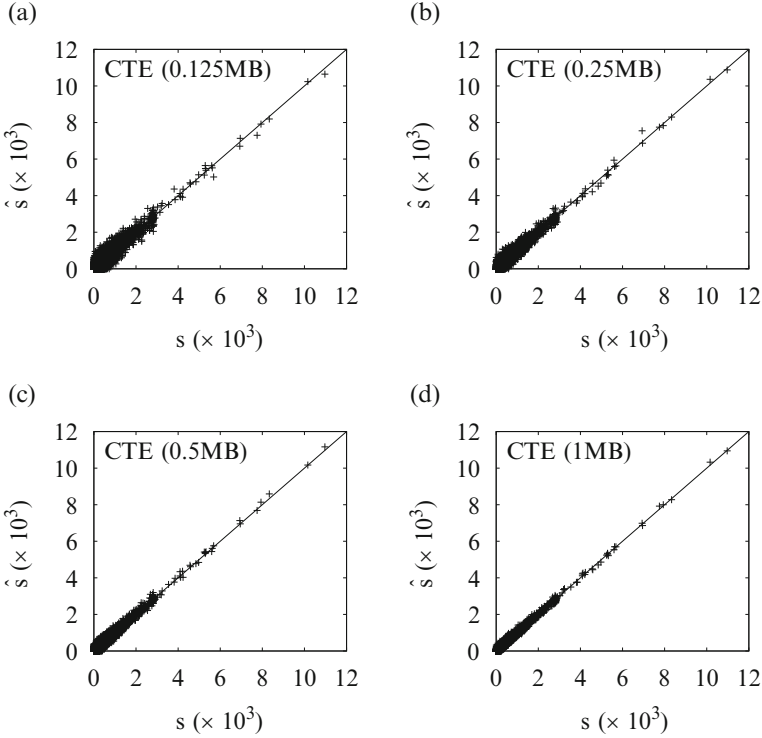


Fig. 2.11 Estimation results by CTE when $M = 0.125$ MB, 0.25 MB, 0.5 MB, and 1 MB (a)–(d), respectively

flow of size 1. If the estimation is 2 (which is in fact a good result because it is off by just 1), the relative error is 100 %. For a large flow of 10,000, if the relative error is 100 %, it will be 20,000, a truly bad estimation. For the same large flow, if the estimation is off by 10, it is a great estimation because the relative error is just 0.1 %. However, if the previous small flow is off by 10, the relative error is 1000 %—even with such a relative error, we would not necessarily say this is a bad estimation because a flow of 11 packets will not be mis-classified as a large flow, in applications of identifying elephant flows or classifying all flows into a few categories based on their sizes.

To numerically evaluate how large flows affect the estimation results of small flows, we only consider small flows that share counters with large flows, and omit the small flows that do not share any counters with large flows. More specifically, we consider all flows with sizes smaller than 100 as small, while those with sizes larger than 1000 as large. However, we observe that all small flows share multiple counters with large flows. The reason is that each flow uses a large number r of counters, which means a small flow will have a good chance to share at least one

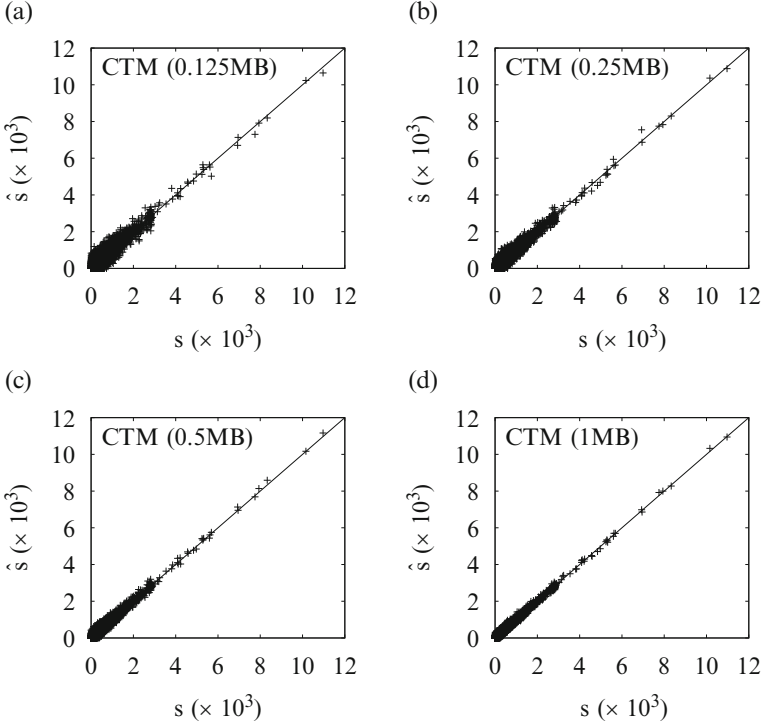


Fig. 2.12 Estimation results by CTM when $M = 0.125$ MB, 0.25 MB, 0.5 MB, and 1 MB (a)–(d), respectively

counter with one of the large flows. We stress that our method is designed to handle such sharing through noise removal and maximum likelihood estimation.

In conclusion, CT works much better than CB and MLM under a tight memory, and it significantly extends the estimation range when compared with MLM.

2.8.4 Impact of b , r , and d

We now vary the system parameters b , r , and d to study their impacts on the performance of CT, while M is fixed to 0.5 MB. The parameters are set as follows:

1. Impact of b : we fix $d = 2$, $r = 100$, $h = 2$ and vary b from 4, 6, to 8.
2. Impact of r : we fix $b = 4$, $d = 2$, $h = 2$ and vary r from 50, 100, to 200.
3. Impact of d : we fix $b = 4$, $r = 100$, $h = 2$ and vary d from 2, 3, to 4.

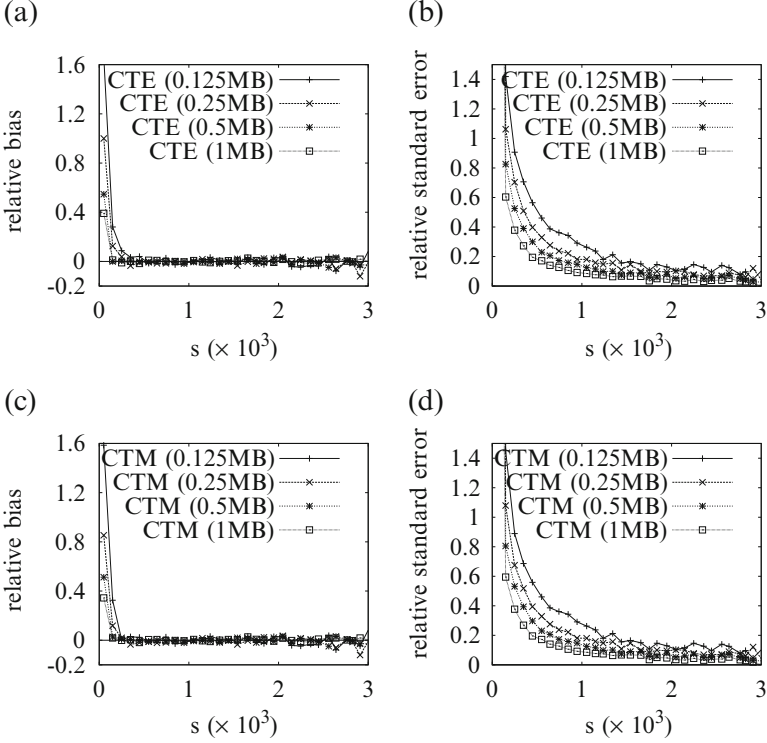


Fig. 2.13 (a) Shows the relative estimation bias $Bias(\hat{\frac{s}{s}})$ of CTE. (b) Shows the relative standard error $StdErr(\hat{\frac{s}{s}})$ of CTE. (c) Shows the relative estimation bias $Bias(\hat{\frac{s}{s}})$ of CTM. (d) Shows the relative standard error $StdErr(\hat{\frac{s}{s}})$ of CTM

We find that those parameters affect CTE and CTM in a similar way. Hence, we only present the estimation results of CTE in Figs. 2.14, 2.15, and 2.16. When we increase the b , r , or d , the estimation range will be increased accordingly. But this does not come for free. Since the increase of b , r , or d makes more flows share each counter, it is expected that the estimation accuracy will degrade due to elevated noise. However, the experimental results show that such degradation is small, and the performance of CTE is not very sensitive to the change of b , r , or d (we will show shortly that the change of h can significantly affect the estimation accuracy of CTE).

2.8.5 Comparison of CTE and E-CTE

Recall that our analysis in Sect. 2.7 points out that the increase of tree height h can degrade the performance of CTE and CTM. To demonstrate this, we run

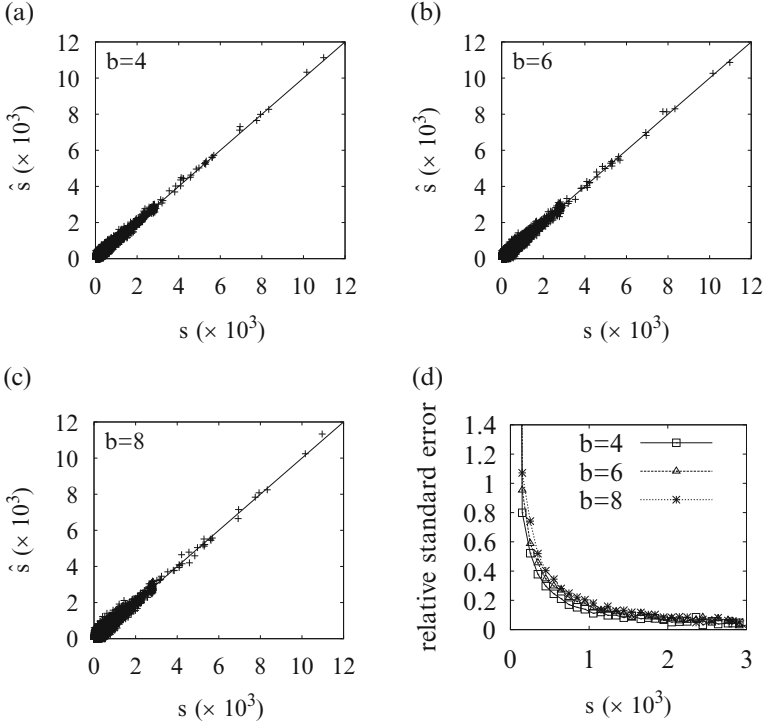


Fig. 2.14 Impact of b on the performance of CTE, where $M = 0.5$ MB, $d = 2$, and $r = 100$. (a) Shows estimation results of CTE when $b = 4$. (b) Shows estimation results of CTE when $b = 6$. (c) Shows estimation results of CTE when $b = 8$. (d) Shows the comparison of relative standard error when $b = 4, 6, 8$

experiments using the CTE method (similar results can be observed if the CTM method is adopted) with different tree heights. We fix $M = 0.5$ MB, $b = 4$, $d = 2$, $r = 100$, and vary h from 2, 4, to 6. The results are presented in Fig. 2.17. As we expect, the estimation accuracy of CTE becomes much worse with the increase of h . The fourth plot in Fig. 2.17 demonstrates that a larger h significantly increases the relative standard error of the estimates. We use E-CTE to stand for the CTE method under the Enhanced Counter Tree architecture which is designed to address this issue. For comparison, we conduct the same experiments on E-CTE, and the results are depicted in Fig. 2.18. Owing to the status bits in E-CTE, the increase of h only slightly degrades the performance of ECT. Therefore, ECT dramatically outperforms CTE when a relatively large h is adopted, e.g., $h = 6$.

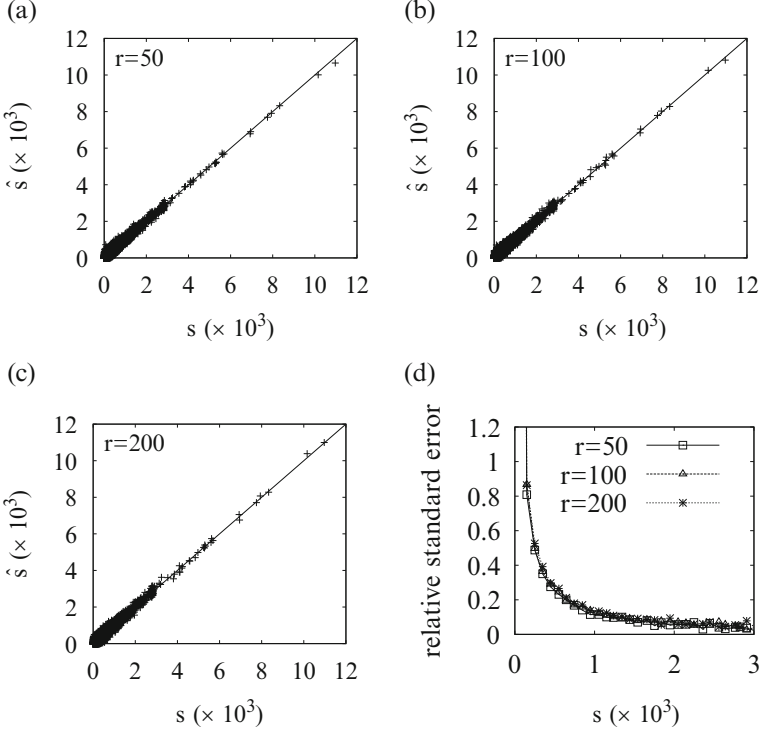


Fig. 2.15 Impact of r on the performance of CTE, where $M = 0.5$ MB, $b = 4$, and $d = 2$. (a) Shows estimation results of CTE when $r = 50$. (b) Shows estimation results of CTE when $r = 100$. (c) Shows estimation results of CTE when $r = 200$. (d) Shows the comparison of relative standard error when $r = 50, 100, 200$

2.8.6 Scalability of Counter Tree

To evaluate the scalability of Counter Tree, we apply it to a larger trace, which contains 8,253,368 TCP flows and 100,000,015 packets, and the minimum, average, and maximum flow size is 1, 12.12, and 295,451 packets, respectively. The available memory space M remains to be 0.125, 0.25, 0.5, and 1 MB, which translate to about 0.125bits/flow, 0.25bits/flow, 0.5bits/flow, and 1bits/flow, respectively. We implement Counter Tree with different system parameters. Figure 2.19 shows the estimation results of E-CTE with $b = 4$, $d = 2$, $r = 100$, and $h = 4$. It is clear that Counter Tree can still yield reasonably accurate estimation results under an extremely tight memory space, e.g., 0.125 bits/flow.

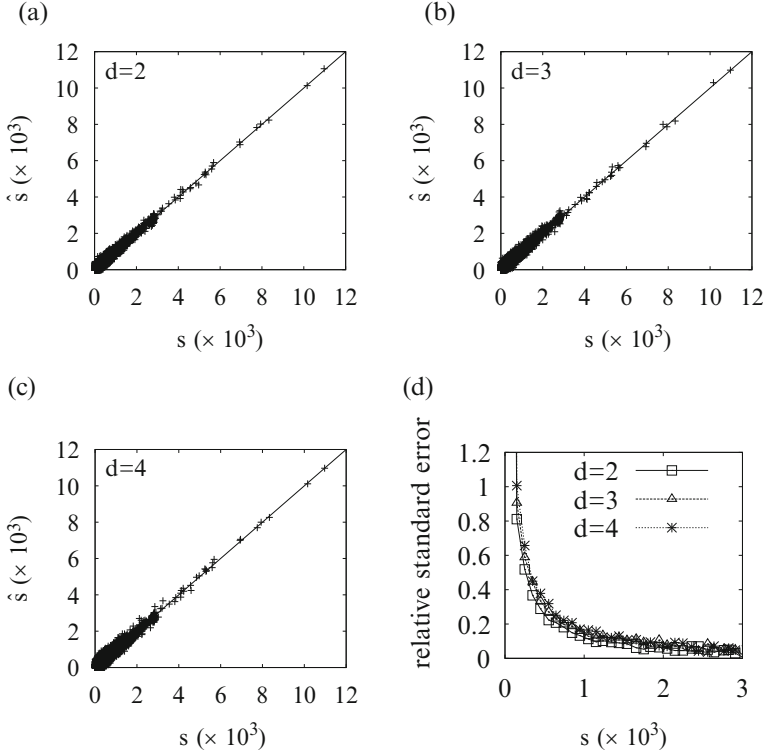


Fig. 2.16 Impact of d on the performance of CTE, where $M = 0.5\text{MB}$, $b = 4$, and $r = 100$. (a) Shows estimation results of CTE when $d = 2$. (b) Shows estimation results of CTE when $d = 3$. (c) Shows estimation results of CTE when $d = 4$. (d) Shows the comparison of relative standard error when $d = 2, 3, 4$

2.9 Summary

This chapter presents a scalable Counter Tree architecture. We develop a two-dimensional sharing technique, where each counter can be shared by multiple virtual counters and each virtual counter can be shared by multiple flows. As a result, Counter Tree significantly reduces memory requirement and extends estimation range. To encode a packet, Counter Tree only requires a little more than 2 memory accesses, which is asymptotically optimal. We use two offline decoding methods to estimate flow sizes. The extensive experiments with real network trace demonstrate that our methods can yield very accurate estimates even under an extremely tight memory space, e.g., 2 bits per flow.

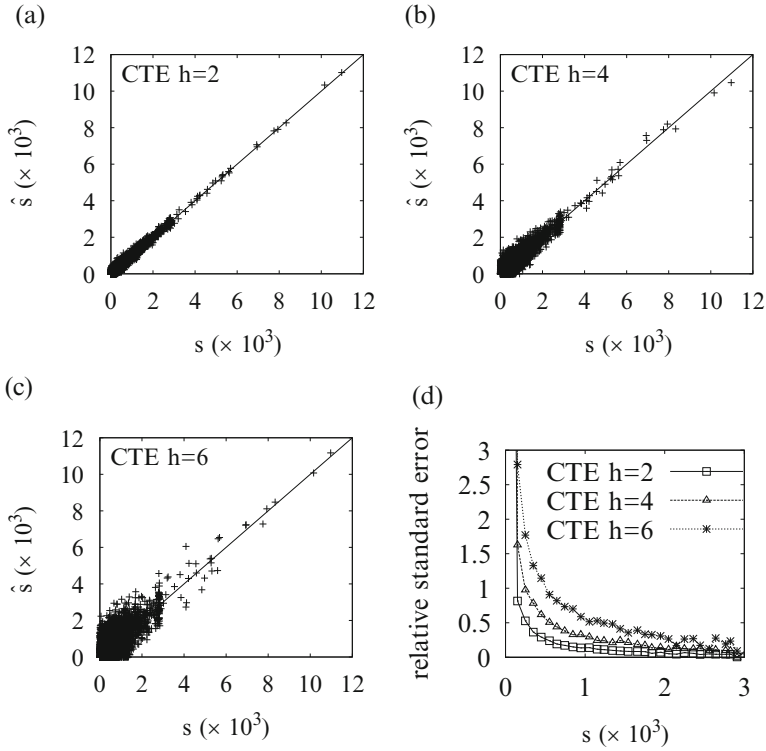


Fig. 2.17 Performance of CTE with different tree height h , where $M = 0.5$ MB, $b = 4$, $d = 2$, and $r = 100$. (a) Shows estimation results of CTE when $h = 2$. (b) Shows estimation results of CTE when $h = 4$. (c) Shows estimation results of CTE when $h = 6$. (d) Shows the comparison of relative standard error when $h = 2, 4, 6$

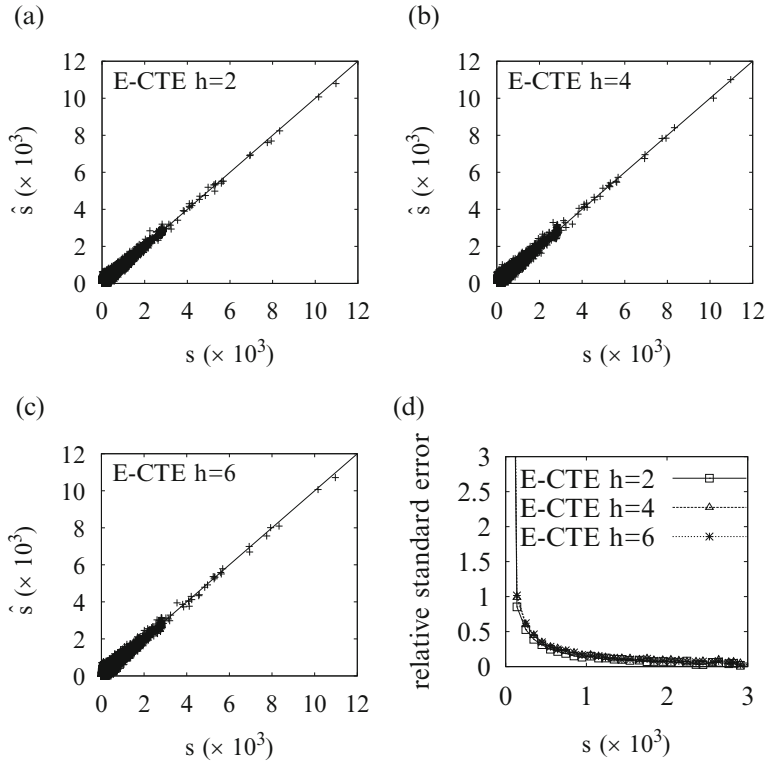


Fig. 2.18 Performance of E-CTE with different tree height h , where $M = 0.5$ MB, $b = 4$, $d = 2$, and $r = 100$. (a) Shows estimation results of E-CTE when $h = 2$. (b) Shows estimation results of E-CTE when $h = 4$. (c) Shows estimation results of E-CTE when $h = 6$. (d) Shows the comparison of relative standard error when $h = 2, 4, 6$

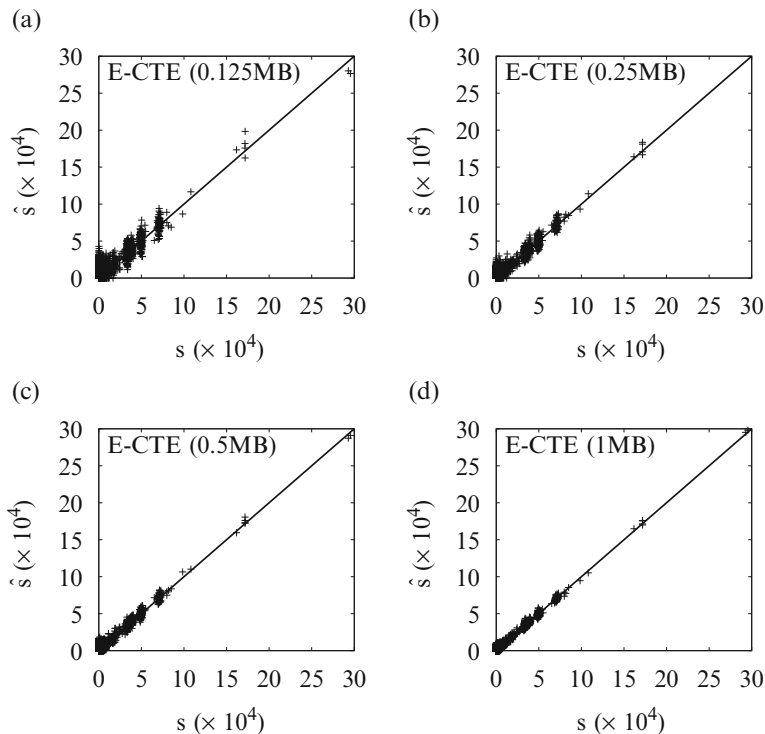


Fig. 2.19 Estimation results of E-CTE for a trace with 100 million packets. We set $b = 4$, $d = 2$, $r = 100$, $h = 4$, and $M = 0.125$ MB, 0.25 MB, 0.5 MB, and 1 MB (a)–(d), respectively

References

1. Chen, M., Chen, S.: Counter tree: a scalable counter architecture for per-flow traffic measurement. In: Proceedings of IEEE ICNP, pp. 111–122 (2015)
2. Cormode, G., Muthukrishnan, S.: An improved data stream summary: the Count-Min sketch and its applications. In: Proceedings of LATIN (2004)
3. Duffield, N., Lund, C., Thorup, M.: Estimating flow distributions from sampled flow statistics. In: Proceedings of ACM SIGCOMM (2003)
4. Einziger, G., Friedman, R.: Counting with tinytable: every bit counts! In: Proceedings of IEEE INFOCOM Workshops, pp. 77–78 (2015)
5. Estan, C., Varghese, G.: New directions in traffic measurement and accounting. In: Proceedings of ACM SIGCOMM (2002)
6. Kamiyama, N., Mori, T.: Simple and accurate identification of high-rate flows by packet sampling. In: Proceedings of IEEE INFOCOM (2006)
7. Kumar, A., Sung, M., Xu, J., Wang, J.: Data streaming algorithms for efficient and accurate estimation of flow size distribution. In: Proceedings of ACM SIGMETRICS (2004)
8. Kumar, A., Xu, J., Wang, J.: Space-code bloom filter for efficient per-flow traffic measurement. *IEEE J. Sel. Areas Commun.* **24**(12), 2327–2339 (2006)
9. Lehmann, E., Casella, G.: Theory of Point Estimation. Springer, New York (1998)

10. Li, T., Chen, S., Ling, Y.: Fast and compact per-flow traffic measurement through randomized counter sharing. In: *Proceedings of IEEE INFOCOM*, pp. 1799–1807 (2011)
11. Li, T., Chen, S., Ling, Y.: Per-flow traffic measurement through randomized counter sharing. *IEEE/ACM Trans. Netw.* **20**(5), 1622–1634 (2012)
12. Lu, Y., Prabhakar, B.: Robust counting via Counter Braids: an error-resilient network measurement architecture. In: *Proceedings of IEEE INFOCOM* (2009)
13. Lu, Y., Montanari, A., Prabhakar, B., Dharmapurikar, S., Kabbani, A.: Counter Braids: a novel counter architecture for per-flow measurement. In: *Proceedings of ACM SIGMETRICS* (2008)
14. Mikians, J., Dhamdhere, A., Dovrolis, C., Barlet-Ros, P., Solé-Pareta, J.: Towards a statistical characterization of the interdomain traffic matrix, pp. 111–123. In: *Networking 2012*, pp. 111–123. Springer, Berlin (2012)
15. Ramabhadran, S., Varghese, G.: Efficient implementation of a statistics counter architecture. In: *ACM SIGMETRICS Performance Evaluation Review*, vol. 31, pp. 261–271 (2003)
16. Shah, D., Iyer, S., Prabhakar, B., McKeown, N.: Analysis of a statistics counter architecture. In: *Hot Interconnects 9*, 2001, pp. 107–111. IEEE, New York (2001)
17. Zhao, Q., Xu, J., Liu, Z.: Design of a novel statistics counter architecture with optimal space and time efficiency. In: *ACM SIGMETRICS Performance Evaluation Review*, vol. 34(1), pp. 323–334 (2006)

Traffic Measurement for Big Network Data

Chen, S.; Chen, M.; Xiao, Q.

2017, VII, 104 p. 45 illus., 2 illus. in color., Hardcover

ISBN: 978-3-319-47339-0