

Fast Processing and Querying of 170TB of Genomics Data via a Repeated And Merged BloOm Filter (RAMBO)

Gaurav Gupta*
Electrical and Computer Engineering
Department, Rice University
gaurav.gupta@rice.edu

Minghao Yan*
Department of Computer Science,
Rice University
my29@rice.edu

Benjamin Coleman
Electrical and Computer Engineering
Department, Rice University
ben.coleman@rice.edu

Bryce Kille
Department of Computer Science,
Rice University
brycekille@gmail.com

R. A. Leo Elworth
Department of Computer Science,
Rice University
chilleo@gmail.com

Tharun Medini
Electrical and Computer Engineering
Department, Rice University
tharun.medini@rice.edu

Todd Treangen
Department of Computer Science,
Rice University
treangen@rice.edu

Anshumali Shrivastava
Department of Computer Science,
Rice University
anshumali@rice.edu

ABSTRACT

DNA sequencing, especially of microbial genomes and metagenomes, has been at the core of recent research advances in large-scale comparative genomics. The data deluge has resulted in exponential growth in genomic datasets over the past years and has shown no sign of slowing down. Several recent attempts have been made to tame the computational burden of sequence search on these terabyte and petabyte-scale datasets, including raw reads and assembled genomes. However, no known implementation provides both fast query and construction time, keeps the low false-positive requirement, and offers cheap storage of the data structure. We propose a data structure for search called RAMBO (Repeated And Merged BloOm Filter) which is significantly faster in query time than state-of-the-art genome indexing methods- COBS (Compact bit-sliced signature index), Sequence Bloom Trees, HowDeSBT, and SSBT. Furthermore, it supports insertion and query process parallelism, cheap updates for streaming inputs, has a zero false-negative rate, a low false-positive rate, and a small index size. RAMBO converts the search problem into set membership testing among K documents. Interestingly, it is a count-min sketch type arrangement of a membership testing utility (Bloom Filter in our case). The simplicity of the algorithm and embarrassingly parallel architecture allows us to stream and index a 170TB whole-genome sequence dataset in a mere 9 hours on a cluster of 100 nodes while competing methods require weeks.

*Equal contribution.

Code available at: https://github.com/gaurav16gupta/RAMBO_MSMT

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than ACM must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.

SIGMOD '21, June 20–25, 2021, Virtual Event, China

© 2021 Association for Computing Machinery.

ACM ISBN 978-1-4503-8343-1/21/06...\$15.00

<https://doi.org/10.1145/3448016.3457333>

CCS CONCEPTS

• **Information systems** → **Search engine indexing**; Distributed retrieval; • **Theory of computation** → **Bloom filters and hashing**; • **Applied computing** → *Computational genomics*.

KEYWORDS

Information retrieval; Genomic sequence search; Bloom filter

ACM Reference Format:

Gaurav Gupta, Minghao Yan, Benjamin Coleman, Bryce Kille, R. A. Leo Elworth, Tharun Medini, Todd Treangen, and Anshumali Shrivastava. 2021. Fast Processing and Querying of 170TB of Genomics Data via a Repeated And Merged BloOm Filter (RAMBO). In *Proceedings of the 2021 International Conference on Management of Data (SIGMOD '21)*, June 20–25, 2021, Virtual Event, China. ACM, New York, NY, USA, 9 pages. <https://doi.org/10.1145/3448016.3457333>

1 INTRODUCTION

The availability of genomic data facilitates necessary biological research like cancer genomics, vaccine development and immunization, infection tracking, early diagnosis and treatments, structural variant analyses, and more [30] [27]. Recent advances in DNA sequencing technologies have both increased the throughput and decreased the cost of reading the DNA sequence of organisms and microbial communities of interest. While this has broadened the horizons of biological research, it poses new challenges for computational biologists. Thanks to these advancements, genome sequence data has doubled in size every 2 years and is likely to grow at an increasing pace [9, 26]. The European Nucleotide Archive (ENA) and NCBI Archive already contain petabytes of data. It has become computationally prohibitive to search these vast archives for DNA sequences of interest. Efficient and frugal search functionality across all available genomic and metagenomic datasets is significant to public health. It would enable quick identification of already-sequenced organisms that are highly similar to an outbreak strain.

Method	Size	Query Time	Comments
Inverted Index	Best case: $\log K \cup_{S \in \mathcal{S}} S $	$O(1)$	Enormous construction time, Impractical for bigger datasets Best case needs MPH and a known k-mer (term) distribution
BIGSI/COBS	$\sum_{S \in \mathcal{S}} S $	$O(K)$	Query time is linear in K , Small size index
Sequence Bloom Trees	$\log K \sum_{S \in \mathcal{S}} S $	Best: $O(\log K)$, Worst: $O(K)$	Sequential query process is bottleneck
RAMBO	$\Gamma \log K \sum_{S \in \mathcal{S}} S $	$O(\sqrt{K} \log K)$	$\Gamma < 1$, Sub-linear query time

Table 1: Theoretical comparison of related algorithms on sequence searching. $S \in \mathcal{S}$ represents a document. K is the total number of documents. Here $\sum_{S \in \mathcal{S}} |S|$ represents total number of terms in K documents and $\cup_{S \in \mathcal{S}} |S|$ is total unique terms. MPH is minimal Perfect Hashing [8]. For the Inverted Index size, the extra $\log K$ comes from the bit precision document IDs. For SBTs, $\log K$ is the height of the tree and Bloom Filters at each level is $O(\sum_{S \in \mathcal{S}} |S|)$ big in total. Refer to Section 4 for the detailed analysis of RAMBO.

The DNA sequence search problem is analogous to Document Retrieval. Given a query gene strand, we are expected to retrieve the whole gene sequence that contains it (Figure 1). The search results are critical for a variety of genomic research and analysis tasks. Its similarities with the problem of web search, in terms of both objective and scale, have triggered a flurry of ideas borrowed from the information retrieval community [17, 23]. In the seminal work BLAST [5], a popular search platform for biological databases, the authors provided the first attempt to search over large databases. However, the method does not scale to large query datasets [5] due to the reliance on computationally expensive local sequence alignment. On the other hand, traditional approaches such as the inverted index [18] cannot quickly index large-scale data without violating memory constraints.

To address this issue, computational biologists and database practitioners have shifted their attention to Bloom Filter based methods [6, 7, 9] and similar bit-signature approaches for gene sequence search due to the sheer scale of genomic data. A recent *Nature Biotechnology* article, BIGSI [9], proposed a method which was successful in indexing the set of 469,654 bacterial, viral and parasitic DNA sequences in the European Nucleotide Archive [3] of December 2016. These sequences come from read archive datasets (FASTQ format) or assembled genomes (FASTA format) consisting of raw sequences from a DNA sequencer or genome assembler, respectively. The average length of each of these half-million sequences is more than 100M characters. This makes the entire archive database about 170TB in size.

To create the index, BIGSI and many other practical indices convert a long gene sequence into a set of length-31 strings (each shifted by 1 character) and compress the strings using a Bloom Filter (sometimes called a Bitsliced signature). These length-31 strings are called k -mers. It is analogous to "terms" in the information retrieval literature. Specifically, a k -mer is a character n -gram where $k = n$. In our experiments $k = 31$, just like most of the state of the art methods. We explain the rationale behind this choice in Section 5. BIGSI essentially creates a Bloom Filter for each document (a set of k -mers for one microbe). This index is simply an array of independent Bloom Filters where the query time grows *linearly* in the number of documents. The **Sequence Bloom Tree (SBT)** [17, 28] is another approach to solve the sequence search problem on the scale of the entire sequence read archive (SRA) [20] using Bloom Filters. To achieve sublinear query time complexity, the SBT uses a tree-like hierarchy of Bloom Filters [28]. However, this introduces

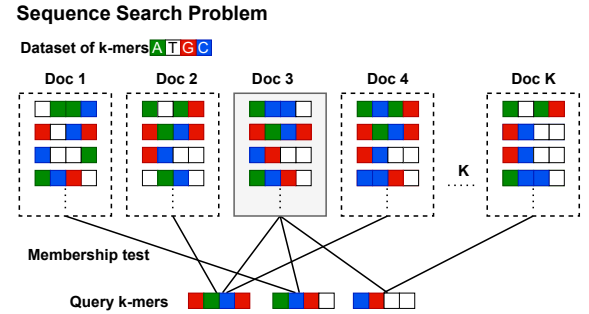


Figure 1: Sequence search problem: First, we convert each of the K documents into a set of k -mers. The k -mers of length 31 are generated using a sliding window on the sequence ($k=4$ in the figure for illustration). Given the k -mers from a query sequence, the task is to determine which of the K documents contain all the k -mers present in the query.

a substantial memory overhead at each node. Moreover, the query process cannot enjoy the parallelism of BIGSI because tree-based traversal is a sequential algorithm. Experimental results from [9] suggest that SBTs become less scalable when the time and evolution of species are factored in, which is the case for bacteria and viruses. Several follow-ups using ideas similar to SBTs also suffer from the same issues [24, 29, 31]. By removing the hierarchy, BIGSI and its recent follow-up **COBS (Compact bit-sliced signature index)** [6] obtained substantial memory savings compared to SBTs. The simplicity of the index, combined with embarrassingly parallel architecture and clever bit manipulation tricks, enables BIGSI and COBS to process and search over 170TB WGS datasets in a record query time. However, with an exponential increase in the number of datasets in the sequence archive, the linear scaling of latency and energy is too expensive.

Our Focus: We propose methods to reduce the query cost of sequence search over the archive of dataset files to address the sheer scale and explosive increase of new sequence files. In particular, unlike BIGSI and COBS, we do not want the number of Bloom Filters used in the query to be of the same order as the number of datasets, which can run into several million. At the same time, we also want an algorithm that maintains all other beneficial BIGSI and COBS features. We are looking for a data structure for sequence search, which has the following properties: **1.** A zero false-negative rate, **2.** A low false-positive rate, **3.** Cheap updates for streaming inputs, **4.** Fast query time, and **5.** A simple, system-friendly data structure

that is straightforward to parallelize. The system should have all these properties with the least possible memory size.

Insights from Computer Science Literature: There is a fundamental algorithmic barrier at the heart of this problem. The classical sub-linear search data structure provides tree-based solutions that mainly implement the SBT [28]. However, trees complicate the query process and have issues with balanced partitions, especially when dimensionality blows up. Fortunately, the Count-Min Sketch (CMS) Algorithm [16] from the data streaming literature provides a workaround. Our proposal for sequence search, Repeated And Merged BloOm Filter (RAMBO) is a CMS using Bloom Filters. It is a simple and intuitive way of creating merges and repetitions of Bloom Filters for membership testing over many sets. RAMBO leads to a better query-time and memory trade-off in practice. It beats the current baselines by achieving a very robust, low memory and ultrafast indexing data structure.

1.1 Contribution

Instead of having separate Bloom Filters for each document, we split the documents into a small number of random partitions. We keep one Bloom Filter for each partition. Inspired by the theory of the Count-Min Sketch [16], if we repeat the partitioning process with different random seeds a small number of times, we can return documents with high accuracy.

Our proposed index RAMBO leads to massive improvements in query cost and would allow effortless scaling to millions of documents. We provide a rigorous experimental and theoretical analysis of the query time. Experimental comparisons show that RAMBO is significantly faster (between **25x** to **2000x** improvement) in query time over the most competitive baselines of gene data indexing while preserving competitive false-positive rates.

RAMBO can be made embarrassingly parallel for insertion and query by an intelligent choice of hash functions and judicious utilization of parallel processors. It can be easily distributed over multiple nodes, cores, and threads for achieving substantial speedup. We show remarkable improvements in construction and query time over the baselines on a 170TB WGS dataset. We reduce the time of offline construction of the index from 6 weeks (1008 hours for BIGSI) to only 9 hours (including the additional download time of 8 hrs). This is attributed to the fact that BIGSI downloads and indexes 460,500 files (170TB) sequentially. Its successor, COBS [6], is much faster and better in all aspects than BIGSI. The RAMBO index has a slightly larger memory requirement than an optimal array of Bloom Filters (COBS), but it keeps a cheap index (1.8 terabytes) for 170TB worth of data. It is important to note that Bloom Filters in RAMBO can be replaced with any other set membership testing method.

2 PRELIMINARIES

2.1 Bloom Filters

The Bloom Filter [7, 15, 22] is an array of m bits which represents a set S of n elements. It is initialized with all bits set to 0. During construction, we apply η universal hash [11] functions $\{h_1, h_2 \dots h_\eta\}$ with range m to the elements of S (η and n are different). We set the bits at the respective locations $\{h_1(x), h_2(x) \dots h_\eta(x)\}$ for each key $x \in S$. Once the construction is done, the Bloom Filter can be used to determine whether a query $q \in S$ by calculating the AND

of the bits at the η locations: $h_1(q), h_2(q) \dots h_\eta(q)$. The output will be True if all η locations are 1 and False otherwise. Bloom Filters have no false negatives as every key $x \in S$ will set all the bits at locations $\{h_1(x), h_2(x) \dots h_\eta(x)\}$. However, there are false positives introduced by hash collisions. The false positive rate of the Bloom Filter, p , is given by: $p = \left(1 - \left[1 - \frac{1}{m}\right]^{\eta n}\right)^\eta \approx \left(1 - e^{-\eta n/m}\right)^\eta$. We should note that this expression makes many simplifying assumptions and is not entirely correct, as we assume independence of the probabilities of each bit being set. A more accurate analysis is given in Christensen et al [13]. However, its deviation from practical numbers is minimal when the Bloom Filter size is large (Figure 2 of [13]). At the scale that we are dealing with, the difference becomes insignificant.

Using the simplified analysis, the false positive rate is minimized when we use $\eta = -\frac{\log p}{\log 2}$ and $m = -n \frac{\log p}{\log 2}$. The size of a Bloom Filter grows linearly in the cardinality n of the set it represents. Bloom Filter has a constant-time query operation.

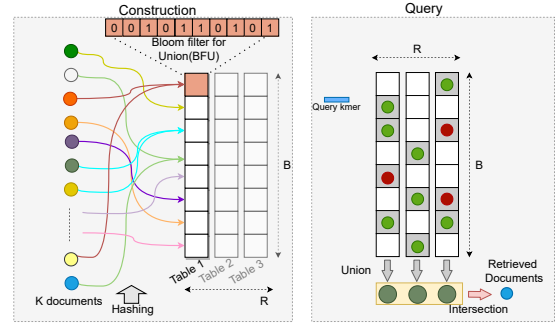


Figure 2: (a) Left: RAMBO architecture and the insertion process. The construction of the first repetition is highlighted. Here the K documents are randomly partitioned (via a 2-universal hash function [11]). Each Bloom Filter (called BFU) is the union of sets equivalent of partitioned documents. (b) Right: For a given query each table, RAMBO returns one or more BFUs (represented by the dots) where the membership is defined. The red dot represents the false positives and green dot represents the true positives. The membership of a query k -mer is defined by the union of the returned documents from each repetition followed by the intersection across R repetitions.

3 RAMBO: REPEATED AND MERGED BLOOM FILTERS

The RAMBO architecture (Figure 2) comprises an array of R tables, each containing B Bloom Filters. We partition the K documents into B groups and compress every group of documents to a Bloom Filter called Bloom Filters for the Union (BFU). Here each document is a set of k -mers. This process is repeated independently for the R tables using R different and independent 2-universal hash functions. Due to the hash functions' universality, every cell of a table in RAMBO contains K/B sets from S in expectation.

3.1 Intuition

We have K documents, partitioned into B partitions, where $2 \leq B \ll K$. Now, we are given a query term q . For simplicity, assume

that this query belongs to a single document. Now, if we query each partition, we can determine which one of them contains q . We refer to this partition as A_1 . Thus, with only B Bloom Filter queries, we have reduced the number of candidate sets from K to $\frac{K}{B}$ in expectation. If we independently repeat this process again, we find another partition A_2 that contains q . Our pool of candidate sets is now the set intersection of A_1 and A_2 , which in expectation has size $\frac{K}{B^2}$. With more repetitions, we progressively rule out more and more options until we are left with only the sets that contain q . The critical insight is that each repetition reduces the number of candidates by a factor of $\frac{1}{B}$, which decreases *exponentially* with the number of repetitions. Since RAMBO is an extension of the Count-Min Sketch (CMS) data structure [16], most theoretical guarantees carry forward. We replace the counters in the CMS with Bloom Filters. Instead of adding counters to construct the CMS, we merge the sets of k -mer terms. The querying procedure of the CMS is replaced with an intersection over the merged sets to determine which sets contain a query term.

3.2 Construction

We assign sets based on a partition hash function $\phi(\cdot)$ that maps the set identity to one of B cells. We use R independent partition hash functions $\{\phi_1, \phi_2, \dots, \phi_R\}$. Suppose we want to add a set of terms in Doc-1 to RAMBO. We first use the partition function $\phi_i(\text{Doc-1})$ to map Doc-1 in repetition $i, \forall i \in \{0, 1, \dots, R\}$. Then we insert the terms (k -mers) of Doc-1 in R assigned BFUs (Algorithm 1). The Bloom Filter insertion process is defined in Section 2.1. We define the size of each BFU based on the expected number of insertions in it. This is further analyzed in Section 5.1.

Our RAMBO data structure is a $B \times R$ CMS of Bloom Filters. Clearly, this structure is conducive to updates to a data stream. Every new term in a document is hashed to unique locations in R tables. The size of the BFU can be predefined or a scalable Bloom Filter [4] can be used for adaptive size.

Algorithm 1 Algorithm for insertion in RAMBO architecture

Input: Set S of K sets
Result: RAMBO (size: $B \times R$)
 Generate R partition hash functions $\phi_1(\cdot), \dots, \phi_R(\cdot)$
 RAMBO $\leftarrow B \times R$ array of Bloom Filters
while Input S_i **do**
 for term $x \in S_i$ **do**
 for $d = 1, \dots, R$ **do**
 Insert(x , RAMBO[$\phi_d(x), d$])
 end for
 end for
end while

3.3 Query

We start the RAMBO query process by performing membership testing for each of the $B \times R$ BFUs. This is followed by taking the union of the sets corresponding to each filter that returns True in each table, and then the intersection of those unions across the R tables. The union and intersection are implemented using fast bitwise operations, and the expected query time is sub-linear (Section 4.2). Algorithm 2 presents the query process. Here, set $A \subseteq S$ is the final set of matched documents.

Note that a k -mer can occur in multiple documents, which we call multiplicity. For this reason, and the fact that Bloom Filters have a nonzero false positive rate, multiple BFUs may return True (Figure 2). If a query k -mer does not exist in any document, RAMBO will most likely return an empty set $A = \emptyset$ or a small set of false positives with low probability.

3.3.1 Large Sequence Query. To query a larger term sequence with length n , we simply use a sliding window of size k to go through the entire sequence. This will create a set of terms Q to query. Then we iterate over the terms in Q and membership test each term. The final output should be the intersection of all returned outputs from each term in Q . Since Bloom Filter does not have any false negatives, we are guaranteed to obtain a valid result. We only need to perform exponentially less (in the cardinality of Q) number of membership tests as the first returned FALSE will be conclusive. It is interesting to note that the final output size is upper bounded by the output size of the rarest k -mer in the query sequence.

Algorithm 2 Algorithm for query using RAMBO architecture

Input: query $q \in \Omega$
Architecture: RAMBO (M) // Size $B \times R$ array of Bloom Filters.
Result: $A \subseteq S$, where $q \in S_i \forall S_i \in A$
 $Q = \text{Terms}(q)$
for $r = 1 : R$ **do**
 $G_r = \{\text{Null}\}$
 for $b = 1 : B$ **do**
 $G_r = G_r \cup \text{DocIDs}(M[b, r])$ **if** $Q \in \text{BFU}(M[b, r])$
 end for
end for
 $A = \cap_i G_i$ {final returned Doc ID's}
Define: $Q \in \text{BFU}(M[b, r])$
 return True if $x \in \text{BFU}(M[b, r]) \forall x \in Q$

4 ANALYSIS

Problem Definition: We are given a set of K documents $S = \{S_1, S_2, \dots, S_K\}$. Each document S_i contains k -mers from a universe Ω of all possible k -mers. Given a query $q \in \Omega$, our goal is to identify all the documents in S that contain q . That is, the task is to return the subset $A_q \subseteq S$, such that $q \in S_i$ if and only if $S_i \in A_q$.

RAMBO has two important parameters, R and B , that control the resource-accuracy trade-off. In this section, we will analyze the false positive rate, query time, and index size to find the optimal values of R and B .

4.1 False-Positives

Our first claim is that RAMBO cannot report false negatives. This follows trivially from our merging procedure and the fact that each BFU cannot produce false negatives [7]. Next, we begin by finding the false positive rate of one document and extend this result to all K documents.

LEMMA 4.1. Per document False Positive Rate

Given the RAMBO data structure with $B \times R$ BFUs, each with false positive rate p and query q , we assume that q belongs to no more than V documents. Under these assumptions, the probability of incorrectly

reporting that $q \in S_i$ when $q \notin S_i$ is

$$F_p = \left(p \left(1 - \frac{1}{B} \right)^V + 1 - \left(1 - \frac{1}{B} \right)^V \right)^R$$

where p is the individual false positive rate of BFUs.

Proof: The probability of selecting a BFU which should return false is $(1 - \frac{1}{B})$ if the multiplicity of the key is 1. If it is V then the probability becomes $(1 - \frac{1}{B})^V$. Since each Bloom Filter has a false positive rate p , the probability of introducing a false positive through a Bloom Filter failure is $p(1 - \frac{1}{B})^V$.

Because each BFU contains multiple documents, ‘True’ documents (containing the query) can occur with ‘False’ documents (not containing the query). Thus, we may also introduce false positives by merging S_i into a BFU that contains the k -mer. The probability for this event is $1 - (1 - \frac{1}{B})^V$. The total per-document false positive rate for R independent repetitions is $F_p = (p(1 - \frac{1}{B})^V + 1 - (1 - \frac{1}{B})^V)^R$. Using this theorem, we can construct the overall false positive rate of RAMBO.

LEMMA 4.2. RAMBO False Positive Rate

Given a RAMBO data structure with $B \times R$ BFUs, each with false positive rate p and query q , we assume that q belongs to no more than V documents. Under this assumption, the probability of reporting an incorrect membership status for any of the K documents, a.k.a. RAMBO False Positive Rate (δ) is upper bounded by

$$\delta \leq K \left(1 - (1 - p) \left(1 - \frac{1}{B} \right)^V \right)^R$$

where p is the individual false positive rate of the BFUs.

This is a direct result of Lemma 4.1 with union bound. A consequence of lemma 4.2 is that we need sub-linear RAMBO repetitions (logarithmic in K) to obtain an overall false positive rate δ . We can state that it is sufficient to keep $R \geq \log K - \log \delta$.

THEOREM 4.3. Number of Repetitions

Given a set of K files, maximum RAMBO false positive rate δ and Bloom Filter for each repetition, we need R repetitions such that-

$$R = O(\log K - \log \delta)$$

4.2 Query Time Analysis

This section demonstrates that RAMBO achieves sublinear query time in expectation. To query the set membership status of an element x , we perform $B \times R$ Bloom Filter look-ups followed by union and intersection operations (Section 3.3).

Since each repetition makes a disjoint partition of the K documents, the union operations do not require any computational overhead. The set intersections between repetitions, however, require $|X_1| + |X_2|$ operations, where X_1 is the set of all active documents in first repetition and X_2 is the set of all active documents in the next repetition. Since there are R repetitions, the total cost for the intersection is $\sum_{r=1}^R |X_r|$. By observing that $\mathbb{E}[|X_r|] \leq V + Bp$, we obtain the following result.

LEMMA 4.4. Expected query time

Given the RAMBO data structure with $B \times R$ BFUs and a query q that

is present in at most V documents, the expected query time is

$$\mathbb{E}[q_t] \leq BR\eta + \frac{K}{B}(V + Bp)R$$

where K is the number of documents, p is the BFU false positive rate, and η is the number of hash functions used in BFUs.

The first term represents the time to query the $B \times R$ BFUs. Note that η ranges from 1 to 6 in practice. The second term is the time required to perform R intersections. We get $B = \sqrt{KV/\eta}$ by minimising the query time (i.e. solving $\nabla_B(\mathbb{E}[q_t]) = 0$). To obtain an expression for the query time in terms of the overall failure probability δ and the number of documents K , we suppose that $p \leq \frac{1}{B}$ and set R according to Theorem 4.3. Our main theorem is a simplified version of this result where we omit lower-order terms.

THEOREM 4.5. RAMBO Query time

Given a RAMBO data structure and a query q that is present in at most V documents, RAMBO performs the search over K documents with false positive rate $\leq \delta$ in query time q_t , where

$$\mathbb{E}[q_t] = O\left(\sqrt{K}(\log K - \log \delta)\right)$$

Note that V is independent of K and δ .

4.3 Memory Analysis

We provide an average case analysis under a simplifying assumption to analyze the expected performance of our method. We assume that every key has a fixed multiplicity V , meaning that every item is present in exactly V documents. Under these assumptions, RAMBO requires the following amount of space.

LEMMA 4.6. Size of RAMBO

For the proposed RAMBO architecture with size $B \times R$ and data with K files, where every key has V number of duplicates, the expected memory requirement is

$$\mathbb{E}_v(M) = \Gamma \log K \log(1/p) \sum_{S \in \mathcal{S}} |S| \quad \text{where } \Gamma < 1$$

Here $\Gamma = \sum_{v=1}^V \frac{1}{v} \frac{(B-1)^{V-2v+1}}{B^{V-1}}$. The expectation is defined over the variable v which takes values from $\{1, 2, \dots, V\}$. This expression of Γ holds if we are hashing document IDs using a universal hash function. If $B = K$, we will have one Bloom Filter per set. In that case, $\Gamma = 1$. We prove the expression of Γ and its variation for any $B < K$ and $V > 1$ in Section 7.

5 EXPERIMENTS

5.1 Parameter Selection and Design Choices

Size of BFU: For each BFU to have a false positive rate p using η hashes, the size of the BFU must be set based on the number of insertions (Section 2.1). One way to determine the BFU size is to preprocess the data by counting the number of terms that will fall into each BFU. In practice, it is sufficient to estimate the average set cardinality from a tiny fraction of the data, and we use this cardinality to set the size for all BFUs. Section 5.2 presents these statistics for our data.

B and R: They are chosen according to $B = O(\sqrt{K})$ and $R = O(\log K)$, where the constants were found empirically.

Bitmap arrays: The intersection may be implemented using either

	Time per query (ms) (CPU time)							Construction time				
	FASTQ				McCortex			FASTQ			McCortex	
#files	HowDe	SSBT	RAMBO	RAMBO ⁺	COBS	RAMBO	RAMBO ⁺	HowDe	SSBT	RAMBO	COBS	RAMBO
100	5.24	8.47	0.018	0.0151	0.19	0.014	0.005	2h30m	52m	35m	1m25s	1m12s
200	10.38	17.12	0.025	0.0202	0.38	0.017	0.011	8h	1h47m	52m	3m58s	2m25s
500	24.15	42.27	0.056	0.0483	1.03	0.04	0.018	21h	4h51m	1h57m	8m28s	6m22s
1000	-	82.32	0.093	0.0747	1.78	0.07	0.031	-	9h16m	4h6m	14m18s	12m32s
2000	-	161.58	0.191	0.149	2.72	0.09	0.059	-	18h22m	8h55m	15m38s	25m41s

Table 2: Performance comparison between RAMBO and baselines on 1000 queries. To ensure a fair comparison, we have selected baseline hyper-parameters from their papers with the target false positive rate range of [0.01, 0.011], where the RAMBO false positive rate always falls in the range [0.0095, 0.01]. HowDeSBT exceeds the available RAM on our platform after 500 files.

#files	Size (FASTQ)			Size (McCortex)	
	HowDe	SSBT	RAMBO	COBS	RAMBO
100	92.5GB	9.5GB	12.8GB	2.4GB	3.5GB
200	182GB	9.5GB	19GB	4.9GB	6.3GB
500	456GB	18GB	42GB	7.5GB	13.9GB
1000	-	36GB	70GB	20GB	23.2GB
2000	-	72GB	140GB	28GB	47 GB

Table 3: Size of index comparison for the same experiment from Table 2. In worst case, RAMBO takes $O(\log K)$ extra space than the optimal Array of Bloom Filter (COBS). HowDeSBT and SSBT uses RRR [25] bitvector compression, however RAMBO does not compress the bitvectors. Any possible compression based optimization is left for the future exploration.

bitmap arrays or sets. For binary operations, the OR operation is $O(1)$, hence intersection is very fast using bitmaps. The complexity of the AND operation depends on the set size; bitmaps are more efficient when 1s occupy > 15% of the bitmap [21]. This is true in our case, so we used bitmaps and found that the AND operations take fewer than 5% of the query process cycles. Extensions such as SIMD-accelerated bitmaps are outside the scope of this work.

Query time speedup: We may avoid querying all $B \times R$ BFUs by analyzing the repetitions sequentially. Specifically, in repetition r we only need to query the BFUs that contain documents returned by previous repetitions $\{1, 2, \dots, r-1\}$. BFUs that do not contain documents identified by previous repetitions cannot change the output, as any documents corresponding to those BFUs will be removed by the intersection operation of Algorithm 2. We obtain a significant speedup using this sparse evaluation process (RAMBO⁺ in Table 2).

k-mer size for sequence indexing: The value of k must be large enough that each gene sequence may be uniquely identified by a set of distinct k -mers. In the laboratory, the gene sequencing of an organism is done in parts by the sequencing machine. Here, each part (called "reads") is around 400 – 600 in length typically. One might be tempted to set k equal to the read length. However, portions of each read are often corrupted by errors, so a smaller k must be used in practice. The work by Chikhi et.al. [12] confirmed $k=31/51/71$ to be optimal for k -mer set size and uniqueness. For the ENA dataset, most of the best methods [10] [6] [28] [19] use $k=31$, partially also because it is small enough to be represented as a 64-bit integer variable with 2-bit encoding. For these reasons and to provide a fair comparison with popular baselines, we use $k=31$.

Dataset: We use the 170TB WGS dataset (containing 460500 files) as described in [9] and [3]. It is the set of all bacterial, viral, and parasitic gene sequence data in the European Nucleotide Archive (ENA) as of December 2016. The data is present in two formats - 1) FASTQ [14] files containing raw, unfiltered sequence reads and 2) McCortex [32] [10] format, which is a filtered set of k -mers that omits low-frequency errors from the sequencing instruments. The k -mer length is 31 and the sequence alphabet (or nucleotides) are A, T, G, and C. Using 1000 random documents we found that the {average, standard deviation} number of k -mers is {377.6M, 354.9M}, number of unique k -mers is {95M, 103.1M} and file size is {145MB, 86.5MB}. The high variation among documents demands either a preprocessing pass or our pooling method from Section 5.1 to set the size of each BFU - we use the pooling procedure.

5.2 Genomic sequence indexing

We start our experiments by indexing only the first 2000 documents (2.4 TB). The results for the subset are shown in Table 2. The details are as follows

Baselines: The COBS (Compact bit-sliced signature index) [6] (Index based on an array of Bloom Filters) prefers McCortex data format and hence gives a very erroneous output on FASTQ. The Bloom Filter tree-based methods, SSBT [29] and HowDeSBT [19] works with FASTQ version but not with McCortex. Hence, we compare with COBS on McCortex and with SSBT and HowDeSBT on FASTQ. The comparison with BIGSI is unnecessary, as COBS is the successor of BIGSI and is better in all aspects. The baseline implementations and RAMBO are in C++.

Parameters: For HowDeSBT, the Bloom Filter size is 7.5×10^9 bits. HowDeSBT only supports 1 hash function and crashes if another value is used. For SSBT, we use 4 hash functions and set Bloom Filter size to 8.5×10^8 bits. For COBS, we use 3 hash function and set the false positive rate to 0.01. These parameters were hand-optimized and hard-coded into the program by the authors of COBS. For RAMBO, we use 2 hash functions, the number of partitions B is 15, 27, 60, 100 and 200 for number of set insertions 100, 200, 500, 1000 and 2000. We set $R = 2$ and the BFU size to 10^9 bits for the McCortex data. For FASTQ, we use $R = 3$ and 2^9 bits. These parameters are optimal for low query time, keeping in mind the allowable (comparable to baselines) index size, false positive rate, and construction time.

Evaluation Metrics: Creating a test set with ground truth requires a very time-consuming procedure of generating inverted indices.

The index’s actual false positive rate can also be assessed by creating an artificial and unseen query-ground-truth set and inserting this set into the index before querying. Therefore, we calculated the false positive rate by creating a test set of 1000 randomly generated 30 length k -mer terms. We used length 30 to ensure that there are no collisions from the existing k -mers already in the RAMBO data structure. These k -mers were assigned to V files (distributed exponentially $(1/\alpha) \exp(-x/\alpha)$ with $\alpha = 100$) randomly. The test set is much much smaller than the dataset’s actual size; hence it makes an insignificant change in the size of RAMBO. To get the index size, we report the maximum resident set size (RSS) in memory as returned by the *time* utility or the serialized index size, whichever is higher. This size includes the main index as well as all auxiliary data structures (like the inverted index mapping B buckets to K documents). Query time is the CPU time on a single thread, and construction time is the wall-clock time on 40 threads (with no other process running on the machine).

System and Platform Details: We ran the experiment on a cluster with multiple 40 core Intel(R) Xeon(R) Gold 6230 CPU @ 2.10GHz processor nodes. Each node has 192 GB of RAM and 960 GB of disk space. The experiments, apart from RAMBO construction on the full dataset, are performed on a single node. We did not use multi-threading for querying.

From Table 2 we can see that RAMBO has much faster query time (from around **25x** to **2000x**) than the baselines. Furthermore, RAMBO achieves a small index size (practically close to the theoretical lower bound - the array of Bloom Filters). The construction of RAMBO is an I/O bound process; hence we see almost linear growth in construction time (with number of files), which is equivalent to COBS and faster than SSBT and HowDeSBT. Insertion from McCortex format is blazing fast and preferred as it has unique and filtered k -mers.

5.3 Smart parallelism- Indexing the full 170TB WGS dataset in 9 hours from scratch

We now address the construction of RAMBO for the entire 170TB dataset. One could theoretically create a single RAMBO data structure with the given R and B parameters on a single machine. However, this is infeasible due to the limited DRAM and compute resources. We could also construct the index over multiple machines using a message passing interface, but this introduces a massive latency overhead due to data transmission over the network. A third way is using a shared-memory cluster of machines, which requires a massive DRAM and is infeasible in current multi-core servers.

We propose a better solution. We parallelize the computation by partitioning the RAMBO data structure over 100 nodes of the cluster. Each node contains a small RAMBO data structure indexing $1/100$ of the whole dataset, which is around 4605 files in our case. In the streaming setting, a file (set of terms) is routed to a BFU of a node randomly. More details about routing are the following-

Routing: We first use a random hash function $\tau(\cdot)$ to assign files to node and then use an independent smaller node-local 2-universal hash function $\phi_i(\cdot)$ to assign the file to the local Bloom Filter (BFU). This process preserves all the mathematical properties and randomness in RAMBO as the final mapping is again 2-universal, i.e., the probability of any two datasets colliding is exactly $1/B$, where B

is the total range (number of partitions in RAMBO). The two-level hash function is given by: $(b \times \tau(D_j)) + \phi_i(D_j)$. For a repetition i , where $i \in \{0..R\}$, b is the number of partitions in RAMBO on a single machine and also the range of $\phi(\cdot)$, D_j is the name ID of j^{th} dataset, and the range of $\tau(\cdot)$ is $\{0..100\}$ in our case. Note that this two-level hash function allows us to divide the insertion process into multiple disjoint parts (100 in our case) without repeating any installation of datasets and internode communications. Effectively, each node will contain a set of 4605 files in expectation. In this way, we eliminate costly transmission of data among the nodes. The data structure on each node has size $B = 500$ and $R = 5$. Stacking them vertically makes the complete RAMBO data structure of size $B = 50000$ and $R = 5$. This process preserves the randomness of set insertion, i.e., the probability of any two sets colliding is exactly $1/B$, where B is the total range (50000 in this case).

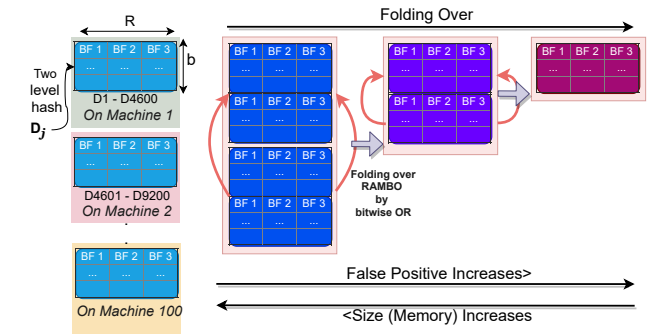


Figure 3: Indexing process of 460K documents over a cluster of 100 nodes. Each machine carries a part of RAMBO with size 500×5 Bloom Filters. The dataset is routed to machine via $\tau(\cdot)$ hash functions followed by insertion using $\phi_i(\cdot)$. The combined direct routing is done by a two-level hash function equivalent $(b \times \tau(\cdot) + \phi(\cdot))$. The stacked view of RAMBO shows the folding process. The folding is done such that number of repetitions R remains the same but B halves, so as the total size. Folding reduces memory progressively by factors of 2, 4, 8... and increases false positive rate super-linearly.

	Query Time (CPU time in ms)	Index size
Fold 2	66.5	7.13 TB
Fold 4	43.5	3.6 TB
Fold 8	26.25	1.78 TB

Table 4: CPU time (in ms) per query of the k -mer averaged over 1000 queries. Each column shows the different number of RAMBO folds. Second column shows the memory size (in TB) of RAMBO for each fold.

This interesting parallel insertion trick results in a fully constructed RAMBO in **about an hour on 100 CPU nodes when using the McCortex file format**. The additional 8 hours are used to download the dataset. It is the round-off time of the highest time taking job. Here we have to ensure that all machines use the same parameters (B, R , Bloom Filter size and hash function $\tau(\cdot)$, $\phi(\cdot)$ and $h(\cdot)$) as well as the random seeds. The consistency of seeds across

machines and larger than required B and R allow us to flexibly reduce the size of RAMBO later by doing bitwise OR between the corresponding BFUs of the first half of RAMBO over the other half (vertically). Each of these processes reduces the index size $B \times R$ to $\frac{B}{2} \times R$. This is called folding over. Refer Figure 3.

Folding Over: The data structures on every machine are independent and disjoint, but they have the same parameters and uses the same hash seeds. Since RAMBO is all made of Bloom Filters, we can perform the bit-wise OR between the first half of RAMBO over the other half to reduce the partitions in RAMBO from B to $\frac{B}{2}$. This operation is depicted in Figure 3. With this folding-over, a one-time processing allows us to create several versions of RAMBO with varying sizes and FP rates (Table 4 and Figure 4).

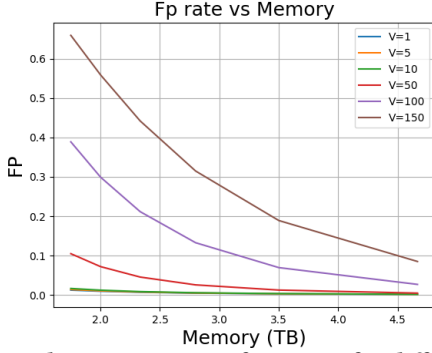


Figure 4: False positive rate of RAMBO for different values of V (k-mer multiplicity per 4605 sets) and memory. Note that the false positive rates are very low if query is rare. For a full sequence search, the returned result depends solely on the rarest k-mer. Hence our method returns very accurate (low false positives) results.

	Wiki-dump (17K)			ClueWeb (50K)		
	QT(ms)	Size	CT	QT(ms)	Size	CT
HowDe	3.781	6.43GB	101m	1.5	8GB	5h
COBS	0.523	157MB	2.71s	0.56	88M	7.6s
RAMBO	0.074	51 MB	1.75s	0.58	62M	5.3s

Table 5: Performance comparison between RAMBO and baselines on wiki-dump data and part ClueWeb data on false positive rate of 0.01. QT is time per query (CPU time) in ms.

5.4 Document indexing

We extend our experiments for web data where each document is represented as a set of English words.

Datasets: We use a sample from Wiki-dump [1] and the popular TREC Category B ClueWeb09 dataset[2]. The Wiki-dump sample has 17618 documents The ClueWeb09 dataset sample has 50K (non-spam) documents of the English language. Both datasets were pre-processed by removing stop words, keeping only alpha-numeric, and tokenizing as word unigrams. Wiki dump is 207 MB, and Clueweb is 98MB after pre-processing.

Parameters: For Wiki dump, RAMBO has $B = 1000$, $R = 2$, and the size of each Bfu is 200000 bits. ClueWeb09, we choose $B = 5000$, $R = 3$, and size of each Bfu = 200000 bits. Clueweb has shorter files (450 terms per file) than Wiki-dump (650 terms per file).

Baseline: We compare with the COBS and Sequence Bloom Tree as in Section 5.

Evaluation Metric: We created a query set of randomly generated terms other than what is present in the data. We inserted them using an exponentially distributed term multiplicity V , similar to the experiment on genomic data. We perform experiments on the same system as in section 5.2. The query is performed sequentially on a single core and thread for a fair comparison. Refer Table 5.

6 DISCUSSION

RAMBO provides a solid trade-off between false positive rate and query time while retaining all desirable properties of Bloom Filter and the bitsliced data structure. Due to cheap updates, RAMBO takes very little time for index creation (Table 4 and Section 5.3). RAMBO performs updates on the stream and is embarrassingly parallel for both insertion and query. The false positive rate of RAMBO is very low for low term multiplicity (Figure 4). This low false positive rate is guaranteed for full sequence/phrase queries, as the rarest of the terms dominates. Therefore, RAMBO can perform a quick and accurate check of an unknown and rare gene sequence. Furthermore, due to sublinear scaling, RAMBO becomes more efficient in memory at a large scale. This property will allow RAMBO to be used as an efficient search engine for extreme-scale applications.

7 APPENDIX

Lemma 4.6 Proof: We want to find the unique insertions in each B Bloom Filter and sum them up to get the size of a single table in RAMBO. If $B = 1$, the unique insertions will be $\frac{N}{V}$ where $N = \sum_{S \in \mathcal{S}} |S|$ is the total number of insertions. If we partition these documents into B bins, every term from the dataset has varying number of duplicates v where $v \in \{0, 1, 2, \dots, V\}$ in a bin. 0 duplicate implies that the term does not exist in the given bin, and 1 duplicate implies that the term has only one copy. Note that each bin corresponds to a Bfu where the terms/kmers are inserted. The expected number of unique terms going in bin b is given by: $|b| = \mathbb{E} \left[\sum_i^{N_b} \frac{1}{v} \right]$, where there are N_b is the number of insertions in b^{th} bucket and $\frac{1}{v}$ is a random variable, $\frac{1}{v} \in \{1, \frac{1}{2}, \frac{1}{3}, \dots, \frac{1}{V}\}$. By the linearity of expectation, we can state that

$$|b| = \sum_i^{N_b} \mathbb{E} \left[\frac{1}{v} \right] = \sum_i^{N_b} \sum_{v=1}^V \frac{1}{v} \times p_v$$

We can view $\frac{1}{v}$ as a multiplicity reduction factor of a term. Here, P_v is the probability of getting v balls in one bucket and $V - v$ in remaining others. Hence we can write $P_v = \frac{1}{B^{v-1}} \times \left(\frac{B-1}{B} \right)^{V-v}$. This gives the expected size of all the bins in a table-

$$\sum_i^N \sum_{v=1}^V \frac{1}{v} \frac{(B-1)^{V-2v+1}}{B^{V-1}} = \sum_{S \in \mathcal{S}} |S| \sum_{v=1}^V \frac{1}{v} \frac{(B-1)^{V-2v+1}}{B^{V-1}}$$

As $B < K$ and $V > 1$, $\Gamma < 1$ always, where $\Gamma = \sum_{v=1}^V \frac{1}{v} \frac{(B-1)^{V-2v+1}}{B^{V-1}}$. The expected size of RAMBO is given by

$$\Gamma \log K \log(1/p) \sum_{S \in \mathcal{S}} |S|$$

REFERENCES

- [1] [n.d.]. Sample wikipedia corpus . Bitfunnel, <http://bitfunnel.org/wikipedia-as-test-corpus-for-bitfunnel>.
- [2] [n.d.]. The ClueWeb09 Dataset. The Lemur Project, <https://www.lemurproject.org/clueweb09.php/>.
- [3] [n.d.]. The European Bioinformatics Institute (EBI): European Nucleotide Archive (ENA) Resource. The European Bioinformatics Institute (EBI) FTP Site, http://ftp.ebi.ac.uk/pub/software/bigsi/nat_biotech_2018/ctx/.
- [4] Paulo Sérgio Almeida, Carlos Baquero, Nuno Preguiça, and David Hutchison. 2007. Scalable bloom filters. *Inform. Process. Lett.* 101, 6 (2007), 255–261.
- [5] Stephen F Altschul, Warren Gish, Webb Miller, Eugene W Myers, and David J Lipman. 1990. Basic local alignment search tool. *Journal of molecular biology* 215, 3 (1990), 403–410.
- [6] Timo Bingmann, Phelim Bradley, Florian Gauger, and Zamin Iqbal. 2019. COBS: a Compact Bit-Sliced Signature Index. In *SPIRE*.
- [7] Burton H Bloom. 1970. Space/time trade-offs in hash coding with allowable errors. *Commun. ACM* 13, 7 (1970), 422–426.
- [8] Fabiano C Botelho, Rasmus Pagh, and Nivio Ziviani. 2007. Simple and space-efficient minimal perfect hash functions. In *Workshop on Algorithms and Data Structures*. Springer, 139–150.
- [9] Phelim Bradley, Henk C den Bakker, Eduardo PC Rocha, Gil McVean, and Zamin Iqbal. 2019. Ultrafast search of all deposited bacterial and viral genomic data. *Nature biotechnology* 37, 2 (2019), 152.
- [10] Phelim Bradley, Henk C den Bakker, Eduardo PC Rocha, Gil McVean, and Zamin Iqbal. 2019. Ultrafast search of all deposited bacterial and viral genomic data. *Nature biotechnology* 37, 2 (2019), 152.
- [11] Larry Carter, Robert Floyd, John Gill, George Markowsky, and Mark Wegman. 1978. Exact and approximate membership testers. In *Proceedings of the tenth annual ACM symposium on Theory of computing*. ACM, 59–65.
- [12] Rayan Chikhi and Paul Medvedev. 2013. Informed and automated k-mer size selection for genome assembly. *Bioinformatics* 30, 1 (06 2013), 31–37. <https://doi.org/10.1093/bioinformatics/btt310> arXiv:<https://academic.oup.com/bioinformatics/article-pdf/30/1/31/643259/btt310.pdf>
- [13] Ken Christensen, Allen Roginsky, and Miguel Jimeno. 2010. A new analysis of the false positive rate of a Bloom filter. *Inform. Process. Lett.* 110, 21 (2010), 944–949. <https://doi.org/10.1016/j.ipl.2010.07.024>
- [14] Peter JA Cock, Christopher J Fields, Naohisa Goto, Michael L Heuer, and Peter M Rice. 2010. The Sanger FASTQ file format for sequences with quality scores, and the Solexa/Illumina FASTQ variants. *Nucleic acids research* 38, 6 (2010), 1767–1771.
- [15] Saar Cohen and Yossi Matias. 2003. Spectral bloom filters. In *Proceedings of the 2003 ACM SIGMOD international conference on Management of data*. ACM, 241–252.
- [16] Graham Cormode and Shan Muthukrishnan. 2005. An improved data stream summary: the count-min sketch and its applications. *Journal of Algorithms* 55, 1 (2005), 58–75.
- [17] Adina Crainiceanu. 2013. Bloofi: a hierarchical Bloom filter index with applications to distributed data provenance. In *Proceedings of the 2nd International Workshop on Cloud Intelligence*. ACM, 4.
- [18] W Bruce Croft, Donald Metzler, and Trevor Strohman. [n.d.]. *Search engines: Information retrieval in practice*. Vol. 520.
- [19] Robert S Harris and Paul Medvedev. 2019. Improved representation of sequence bloom trees. *Bioinformatics* (08 2019).
- [20] Yuichi Kodama, Martin Shumway, and Rasko Leinonen. 2011. The Sequence Read Archive: explosive growth of sequencing data. *Nucleic acids research* 40, D1 (2011), D54–D56.
- [21] Daniel Lemire. 2012. When is a bitmap faster than an integer list? <https://lemire.me/blog/2012/10/23/when-is-a-bitmap-faster-than-an-integer-list/>
- [22] Michael Mitzenmacher. 2002. Compressed bloom filters. *IEEE/ACM Transactions on Networking (TON)* 10, 5 (2002), 604–612.
- [23] Brian D Ondov, Todd J Treangen, Páll Melsted, Adam B Mallonee, Nicholas H Bergman, Sergey Koren, and Adam M Phillippy. 2016. Mash: fast genome and metagenome distance estimation using MinHash. *Genome biology* 17, 1 (2016), 132.
- [24] Prashant Pandey, Fatemeh Almodaresi, Michael A Bender, Michael Ferdman, Rob Johnson, and Rob Patro. 2018. Mantis: A fast, small, and exact large-scale sequence-search index. *Cell systems* 7, 2 (2018), 201–207.
- [25] Rajeev Raman, Venkatesh Raman, and Srinivasa Rao Satti. 2007. Succinct indexable dictionaries with applications to encoding k-ary trees, prefix sums and multisets. *ACM Transactions on Algorithms (TALG)* 3, 4 (2007), 43–es.
- [26] Michael C Schatz and Ben Langmead. 2013. The DNA data deluge: fast, efficient genome sequencing machines are spewing out more data than geneticists can analyze. *IEEE Spectrum* 50, 7 (2013), 26.
- [27] Evan S Snitkin, Adrian M Zelazny, Pamela J Thomas, Frida Stock, David K Henderson, Tara N Palmore, Julia A Segre, NISC Comparative Sequencing Program, et al. 2012. Tracking a hospital outbreak of carbapenem-resistant *Klebsiella pneumoniae* with whole-genome sequencing. *Science translational medicine* 4, 148 (2012), 148ra116–148ra116.
- [28] Brad Solomon and Carl Kingsford. 2016. Fast search of thousands of short-read sequencing experiments. *Nature biotechnology* 34, 3 (2016), 300.
- [29] Brad Solomon and Carl Kingsford. 2017. Improved search of large transcriptomic sequencing databases using split sequence bloom trees. In *International Conference on Research in Computational Molecular Biology*. Springer, 257–271.
- [30] Eric L Stevens, Ruth Timme, Eric W Brown, Marc W Allard, Errol Strain, Kelly Bunning, and Steven Musser. 2017. The public health impact of a publically available, environmental database of microbial genomes. *Frontiers in microbiology* 8 (2017), 808.
- [31] Chen Sun, Robert S Harris, Rayan Chikhi, and Paul Medvedev. 2018. Allsome sequence bloom trees. *Journal of Computational Biology* 25, 5 (2018), 467–479.
- [32] Isaac Turner, Kiran V Garimella, Zamin Iqbal, and Gil McVean. 2018. Integrating long-range connectivity information into de Bruijn graphs. *Bioinformatics* 34, 15 (2018), 2556–2565.