



Assembler artifacts include misassembly because of unsafe unitigs and under-assembly because of bidirected graphs

Amatur Rahman and Paul Medvedev

Genome Res. published online July 27, 2022

Access the most recent version at doi:[10.1101/gr.276601.122](https://doi.org/10.1101/gr.276601.122)

P<P Published online July 27, 2022 in advance of the print journal.

Accepted Manuscript Peer-reviewed and accepted for publication but not copyedited or typeset; accepted manuscript is likely to differ from the final, published version.

Creative Commons License This article is distributed exclusively by Cold Spring Harbor Laboratory Press for the first six months after the full-issue publication date (see <https://genome.cshlp.org/site/misc/terms.xhtml>). After six months, it is available under a Creative Commons License (Attribution-NonCommercial 4.0 International), as described at <http://creativecommons.org/licenses/by-nc/4.0/>.

Email Alerting Service Receive free email alerts when new articles cite this article - sign up in the box at the top right corner of the article or [click here](#).



To subscribe to *Genome Research* go to:
<https://genome.cshlp.org/subscriptions>

Published by Cold Spring Harbor Laboratory Press

Assembler artifacts include mis-assembly because of unsafe unitigs and under-assembly because of bidirected graphs

Amatur Rahman¹Paul Medvedev^{1,2,3,†}¹ Department of Computer Science and Engineering, The Pennsylvania State University² Department of Biochemistry and Molecular Biology, The Pennsylvania State University³ Huck Institutes of the Life Sciences, The Pennsylvania State University† Corresponding author, pzm11@psu.edu

Abstract

Recent assemblies by the T2T and VGP consortia have achieved significant accuracy but required a tremendous amount of effort and resources. More typical assembly efforts, on the other hand, still suffer both from mis-assemblies (joining sequences that should not be adjacent) and from under-assemblies (not joining sequences that should be adjacent). To better understand the common algorithm-driven causes of these limitations, we investigated the unitig algorithm, which is a core algorithm at the heart of most assemblers. We prove that, contrary to popular belief, even when there are no sequencing errors, unitigs are not always safe (i.e. they are not guaranteed to be substrings of the sequenced genome). We also prove that the unitigs of a bidirected de Bruijn graph are different from those of a doubled de Bruijn graph and, contrary to our expectations, result in under-assembly. Using experimental simulations, we then confirm that these two artifacts exist not only in theory but also in the output of widely used assemblers. In particular, when coverage is low then even error-free data results in unsafe unitigs; also, unitigs may unnecessarily split palindromes in half if special care is not taken. To the best of our knowledge, this paper is the first to theoretically predict the existence of these assembler artifacts and confirm and measure the extent of their occurrence in practice.

Introduction

Reconstructing the full sequence of a genome from its sequencing data remains one of the most challenging problems in bioinformatics. Assemblers have suffered both from mis-assemblies (putting together sequences that should not be adjacent) and under-assemblies (not putting together sequences whose adjacency should be apparent from the data) (Alkan et al 2011; Simpson and Pop 2015). Recent efforts by the Telomere-to-Telomere consortium (Miga et al 2020; Nurk, Koren, et al 2022) and the Vertebrate Genome Project (Rhie et al 2021) demonstrated how long read technologies, long-range contact mapping, and manual curation can alleviate these errors. However, the time and cost of those efforts remain prohibitive for most biology labs. In such cases, mis- and under-assemblies continue to be a major limitation (e.g. (Yang et al 2021)).

Understanding the common algorithm-driven causes of these limitations is made complicated by the diversity and complexity of assembly algorithms. We can start by focusing on assemblers that use de Bruijn graphs (DBGs) (Idury and Waterman 1995), which continue to be popular even for long-read data (Bankevich, Bzikadze, et al 2022). But even DBG-based assemblers differ on how they handle complexities arising from sequencing errors, heterogeneity, or DNA double strandedness. Nevertheless, most assemblers are built on top of the unitig algorithm, which returns all the maximal unitigs in an assembly graph (Simpson and Pop 2015); a unitig is a path whose vertices have exactly one incoming and one outgoing edge, with the exception that the first and last vertex can have any number of incoming and outgoing edges, respectively. Being a common denominator of most assemblers, the unitig algorithm is a good target for investigating shared sources of mis- and under-assemblies.

It is already known that the unitig algorithm contributes to under-assembly (e.g. see the safe and complete framework of (Cairo et al 2020; Tomescu and Medvedev 2016)) and can trivially create mis-assemblies when there are sequencing errors. The effect of sequencing errors on assembly errors has even been theoretically studied more broadly in (Shomorony, Courtade, et al 2015; Shomorony, Courtade, et al 2016; Shomorony, Kim, et al 2016). However, it is widely assumed that if it were not for sequencing errors, unitigs would always be safe (i.e. substrings of the sequenced genome). In an earlier work (Medvedev 2019), we attempted to formally prove this but could only do so by assuming perfect coverage. This assumption was also necessary in another earlier work (Tomescu and Medvedev 2016), where it was suggested that without it, unitigs may not be safe. Unitigs were also implied to be unsafe in certain models of the assembly problem (Cairo et al 2020). We therefore hypothesize that, contrary to popular belief, there are non-contrived conditions which lead to unsafe unitigs on error-free data.

The unitig algorithm also needs to account for the fact that the strand from which a read is sequenced is unknown. Most assemblers do so via two common approaches to constructing the DBG. In one, every k -mer is “doubled” prior to constructing the DBG, i.e. for every k -mer in the input, both it and its reverse complement is added to the DBG (e.g. SPAdes (Bankevich, Nurk, et al 2012)). In the other approach, edges are given two instead of one orientation, thereby capturing the way that double-stranded strings can overlap. This results in a bidirected DBG (Medvedev, Georgiou, et al 2007), used in assemblers such as ABySS (Jackman et al 2017; Simpson, Wong, et al 2009)). Since this is a more elegant construction for capturing the double-stranded nature of the data, one would hypothesize that it should not hurt assembly accuracy. In this paper, we will perform a theoretical and empirical study to validate our two hypothesis about common algorithm-driven sources of mis- and under-assemblies.

Results

Summary of findings: First, despite widespread belief, we show that even on error-free data, unitigs do not always appear in the sequenced genome (i.e. they are unsafe). Our experimental results confirm that at least two different assemblers exhibit this behavior in practice. Second, we establish that there is a bijection between maximal unitigs in the doubled and bidirected DBGs, except that palindromic unitigs in the doubled DBG are split in half in the bidirected DBG. This shows that, contrary to intuition, naively using the bidirected graph actually contributes to under-assembly compared to the doubled graph. Our experimental results confirm that this artifact appears in some assemblers but not in others. Nevertheless, we also find that the extent of these two artifacts is limited. To the best of our knowledge, this paper is the first to theoretically predict the existence of these assembler artifacts and confirm and measure the extent of their occurrence in practice.

Theoretical results: The main theoretical results in this paper are two theorems, whose precise statement is given in the Methods. Theorem 1 gives a characterization of unitigs that are unsafe, i.e. they are not present in the sequenced parts of the genome. Theorem 2 breaks down the categories of maximal unitigs in the doubled DBG and the bidirected DBG, and gives a relationship between the two. For the doubled graph, the set of maximal unitigs is partitioned into non-palindromic strings ($D_{\text{non-pal}}$) and palindromic strings (D_{pal}). For the bidirected graph, the maximal unitigs is partitioned into three sets: $B_{\text{no-loop}}$, $B_{\text{first-loop}}$, and $B_{\text{last-loop}}$. Theorem 2 states that there is a one-to-one correspondence between $D_{\text{non-pal}}$ and $B_{\text{no-loop}}$. However, each $D_{\text{non-pal}}$ unitig is split in half in the bidirected graph, with one half appear in $B_{\text{first-loop}}$ and the other half appearing in $B_{\text{last-loop}}$.

Occurrence of unsafe unitigs in real genomes: Theorem 1 predicts the possibility of unsafe unitigs. To verify the extent to which this happens with real genomes, we use T2T human reference Chromosome 1 (Nurk, Koren, et al 2022). We simulated error-free reads of length 100 with varying target coverages and varying k . Note that for this experiment, we want to test if mis-assemblies occur even when the data is perfect, so making the reads error-free is necessary. The sequenced read intervals correspond to the source location of each simulated read, and the sequenced segments are defined as in previous section. From these reads, we constructed the basic de Bruijn graph and output its maximal unitigs, using a version of BCALM (Chikhi, Limasset, Jackman, et al 2014; Chikhi, Limasset, and Medvedev 2016) modified to ignore reverse-complementary. We confirmed that the unitigs that were unsafe (i.e. not a substring of the sequenced segments) were exactly the unitigs that satisfied the conditions of Theorem 1.

Table 1 shows the number of unsafe unitigs, as a function of the coverage and of k . There are as many as 17,635 unsafe unitigs (at coverage $2\times$ and $k = 71$). The best indicator for the number of unsafe unitigs is the percent of k -mers sampled (or the number of sequenced segments), i.e. the number of unsafe unitigs goes down as the percent of sampled k -mers goes up. This trend is in line with the prediction of Corollary 1, which states that once the coverage is perfect, we expect to see at most one unsafe unitig. Our results indicate that the artifacts identified by Theorem 1 do occur in real genomes, though they become less common as more of the genomic k -mers are sampled.

An unsafe unitig is not necessarily a mis-assembly, as it may be a substring of the unsequenced genome by luck. We define an unitig to be *mis-assembled* if its spelling is not a substring of the reference. Table 1 shows that the number of mis-assembled unitigs is substantially lower than the unsafe unitigs, e.g. with 708 mis-assembled unitigs at $2\times$ coverage and $k = 71$. Thus the potential

Table 1: The presence of unsafe and mis-assembled unitigs in human Chromosome 1, using simulated error-free reads.

coverage	k	% k -mers sampled	n. sequenced segments	n. unitigs	n. unsafe	n. mis-assembled	n. Figure 3 cases
1×	71	26.47	1,838,685	1,747,456	12,396	449	383
2×		45.94	2,710,240	2,582,737	17,635	708	628
10×		95.82	1,051,772	1,243,975	2,758	36	36
20×		99.88	61,358	373,335	32	1	0
2×	21	80.76	987,257	4,545,450	2,924	260	224
	31	76.25	1,208,314	2,823,743	4,489	379	352
	41	70.78	1,478,524	2,251,762	6,460	480	447
	51	64.09	1,808,896	2,093,319	9,115	578	532
	61	55.93	2,213,535	2,230,578	12,838	709	645
	71	45.94	2,710,240	2,582,737	17,635	708	628

Table 2: The presence of unsafe and mis-assembled unitigs in the CAMI dataset, using simulated error-free reads.

k	% k -mers sampled	n. sequenced segments	n. unitigs	n. unsafe	n. mis-assembled
55	65	171,682	174,954	593	37
75	61	207,429	207,402	656	33

for mis-assembly does not usually translate into a real mis-assembly, though many mis-assemblies remain.

We further check how many of these mis-assembled unitigs fit the example in Fig. 3. A formal definition to capture this example is included in Appendix A for reference. Table 1 shows that the vast majority of mis-assembled cases are in fact caused by this situation, where a repeat has an occurrence in which its start is unsequenced and another occurrence in which its end is unsequenced.

The simulations in Table 1 suggest that the mis-assembly artifact can be removed by simply increasing coverage. In a metagenome experiment, however, this is not always possible. Even when one increases the number of reads, there will continue to be genomes in the sample whose abundance is low enough that their coverage is low. To verify this intuition, we used a standard benchmark dataset generated by the CAMI competition (Sczyrba et al 2017), containing 70 million synthetic reads from 30 genomes. Table 2 shows there are 33-37 mis-assembled unitigs, indicating that this artifact remains under realistic coverage of a metagenomic dataset. The section “CAMI dataset” in Appendix C contains more details about the experiment, including Table S1 which shows the details of the dataset.

Presence of unsafe unitigs in the contig output of real assemblers: We investigated the extent to which the artifact predicted by Theorem 1 appears in output of real assemblers. Assemblers do not simply output the unitigs of a graph but perform many other steps, hence it was not clear if this artifact would appear in the output contigs. It is not clear how to verify this artifact with real data, as sequencing errors make it difficult to know which of the mis-assembled contigs are caused by the conditions of Theorem 1. We therefore again used a simulated error-free dataset from the T2T Chromosome 1, using the ART simulator (Huang et al 2012), with read length of 250 and varying coverages. This time, we simulated reads from either strand, since assemblers are not typically run in single-stranded mode. We also used the CAMI dataset, but simulating reads in double-stranded mode. We then constructed the doubled de Bruijn graph using $k = 74$ and output its maximal unitigs (note that Theorem 1 holds for even k). We also ran SPAdes (Bankevich, Nurk, et al 2012) and MEGAHIT (Li et al 2015) to assemble the reads (see Appendix C for parameter

Table 3: The extent to which mis-assembled unitigs contribute to mis-assembled contigs of real assemblers. Here, U is the set of mis-assembled unitigs in G_{dbl} , S is the set of mis-assembled contigs of SPAdes, and M is the set of mis-assembled contigs of MEGAHIT. We use $A \sqsubset_t B$ to indicate the subset of A that matches at least one element of B at a threshold of t .

	G_{dbl}	SPAdes			MEGAHIT		
coverage	$ U $	$ S $	$ S \sqsubset_1 U $	$ U \sqsubset_1 S $	$ M $	$ M \sqsubset_{0.5} U $	$ U \sqsubset_{0.5} M $
$1\times$	234	3,366	233	209	3,070	111	179
$2\times$	129	4,423	128	87	2,677	75	119
$3\times$	44	8,329	44	40	1832	21	39
$4\times$	13	7,365	13	13	1240	5	13
$5\times$	5	6,526	5	5	986	0	5
$6\times$	1	5,795	1	1	832	0	1

Table 4: The extent to which mis-assembled unitigs contribute to mis-assembled contigs of real assemblers for the CAMI metagenomic dataset.

	G_{dbl}	metaSPAdes			MEGAHIT		
k	$ U $	$ S $	$ S \sqsubset_1 U $	$ U \sqsubset_1 S $	$ M $	$ M \sqsubset_{0.5} U $	$ U \sqsubset_{0.5} M $
55	37	384	36	36	255	34	32
75	33	208	30	30	144	30	26

details). We then identified unitigs and the assembler contigs that were mis-assembled, but allowing for reverse complements. We will say that a string x *matches* a string y with a threshold of t if a fraction t of the k -mers of x occur in y .

Tables 3 and 4 show that nearly all of the mis-assembled unitigs matched at least one mis-assembled SPAdes contig with a threshold of 1. For MEGAHIT, the threshold of 1 turned out to be stringent; this is not surprising, since assemblers have many steps that may add or remove k -mers from the graph; additionally, MEGAHIT varies the value of k internally and may therefore join k -mers that do not have an overlap of length $k-1$. Using a threshold of 0.5, however, we found that, similarly to SPAdes, most mis-assembled unitigs matched a mis-assembled contig of MEGAHIT. These results indicate that the artifact predicted by Theorem 1 not only appears in unitigs of the raw graph but also in the output of widely used assemblers like SPAdes and MEGAHIT.

Presence of palindrome splitting in a real genome: To measure the extent of the “palindrome splitting” artifact predicted by Theorem 2, we let K be the set of all constituent k -mers in human Chromosome 21 (GRCh38.p13), after excising the Ns. We confirmed the correctness of Theorem 2 by verifying that the spellings of $D_{\text{non-pal}}$ are equal to the spellings of $B_{\text{no-loop}}$ and that the spellings of $B_{\text{last-loop}}$ are equal to the spellings of D_{pal} and are the reverse complements of the spellings of $B_{\text{first-loop}}$. Table 5 shows that the splitting artifact is present but rare, e.g. for $k = 15$, there were 186 palindromic maximal unitigs in $G_{\text{dbl}}(K)$ which were split in $G_{\text{bid}}(K)$. The artifact becomes rarer with increasing k (e.g. for $k = 43$, there were only 3 split palindromes), which is expected since palindrome frequency in real genomes decreases with length.

Presence of palindrome splitting in real assemblers: Most assembler papers do not contain enough detail to ascertain what kind of de Bruijn graph they use to handle reverse complements nor what modifications, if any, they make to the unitig algorithm used for the final output. Looking at MEGAHIT (Li et al 2015), SPAdes (Bankevich, Nurk, et al 2012), ABySS (Jackman et al 2017; Simpson, Wong, et al 2009), and minia (Chikhi and Rizk 2013), only the SPAdes paper is unambiguously clear in saying how it handled reverse complements (it used the doubled DBG). Fur-

Table 5: Extent of the palindrome splitting artifact predicted by Theorem 2 in Chr21.

k	$ D $	$ B $	$ D_{\text{non-pal}} $	$ D_{\text{pal}} $	$ B_{\text{last-loop}} $	$ B_{\text{first-loop}} $	$ B_{\text{no-loop}} $
15	1,465,800	1,465,986	1,465,614	186	186	186	1,465,614
29	60,849	60,866	60,832	17	17	17	60,832
35	36,542	36,552	36,532	10	10	10	36,532
43	18,459	18,462	18,456	3	3	3	18,456

Table 6: Presence of the palindrome splitting artifact in real assemblers on a synthetic genome. We used $k = 31$ for all the assemblers (see Appendix C for details). We filtered out unitigs shorter than 500 bp, amounting to 440 palindromic strings in ABySS and minia and 433 palindromic strings in SPAdes and MEGAHIT.

Contigs		Unitigs in D_{pal}		
Assembler	n. contig	n. fully-covered	n. split	n. ambiguous
MEGAHIT	4,882	427	0	13
SPAdes	7,209	423	0	17
ABySS	53,046	0	66	367
minia	23,318	383	34	16

thermore, since these assemblers implement many heuristics, the splitting artifact may be absent (respectively, present) even if they did (respectively, did not) use bidirected graphs. We therefore tested the behavior of these assemblers by looking for evidence of palindrome splitting in their output, rather than in their technical descriptions.

Since large exact palindromes are uncommon in typical genomes, we created a synthetic genome by modifying a ~ 7 mil bp long contig from human Chromosome 4 (GRCh38.p13) as follows. We randomly sampled a 1,000bp-long region and replaced the last 500bp by the reverse complement of first 500 bp; we then repeated the sampling process 700,000 times. We then simulated error-free Illumina reads with ART. We used a read length of 100bp so that assemblers will not be able to supplement the DBG with read information in a way that hides the palindrome splitting artifact. We used $10\times$ coverage so that most k -mers would be sampled.

First, we find the reference location of each unitig w in D_{pal} . Then, we find all exact alignments of the assembler contigs to the reference. We say that w is *fully-covered* if there exists a contig whose alignment spans w 's. Otherwise, we say w is *split* if one half of w 's region does not overlap with any contig alignments while the other half has a contig aligned that ends precisely in the middle of w at one end and extends past w at the other end. A unitig is *ambiguous* if it does fall into either category. Appendix C contains a more precise definition of these cases.

Table 6 shows that ABySS clearly exhibits the palindrome splitting artifact, with all non-ambiguous unitigs being split and none fully-covered. In fact, this is due to a heuristic that breaks unitigs at any palindromic edges or vertices [see relevant code in <https://github.com/bcgsc/abyss/blob/master/Common/Kmer.cpp>]. The opposite was true for SPAdes and MEGAHIT, with all non-ambiguous unitigs being fully-covered and none split. minia on the other hand exhibited mixed behavior. Of the 417 non-ambiguous cases, 34 were split and 383 were fully-covered. These results indicate that the palindrome splitting artifact of Theorem 2 does persist all the way to the contig output stage in some assemblers. However, this artifact requires the presence of long exact palindromes in the reference, which is uncommon in most genomes.

Discussion

Our theoretical study uncovered two artifacts of the unitig algorithm for genome assembly. The first is that even without sequencing errors, it can create mis-assemblies in places of imperfect coverage. The second is that when the bidirected graph is used to model double-strandedness, the unitig algorithm under-assembles by failing to merge the two halves of palindromes. Our experiments confirmed the presence of these theoretically-predicted artifacts in real genomes and popular assemblers. Fortunately, the impact of these artifacts is not large and can be addressed. Mis-assembly issues due to the first artifact can be resolved by increasing coverage or, potentially, breaking unitigs at places where the coverage along them is uneven. Under-assembly issues due to the palindrome artifact are rare in real genomes and, moreover, can sometimes be fixed by forcing the unitigs to “push their way through” lonely inverted loops (however, it is not always possible, e.g. (Bushmanova et al 2019; Meleshko et al 2022)).

One of the tangential outcomes of this paper is that we have given proper definitions for things like walks and unitigs in the context of bidirected graphs. Previous papers used these concepts somewhat informally; when definitions were given, they worked in the context of that paper but failed to have more general desired properties. For example, our previous work had an inconsistency in the way that a walk was defined on a single vertex versus on many vertices (Rahman, Chikhi, et al 2021). One key takeaway is that as a rule thumb, when working with bidirected graphs one should avoid thinking in terms of vertices but think instead of vertex-sides. The definitions we have provided in this paper generalize further than previous ones and are able to form the basis for the type of analysis we have done in this paper. For example, we are the first to prove the bijection between walks in the doubled and bidirected DBGs. We hope that these definitions will facilitate future attempts to formally study questions in bidirected graphs.

Bidirected graphs give an elegant way to capture the double-stranded nature of DNA in a DBG, but our results here indicate that, for the unitig algorithm, they do not give any theoretical advantage. One of the claimed advantages of using the bidirected graph framework in assembly is that it allows one to take advantage of results from graph theory that may otherwise be hidden. The primary example of this is a result (involving one of the authors) in (Medvedev, Georgiou, et al 2007) where a variant of the assembly problem was theoretically solved in polynomial time by relying on a reduction to the flow problem in bidirected graphs (Gabow 1983). When viewed in retrospect, however, it is not clear that this connection was necessary. The algorithm being reduced to (Gabow 1983) was too cumbersome to implement and, when the assembly problem later necessitated a software solution, an approximation algorithm was used instead (Medvedev and Brudno 2008; Medvedev, Fiume, et al 2010). But the approximation algorithm worked on the doubled graph, erasing the advantage of having initially formulated the problem on bidirected graphs. Therefore, it remains to be seen if there are situations where the connection to graph theoretical results on bidirected graphs can prove useful for genome assembly. Alternatively, using a different setting may better help identify the advantages of bidirected graphs, e.g. pangenomics (Paten et al 2017), rearrangement analysis (Bergeron et al 2006), or compression (Rahman and Medvedev 2021). Quantifying these advantages would be an exciting future direction.

Methods

Preliminaries

In this section we give the formal definitions for our paper. The reader may wish to delay reading the last three paragraphs (relating to bidirected graphs) until they are used later in the text.

Strings: In this paper, we assume all strings are over the four-letter DNA alphabet. A string of length k is called a k -mer. We define $\text{suf}_k(x)$ (respectively, $\text{pre}_k(x)$) to be the last (respectively, first) k characters of x . When the subscript is omitted from pre , and suf , we assume it is $k - 1$. For x and y with $\text{suf}(x) = \text{pre}(y)$, we define *gluing* x and y , denoted by $x \odot y$, as x concatenated with the last $|y| - k + 1$ characters of y . Given two strings x and y , we define $\text{occ}_y(x)$ as the number of times that x occurs in y . The reverse complement of x is denoted as $\bar{x}$. For a set of strings $\mathcal{S}$, $\bar{\mathcal{S}}$ denotes the set of the reverse complements of all strings of $\mathcal{S}$. A string x is a *palindrome* iff $x = \bar{x}$. A string x is *canonical* if it is the lexicographically smaller of x and $\bar{x}$. For $s \in \{0, 1\}$, we define $\text{orient}(x, s)$ to be x if $s = 0$ and to be $\bar{x}$ if $s = 1$. To *canonicalize* x is to replace it by its canonical version, $\text{canon}(x) = \min_i(\text{orient}(x, i))$. We say that x_0 and x_1 have a (s_0, s_1) -oriented-overlap if $\text{suf}(\text{orient}(x_0, 1 - s_0)) = \text{pre}(\text{orient}(x_1, s_1))$. Informally, such an overlap exists between two strings if we can orient them in such a way that they are glueable. For example, GTT and TTG has a $(1, 0)$ -oriented overlap, and AAC and TTG have a $(0, 0)$ -oriented overlap. We define the *non-canonical k -spectrum* $\text{sp}^k(x)$ as the set of all k -mer substrings of x .

Directed de Bruijn graphs: Given a set of k -mers K , the *basic node-centric directed de Bruijn Graph*, $G_{\text{basic}}(K)$, is directed graph where nodes are the k -mers of K , and an edge exists from k -mer x to k -mer y iff $\text{suf}(x) = \text{pre}(y)$. A *double directed de Bruijn graph* on K , $G_{\text{dbl}}(K)$ is a basic de Bruijn graph on the set of k -mers $K \cup \bar{K}$, i.e. $G_{\text{dbl}}(K) = G_{\text{basic}}(K \cup \bar{K})$. Observe that for any k -mer x such that $\text{suf}(x) \neq \text{pre}(\bar{x})$, the existence of the edge from x to y in $G_{\text{dbl}}(K)$ implies the existence of a different edge from $\bar{y}$ to $\bar{x}$. We refer to such a pair of edges as *mirrors*. For a k -mer x such that $\text{suf}(x) = \text{pre}(\bar{x})$, the $G_{\text{dbl}}(K)$ will contain an edge from x to $\bar{x}$; we call this edge a *self-mirror*.

Walks and unitigs in directed graphs For a vertex x in a directed graph, the *in-degree* $d^-(x)$ (respectively, *out-degree* $d^+(x)$) is the number of edges incoming to (respectively, outgoing from) it. The sequence of vertices $w = (x_0, \dots, x_n)$, for $n \geq 0$, is a *walk* iff for all $1 \leq i \leq n$ there exists an edge from x_{i-1} to x_i . Vertices x_0 and x_n are called *endpoints*, and a walk sometimes has one endpoint. The *spelling* of a walk is defined as $\text{spell}(w) = x_0 \odot \dots \odot x_n$. A walk is said to be *circular* iff $n \geq 1$ and $x_0 = x_n$, and a simple *cycle* if for all i and j such that $0 \leq i < j \leq n$, $x_i = x_j$ implies $i = 0$ and $j = n$. A simple periodic cycle is a walk that starts with a simple cycle and then keeps on looping around it without ever exiting; formally, a walk is a *simple periodic cycle* if there exists $0 \leq i \leq n - 1$ such that $(x_0, \dots, x_i)$ is a simple cycle and $x_{i+1}, \dots, x_n$ is a repetition of $x_0, \dots, x_i$, except the last repetition may be partial. A walk is a *unitig* if it is not a periodic cycle and for all $1 \leq i \leq n$, $d^-(x_i) = 1$ and for all $0 \leq i \leq n - 1$, $d^+(x_i) = 1$. A unitig is *maximal* if it is not a proper subwalk of another unitig.

Bidirected de Bruijn graph: A *bidirected graph* G is a pair (V, E) where the set V are called vertices and E is a set of edges. Informally, every vertex has two sides; formally, a *vertex-side* is a pair (u, s) , where $u \in V$ and $s \in \{0, 1\}$. An edge e is a set of two vertex-sides $\{(u_0, s_0), (u_1, s_1)\}$, where $u_i \in V$ and $s_i \in \{0, 1\}$, for $i \in \{0, 1\}$. Informally, an edge is an undirected connection between two (not-necessarily distinct) vertex-sides. We say that an edge e is incident to each of the two vertex-sides. Note that there can be multiple edges between two vertices, but only one edge once the sides are fixed. A *labeled bidirected graph* is a bidirected graph G where every vertex u has a string label $\text{lab}(u)$, and for every edge $e = \{(u_0, s_0), (u_1, s_1)\}$, there is a (s_0, s_1) -oriented-overlap between $\text{lab}(u_0)$ and $\text{lab}(u_1)$. G is said to be *overlap-closed* if there is an edge for every such overlap. Let K be a set of canonical k -mers. The node-centric *bidirected de Bruijn graph*, denoted by $G_{\text{bid}}(K)$, is the overlap-closed labeled bidirected graph where the vertices and their labels correspond to K . Figure S1A shows an example of a bidirected graph.

Walks and unitigs in bidirected graphs: An edge in a bidirected graph is an *inverted loop* if its two vertex-sides are equal. An inverted loop $\{(u, s), (u, s)\}$ is *lonely* if it is the only edge incident to (u, s) . We define the *degree* of a vertex-side $d(u, s)$ to be the number of edges incident to it, but with an inverted loop contributing two to the degree. A sequence $t = (u_0, s_0, u_1, s_1, \dots, u_n, s_n)$ with $n \geq 0$ is a *walk* if for all $1 \leq i \leq n$, there exists an edge $e_i = \{(u_{i-1}, 1 - s_{i-1}), (u_i, s_i)\}$. The vertex-sides (u_0, s_0) and $(u_n, 1 - s_n)$ are called the first and last *endpoint sides*, respectively. Note that even when $n = 0$, there are two endpoint sides. The *spelling* of a walk is defined as $spell(w) = orient(lab(u_0), s_0) \odot \dots \odot orient(lab(u_n), s_n)$. The reverse of t is $rev(t) = (u_n, 1 - s_n, \dots, u_0, 1 - s_0)$. Note that, as expected, $spell(t) = spell(rev(t))$. Note that if t' is a subwalk of t , then $rev(t')$ is a subwalk of $rev(t)$ and $spell(t')$ is a substring of $spell(t)$ (the converse is not necessarily true when k is even). Figure S1BC gives an example illustrating a walk in a bidirected graph and Figure S1D shows a corresponding walk in a doubled directed DBG.

A walk $w = (u_0, s_0, \dots, u_n, s_n)$ is said to be *circular* iff $n \geq 1$ and $(u_0, s_0) = (u_n, s_n)$, and a *simple cycle* if for all i and j such that $0 \leq i < j \leq n$, $(u_i, s_i) = (u_j, s_j)$ implies $i = 0$ and $j = n$. A *simple periodic cycle* is a walk that starts with a simple cycle and then keeps on looping around it without ever exiting; formally, w is a *simple periodic cycle* if there exists $0 \leq i \leq n - 1$ such that $(u_0, s_0, \dots, u_i, s_i)$ is a simple cycle and $(u_{i+1}, s_{i+1}, \dots, u_n, s_n)$ is a repetition of $(u_0, s_0, \dots, u_i, s_i)$, except the last repetition may be partial. A walk is a *unitig* if it is not a periodic cycle and for all $1 \leq i \leq n$, $d^-(x_i) = 1$ and for all $0 \leq i \leq n - 1$, $d^+(x_i) = 1$. A walk $(u_0, s_0, \dots, u_n, s_n)$ is a *unitig* if it is not a periodic cycle and for all $0 \leq i < n$, $d(u_i, 1 - s_i) = 1$ and, for all $0 < i \leq n$, $d(u_i, s_i) = 1$. A unitig is said to be *maximal* if it is not a proper subwalk of another unitig. Note that all the subwalks of a unitig must themselves be unitigs.

Safety of unitigs

In this subsection, we will give necessary and sufficient conditions for a unitig to be unsafe in the basic DBG constructed from error-free reads. To properly formulate this question, we define a *sequenced read interval* as a genomic interval that generated a read, i.e. from which a read was sequenced. A sequencing experiment then corresponds to a set of sequenced read intervals. A *sequenced interval* is then defined as a maximal interval which is covered by sequenced read intervals, with the additional constraint that any two consecutive sequenced intervals overlap by at least $k - 1$. We define a *sequenced segment* as the string corresponding to a sequenced interval.

Observe that the sequenced intervals do not overlap by more than $k - 2$ bases (otherwise they would not be maximal), but the sequenced segment may have longer overlaps due to repeats. A set of reads then induces a set $\mathcal{S} = \{S_1, \dots, S_{|\mathcal{S}|}\}$ of sequenced segments. Figure 1 illustrates this formulation. In this subsection, we do not explicitly account for reverse complements, since they will be considered in the next subsection.

Given a set of sequenced segments $\mathcal{S}$, we say that a unitig w in $G_{\text{basic}}(\text{sp}^k(\mathcal{S}))$ is *unsafe* iff $spell(w)$ is not a substring of a string in $\mathcal{S}$. Equivalently, w is unsafe iff it is not a subwalk of a walk that corresponds to a string in $\mathcal{S}$. Our definition of unsafe captures the notion of a potential mis-assembly, as the unitig is not present in the sequenced part of the genome.¹ Observe that in formulating the problem, we start with the set of sequenced segments themselves; the read set that induced them is irrelevant. We can now state the main result of this subsection, which gives the

¹The safety of unitigs has been previously studied for other notions of “safety” by (Cairo et al 2020). While the authors did not make the explicit conclusion and did not verify it in practice, their Theorem 6.1(d) implies that unitigs are not guaranteed to be safe in the model of assembly they consider. Concretely, while a suffix or prefix of the unitig may be present at the starts and ends of parts of the genome, the whole unitig might never be contained as a contiguous sequence.

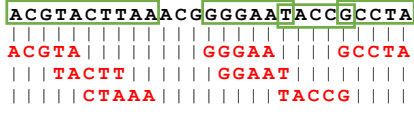


Figure 1: Illustration of sequenced segments. The black text on top shows the reference genome of length 26. The seven sequences in red are reads aligned to the reference. The green boxes highlight the resulting sequenced segments when $k = 3$. Note that the reads TACCG and GCCTA form two separate segments as the k -mer CGC is not present in K .

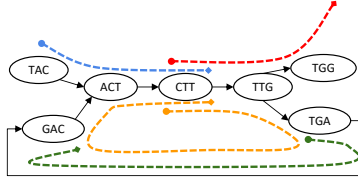


Figure 2: Illustration of the cases in Theorem 1. The graph in the figure represents $G_{\text{basic}}(\text{sp}^k(\mathcal{S}))$, where $k = 3$ and $\mathcal{S} = \{\text{CTTGG}, \text{CTTGACTT}, \text{TACTT}, \text{TGAC}\}$. The segments are marked by dashed lines with their starts marked with a dot and their ends marked with a diamond. The unitig $w = \{\text{ACT}, \text{CTT}, \text{TTG}\}$ is unsafe because for each of the segment, one of the cases in Theorem 1 is true. For segment colored in green, case (i) holds. For red, case (ii), for blue, case (iii) and for orange, case (iv) holds.

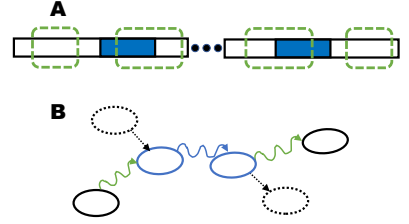


Figure 3: Illustration of an unsafe unitig. Panel A shows two parts of a sequenced genome X . Regions surrounded by green dashed boxes are the sequenced segments $\mathcal{S}$. The solid blue boxes represent two copies of a repeat. Panel B shows the resulting $G_{\text{basic}}(\text{sp}^k(\mathcal{S}))$, with dashed vertices and edges representing vertices that are in $G_{\text{basic}}(\text{sp}^k(X))$ but not sequenced.

necessary and sufficient conditions for a unitig to be unsafe. The proof of this theorem, along with the necessary Lemmas, is left for Appendix A due to space constraints.

Theorem 1. *Let $\mathcal{S}$ be a set of sequenced segments and let $w = (x_0, \dots, x_m)$ be a unitig in $G_{\text{basic}}(\text{sp}^k(\mathcal{S}))$. Then w is unsafe if and only if for all $S \in \mathcal{S}$, one of the following holds:*

- (i) S does not contain any k -mer of w ,
- (ii) $\text{occ}_S(\text{pre}_k(S)) = 1$ and $\text{pre}_k(S) = x_i$ for some $1 \leq i \leq m$,
- (iii) $\text{occ}_S(\text{suf}_k(S)) = 1$ and $\text{suf}_k(S) = x_j$ for some $0 \leq j \leq m-1$, or
- (iv) $\text{occ}_S(\text{pre}_k(S)) = \text{occ}_S(\text{suf}_k(S)) = 2$ and there exists $1 \leq i \leq j \leq m-1$ such that $\text{pre}_k(S) = x_i$ and $\text{suf}_k(S) = x_j$.

The cases of Theorem 1 are illustrated in Figure 2 and can be understood informally as follows. Since every k -mer of $G_{\text{basic}}(\text{sp}^k(\mathcal{S}))$ is in $\mathcal{S}$, every k -mer of w must be touched by some $S \in \mathcal{S}$. Then, consider a walk g corresponding to such a string S . If g starts in the middle of w and does not visit its own starting vertex again, then g does not fully contain w (case (ii)). Similarly, if g ends in the middle of w and did not visit its own ending vertex previously, then g does not fully contain w (case (iii)). If g starts and ends in the middle of w , with the ending vertex to the right of the starting vertex, and contains each of those vertices exactly twice, then g does not fully contain w (case (iv)). This is the “if” direction of Theorem 1, with the “only if” direction further stating that under all other conditions, g fully contains w .

When the genome is a single chromosome and the coverage is high enough so that every k -mer is sequenced, the whole genome becomes one sequenced segment. In this case, Theorem 1 simplifies because the genome has only one starting and ending vertex and, for a unitig w to be unsafe, the genome must somehow contain every vertex of w without containing w as a subwalk.

Corollary 1. *Let X be a string and let $w = (x_0, \dots, x_m)$ be a unitig in $G_{\text{basic}}(\text{sp}^k(X))$. Then $\text{spell}(w)$ is not a substring of X iff one of the following holds:*

1. $occ_X(pre_k(X)) = occ_X(suf_k(X)) = 1$, $pre_k(X) = x_i$, $suf_k(X) = x_{i-1}$ for some $1 \leq i \leq m$.
2. $occ_X(pre_k(X)) = occ_X(suf_k(X)) = 2$, $pre_k(X) = x_i$, $suf_k(X) = x_j$ for some $0 < i \leq j < m$.

Moreover, this can hold for at most one unitig in $G_{basic}(sp^k(X))$.

This corollary tells us that with perfect coverage, all unitigs, except possibly one, are safe. Note that this is a stronger version of the perfect coverage case than the one given in (Medvedev 2019), which made an assumption that the starting vertex of X is a source and the ending vertex of X is a sink.

A natural question is how a scenario which gives an unsafe unitig looks like in terms of the original genome. Figure 3A-B visualized the following natural possibility. Suppose that the sequenced genome X has a repeat that appears as a maximal unitig ψ in $G_{basic}(sp^k(X))$. Then, suppose that the region encompassing the start of one copy and the region encompassing the end of the other copy is not sequenced. Then ψ loses its maximality in $G_{basic}(sp^k(\mathcal{S}))$ and becomes a subwalk of a bigger unitig w . Though w is a unitig in the graph from the sequencing data, it would not be a unitig if all the k -mers of X were included in the graph. In the Results section, we will show that this situation accounts for the majority of our experimental observations.

The relationship between the doubled dBG ($G_{dbl}(K)$) and the bidirected dBG ($G_{bid}(K)$)

In this subsection, we will characterize the relationship between the maximal unitigs of $G_{dbl}(K)$ and the maximal unitigs of $G_{bid}(K)$ (Theorem 2). Due to space constraints, the lemmas and proofs needed to prove Theorem 2 are in Appendix B. Here, we will instead give an informal walk-through to elucidate the relationship between the two graphs. We will incrementally show the relationship between objects in the doubled graph and the bidirected graph — first between vertices and vertex-sides, then between edges, then between walks, and finally between maximal unitigs.

Let K be a set of canonical k -mers, with k odd. We only consider the case of odd k ; when k is even, there may be palindrome k -mers, which create special cases to handle both in the practical assembler implementation and in the theoretical analysis. Since most assemblers anyway restrict k to be odd, we limit ourselves to this case as well.

There is a natural mapping between vertices of $G_{dbl}(K)$ and vertex-sides of $G_{bid}(K)$. For a vertex x in $G_{dbl}(K)$, define $F_V(x) = (u, s)$, where u is a vertex in $G_{bid}(K)$ and $s \in \{0, 1\}$ such that $lab(u) = orient(x, s)$. By the definition of $G_{bid}(K)$, there exists a unique u and unique s that satisfy this condition. The uniqueness of s is guaranteed by the fact that x cannot be a palindrome. Formally, F_V is a bijection between vertices of $G_{dbl}(K)$ and vertex-sides of $G_{bid}(K)$ (Lemma B.10).

There is also a natural mapping between edges in $G_{dbl}(K)$ and $G_{bid}(K)$. Let x_1 and x_2 be two k -mers in $G_{dbl}(K)$ and let $(u_1, s_1) = F_V(x_1)$ and $(u_2, s_2) = F_V(x_2)$. We define the mapping $F_E(x_1, x_2) = \{(u_1, 1 - s_1), (u_2, s_2)\}$ such that (x_1, x_2) is an edge in $G_{dbl}(K)$ if and only if $F_E(x_1, x_2)$ is an edge in $G_{bid}(K)$ (Lemma B.11). Note, however, that F_E is not a bijection, since a pair of mirror edges (x, y) and $(\bar{y}, \bar{x})$ map to the same bidirected edge, i.e. $F_E(x, y) = F_E(\bar{y}, \bar{x})$.

The F_V and F_E mappings allow us to naturally define a mapping from walks in $G_{dbl}(K)$ to walks in $G_{bid}(K)$. Let $w = (x_0, \dots, x_n)$ be a walk in $G_{dbl}(K)$. For each $0 \leq i \leq n$, let $(u_i, s_i) = F_V(x_i)$ and define $F_W(w) \triangleq (u_0, s_0, \dots, u_n, s_n)$. F_W is a spell-preserving bijection between the set of walks in $G_{dbl}(K)$ and the set of walks in $G_{bid}(K)$ (Lemma B.12).

One might hypothesize that F_W is also a bijection between the maximal unitigs of $G_{dbl}(K)$ and the maximal unitigs of $G_{bid}(K)$. It turns out to not be the case, though the following more careful analysis reveals a close relationship. For $G_{dbl}(K)$, let us partition the set of maximal unitigs into

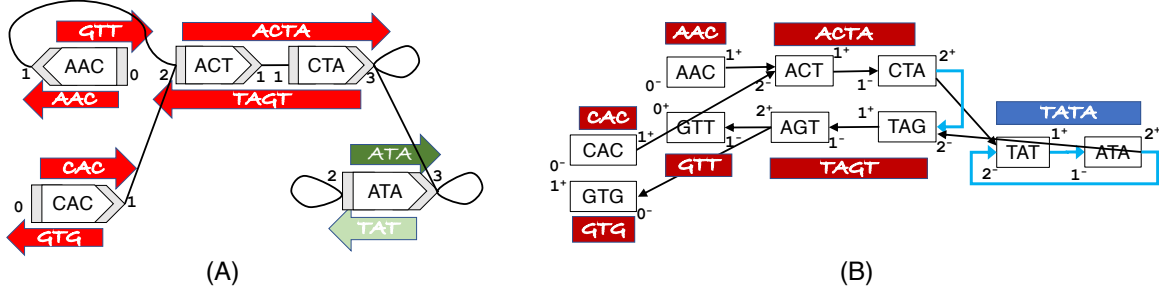


Figure 4: Example of a bidirected dBG ($G_{\text{bid}}(K)$) (panel A) and a doubled dBG ($G_{\text{dbl}}(K)$) (panel B) on the same underlying set of k -mers $K = \{CAC, AAC, ACT, CTA, ATA\}$. Each vertex side in $G_{\text{bid}}(K)$ and each in- and out-side of a vertex in $G_{\text{dbl}}(K)$ is numbered with the corresponding degree. All maximal unitigs are shown using a long filled rectangle with an arrow. The maximal unitigs of $G_{\text{bid}}(K)$ are color coded so that red is $B_{\text{no-loop}}$, dark green is $B_{\text{last-loop}}$, and light green is $B_{\text{first-loop}}$. The maximal unitigs of $G_{\text{dbl}}(K)$ are color coded so that dark red is $D_{\text{non-pal}}$ and blue is D_{pal} . Self-mirror edges in $G_{\text{dbl}}(K)$ are shown in blue.

non-palindromic strings $D_{\text{non-pal}}$ and palindromic strings D_{pal} . For $G_{\text{bid}}(K)$, let $B_{\text{no-loop}}$ be the set of maximal unitigs where neither endpoint side has an incident lonely inverted loop, let $B_{\text{first-loop}}$ be the set of maximal unitigs where the only endpoint side with a lonely inverted loop is the first one, and let $B_{\text{last-loop}}$ be the set of maximal unitigs where the only endpoint side with a lonely inverted loop is the last one. To avoid corner cases, let us further assume that there are no circular unitigs in $G_{\text{dbl}}(K)$, which eliminates the possibility of a maximal unitig having lonely inverted loops at both endpoint sides and implies that $B_{\text{no-loop}}$, $B_{\text{first-loop}}$, and $B_{\text{last-loop}}$ are a partition of the maximal unitigs of $G_{\text{bid}}(K)$ (Lemma B.16). Figure 4A-B shows an example.

We also need to define a function HEAD which, informally, takes a maximal palindromic unitig in $G_{\text{dbl}}(K)$, extracts the first half of it, and maps it to $G_{\text{bid}}(K)$. Formally, $\text{HEAD}(w)$ maps a walk $w = (x_0, \dots, x_n)$ in D_{pal} to the walk $F_W((x_0, \dots, x_{\frac{n-1}{2}}))$ in $G_{\text{bid}}(K)$. Note that $\frac{n-1}{2}$ is necessarily an integer since w is a palindrome and hence n must be odd (Lemma B.1). We can now state the main theorem of this subsection.

Theorem 2. *Let K be a set of canonical k -mers where k is odd and $G_{\text{dbl}}(K)$ does not contain a circular unitig.*

- (i) *The function F_W is a bijection from $D_{\text{non-pal}}$ to $B_{\text{no-loop}}$.*
- (ii) *The function rev is a bijection between $B_{\text{last-loop}}$ and $B_{\text{first-loop}}$.*
- (iii) *HEAD is a bijection from D_{pal} and $B_{\text{last-loop}}$*

Figure 5 schematically illustrates the relationship captured by Theorem 2. The theorem says that for maximal unitigs that are non-palindromic in $G_{\text{dbl}}(K)$ and do not have inverted self loops incident at the endpoint sides in $G_{\text{bid}}(K)$, F_W is in fact a bijection. However, every maximal unitig w that is palindromic in $G_{\text{dbl}}(K)$ is split into two maximal unitigs in $G_{\text{bid}}(K)$: one that spells the first half of w and has a self loop incident at the last endpoint side, and one that spells the second half of w and has a self loop at the first endpoint side. These are necessarily reverses of each other.

Inverted loops are caused by k -mers x where $\text{suf}(x) = \overline{\text{suf}(x)}$ (e.g. GTA). When these type of k -mers are not present in K , there are no inverted loops in $G_{\text{bid}}(K)$ or palindromic unitigs in $G_{\text{dbl}}(K)$. Hence, $D_{\text{pal}} = B_{\text{first-loop}} = B_{\text{last-loop}} = \emptyset$, and Theorem 2 immediately simplifies.

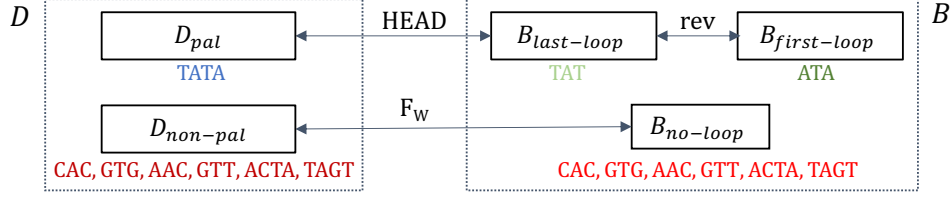


Figure 5: Overview of relationship between maximal units in double and bidirected graph for odd k . We use the example from Figure 4, where $K = \{AAC, ACT, CTA, CAC, ATA\}$. The set of maximal unitigs from $G_{dbl}(K)$, D is partitioned into D_{pal} and $D_{non-pal}$. The set of maximal unitigs from $G_{bid}(K)$, B is partitioned into $B_{last-loop}$, $B_{first-loop}$ and $B_{no-loop}$. The arrows between different subsets of D and B denote bijections.

Corollary 2. *Let K be a set of k -mers, with odd k , which does not contain any x such that $suf(x) = suf(\overline{x})$. Then F_W is a bijection from the maximal unitigs in $G_{dbl}(K)$ to the maximal unitigs in $G_{bid}(K)$.*

Software availability

Scripts for the experimental evaluations are available on GitHub [<https://github.com/medvedevgroup/assembly-artifacts-paper-experiments>].

Competing interest statement

The authors declare that they have no competing interests.

Acknowledgements

Both PM and AR took part in all aspects of the project, except that AR was the one who performed all the experiments. PM thanks Rayan Chikhi, Alexandru Tomescu, and Mihai Pop for useful discussions. This material is based upon work supported by the National Science Foundation under Grant No. 1453527 and 1931531. AR was supported by NIH Computation, Bioinformatics, and Statistics training program.

References

- Alkan, C., Sajjadian, S., & Eichler, E. E. 2011. Limitations of next-generation genome sequence assembly. *Nature methods*, 8. 1., 61–65.
- Bankevich, A., Bzikadze, A. V., Kolmogorov, M., Antipov, D., & Pevzner, P. A. 2022. Multiplex de bruijn graphs enable genome assembly from long, high-fidelity reads. *Nature biotechnology*, 1–7.
- Bankevich, A., Nurk, S., Antipov, D., Gurevich, A. A., Dvorkin, M., Kulikov, A. S., Lesin, V. M., Nikolenko, S. I., Pham, S., Prjibelski, A. D. et al. 2012. Spades: A new genome assembly algorithm and its applications to single-cell sequencing. *Journal of computational biology*, 19. 5., 455–477.
- Bergeron, A., Mixtacki, J., & Stoye, J. 2006. A unifying view of genome rearrangements. *International Workshop on Algorithms in Bioinformatics*, 163–173.
- Bushmanova, E., Antipov, D., Lapidus, A., & Prjibelski, A. D. 2019. Rnaspades: A de novo transcriptome assembler and its application to rna-seq data. *GigaScience*, 8. 9., giz100.
- Cairo, M., Khan, S., Rizzi, R., Schmidt, S., Tomescu, A. I., & Zironde, E. C. 2020. The Hydrostructure: a Universal Framework for Safe and Complete Algorithms for Genome Assembly. *arXiv preprint arXiv:2011.12635*.
- Chikhi, R., Limasset, A., Jackman, S., Simpson, J. T., & Medvedev, P. 2014. On the representation of de Bruijn graphs. *Research in Computational Molecular Biology, RECOMB 2014*, 8394, 35–55.
- Chikhi, R., Limasset, A., & Medvedev, P. 2016. Compacting de Bruijn graphs from sequencing data quickly and in low memory. *Bioinformatics*, 32. 12., i201–i208.
- Chikhi, R., & Rizk, G. 2013. Space-efficient and exact de bruijn graph representation based on a bloom filter. *Algorithms for Molecular Biology*, 8. 1., 1–9.
- Fritz, A., Hofmann, P., Majda, S., Dahms, E., Dröge, J., Fiedler, J., Lesker, T. R., Belmann, P., DeMaere, M. Z., Darling, A. E. et al. 2019. Camisim: Simulating metagenomes and microbial communities. *Microbiome*, 7. 1., 1–12.
- Gabow, H. N. 1983. An efficient reduction technique for degree-constrained subgraph and bidirected network flow problems. *STOC*, 448–456.
- Huang, W., Li, L., Myers, J. R., & Marth, G. T. 2012. Art: A next-generation sequencing read simulator. *Bioinformatics*, 28. 4., 593–594.
- Idury, R. M., & Waterman, M. S. 1995. A new algorithm for dna sequence assembly. *Journal of computational biology*, 2. 2., 291–306.
- Jackman, S. D., Vandervalk, B. P., Mohamadi, H., Chu, J., Yeo, S., Hammond, S. A., Jahesh, G., Khan, H., Coombe, L., Warren, R. L. et al. 2017. Abyss 2.0: Resource-efficient assembly of large genomes using a bloom filter. *Genome research*, 27. 5., 768–777.
- Li, D., Liu, C.-M., Luo, R., Sadakane, K., & Lam, T.-W. 2015. Megahit: An ultra-fast single-node solution for large and complex metagenomics assembly via succinct de bruijn graph. *Bioinformatics*, 31. 10., 1674–1676.
- Medvedev, P. 2019. Modeling biological problems in computer science: A case study in genome assembly. *Briefings in bioinformatics*, 20. 4., 1376–1383.
- Medvedev, P., & Brudno, M. 2008. Ab initio whole genome shotgun assembly with mated short reads. *RECOMB*, 50–64.
- Medvedev, P., Fiume, M., Dzamba, M., Smith, T., & Brudno, M. 2010. Detecting copy number variation with mated short reads. *Genome Research*, 20. 11., 1613–1622.
- Medvedev, P., Georgiou, K., Myers, G., & Brudno, M. 2007. Computability of models for sequence assembly. *WABI*, 289–301.

- Meleshko, D., Hajirasouliha, I., & Korobeynikov, A. 2022. Coronaspades: From biosynthetic gene clusters to rna viral assemblies. *Bioinformatics*, 38. 1., 1–8.
- Miga, K. H., Koren, S., Rhie, A., Vollger, M. R., Gershman, A., Bzikadze, A., Brooks, S., Howe, E., Porubsky, D., Logsdon, G. A. et al. 2020. Telomere-to-telomere assembly of a complete human x chromosome. *Nature*, 585. 7823., 79–84.
- Nurk, S., Koren, S., Rhie, A., Rautiainen, M., Bzikadze, A. V., Mikheenko, A., Vollger, M. R., Altemose, N., Uralsky, L., Gershman, A. et al. 2022. The complete sequence of a human genome. *Science*, 376. 6588., 44–53.
- Nurk, S., Meleshko, D., Korobeynikov, A., & Pevzner, P. A. 2017. Metaspades: A new versatile metagenomic assembler. *Genome research*, 27. 5., 824–834.
- Paten, B., Novak, A. M., Eizenga, J. M., & Garrison, E. 2017. Genome graphs and the evolution of genome inference. *Genome research*, 27. 5., 665–676.
- Rahman, A., Chikhi, R., & Medvedev, P. 2021. Disk compression of k-mer sets. *Algorithms for Molecular Biology*, 16. 1., 1–14.
- Rahman, A., & Medvedev, P. 2021. Representation of k-mer sets using spectrum-preserving string sets. *Journal of Computational Biology*, 28. 4., 381–394.
- Rhie, A., McCarthy, S. A., Fedrigo, O., Damas, J., Formenti, G., Koren, S., Uliano-Silva, M., Chow, W., Fungtammasan, A., Kim, J. et al. 2021. Towards complete and error-free genome assemblies of all vertebrate species. *Nature*, 592. 7856., 737–746.
- Sczyrba, A., Hofmann, P., Belmann, P., Koslicki, D., Janssen, S., Dröge, J., Gregor, I., Majda, S., Fiedler, J., Dahms, E. et al. 2017. Critical assessment of metagenome interpretation—a benchmark of metagenomics software. *Nature methods*, 14. 11., 1063–1071.
- Shomorony, I., Courtade, T., & Tse, D. 2015. Do read errors matter for genome assembly? *2015 IEEE International Symposium on Information Theory (ISIT)*, 919–923.
- Shomorony, I., Courtade, T. A., & Tse, D. 2016. Fundamental limits of genome assembly under an adversarial erasure model. *IEEE Transactions on Molecular, Biological and Multi-Scale Communications*, 2. 2., 199–208.
- Shomorony, I., Kim, S. H., Courtade, T. A., & Tse, D. N. 2016. Information-optimal genome assembly via sparse read-overlap graphs. *Bioinformatics*, 32. 17., i494–i502.
- Simpson, J. T., & Pop, M. 2015. The theory and practice of genome sequence assembly. *Annual review of genomics and human genetics*, 16, 153–172.
- Simpson, J. T., Wong, K., Jackman, S. D., Schein, J. E., Jones, S. J., & Birol, I. 2009. Abyss: A parallel assembler for short read sequence data. *Genome research*, 19. 6., 1117–1123.
- Tomescu, A. I., & Medvedev, P. 2016. Safe and complete contig assembly via omnitigs. *International Conference on Research in Computational Molecular Biology*, 152–163.
- Yang, L., Malhotra, R., Chikhi, R., Elleder, D., Kaiser, T., Rong, J., Medvedev, P., & Poss, M. 2021. Recombination Marks the Evolutionary Dynamics of a Recently Endogenized Retrovirus. *Molecular Biology and Evolution*.