

A Linear-Time Algorithm for Computing Inversion Distance between Signed Permutations with an Experimental Study

DAVID A. BADER,¹ BERNARD M.E. MORET,² and MI YAN¹

ABSTRACT

Hannenhalli and Pevzner gave the first polynomial-time algorithm for computing the inversion distance between two signed permutations, as part of the larger task of determining the shortest sequence of inversions needed to transform one permutation into the other. Their algorithm (restricted to distance calculation) proceeds in two stages: in the first stage, the overlap graph induced by the permutation is decomposed into connected components; then, in the second stage, certain graph structures (hurdles and others) are identified. Berman and Hannenhalli avoided the explicit computation of the overlap graph and gave an $O(n\alpha(n))$ algorithm, based on a Union-Find structure, to find its connected components, where α is the inverse Ackerman function. Since for all practical purposes $\alpha(n)$ is a constant no larger than four, this algorithm has been the fastest practical algorithm to date. In this paper, we present a new linear-time algorithm for computing the connected components, which is more efficient than that of Berman and Hannenhalli in both theory and practice. Our algorithm uses only a stack and is very easy to implement. We give the results of computational experiments over a large range of permutation pairs produced through simulated evolution; our experiments show a speed-up by a factor of 2 to 5 in the computation of the connected components and by a factor of 1.3 to 2 in the overall distance computation.

Key words: signed permutation, sorting by reversals, inversion distance, overlap graph, cycle graph, overlap forest.

1. INTRODUCTION

SOME ORGANISMS HAVE A SINGLE CHROMOSOME or contain single-chromosome organelles (such as mitochondria or chloroplasts), the evolution of which is largely independent on the evolution of the nuclear genome. Given a particular strand from a single chromosome, whether linear or circular, we can infer the ordering and directionality of the genes, thus representing each chromosome by an ordering of oriented genes. In many cases, the evolutionary process that operates on such single-chromosome organisms consists mostly of inversions of portions of the chromosome; this finding has led many biologists to

¹Department of Electrical and Computer Engineering, University of New Mexico, Albuquerque, NM 87131.

²Department of Computer Science, University of New Mexico, Albuquerque, NM 87131.

reconstruct phylogenies based on gene orders, using as a measure of evolutionary distance between two genomes the inversion distance, i.e., the smallest number of inversions needed to transform one signed permutation into the other (Olmstead and Palmer, 1994; Palmer, 1992; Raubeson and Jansen, 1992).

Both inversion distance and the closely related transposition distance are difficult computational problems that have been studied intensively over the last five years (Bafna and Pevzner, 1995; Bafna and Pevzner, 1996; Bafna and Pevzner, 1998; Berman and Hannenhalli, 1996; Caprara, 1997; Caprara, 1999; Hannenhalli and Pevzner, 1995; Kaplan, Shamir, and Tarjan, 1999). Finding the inversion distance between unsigned permutations is NP-hard (Caprara, 1997), but with signed ones, it can be done in polynomial time (Hannenhalli and Pevzner, 1995). The fastest published algorithm for the computation of inversion distance between two signed permutations has been that of Berman and Hannenhalli (1996), which uses a Union-Find data structure and runs in $O(n\alpha(n))$ time, where $\alpha(n)$ is the inverse Ackerman function. (The later KST algorithm (Kaplan, Shamir, and Tarjan, 1999) reduces the time needed to compute the shortest sequence of inversions, but uses the same algorithm for computing the length of that sequence.) We have found only two implementations on the web, both designed to compute the shortest sequence of inversions as well as its length; one, due to Hannenhalli (Hannenhalli, URL), implements his first algorithm (Hannenhalli and Pevzner, 1995), which runs in quadratic time when computing distances, while the other, a Java applet written by Mantin (Mantin and Shamir, 1999), a student of Shamir, implements the KST algorithm (Kaplan, Shamir, and Tarjan, 1999), but uses an explicit representation of the overlap graph and thus also takes quadratic time.

We present a simple and practical, worst-case linear-time algorithm to compute the connected components of the overlap graph, which results in a simple linear-time algorithm for computing the inversion distance between two signed permutations. We also provide ample experimental evidence that our linear-time algorithm is efficient in practice as well as in theory: we coded it, as well as the algorithm of Berman and Hannenhalli, using the best principles of algorithm engineering (McGeoch, 1996; Moret and Shapiro, 2001) to ensure that both implementations would be as efficient as possible and compared their running times on a large range of instances generated through simulated evolution. (The two implementations on the web are naturally far slower).

The paper is organized as follows. We begin by recalling some definitions, briefly review past work on sorting by reversals, then introduce the concepts that we will need in our algorithm including the fundamental theorem that makes it possible. We then describe and analyze our algorithm, discuss our experimental setup, present and comment on our results, and briefly mention an application of our distance computation in a whole-genome phylogeny study.

2. INVERSIONS ON SIGNED PERMUTATIONS

We assume a fixed set of genes $\{g_1, g_2, \dots, g_n\}$. Each genome is then an ordering (circular or linear) of these genes, each gene given with an orientation that is either positive (g_i) or negative ($-g_i$). The ordering $g_1, g_2, \dots, g_n$, whether linear or circular, is considered equivalent to that obtained by considering the complementary strand, i.e., the order $-g_n, -g_{n-1}, \dots, -g_1$.

Let G be the genome with signed ordering (linear or circular) $g_1, g_2, \dots, g_n$. An *inversion* between indices i and j , for $i \leq j$, produces the genome with linear ordering

$$g_1, g_2, \dots, g_{i-1}, -g_j, -g_{j-1}, \dots, -g_i, g_{j+1}, \dots, g_n.$$

If we have $j < i$, we can still apply an inversion to a circular (but not linear) genome by simply rotating the circular ordering until the two indices are in the proper relationship—recall that we consider all rotations of the complete circular ordering of a circular genome as equivalent.

The inversion distance between two genomes (two signed permutations of the same set) is then the minimum number of inversions that must be applied to one genome in order to produce the other. (This measure is easily seen to be a true metric.) Computing the shortest sequence of inversions that gives rise to this distance is also known as *sorting by reversals*—we shall shortly see why it can be regarded as a sorting problem.

3. PREVIOUS WORK

Bafna and Pevzner (1993) introduced the cycle graph of a permutation, thereby providing the basic data structure for inversion distance computations. Hannenhalli and Pevzner then developed the basic theory for expressing the inversion distance in easily computable terms (number of breakpoints minus number of cycles plus number of hurdles plus a correction factor for a fortress (Bafna and Pevzner, 1993; Setubal and Meidanis, 1997)—hurdles and fortresses are easily detectable from a connected component analysis). They also gave the first polynomial-time algorithm for sorting signed permutations by reversals (Hannenhalli and Pevzner, 1995); they also proposed a $O(n^4)$ implementation of their algorithm (Hannenhalli, URL), which runs in quadratic time when restricted to distance computation. Their algorithm requires the computation of the connected components of the overlap graph, which is the bottleneck for the distance computation. Berman and Hannenhalli (1996) later exploited some combinatorial properties of the cycle graph to give a $O(n\alpha(n))$ algorithm to compute the connected components, leading to a $O(n^2\alpha(n))$ implementation of the sorting algorithm. (We will refer to this approach as the UF approach.) Algorithms for finding the connected components of interval graphs (a class of graphs that include the more specialized overlap graphs used in sorting by reversals) that run in linear time are known, but they use range minima and lowest common ancestor data structures and algorithms so that, in addition to being complex and hard to implement, they suffer from high overhead—high enough, in fact, that the UF approach would remain the faster solution in practice.

4. OVERLAP GRAPH AND FOREST

Given a signed permutation of $\{1, \dots, n\}$, we transform it into an unsigned permutation π of $\{1, \dots, 2n\}$ by substituting the ordered pair $(2x - 1, 2x)$ for the positive element x and the ordered pair $(2x, 2x - 1)$ for the negative elements $-x$ then extend π to the set $\{0, 1, \dots, 2n, 2n + 1\}$ by setting $\pi(0) = 0$ and $\pi(2n + 1) = 2n + 1$. By convention, we assume that the two signed permutations for which we must compute a distance have been turned in this manner into unsigned permutations and then both have been permuted (by the same transformation) so that the first permutation becomes the linear ordering $(0, 1, \dots, 2n, 2n + 1)$; these manipulations do not affect the distance value. (This is the reason why transforming one permutation into the other can be viewed as sorting—we want to find out how many inversions are needed to produce the identity permutation from the given one.) We represent an extended unsigned permutation with an edge-colored graph, the *cycle graph* of the permutation. The graph has $2n + 2$ vertices; for each i , $0 \leq i \leq n$, we join vertices $\pi(2i)$ and $\pi(2i + 1)$ by a *gray* edge and vertices $2i$ and $2i + 1$ by a *black* edge, as illustrated in Fig. 1(a). The resulting graph consists of disjoint cycles in which

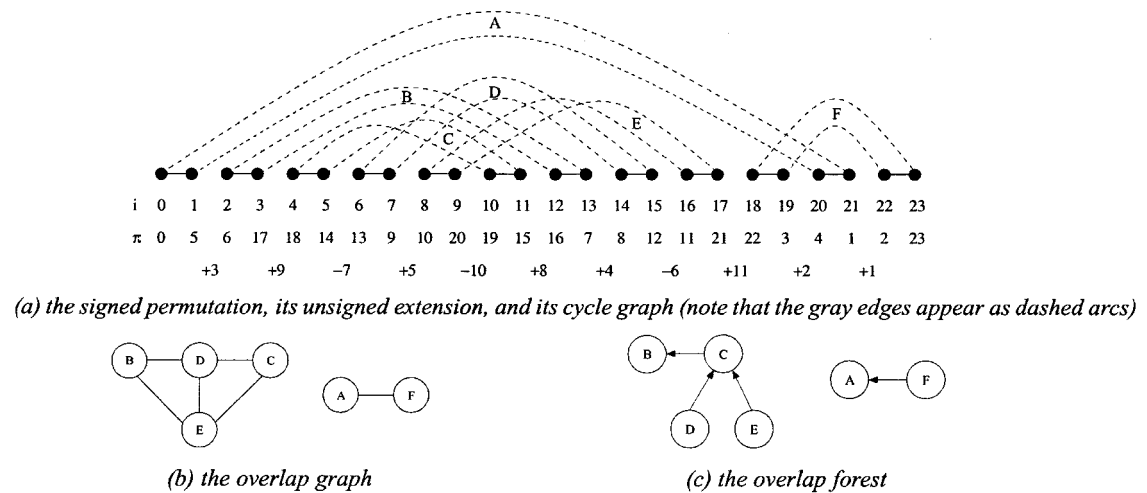


FIG. 1. The signed permutation $(+3, +9, -7, +5, -10, +8, +4, -6, +11, +2, +1)$ and its various representations.

edges alternate colors; we remove from it all 2-cycles (because these cycles correspond to portions of the permutation that are already sorted and cannot intersect with any other cycles).

We say that gray edges $(\pi(i), \pi(j))$ and $(\pi(k), \pi(t))$ *overlap* whenever the two intervals $[i, j]$ and $[k, t]$ overlap, but neither contains the other. Similarly, we say that cycles C_1 and C_2 *overlap* if there exist overlapping gray edges $e_1 \in C_1$ and $e_2 \in C_2$.

Definition 1. The overlap graph of permutation π has one vertex for each cycle in the cycle graph and an edge between any two vertices that correspond to overlapping cycles.

Figure 1 illustrates the concept. The *extent* of cycle C is the interval $[C.B, C.E]$, where we have $C.B = \min\{i \mid \pi(i) \in C\}$ and $C.E = \max\{i \mid \pi(i) \in C\}$. The extent of a set of cycles $\{C_1, \dots, C_k\}$ is $[B, E]$, with $B = \min_{i=1}^k C_i.B$ and $E = \max_{i=1}^k C_i.E$. In Fig. 1, the extent of cycle A is $[0, 21]$, that of cycle F is $[18, 23]$, and that of the set $\{A, F\}$ is $[0, 23]$.

No algorithm that actually builds the overlap graph can run in linear time, since that graph can be of quadratic size. Thus, our goal is to construct an *overlap forest* such that two vertices f and g belong to the same tree in the forest exactly when they belong to the same connected component in the overlap graph. An overlap forest (the composition of its trees is unique, but their structure is arbitrary) has exactly one tree per connected component of the overlap graph and is thus of linear size.

5. THE LINEAR-TIME ALGORITHM FOR CONNECTED COMPONENTS

Our algorithm for computing the connected components scans the permutation twice. The first scan sets up a trivial forest in which each node is its own tree, labeled with the beginning of its cycle. The second scan carries out an iterative refinement of this first forest, by adding edges and so merging trees in the forest; unlike a Union-Find, however, our algorithm does not attempt to maintain the trees within certain shape parameters.

Recall that a node in the overlap graph (or forest) corresponds to a cycle in the cycle graph. The extent $[f.B, f.E]$ of a node f of the overlap forest is the extent of the set of nodes in the subtree rooted at f . Let F_0 be the trivial forest set up in the first scan and assume that the algorithm has processed elements 0 through $j - 1$ of the permutation, producing forest F_{j-1} . We construct F_j from F_{j-1} as follows. Let f be the cycle containing element j of the permutation. If j is the beginning of its own cycle f , then it must be the root of a single-node tree; otherwise, if f overlaps with another cycle g , then we add a new arc (g, f) and compute the combined extent of g and of the tree rooted at f . We say that a tree rooted at f is *active* at stage j whenever j lies properly within the extent of f ; we shall store the extent of the active trees in a stack.

Figure 2 summarizes our algorithm for constructing the overlap forest; in the algorithm, *top* denotes the top element of the stack. The conversion of a forest of up-trees into connected component labels is accomplished in linear time by a simple sweep of the array, taking advantage of the fact that the parent of i must appear before i in the array.

Lemma 1. At iteration i of Step (3) of the algorithm, if the tree rooted at *top* is active and i lies on cycle f and we have $f.B < \text{top}.B$, then there exists h in the tree rooted at *top* such that h overlaps with f .

Proof. Since *top* is active, it must have been pushed onto the stack before the current iteration ($\text{top}.B < i$) and we must not have reached the end of *top*'s extent ($i < \text{top}.E$). Hence, i must be contained in *top*'s extent ($\text{top}.B < i < \text{top}.E$). Since i lies on the cycle f that begins before *top* ($f.B < \text{top}.B$), there must be an edge from cycle f that overlaps with *top*. ■

Theorem 1. The algorithm produces a forest in which each tree is composed of exactly those nodes that form a connected component.

Input: permutation

Output: $\text{parent}[i]$, the parent of i in the overlap forest

Begin

1. scan the permutation, label each position i with $C[i].B$, and set up $[C[i].B, C[i].E]$
2. initialize empty stack
3. for $i \leftarrow 0$ to $2n + 1$
 - (a) if $i = C[i].B$
then push $C[i]$
 - (b) $\text{extent} \leftarrow C[i]$
while ($\text{top}.B > C[i].B$)
 $\text{extent}.B \leftarrow \min\{\text{extent}.B, \text{top}.B\}$
 $\text{extent}.E \leftarrow \max\{\text{extent}.E, \text{top}.E\}$
 pop top
 $\text{parent}[\text{top}.B] \leftarrow C[i].B$
endwhile
 $\text{top}.B \leftarrow \min\{\text{extent}.B, \text{top}.B\}$
 $\text{top}.E \leftarrow \max\{\text{extent}.E, \text{top}.E\}$
 - (c) if $i = \text{top}.E$
then pop top
4. convert each tree into a labeling of its vertices

End

FIG. 2. Constructing the interleaving forest in linear time.

Proof. It suffices to show that, after each iteration, the trees in the forest correspond exactly to the connected components determined by the permutation values scanned up to that point. We prove this invariant by induction on the number of applications of Step (3) of the algorithm.

The base case is trivial: each tree of F_0 has a single node and no two nodes belong to the same connected component since we have not yet processed any element of the permutation.

Assume that the invariant holds after the $(i - 1)$ st iteration and let i lie on cycle f . We prove that the nodes of the tree containing i form the same set as the nodes of the connected component containing i —other trees and connected components are unaffected and so still obey the invariant.

- We prove that a node in the tree containing i must be in the same connected component as i . If we have $i = f.B$, then, as we remarked earlier, nothing changes in the overlap graph (and thus in the connected components); from Step (3), it is also clear that the forest remains unchanged, so that the invariant is preserved. On the other hand, if we have $i > f.B$, then at Step (3) the edge (top, f) will be added to the forest whenever $f.B < \text{top}.B$ holds. This edge will join the subtree rooted at f with that rooted at top into a single subtree. From Lemma 1, we also know that, whenever $f.B < \text{top}.B$ holds, there must exist h in the tree rooted at top such that h and f overlap, so that edge (h, f) must belong to the overlap graph, thereby connecting the component containing f with that containing top and merging them into a single connected component, which maintains the invariant.
- We prove that a node in the same connected component as i must be in the tree containing i . Whenever (j, i) and (k, l) , with $j < k < i < l$, are gray edges on cycles f and h , respectively, then edge (f, h) must belong to the overlap graph built from the first i entries of the permutation. In such a case, our algorithm ensures that edge (h, f) belong to the overlap forest. Our conclusion follows. ■

Obviously, each step of the algorithm takes linear time, so the entire algorithm runs in worst-case linear time.

6. EXPERIMENTS

6.1. Programs

The implementation due to Hannenhalli is very slow and implements the original method of Hannenhalli and Pevzner and not the faster one of Berman and Hannenhalli. The KST applet is very slow as well since it explicitly constructs the overlap graph; it is also written in Java which makes it difficult to compare with C code. For these reasons we wrote our own implementation of the Berman and Hannenhalli algorithm (just the part handling the distance computation) with a view to efficiency. Thus, we not only have an efficient implementation to compare to our linear-time algorithm, but also we have ensured that the two implementations are truly comparable because they share much of their code (hurdles, fortresses, breakpoints), were written by the same person, and used the same algorithmic engineering techniques.

6.2. Experimental setup

We ran experiments on signed permutations of length 10, 20, 40, 80, 160, 320, and 640, in order to verify rate of growth as a function of the number of genes and also to cover the full range of biological applications. We generated groups of three signed permutations from the identity permutation using the evolutionary model of Nadeau and Taylor (1984); in this model, randomly chosen inversions are applied to the permutation at a node to generate the permutations labeling its children, repeating the process until all nodes have been assigned a permutation. The expected number of inversions per edge, r , is fixed in advance, reflecting assumptions about the evolutionary rate in the model. We use five evolutionary rates: 4, 16, 64, 256, and 1024 inversions per edge and generated 10 groups of three-leaf trees—or 10 groups of three genomes each—at each of the six selected lengths. We also generated 10 groups of three random permutations (from a uniform distribution) at each length to provide an extreme test case. For each of these 42 test suites, we computed the three distances among the three genomes in each group 20,000 times in a tight loop, in order to provide accurate timing values for a single computation, then averaged the values over the 10 groups and computed the standard deviation. The computed inversion distances are expected to be at most twice the evolutionary rate since there are two tree edges between each pair of genomes. Our linear algorithm exhibited very consistent behavior throughout, with standard deviations never exceeding 2% of the mean; the UF algorithm showed more variation for $r = 4$ and $r = 16$.

We ran all our tests on a 300MHz Pentium II with 16KB of L1 data cache, 16KB of L1 instruction cache, and 512KB of L2 cache running at half clockspeed; our codes were compiled under Linux with the GNU gcc compiler with options `-O3 -mpentiumpro`. Our code also runs on other systems and machines (e.g., Solaris and Microsoft), where we observed the same behavior.

6.3. Experimental results

We present our results in four plots. The first two (Fig. 3) show the actual running time of our linear-time algorithm for the computation of the inversion distance between two permutations as a function of the size of the permutation, with one plot for the computation of the connected components alone and the other for the complete distance computation. Each plot shows one curve each for the various evolutionary rates and one for the random permutations. We added a third plot showing the average inversion distance; not the very close correlation between the distance and the running time.

For small permutation sizes (10 or less), the L1 data cache holds all of the data without any cache misses, but, as the permutation size grows, the hit rate in the direct-mapped L1 cache steadily decreases until, for permutations of size 100 and larger, execution has slowed down to the speed of the L2 cache (a ratio of two). From that point on, it is clear that the rate of growth is linear, as predicted. It is also clear that $r = 1,024$ is as high a rate of evolution as we need to test, since the number of connected components and inversion distance are nearly indistinguishable from those of the random permutations (see the side plot in Fig. 3 that plots inversion distance as a function of the permutation size). The speed is remarkable: for a typical genome of 100 gene fragments (as found in chloroplast data, for instance [Cosner *et al.*, 2000]), well over 20,000 distance computations can be carried out every second on our rather slow workstation.

Our second two plots (Fig. 4) compare the speed of our linear-time algorithm and that of the UF approach. We plot speed-up ratios, i.e., the ratio of the running time of the UF approach to that of our linear-time

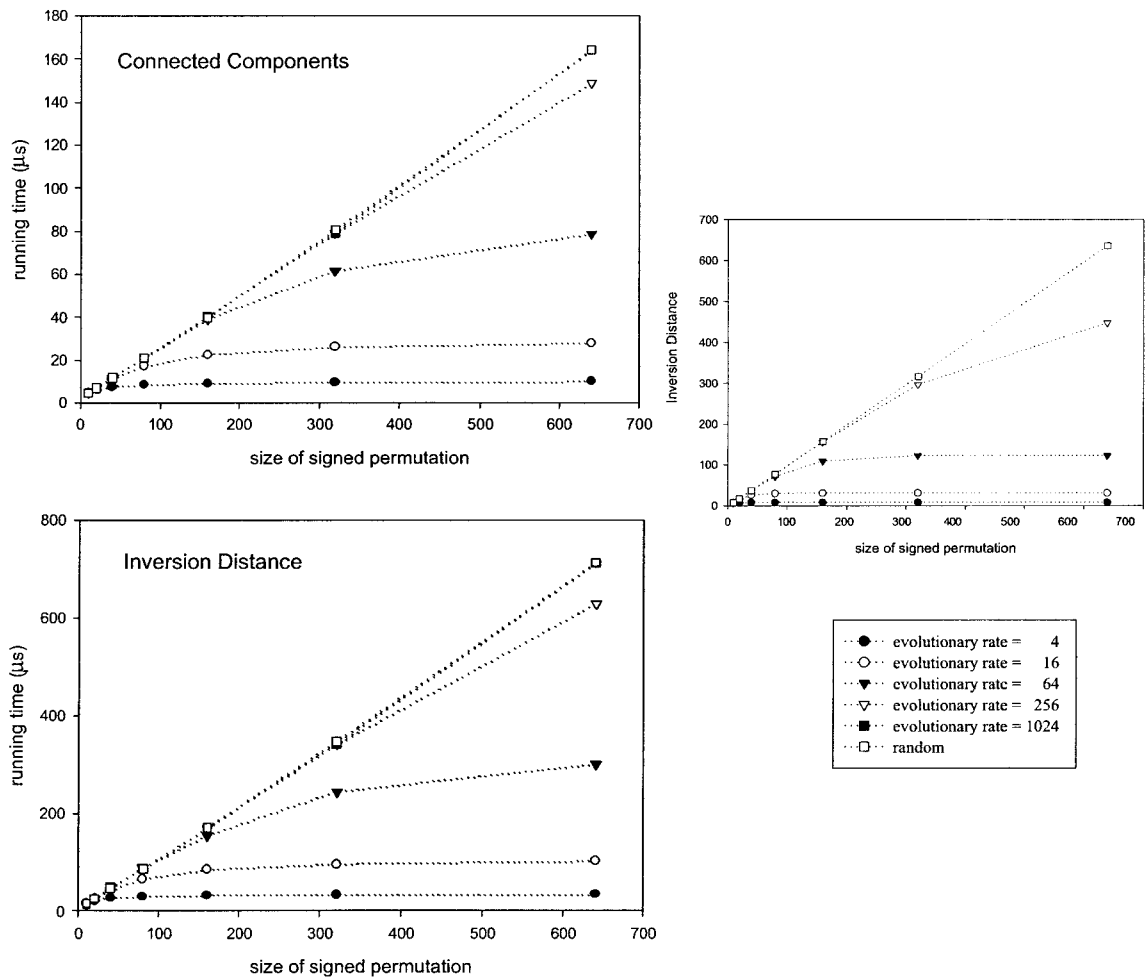


FIG. 3. The running time of our linear-time algorithms as a function of the size of the signed permutation.

algorithm. Again, the first plot addresses only the connected components part of the computation, while the second captures the complete distance computation.

Since the two approaches use a different amount of memory and give rise to a very different pattern of addressing (and thus different cache conflicts), the ratios up to permutations of size 100 vary quite a bit as the size increases—reflecting a transition from the speed of the L1 cache to that of the L2 cache at different permutation sizes for the two algorithms. Beyond that point, however, the ratios stabilize and clearly demonstrate the gain of our algorithm, a gain that increases with increasing permutation size as well as with increasing evolutionary rate.

7. CONCLUDING REMARKS

We have presented a new, very simple, practical, linear-time algorithm for computing the inversion distance between two signed permutations, along with a detailed experimental study comparing the running time of our algorithm with that of Berman and Hannenhalli. Our code is available from the web page www.cs.unm.edu/~moret/GRAPPA under the terms of the GNU public license (GPL); it has been tested under Linux (including the parallel version), FreeBSD, Solaris, and Windows NT. This code includes inversion distance as part of a much larger context, which provides means of reconstructing phylogenies based on gene order data. We found that using our inversion distance computation in lieu of the surrogate breakpoint distance (which was used by previous researchers in an attempt to speed up computation

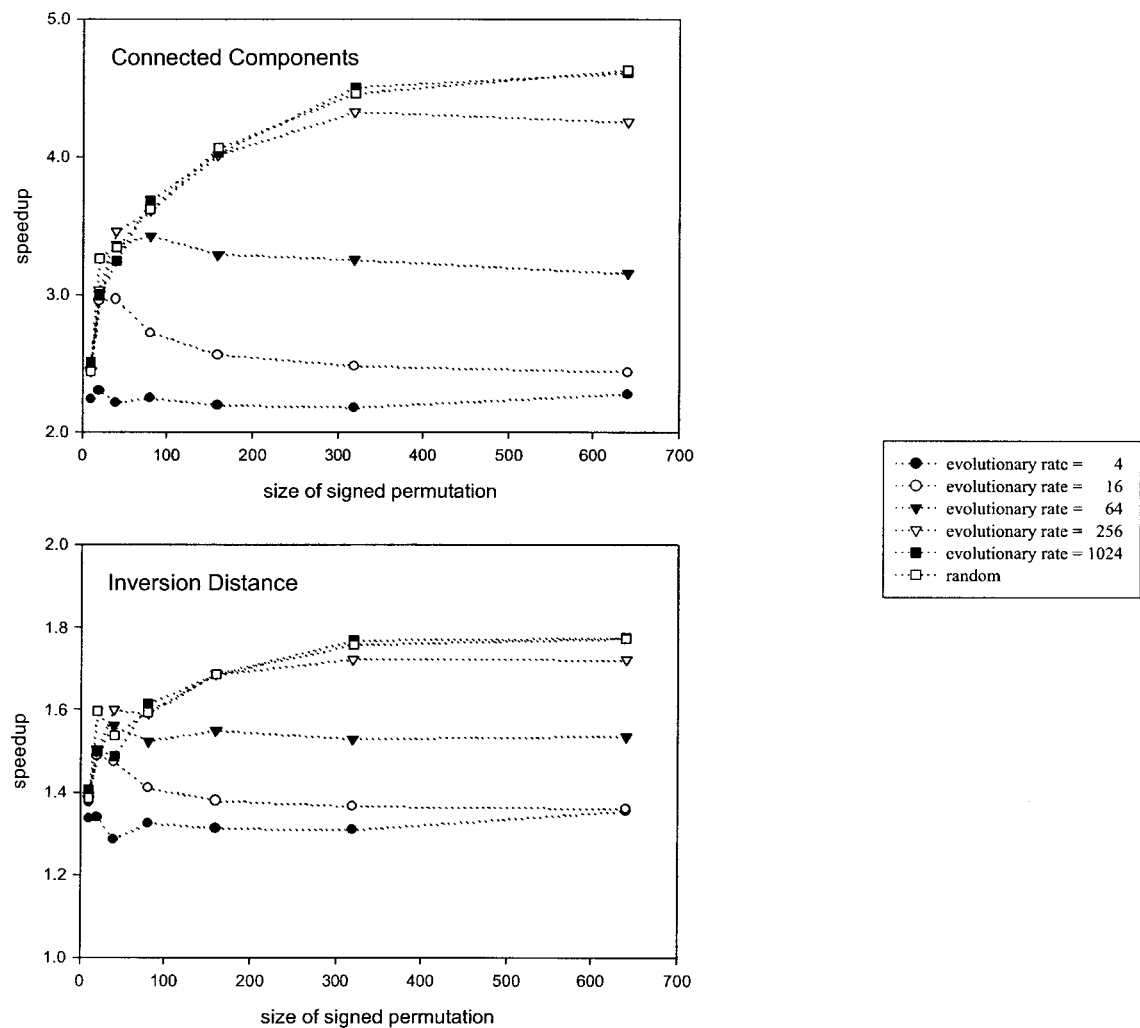


FIG. 4. The speedup of our linear-time algorithms over the UF approach as a function of the size of the signed permutation.

[Blanchette *et al.*, 1997; Sankoff and Blanchette, 1998]) only slowed down the reconstruction algorithm by about 30%, enabling us to extend work on breakpoint analysis (as reported by Cosner *et al.* [2000] and by Moret *et al.* [2001]) to similar work on inversion phylogeny.

ACKNOWLEDGMENTS

David A. Bader is supported in part by NSF grants CAREER 00-93039, ITR 00-81404, and DEB 99-10123; and by DOE CSRI 14968. Bernard M.E. Moret is supported in part by NSF grant ITR 00-81404. Mi Yan acknowledges her support by the UNM Albuquerque High Performance Computing Center.

REFERENCES

Bafna, V., and Pevzner, P. 1995. Sorting permutations by transpositions. *Proc. 6th Ann. Symposium on Discrete Algorithms*, 614–623, ACM Press, New York.

Bafna, V., and Pevzner, P.A. 1993. Genome rearrangements and sorting by reversals. *Proc. 34th Ann. IEEE Symposium on Foundations of Computer Science*, 148–157, IEEE Press.

- Bafna, V., and Pevzner, P.A. 1996. Genome rearrangements and sorting by reversals. *SIAM J. Computing* 25, 272–289.
- Bafna, V., and Pevzner, P.A. 1998. Sorting by transpositions. *SIAM J. Discrete Math.* 11, 2, 224–240.
- Berman, P., and Hannenhalli, S. 1996. Fast sorting by reversal. In D. Hirschberg and E. Myers, eds., *Proc. 7th Ann. Symposium on Combinatorial Pattern Matching*, Laguna Beach, CA, 168–185, *Lecture Notes in Computer Science*, 1075, Springer-Verlag, New York.
- Blanchette, M., Bourque, G., and Sankoff, D. 1997. Breakpoint phylogenies. In S. Miyano and T. Takagi, eds., *Genome Informatics*, 25–34. University Academy Press, Tokyo.
- Caprara, A. 1997. Sorting by reversals is difficult. In *Proc. 1st Conf. Computational Molecular Biology*, Santa Fe, NM, 78–83. ACM Press, New York.
- Caprara, A. 1999. Sorting permutations by reversals and Eulerian cycle decompositions. *SIAM J. Discrete Mathematics* 12, 1, 91–110.
- Cosner, M., Jansen, R., Moret, B., Raubeson, L., Wang, L., Warnow, T., and Wyman, S. 2000. A new fast heuristic for computing the breakpoint phylogeny and a phylogenetic analysis of a group of highly rearranged chloroplast genomes. In *Proc. 8th Int. Conf. Intelligent Systems for Molecular Biology*, San Diego, CA, 104–115, AAAI.
- Hannenhalli, S. URL. Software for computing inversion distances between signed gene orders. Department of Mathematics. University of Southern California. www-hto.usc.edu/plain/people/Hannenhalli.html.
- Hannenhalli, S., and Pevzner, P. 1995. Transforming cabbage into turnip (polynomial algorithm for sorting signed permutations by reversals). In *Proc. 27th Ann. Symposium on Theory of Computing*, Las Vegas, NV, 178–189. ACM Press, New York.
- Kaplan, H., Shamir, R., and Tarjan, R. 1999. A faster and simpler algorithm for sorting signed permutations by reversals. *SIAM J. Computing* 29, 3, 880–892. First appeared in *Proc. 8th Ann. Symposium on Discrete Algorithms*, 344–351, New Orleans, LA. ACM Press, New York.
- Mantini, I., and Shamir, R. 1999. Genome rearrangement algorithm applet: An algorithm for sorting signed permutations by reversals. www.math.tau.ac.il/~rshamir/GR/.
- McGeoch, C. 1996. Toward an experimental method for algorithm simulation. *INFORMS J. Computing* 8, 1–15.
- Moret, B.M., and Shapiro, H.D. 2001. Algorithms and experiments: The new (and old) methodology. *J. Universal Comp. Sci.* 7, 5, 434–446. URL: http://www.jucs.org/juc_7_5/algorithms_and_experiments_the.
- Moret, B.M., Wyman, S., Bader, D., Warnow, T., and Yan, M. 2001. A new implementation and detailed study of breakpoint analysis. *Proc. 6th Pacific Symposium on Biocomputing*, Hawaii, 583–594, World Scientific.
- Nadeau, J., and Taylor, B. 1984. Lengths of chromosome segments conserved since divergence of man and mouse. *Proc. Natl. Acad. Sci USA* 81, 814–818.
- Olmstead, R., and Palmer, J. 1994. Chloroplast DNA systematics: A review of methods and data analysis. *Am. J. Botany* 81, 1205–1224.
- Palmer, J. 1992. Chloroplast and mitochondrial genome evolution in land plants. In R. Herrmann, ed., *Cell Organelles*, 99–133. Springer Verlag, New York.
- Raubeson, L., and Jansen, R. 1992. Chloroplast DNA evidence on the ancient evolutionary split in vascular land plants. *Science* 255, 1697–1699.
- Sankoff, D., and Blanchette, M. 1998. Multiple genome rearrangement and breakpoint phylogeny. *J. Comp. Biol.* 5, 555–570.
- Setubal, J., and Meidanis, J. 1997. *Introduction to Computational Molecular Biology*. PWS, Boston.

Address correspondence to:

David A. Bader

The University of New Mexico

Department of Electrical and Computer Engineering

Albuquerque, NM 87131-1356

E-mail: dbader@ece.unm.edu