

Subject Section

The Effect of Genome Graph Expressiveness on the Discrepancy Between Genome Graph Distance and String Set Distance

Yutong Qiu¹ and Carl Kingsford^{1*}

¹ Computational Biology Department, Carnegie Mellon University, Pittsburgh, 15232, United States.

*To whom correspondence should be addressed. Email: carlk@cs.cmu.edu

Associate Editor: XXXXXXX

Received on XXXXX; revised on XXXXX; accepted on XXXXX

Abstract

Motivation: Intra-sample heterogeneity describes the phenomenon where a genomic sample contains a diverse set of genomic sequences. In practice, the true string sets in a sample are often unknown due to limitations in sequencing technology. In order to compare heterogeneous samples, genome graphs can be used to represent such sets of strings. However, a genome graph is generally able to represent a string set universe that contains multiple sets of strings in addition to the true string set. This difference between genome graphs and string sets is not well characterized. As a result, a distance metric between genome graphs may not match the distance between true string sets.

Results: We extend a genome graph distance metric, Graph Traversal Edit Distance (GTED) proposed by Ebrahimpour Boroojeny et al., to FGTEd to model the distance between heterogeneous string sets and show that GTED and FGTEd always underestimate the Earth Mover's Edit Distance (EMED) between string sets. We introduce the notion of string set universe diameter of a genome graph. Using the diameter, we are able to upper-bound the deviation of FGTEd from EMED and to improve FGTEd so that it reduces the average error in empirically estimating the similarity between true string sets. On simulated TCR sequences and actual Hepatitis B virus genomes, we show that the diameter-corrected FGTEd reduces the average deviation of the estimated distance from the true string set distances by more than 250%.

Availability: Data and source code for reproducing the experiments are available at:

<https://github.com/Kingsford-Group/gtedemtest/>

Contact: carlk@cs.cmu.edu

1 Introduction

Intra-sample heterogeneity describes the phenomenon where a genomic sample contains a diverse set of genomic sequences. A heterogeneous string set is a set of strings where each string is assigned a weight representing its abundance in the set. Computing the distance between heterogeneous string sets is essentially computing the distance between two distributions of strings. We formulate the problem of heterogeneous sample comparison as the heterogeneous string set comparison problem.

This problem can be used to compare samples where differences can be traced to the differences between sets of genomic sequences. For example, cancer samples are clustered based on differences in their genomic and transcriptomic features (Morris *et al.*, 2016; Zhao *et al.*,

2019) into cancer subtypes that correlate with patient survival rates. The dissimilarities between T-cell receptor (TCR) sequences are computed between individuals to study immune responses (Bolen *et al.*, 2017). Different compositions of these sequences result in different clinical outcomes such as response to treatment.

We point out that the Earth Mover's Distance (EMD) (Rubner *et al.*, 2000), or the Wasserstein distance (Wasserstein *et al.*, 1969), with edit distance as the ground metric is an elegant metric to compare a pair of heterogeneous string sets. Given two distributions of items and a cost to transform one item into another, EMD computes the total cost of transforming one distribution into another. The EMD was initially used in computer vision to compare distributions of pixel values in images (Levina and Bickel, 2001) and later adapted to natural language processing (Kusner *et al.*, 2015). It has also been used to approximate the distance between two

genomes (Mangul and Koslicki, 2016) by computing the distance between two distributions of k-mers. To compare heterogeneous string sets, when the strings and their distributions are known, we use edit distance as the cost to transform one string to another. We refer to this as the Earth Mover’s Edit Distance (EMED).

In practice, the complete strings of interest and their abundances are often unknown, and these strings are only observed as fragmented sequencing reads. It is impossible to exactly compute EMED between the true sets of complete strings from the sequencing reads only.

The challenges posed by incomplete observed sequences can be alleviated by representing the string set using a graph structure. Multiple types of genome graphs have been introduced (Holley and Melsted, 2020; Almodaresi et al., 2018; Li et al., 2020; Iqbal et al., 2012; Dilthey et al., 2015; Garrison et al., 2018; Minkin et al., 2017; Paten et al., 2017, 2011; Lee and Kingsford, 2018). For our purposes, a genome graph is a directed multigraph with labeled nodes and weighted edges, along with a source and a sink node. A string is spelled by a source-to-sink path, or $s - t$ path, if it is equal to the concatenation of node labels on the path. We say that a genome graph represents a string set if the union of paths that spells each string in the set is equal to the graph. In other words, a string set can be spelled by a decomposition of the genome graph.

There are several methods that compute the distance between genome graphs (Minkin and Medvedev, 2020; Polevikov and Kolmogorov, 2019; Ebrahimpour Boroojeny et al., 2020). Among those, Graph Traversal Edit Distance (GTED) (Ebrahimpour Boroojeny et al., 2020) is a general measure that can be applied to genome graphs and does not rely on the type of genome graphs nor the knowledge of the true string sets. Given two genome graphs, GTED finds an Eulerian cycle in each graph that minimizes the edit distance between the strings spelled by each cycle.

However, applying GTED on genome graphs representing heterogeneous string sets may overestimate the similarity between these string sets for two reasons. First, since GTED computes the distance between Eulerian cycles in genome graphs, it may align the prefix of a string to the suffix of another string with no additional penalties. We address this challenge by proposing an extension of GTED, called FGTED, that penalizes direct alignment of prefixes of a string with suffixes of other strings.

Second, and more significantly, both FGTED and GTED compute the edit distance between the two string sets represented by each genome graph that are most similar to each other. However, a genome graph that is constructed from sequencing fragments typically is able to represent more than one set of strings (Kingsford et al., 2010; Paten et al., 2018). As a genome graph merges shared sequences into the same node, it creates chains of bubble structures (Zerbino and Birney, 2008) that result in exponential number of possible paths, and these paths spell a much more diverse collection of strings than the original set. We call the degree to which a genome graph encodes a larger set of strings than the true underlying set the “expressiveness” of a genome graph. Due to the expressiveness of a genome graph, the Eulerian cycles found by GTED may not spell the true set of strings and the computed distance may be far from the true distance between string sets used to construct the graphs (Figure 1(a)).

We prove both that FGTED always produces a distance that is larger than or equal to GTED, and that FGTED computes a metric that is always less than or equal to the EMED between true sets of strings.

However, FGTED and GTED can be quite far from the EMED. To resolve this discrepancy between FGTED and EMED, we define the collection of strings that can be represented by the genome graph as its string set universe, and genome graph expressiveness as the diameter of its string set universe (SUD), which is the maximum EMED between two string sets that can be represented by the graph (Figure 1(b)).

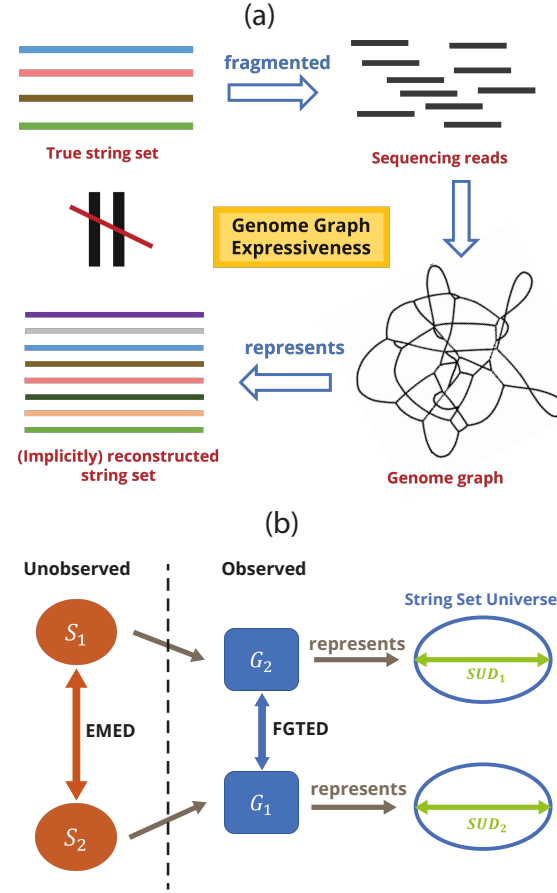


Fig. 1. (a) Genome graph expressiveness results in inexact representations of true string sets. (b) Overview of part of theoretical contributions.

Using diameters, we are able to upper-bound the deviation of FGTED from EMED. Additionally, we are able to correct FGTED and more accurately estimate the true string set distance empirically. On simulated TCR sequences, we reduce the average deviation of FGTED from EMED by more than 300%, and increase the correlation between the true and estimated string set distances by 20%. On Hepatitis B virus genomes, we reduce the average deviation by more than 250%.

These results provide the first connection between comparisons of genome graphs that encode multiple sequences and a natural string distance and provide the first formalization of the expressiveness of genome graphs. Additionally, they provide a practical method to estimate and reduce discrepancy between genome graph distances and string set distances.

2 Preliminary Concepts

2.1 Strings

Definition 1 (Heterogeneous string set). A heterogeneous string set $\mathcal{S} = \{(w_1, s_1), \dots, (w_n, s_n)\}$ contains a set of strings, where each string s_i is assigned a weight $w_i \in [0, 1]$ that indicates the abundance of s_i in $\mathcal{S}$. We say that the total weight of $\mathcal{S}$ is $\sum_{i \in [1, n]} w_i = 1$.

$ED(s_1, s_2)$ is the minimum cost to transform s_1 into s_2 under edit distance (Levenshtein et al., 1966). The set of operations that transforms s_1 to s_2 can be written as an alignment between s_1 and s_2 , or $A =$

$align(s_1, s_2)$. The i -th position in A is denoted by $A[i] = \begin{bmatrix} c_{(1,i)} \\ c_{(2,i)} \end{bmatrix}$, where $c_{(a,i)}$ is either a gap character “-” or a character in s_a .

2.2 Earth Mover’s Edit Distance

To find a distance between two heterogeneous string sets, we need to take into account not only the distance between pairs of strings, but also the weight, or abundance of each string in the set. When we are comparing two heterogeneous string sets, we are essentially comparing two distributions of strings. Therefore, we propose using the Earth Mover’s Distance (EMD) as a natural distance measure.

Given two distributions of items (here, strings) and a cost function that quantifies the cost of transforming one item into another, the EMD between the two distributions is the minimum cost to transform one distribution into another. Computing EMD can be viewed as a transportation problem that finds a many-to-many mapping between two sets of items and minimizes the total cost of the mapping (Rubner *et al.*, 2000; Wasserstein *et al.*, 1969).

Given two heterogeneous string sets $\mathcal{S}_1 = \{(w_1, s_1), \dots, (w_n, s_n)\}$ and $\mathcal{S}_2 = \{(w_{n+1}, s_{n+1}), \dots, (w_m, s_m)\}$, to compute the Earth Mover’s Edit Distance (EMED), we use the edit distance between s_i and s_j as the cost of transforming one string to another. Following procedures to compute EMD (Rubner *et al.*, 2000) as a min-cost max-flow problem, we find a mapping M , where $M(s_i, s_j)$ is the amount of $s_i \in \mathcal{S}_1$ to be transformed into $s_j \in \mathcal{S}_2$, that minimizes $cost(M) = \sum_{s_i \in \mathcal{S}_1} \sum_{s_j \in \mathcal{S}_2} M(s_i, s_j) \cdot ED(s_i, s_j)$. We define that $EMED(\mathcal{S}_1, \mathcal{S}_2) = \min_M cost(M)$.

2.3 Flow Networks

Definition 2 (Valid flow network). A directed graph $\mathcal{G} = (V, E, w)$, where $w(e)$ is the weight of each edge, is a valid flow network if there exists a source s and sink node t such that:

$$(Flow\ conservation) \quad \sum_{(u,v) \in E} w(u,v) = \sum_{(v,w) \in E} w(v,w)$$

$$\forall v \in V, v \neq s, v \neq t,$$

$$(Total\ capacity) \quad \sum_{(s,u) \in E} w(s,u) = \sum_{(v,t) \in E} w(v,t).$$

Definition 3 (Flow decomposition). A flow decomposition of a valid flow graph $\mathcal{G}$, denoted as $D(\mathcal{G})$, is a collection of paths and their weights $\mathcal{P} = \{(w_1, p_1), \dots, (w_n, p_n)\}$, where $p_i = ((s, u_1), \dots, (u_m, t))$ is an ordered sequence of edges in $\mathcal{G}$, such that:

$$(Flow\ coverage) \quad \sum_{p_i \in \mathcal{P}} O(e, i) \cdot w_i = w(e) \quad \forall e \in \mathcal{G},$$

where $O(e, i)$ is equal to the number of occurrences of edge e in path p_i .

A valid flow network typically has more than one flow decomposition. Let the set of possible flow decompositions of $\mathcal{G}$ be $\mathcal{D}_{\mathcal{G}}$.

2.4 Genome Graphs

There are many variants of genome graphs used for various purposes and in various settings. Here, we introduce the definition of genome graphs we will use.

Definition 4 (Genome graph). A genome graph $\mathcal{G} = (V, E, l, w)$ is a valid flow network with node set V , edge set E , node labels $l(u)$ for each $u \in V$ and edge weights $w(e)$ for each $e \in E$. A genome graph contains a source node s and a sink node t , and $l(s) = “\$”$, $l(t) = “\#”$, where $\$$ and $\#$ are special characters that do not appear in any string set considered in the scope of this manuscript.

Define operator $S(\cdot)$ that transforms a set of paths in a genome graph $\mathcal{G}$ to a set of strings by concatenating the node labels on each path. $S(\mathcal{P}) = \{concat(p), w(p) \mid p \in \mathcal{P}\}$ is a heterogeneous string set where the weight of each string is equal to the weight of the path that spells the string.

Definition 5 (String set represented by a genome graph). A genome graph $\mathcal{G}$ represents a string set $\mathcal{S}$ if there exists a decomposition $D(\mathcal{G}) \in \mathcal{D}_{\mathcal{G}}$, such that $S(D(\mathcal{G})) = \mathcal{S}$.

We use $\mathcal{G} = G(\mathcal{S})$ to denote when $\mathcal{G}$ represents $\mathcal{S}$.

Definition 6 (String set universe represented by a genome graph). The string set universe $SU(\mathcal{G})$ of a genome graph $\mathcal{G}$ is the collection of heterogeneous string sets that can be represented by $\mathcal{G}$. Formally, $SU(\mathcal{G}) = \{S(D) \mid D \in \mathcal{D}_{\mathcal{G}}\}$.

2.5 Alignment Graph

An alignment graph is used to align two genome graphs (Ebrahimpour Boroojeni *et al.*, 2020) and can be viewed as a graph product between two genome graphs. A special case of the alignment graph (Jain *et al.*, 2020) is used to align a string to a graph where the string is represented as a graph with only one path. We assume that the genome graphs to be aligned are transformed so that the label of each node contains only one character.

Definition 7 (Alignment graph). Given genome graphs $\mathcal{G}_1 = (V_1, E_1, l_1, w_1)$ and $\mathcal{G}_2 = (V_2, E_2, l_2, w_2)$, an alignment graph $AG(\mathcal{G}_1, \mathcal{G}_2) = (V_A, E_A, cost, w)$ is a directed graph with node set V_A , edge set E_A , edge cost $cost(e)$ and edge weight $w(e)$ for each edge $e \in E_A$. The alignment graph is defined following the steps:

- V_A is constructed by adding pairings of nodes in V_1 and V_2 ; that is $V_A = \{(u_1, u_2) \mid u_1 \in V_1, u_2 \in V_2\}$.
- For each edge $(u_1, v_1) \in E_1$ and $(u_2, v_2) \in E_2$, where $(u_1, u_2) \in V_A$ and $(v_1, v_2) \in V_A$, create three types of edges:
 1. A match/mismatch edge $e = ((u_1, u_2), (v_1, v_2))$ with $w(e) = \min\{w_1(u_1, v_1), w_2(u_2, v_2)\}$.
 2. An insertion (in) edge $e = ((u_1, u_2), (u_1, v_2))$ with $w(e) = w_2(u_2, v_2)$.
 3. A deletion (del) edge $e = ((u_1, u_2), (v_1, u_2))$ with $w(e) = w_1(u_1, v_1)$.

The cost of an in/del edge and a mismatch edge is equal to a customized penalty. The cost of a match edge is equal to zero. A match/mismatch edge should be distinguished with an in/del edge if the corresponding edge in one of the input graphs is a self-loop.

Each edge $e = ((u_1, u_2), (v_1, v_2))$ in an alignment graph can be projected onto one edge in each of the input graphs. An edge in each of the input graphs can also be projected onto a set of edges in AG . The size of AG is $O(|E_1| \cdot |E_2|)$ and thus the time to construct AG is quadratic in the sizes of the input genome graphs.

Definition 8 (Projection function). Define the projection function as $P_{(\mathcal{G}, \mathcal{H})}(e) = E'$ that maps an edge e from graph $\mathcal{G}$ to a set of edges E' in graph $\mathcal{H}$. The projection function maps an edge in the alignment graph to the edges in the input graphs that are matched together by that edge. It also maps an edge in one of the input graphs to a set of edges in the alignment graph where it is matched with other edges in another input graph. Specifically:

Projection from alignment graph to one of the input graphs is defined by

$$P_{(AG, \mathcal{G}_i)}((u_1, u_2), (v_1, v_2)) = \{(u_i, v_i)\}, i \in \{1, 2\}.$$

Projection from one of the input graphs to alignment graph is defined by

$$P_{(\mathcal{G}_1, AG)}((u_1, v_1)) = \{e \mid e = ((u_1, u_2), (v_1, v_2)) \in E_{AG}\},$$

$$P_{(\mathcal{G}_2, AG)}((u_2, v_2)) = \{e \mid e = ((u_1, u_2), (v_1, v_2)) \in E_{AG}\}.$$

Given a set of paths $\mathcal{P}$ in AG , we use $P_{(AG, \mathcal{G}_i)}(\mathcal{P})$ to denote the projection of $\mathcal{P}$ onto $\mathcal{G}_i$, where $P_{(AG, \mathcal{G}_i)}(\mathcal{P}) = \{(P_{(AG, \mathcal{G}_i)}(e) \mid e \in p) \mid p \in \mathcal{P}\}$.

For convenience, we define that $f_i(D(AG)) = S(P_{(AG, \mathcal{G}_i)}(D(AG)))$, which is the set of strings spelled by a path decomposition in AG that is projected onto $\mathcal{G}_i$.

2.6 Graph Traversal Edit Distance (GTED)

Graph Traversal Edit Distance (GTED), proposed by Ebrahimpour Boroojeny *et al.* (2020), is a distance between two labeled graphs which are assumed to be Eulerian graphs. Given a genome graph in our definition, we add an edge directing from sink to source with weight equal the sum of edge weights that are directing from the source node in order to make an Eulerian graph.

Let the language of $\mathcal{G}$, $L(\mathcal{G})$, be the set of strings spelled by Eulerian cycles in $\mathcal{G}$. Formally, $L(\mathcal{G}) = \{S(c) \mid c \text{ is an Eulerian cycle in } \mathcal{G}\}$.

Definition 9 (Graph Traversal Edit Distance (Ebrahimpour Boroojeny *et al.*, 2020)). *Let $\mathcal{G}_1$ and $\mathcal{G}_2$ be two Eulerian graphs, where the weights on the edges are seen as the number of times an edge must be visited in each Eulerian cycle. Then,*

$$GTED(\mathcal{G}_1, \mathcal{G}_2) = \min_{\substack{s_1 \in L(\mathcal{G}_1) \\ s_2 \in L(\mathcal{G}_2)}} ED(s_1, s_2).$$

GTED finds one Eulerian cycle in each genome graph such that the edit distance between the strings spelled by the Eulerian cycles is minimized. GTED is computed by solving a linear programming (LP) formulation (Equations (1)-(4)) on the alignment graph $AG(\mathcal{G}_1, \mathcal{G}_2)$, which minimizes the cost of a flow in the graph with the flow conservation (Equation (4)) and flow coverage constraints (Equations (2)-(3)). The LP formulation is as follows:

$$\min_{x \in \mathbb{R}^{|E_A|}} \sum_{e \in E_A} cost(e) \cdot x_e \quad (1)$$

$$s.t. \sum_{j,l} x_{((i,j),(k,l))} = w_1(i, k) \quad \forall (i, k) \in E_1 \quad (2)$$

$$\sum_{i,k} x_{((i,j),(k,l))} = w_2(j, l) \quad \forall (j, l) \in E_2 \quad (3)$$

$$\sum_{(u,v) \in E_A} x_{(u,v)} = \sum_{(v,w) \in E_A} x_{(v,w)} \quad \forall v \in V_A \quad (4)$$

Ebrahimpour Boroojeny *et al.* (2020) prove that GTED is equal to the optimal solution of this LP formulation, and thus GTED is computable in polynomial time. The number of constraints in the above LP is linear in the size of the alignment graph and thus quadratic in the size of the input genome graphs.

3 An Extension of GTED

GTED was originally used to compare genome graphs that are assumed to contain single genomes. It is therefore intuitive that each string represented by the genome graph is spelled with an Eulerian cycle. This property follows the property of assembly graphs (Pevzner *et al.*, 2001). When the genome graph represents more than one string, finding a string spelled by

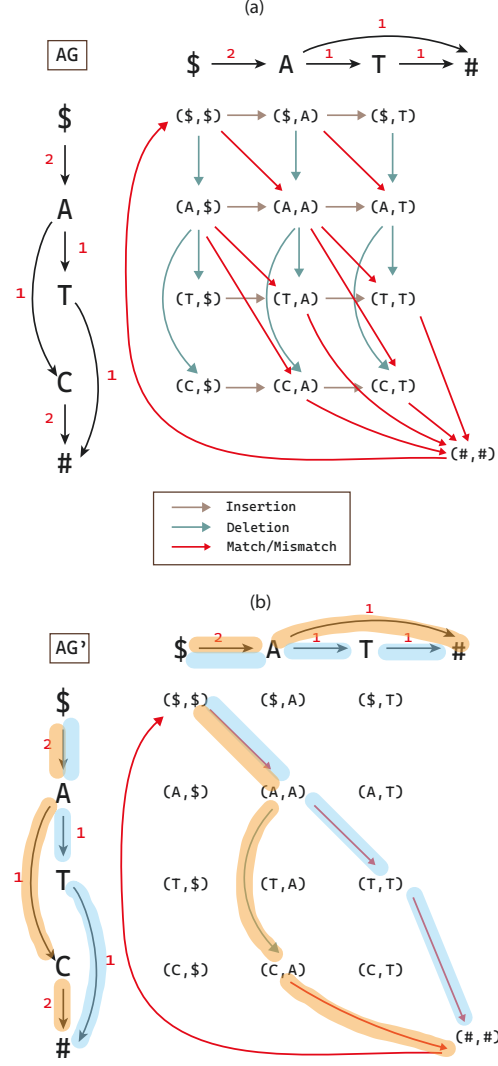


Fig. 2. (a) An alignment graph AG between $\mathcal{G}_1$ (vertical) and $\mathcal{G}_2$ (horizontal). Insertion, deletion and match/mismatch edges are labeled with different colors. (b) AG' after removing all the edges with zero flow in a solution to FGTEd($\mathcal{G}_1, \mathcal{G}_2$). Edges in $\mathcal{G}_1$ and $\mathcal{G}_2$ that are highlighted with matching colors are projections from edges in AG' to $\mathcal{G}_1$ and $\mathcal{G}_2$, respectively. Path $(\$, A, T, \#) \in \mathcal{G}_1$ is aligned to $(\$, A, T, \#) \in \mathcal{G}_2$ and path $(\$, A, C, \#) \in \mathcal{G}_1$ is aligned to $(\$, A, \#) \in \mathcal{G}_2$. The weights on AG and AG' edges are omitted for simplicity.

an Eulerian cycle c in the graph is equivalent to finding a concatenation of a permutation of strings in a string set. When aligning two Eulerian cycles, c_1 and c_2 , from input graphs, the boundaries between strings are ignored and the prefix of one string may be aligned to the suffix of another string with no cost. However, such alignment is not allowed when we align sets of strings using EMED.

We propose an extension of GTED with a modified cost function in edit distance computation so that the cost of aligning the sink character $\#$ with any other character is infinity.

Figure 2(a) shows an example of the alignment graph built from two input graphs using the proposed cost function. Let the sink nodes in $\mathcal{G}_1$ and $\mathcal{G}_2$ be t_1 and t_2 , and the source nodes be s_1 and s_2 , respectively. After removing all the alignment edges with infinite costs, there is an edge to the alignment node (t_1, t_2) in AG if and only if there exists an edge (u_1, t_1) in $\mathcal{G}_1$ and an edge (u_2, t_2) in $\mathcal{G}_2$. The only incoming edge to (s_1, s_2) is $((t_1, t_2), (s_1, s_2))$. We refer to the edge $((t_1, t_2), (s_1, s_2))$

as the sink-to-source edge, or $t - s$ edge in alignment graph in the rest of the manuscript.

We let Flow-GTED, or FGTED, denote the distance computed using the alignment graph after removing all infinity cost edges that forbids aligning the sink with any nodes other than the source node. FGTED assumes that the input genome graphs are flow networks that represent string sets, which can be seen as an analog to Eulerian tours in the graphs that are used as input for GTED. Sink-to-source edges are added to transform flow networks into Eulerian graphs such that FGTED can be reduced to GTED. As FGTED solves a similar LP formulation as GTED that is constructed on a slightly smaller alignment graph, FGTED is also solvable in polynomial time.

Theorem 1. $GTED(\mathcal{G}_1, \mathcal{G}_2) \leq FGTED(\mathcal{G}_1, \mathcal{G}_2)$ for any pair of genome graphs $\mathcal{G}_1, \mathcal{G}_2$.

Proof. Since FGTED is computed on a smaller alignment graph that contains fewer edges than that for computing GTED, FGTED explores a smaller solution space than GTED in solving the LP formulation. Therefore, any feasible solution to the LP formulation for FGTED($\mathcal{G}_1, \mathcal{G}_2$) is a feasible solution to the LP formulation for GTED($\mathcal{G}_1, \mathcal{G}_2$). Since $GTED(\mathcal{G}_1, \mathcal{G}_2)$ minimizes the objective, the theorem is true.

4 The Relationship Between GTED, FGTED and EMED

Let AG^* be the alignment graph after removing the $t - s$ edge and all the edges from $\{e \mid x_e = 0, e \in E_A\}$ from the LP solution to Equations (1)–(4). We say that AG^* is a solution of FGTED. Due to constraints (2)–(4), AG^* is a valid flow network. Let $D(AG^*)$ be a flow decomposition in AG^* . Similar to the Eulerian cycles found during the GTED computation, each path in $D(AG^*)$ can be projected to a path in $\mathcal{G}_1$ and a path in $\mathcal{G}_2$.

Denote $\mathcal{S}_1^* = f_1(D(AG^*)) = S(P_{(AG^*, \mathcal{G}_1)}(D(AG^*)))$ as the set of strings spelled by the set of projected paths from a decomposition $D(AG^*)$ to $\mathcal{G}_1$. Similarly, $\mathcal{S}_2^* = f_2(D(AG^*)) = S(P_{(AG^*, \mathcal{G}_2)}(D(AG^*)))$. We show an example of path projections in Figure 2(b).

Observing that we can do flow decomposition in both the FGTED solution and input genome graphs, we will show in this section that FGTED can be bounded by EMED between decompositions in input genome graphs and in the alignment graph solutions.

Theorem 2. Given two sets of strings $\mathcal{S}_1$ and $\mathcal{S}_2$, and genome graphs representing these string sets, $\mathcal{G}_1 = G(\mathcal{S}_1)$ and $\mathcal{G}_2 = G(\mathcal{S}_2)$,

$$\begin{aligned} 0 &\leq EMED(\mathcal{S}_1, \mathcal{S}_2) - FGTED(\mathcal{G}_1, \mathcal{G}_2) \\ &\leq \min_{\substack{D(AG^*) \in \mathcal{D}_{AG^*} \\ \mathcal{S}_1^* = f_1(D(AG^*)) \\ \mathcal{S}_2^* = f_2(D(AG^*))}} (EMED(\mathcal{S}_1, \mathcal{S}_1^*) + EMED(\mathcal{S}_2, \mathcal{S}_2^*)) \end{aligned}$$

where AG^* is the solution obtained from FGTED($\mathcal{G}_1, \mathcal{G}_2$).

The proof of this theorem is completed in two parts. The first inequality is proven in Section 4.1 and the second is proven in Section 4.2. Since FGTED computes a distance that is larger than GTED between the same pair of genome graphs (Theorem 1), Theorem 2 also shows that FGTED always estimates the distance between true string sets more accurately than GTED.

4.1 FGTED is Always Less Than or Equal to EMED

We show in this section that FGTED can be expressed in terms of EMED between string sets constructed from decomposing AG^* . In other words,

while GTED finds an Eulerian tour in each input graph, FGTED finds a flow decomposition in $\mathcal{G}_1$ and $\mathcal{G}_2$, respectively, that minimizes the EMED between them. Analogous to Definition 9, we have:

Theorem 3. Given two genome graphs $\mathcal{G}_1$ and $\mathcal{G}_2$,

$$FGTED(\mathcal{G}_1, \mathcal{G}_2) = \min_{\substack{D(\mathcal{G}_1) \in \mathcal{D}_{\mathcal{G}_1}, \\ D(\mathcal{G}_2) \in \mathcal{D}_{\mathcal{G}_2}}} EMED(S(D(\mathcal{G}_1)), S(D(\mathcal{G}_2)))$$

Theorem 3 allows us to define FGTED as the minimum EMED between flow decompositions in input graphs. To prove Theorem 3, we first explore the relationship between an $s - t$ path in AG^* and the strings spelled by the projections of this path onto $\mathcal{G}_1$ and $\mathcal{G}_2$.

Lemma 1. Given an s - t path $p \in D(AG^*)$, let $s_1 = S(P_{(AG^*, \mathcal{G}_1)}(p))$ be the string spelled by projecting p onto $\mathcal{G}_1$, and $s_2 = S(P_{(AG^*, \mathcal{G}_2)}(p))$. Then for any $p \in D(AG^*)$,

$$\sum_{e \in p} cost(e) = ED(s_1, s_2).$$

Proof. We prove in two directions.

($\geq$ direction) We construct $A = align(s_1, s_2)$ from p . For each $e = ((u_1, u_2), (v_1, v_2)) \in p$:

(1) if $u_1 = v_1$, add $c = [l_{(v_2)}^-]$ to A , (2) if $u_2 = v_2$, add $c = [l_{(v_1)}^+]$ to A , (3) else, add $c = [l_{(v_2)}^+]$ to A .

By definition of an alignment graph, $cost(e) = cost(c)$ in for all e , and therefore $cost(A) = \sum_{c \in A} cost(c) = \sum_{e \in p} cost(e)$. Since edit distance minimizes the cost of edit operations, $cost(A) = cost(p) \geq ED(s_1, s_2)$.

($\leq$ direction) We construct p' from $A^* = align(s_1, s_2)$ such that $cost(A^*) = ED(s_1, s_2)$. The procedure is similar as above — for each pair of adjacent entries in A^* , add corresponding edge to p' . Then $cost(p') = cost(A^*) = ED(s_1, s_2)$.

Let $AG' = AG^* \setminus p \cup p'$. Both p and p' can be found in AG , and both p and p' can be constructed by the alignment of the same pair of strings. Therefore, AG' is also a valid flow network and a feasible solution to FGTED. Since AG^* is the optimal solution to FGTED, $cost(AG^*) \leq cost(AG')$, and

$$\begin{aligned} cost(AG^*) - cost(AG') &\leq 0 \\ \Rightarrow w(p) \cdot (cost(p) - cost(p')) &\leq 0 \\ \Rightarrow w(p) \cdot (cost(p) - ED(s_1, s_2)) &\leq 0 \\ \Rightarrow cost(p) &\leq ED(s_1, s_2). \end{aligned}$$

We have shown that the cost of an $s - t$ path in AG^* is equal to the edit distance between its projections onto input graphs. Using this lemma, we can transform an optimal FGTED solution into an EMED solution.

Given an optimal FGTED solution, AG^* , let the set of possible flow decompositions of AG^* be $\mathcal{D}_{AG^*}$. Let $D(AG^*)$ be one of the flow decompositions that is a set of weighted $s - t$ paths. We can construct heterogeneous string sets $\mathcal{S}_1^*$ and $\mathcal{S}_2^*$ by projecting paths in $D(AG^*)$ to $\mathcal{G}_1$ and $\mathcal{G}_2$. Formally, $\mathcal{S}_1^* = \{S(P_{(AG^*, \mathcal{G}_1)}(p)) \mid p \in D(AG^*)\}$ and $\mathcal{S}_2^* = \{S(P_{(AG^*, \mathcal{G}_2)}(p)) \mid p \in D(AG^*)\}$.

Lemma 2. Given $\mathcal{S}_1^*$ and $\mathcal{S}_2^*$ obtained from any decomposition $D(AG^*) \in \mathcal{D}_{AG^*}$,

$$EMED(\mathcal{S}_1^*, \mathcal{S}_2^*) = cost(AG^*),$$

where $cost(AG^*)$ is sum of edge costs in the solution alignment graph to FGTED.

Proof. We prove in two directions.

($\leq$ direction) We construct a mapping M between strings in $\mathcal{S}_1^*$ and $\mathcal{S}_2^*$ from the decomposition $D(AG^*)$, where $M(s_i, s_j)$ is the portion of $s_i \in \mathcal{S}_1$ and $s_j \in \mathcal{S}_2$ that are aligned. For each $p \in D(AG^*)$, we obtain s_1 and s_2 as strings constructed from projections of p onto $\mathcal{G}_1$ and $\mathcal{G}_2$ and increment the weight of mapping $M(s_1, s_2)$ by $w(p)$. After iterating through all paths in $D(AG^*)$, the cost of M is

$$\begin{aligned} \text{cost}(M) &= \sum_{(s_1, s_2) \in M} M(s_1, s_2) \cdot \text{ED}(s_1, s_2) \\ &= \sum_{p \in D(AG^*)} w(p) \cdot \text{cost}(p) = \text{cost}(AG^*). \end{aligned}$$

M is also a feasible solution to the LP formulation of EMED. Since EMED minimizes the cost of mapping between $\mathcal{S}_1^*$ and $\mathcal{S}_2^*$, $\text{EMED}(\mathcal{S}_1^*, \mathcal{S}_2^*) \leq \text{cost}(AG^*)$.

($\geq$ direction) We construct a valid flow network, AG' using an optimal solution to $\text{EMED}(\mathcal{S}_1^*, \mathcal{S}_2^*)$. For each pairing (s_i, s_j) for $s_i \in \mathcal{S}_1$ and $s_j \in \mathcal{S}_2$, we obtain its weight w and cost c from the EMED solution. Let $A = \text{align}(s_i, s_j)$ be an optimal alignment under edit distance, and $\text{cost}(A) = c$. We then add a path corresponding to A with weight w in AG' . This follows the same procedure in the proof of Lemma 1. After adding all paths, we obtain AG' with $\text{cost}(AG') = \text{EMED}(\mathcal{S}_1^*, \mathcal{S}_2^*)$. Since $\text{cost}(AG^*)$ is minimized by FGTEd, $\text{cost}(AG') \geq \text{cost}(AG^*) \Rightarrow \text{EMED}(\mathcal{S}_1^*, \mathcal{S}_2^*) \leq \text{cost}(AG^*)$.

Lemma 2 provides a transformation algorithm between optimal solutions to EMED and solutions to FGTEd. Using Lemma 2, we can show that the EMED between $\mathcal{S}_1^*$ and $\mathcal{S}_2^*$ constructed from any decomposition in AG^* is equal to the decompositions of $\mathcal{G}_1$ and $\mathcal{G}_2$ that are closest in terms of EMED.

Lemma 3. Given $\mathcal{S}_1^*$ and $\mathcal{S}_2^*$ obtained from any decomposition $D(AG^*) \in \mathcal{D}_{AG^*}$,

$$\text{EMED}(\mathcal{S}_1^*, \mathcal{S}_2^*) = \min_{\substack{D(\mathcal{G}_1) \in \mathcal{D}_{\mathcal{G}_1}, \\ D(\mathcal{G}_2) \in \mathcal{D}_{\mathcal{G}_2}}} \text{EMED}(S(D(\mathcal{G}_1)), S(D(\mathcal{G}_2)))$$

Proof. In Lemma 2, $\mathcal{S}_1^*$ and $\mathcal{S}_2^*$ can be constructed from decomposing $\mathcal{G}_1$ and $\mathcal{G}_2$. Suppose for contradiction that there exists a decomposition that constructs string sets $\mathcal{S}_1'$ and $\mathcal{S}_2'$, such that $\text{EMED}(\mathcal{S}_1', \mathcal{S}_2') > \text{EMED}(\mathcal{S}_1^*, \mathcal{S}_2^*)$. Following the procedure in the proof of Lemma 2, we can construct a feasible solution to FGTEd with cost equal to $\text{EMED}(\mathcal{S}_1', \mathcal{S}_2')$, which is less than $\text{cost}(AG^*) = \text{EMED}(\mathcal{S}_1^*, \mathcal{S}_2^*)$. This contradicts with the assumption that FGTEd minimizes $\text{cost}(AG^*)$.

Theorem 3 is therefore true because of Lemma 2 and 3. Using Theorem 3, we are able to prove the first inequality in Theorem 2 with Lemma 4.

Lemma 4. Given heterogeneous string sets $\mathcal{S}_1$ and $\mathcal{S}_2$ and genome graphs representing these string sets, $\mathcal{G}_1 = G(\mathcal{S}_1)$ and $\mathcal{G}_2 = G(\mathcal{S}_2)$, $\text{FGTEd}(\mathcal{G}_1, \mathcal{G}_2) \leq \text{EMED}(\mathcal{S}_1, \mathcal{S}_2)$.

Proof. Given Theorem 3, FGTEd finds flow decomposition in $\mathcal{D}_{\mathcal{G}_1}$ and $\mathcal{D}_{\mathcal{G}_2}$ that minimizes the EMED between them. Since $\mathcal{S}_1$ and $\mathcal{S}_2$ can be constructed from a flow decomposition in $\mathcal{D}_{\mathcal{G}_1}$ and $\mathcal{D}_{\mathcal{G}_2}$, respectively, this lemma is true.

4.2 Genome Graph Expressiveness

A genome graph typically can represent more than one set of strings. We name the collection of string sets representable by a genome graph the *string set universe* of that genome graph, or $SU(\mathcal{G})$. Using Theorem 3, we

can say that FGTEd finds two sets of strings in the string set universe of $\mathcal{G}$ that are closest in the metric space of EMED. We define the expressiveness of a genome graph as the diameter of its string set universe, which is the maximum EMED between the string sets in $SU(\mathcal{G})$.

Definition 10 (String Set Universe Diameter (SUD)). Given a genome graph $\mathcal{G}$,

$$\text{SUD}(\mathcal{G}) = \max_{\mathcal{S}_a, \mathcal{S}_b \in SU(\mathcal{G})} \text{EMED}(\mathcal{S}_a, \mathcal{S}_b)$$

4.2.1 String Set Universe Diameter as an Upper Bound on Deviation of FGTEd from EMED

The string set universe diameter gives one measure of the size of $SU(\mathcal{G})$, and it can also be used to characterize the deviation of FGTEd from EMED.

Recall that $\mathcal{S}_1^*$ and $\mathcal{S}_2^*$ are string sets obtained from a decomposition $D(AG^*)$, and that $\text{EMED}(\mathcal{S}_1^*, \mathcal{S}_2^*) = \text{FGTEd}(G(\mathcal{S}_1), G(\mathcal{S}_2))$, where $\mathcal{S}_1$ and $\mathcal{S}_2$ are true string sets. From Theorem 2, we have that $\text{EMED}(\mathcal{S}_1, \mathcal{S}_2) \geq \text{EMED}(\mathcal{S}_1^*, \mathcal{S}_2^*)$. We can bound the deviation of $\text{EMED}(\mathcal{S}_1^*, \mathcal{S}_2^*)$ from $\text{EMED}(\mathcal{S}_1, \mathcal{S}_2)$ using triangle inequalities.

Lemma 5. Given string sets $\mathcal{S}_1$ and $\mathcal{S}_2$ and genome graphs $\mathcal{G}_1 = G(\mathcal{S}_1)$ and $\mathcal{G}_2 = G(\mathcal{S}_2)$,

$$\begin{aligned} &\text{EMED}(\mathcal{S}_1, \mathcal{S}_2) - \text{FGTEd}(\mathcal{G}_1, \mathcal{G}_2) \\ &\leq \min_{\substack{D(AG^*) \in \mathcal{D}_{AG^*} \\ \mathcal{S}_1^* = f_1(D(AG^*)) \\ \mathcal{S}_2^* = f_2(D(AG^*))}} (\text{EMED}(\mathcal{S}_1, \mathcal{S}_1^*) + \text{EMED}(\mathcal{S}_2, \mathcal{S}_2^*)), \end{aligned} \quad (5)$$

where AG^* is the solution obtained from FGTEd($\mathcal{G}_1, \mathcal{G}_2$).

Proof. Both edit distance and EMD are metrics (Rubner et al., 2000; Levenshtein et al., 1966), which means that triangle inequality holds for EMED between strings. Therefore, for any string sets $\mathcal{S}_1^*$ and $\mathcal{S}_2^*$,

$$\begin{aligned} \text{EMED}(\mathcal{S}_1, \mathcal{S}_1^*) + \text{EMED}(\mathcal{S}_1^*, \mathcal{S}_2) &\geq \text{EMED}(\mathcal{S}_1, \mathcal{S}_2) \\ \text{EMED}(\mathcal{S}_2, \mathcal{S}_2^*) + \text{EMED}(\mathcal{S}_1^*, \mathcal{S}_2^*) &\geq \text{EMED}(\mathcal{S}_1^*, \mathcal{S}_2) \end{aligned}$$

Combining two inequalities, we have

$$\begin{aligned} &\text{EMED}(\mathcal{S}_1^*, \mathcal{S}_2^*) = \text{FGTEd}(\mathcal{G}_1, \mathcal{G}_2) \\ &\geq \text{EMED}(\mathcal{S}_1, \mathcal{S}_2) - (\text{EMED}(\mathcal{S}_1, \mathcal{S}_1^*) + \text{EMED}(\mathcal{S}_2, \mathcal{S}_2^*)) \\ &\Rightarrow \text{EMED}(\mathcal{S}_1, \mathcal{S}_2) - \text{FGTEd}(\mathcal{G}_1, \mathcal{G}_2) \\ &\leq \text{EMED}(\mathcal{S}_1, \mathcal{S}_1^*) + \text{EMED}(\mathcal{S}_2, \mathcal{S}_2^*). \end{aligned} \quad (6)$$

The above inequality (6) holds for any string sets $\mathcal{S}_1^*$ and $\mathcal{S}_2^*$. To give a tight upper bound on the deviation, we take the minimum over all possible pairs of string sets constructed from decomposing AG^* that yields inequality (5).

Lemma 5 proves the second inequality of Theorem 2 thus completing the proof for Theorem 2 with Lemma 4.

The upper-bound found in Lemma 5 can be used as a factor that evaluates the pair-wise expressiveness of two genome graphs. While a genome graph may represent a large universe of string sets, as long as the true string set is close to the “best” string set in the pair-wise comparison, the deviation of FGTEd from EMED is small. We define this upper bound as the String Universe Co-Expansion Factor (SUCEF), which can be used to evaluate the discrepancy between FGTEd and EMED.

Definition 11 (String Universe Co-Expansion Factor (SUCEF)).

$$\begin{aligned} \text{SUCEF}(\mathcal{S}_1, \mathcal{S}_2, \mathcal{G}_1, \mathcal{G}_2) \\ = \min_{\substack{D(AG^*) \in \mathcal{D}_{AG^*} \\ \mathcal{S}_1^* = f_1(D(AG^*)) \\ \mathcal{S}_2^* = f_2(D(AG^*))}} (\text{EMED}(\mathcal{S}_1, \mathcal{S}_1^*) + \text{EMED}(\mathcal{S}_2, \mathcal{S}_2^*)), \end{aligned}$$

where AG^* is the solution to $\text{FGTED}(\mathcal{G}_1, \mathcal{G}_2)$.

On the other hand, finding SUCEF not only requires knowledge of true string sets $\mathcal{S}_1$ and $\mathcal{S}_2$, but SUCEF is also a pair-dependent measure that needs to be calculated for every pair of string sets and corresponding genome graphs. In order to characterize the effect of the expressiveness of individual genome graphs, we derive another upper bound on the deviation of FGTED from EMED using the string set universe diameters.

The sum of string set universe diameters of two genome graphs is an upper bound on SUCEF of these graphs and any two sets of strings they represent.

Lemma 6. *Given two genome graphs $\mathcal{G}_1$ and $\mathcal{G}_2$ and two sets of strings $\mathcal{S}_1$ and $\mathcal{S}_2$ they represent,*

$$\begin{aligned} \text{EMED}(\mathcal{S}_1, \mathcal{S}_2) - \text{FGTED}(\mathcal{G}_1, \mathcal{G}_2) &\leq \text{SUCEF}(\mathcal{S}_1, \mathcal{S}_2, \mathcal{G}_1, \mathcal{G}_2) \\ &\leq \text{SUD}(\mathcal{G}_1) + \text{SUD}(\mathcal{G}_2). \end{aligned}$$

Proof. Both $\mathcal{S}_1$ and $\mathcal{S}_1^*$ are represented by $\mathcal{G}_1$ and belong to $\text{SU}(\mathcal{G}_1)$. Therefore, by definition of string set universe diameter, $\text{EMED}(\mathcal{S}_1, \mathcal{S}_1^*) \leq \text{SUD}(\mathcal{G}_1)$ as the diameter maximizes the distance between any pair of strings represented by the genome graph. The same holds for $\text{EMED}(\mathcal{S}_2, \mathcal{S}_2^*) \leq \text{SUD}(\mathcal{G}_2)$.

Using Lemma 6, we can bound the deviation of FGTED from EMED using the expressiveness of individual genome graphs even when we do not have the knowledge of ground truth string sets. In practice, we can construct genome graphs using known sequences from the species of interest and form a training set. Using the training set, we can learn the relationship between SUDs and the deviation of FGTED from EMED, and then empirically estimate the anticipated discrepancy between FGTED and EMED. In the following sections, we show that we can improve FGTED using SUDs to obtain reduced anticipated deviation from EMED and stronger correlation with EMED.

5 Practically Correcting the Discrepancy between FGTED and EMED

5.1 Estimating String Set Universe Diameters

The string set universe diameter of a genome graph can be estimated by sampling flow decompositions of the graph. To sample a flow decomposition, we first sample one $s - t$ path. At each node u , we choose the neighbor v with the highest edge weight $w(u, v)$ with probability 0.5 and randomly choose a neighbor otherwise. After sampling a path, we send flow that is equal to the minimum edge weight on that path and produce the residual graph by subtracting the flow from edge weights on that $s - t$ path. We repeat this process on the residual graph until all edge weights are zero. This process assumes that the input genome graphs are acyclic to ensure all edge capacities (weights) are satisfied. If a genome graph is cyclic, e.g. de Bruijn graphs, string sets from $\text{SU}(\mathcal{G}_1)$ can be obtained by sampling Eulerian cycles in the genome graph, and each string in the string set is obtained by segmenting the sampled Eulerian cycle at source and sink nodes. After sampling 50 pairs of flow decompositions, we construct string sets from sampled flow decompositions and calculate pairwise EMED. We then obtain the highest pairwise EMED and use it as the estimated diameter.

5.2 Correcting FGTED Using String Set Universe Diameters

Using the sum of SUDs, we empirically estimate the deviation of FGTED from EMED with a linear regression model. We denote the deviation of FGTED from EMED by $\text{deviation}(\mathcal{S}_1, \mathcal{S}_2, \mathcal{G}_1, \mathcal{G}_2)$, which is computed as $|\text{EMED}(\mathcal{S}_1, \mathcal{S}_2) - \text{FGTED}(\mathcal{G}_1, \mathcal{G}_2)|$. The linear regression model, LR , has the following form

$$\begin{aligned} \text{deviation}(\mathcal{S}_1, \mathcal{S}_2, \mathcal{G}_1, \mathcal{G}_2) \\ = a \cdot (\text{SUD}(\mathcal{G}_1) + \text{SUD}(\mathcal{G}_2)) + b = LR(\text{SUD}(\mathcal{G}_1) + \text{SUD}(\mathcal{G}_2)), \end{aligned}$$

where a is the coefficient of the model and b is the intercept. The fitted model will minimize the mean squared error between predicted deviation and true deviation in the training set.

The corrected FGTED for each pair of graphs is calculated using the learned linear regression model as follows.

$$\begin{aligned} \text{correctedFGTED}(\mathcal{G}_1, \mathcal{G}_2) \\ = \text{FGTED}(\mathcal{G}_1, \mathcal{G}_2) + LR(\text{SUD}(\mathcal{G}_1) + \text{SUD}(\mathcal{G}_2)) \end{aligned}$$

The deviation of corrected FGTED from EMED has the same form as the deviation of uncorrected FGTED from EMED.

5.3 Data

We evaluate the use of string set universe diameters on two sequence sets:

1. **Simulated T-Cell Receptor (TCR) Repertoire.** We simulate 50 sets of TCR sequences and assign weights to each sequence using reference gene sequences of V, D and J genes from Immunogenetics (IMGT) V-Quest sequence directory (Lefranc and Lefranc, 2001). The number of sequences in each set varies from 2 to 5. We then generate 225 pairs of TCR string sets. Each TCR sequence is about 300 base pairs long. See Supplementary Materials for detailed simulation process.
2. **Hepatitis B Virus (HBV) Genomes.** We collect 9 sets of HBV genomes from three hosts — humans, bats and ducks — from the NCBI virus database (Hatcher *et al.*, 2017). We build 36 pairs of HBV string sets. See Supplementary Materials for detailed string set construction process.

We construct a partial order MSA graph on each string set (Lee *et al.*, 2002). We first conduct multiple sequence alignment (MSA) for each string set using Clustal Omega (Sievers *et al.*, 2011). Then for each column of the MSA, we create a node for each unique character and add an edge between two nodes if the characters in node labels are adjacent in the input strings at that column. For each consecutive stretch of gap characters, no nodes are created, but an edge is added between flanking columns of the stretch of gaps. We also create a source node and a sink node that are connected to nodes representing the first and last characters of the input strings. The MSA graphs created in this process are all acyclic. We compute FGTED on MSA graphs by adding sink-to-source edges.

We also construct a de Bruijn graph (Pevzner *et al.*, 2001) with k-mer size equal to 4 on TCR sequence sets, which we refer to as dBG4 in the following sections. This k-mer size is reasonable as compared to the average lengths of TCR sequences which is 350 base pairs and allows us to experiment with graphs that are expected to have higher expressiveness. In dBG4, each node corresponds to a k-mer, $S[i : i + k]$, where S is a string from the ground truth string set, $\mathcal{S}$. Each edge corresponds to the overlap between two k-mers, $S[i : i + k]$ and $S[i + 1 : i + k + 1]$ for any $S \in \mathcal{S}$. In order to construct the alignment graph, we process the de Bruijn graphs such that each node represents one character. We add a source node and a

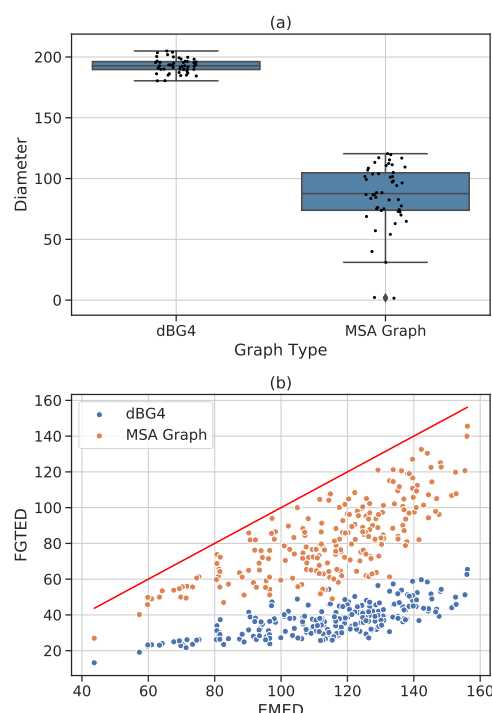


Fig. 3. Comparison between de Bruijn graphs and MSA graphs constructed with TCR sequence sets. (a) The distribution of diameters sampled in both types of graphs. Each box shows the quartiles of the distribution, and the whiskers show the rest of the distribution. Each black dot represent the diameter of one graph. (b) The correlation between FGTED and EMED with different types of graphs. The red line denotes equality between FGTED and EMED.

sink node to each dBG4 and connect them to nodes that represent the first and the last character in each string, respectively.

5.4 FGTED Deviates from EMED as the Expressiveness of the Genome Graph Increases

We compute EMED and FGTED on string set pairs and genome graph pairs. The alignment graphs are constructed using one thread, which on average takes 6 seconds for dBG4s, 8 seconds for each MSA graph on TCR sequences, 9.43 minutes for each MSA graph on HBV genomes. Optimization for LP with 10 threads takes on average 601 seconds for each dBG4, 1 hour for each MSA graph of TCR sequence sets and 4 hours for each MSA graph on HBV genomes (Figure S1).

We show that the deviation of FGTED from EMED is higher on genome graphs that are more expressive. We compare the FGTED computed on dBG4s and MSA graphs constructed with TCR sequences and the diameters of two types of graphs. dBG4 represents all sequences with the same 5-mer distributions as the ground truth sequences. Therefore, as expected, we observe larger sampled SUDs from dBG4 than the MSA graphs (Figure 3(a)). The deviation of FGTED from EMED is also larger with dBG4s than the MSA graphs (Figure 3(b)). This further illustrates the effect of graph construction approaches on the resulting expressiveness.

5.5 Corrected FGTED More Accurately Estimates Distance Between Unseen String Sets Encoded With Genome Graphs

For each pair of string sets, we obtain the deviation of FGTED from EMED and sum of estimated SUDs. We fit three linear regression models,

LR_{dBG4} , LR_{TCR} and LR_{HBV} , to predict deviation from sum of SUDs on simulated TCR sequences and HBV genomes of different types of graphs separately.

We evaluate the corrected and uncorrected FGTED by performing Pearson correlation experiments. We fit LR models on half of the data and compute the corrected FGTED on the other half as the test set. We evaluate the correlation between corrected and uncorrected FGTED and EMED on the test set. Two-tail P-values are calculated for each correlation experiment to test for non-correlation.

The LR models are evaluated with 10-fold cross validation. We randomly permute and split data into 10 equal parts. In each of the 10 iterations, we use one part as the test set and the rest as the training set. An average deviation is calculated across all iterations.

FGTED	Pearson Correlation		
	TCR (dBG4)	TCR (MSA Graph)	HBV (MSA Graph)
Uncorrected	0.75	0.74	0.99
Corrected	0.68	0.90	0.99

Table 1. Pearson correlation between EMED and corrected and uncorrected FGTED on simulated TCR and HBV sequences. Pearson correlation is calculated on a held-out set of data for both simulated TCR and HBV that consist of 50% of data, and LR model is fit on the other half.

In Table 1 and Table 2, we show that using string set universe diameters, we are able to improve the correlation between FGTED and EMED on MSA graphs of both the simulated TCR sequences and HBV genomes. On dBG4s, the correlation is reduced slightly by the correction. All Pearson correlation experiments are statistically significant with P-values < 0.01 . On HBV genomes, since the correlation between uncorrected FGTED and EMED is approaching 1, no significant improvement is observed. On the other hand, significant reduction in average deviation is observed on both types of data. We are able to reduce the average deviation from 77.29 to 19.08 on de Bruijn Graphs with TCR sequences, from 32.74 to 9.13 on MSA graphs containing simulated TCR sequences and from 140.12 to 54.87 on HBV genomes.

FGTED	Average Deviation		
	TCR (dBG4)	TCR (MSA Graph)	HBV (MSA Graph)
Uncorrected	77.29	32.74	140.12
Corrected	19.08	9.13	54.87

Table 2. Average deviation of corrected and uncorrected FGTED from EMED on simulated TCR and HBV sequences. The average deviation is calculated over a 10-fold cross validation of the LR model.

One caveat of using SUDs for correcting distances between genome graphs is that this correction is not guaranteed to always improve the distance. Given two string sets, there is usually an adversarial worst case where adjusting the distance using this approach reduces the accuracy in estimating string sets distances. When EMED between true string sets are small, the corrected FGTED may overestimate the EMED and result in a larger deviation. Nevertheless, we show that corrected FGTED reduces the anticipated deviation from EMED.

6 Discussion

A genome graph's string set universe diameter (SUD) provides information on the size and diversity of the represented string sets. We show that we

can use SUDs to practically characterize the discrepancy between FGTEd and EMED and to obtain a more accurate distance between unseen string sets encoded in genome graphs on average. While the results are obtained on short genomic sequences due to the high computational cost of FGTEd and GTED, this result is encouraging.

The corrected FGTEd can be used to compute a more accurate distance between heterogeneous samples represented by genome graphs in applications such as immune repertoire analysis and cancer subtyping. This opens up avenues for more comprehensive heterogeneous sample comparison methods. However, FGTEd, as well as GTED, is not scalable to mammalian genomes due to the quadratic size of the alignment graph and time it takes to solve the LP formulations. Algorithms that compute FGTEd faster or efficient approximation genome graph comparison methods (Minkin and Medvedev, 2020; Polevikov and Kolmogorov, 2019) are needed for comparing large heterogeneous string sets.

SUDs may also be used to characterize the diversity of strings represented by reference genome graphs that are used in sequence-to-graph alignment (Rautiainen and Marschall, 2020; Sirén *et al.*, 2020). In sequence-to-graph alignment, it is often desired that a more diverse set of strings than the original reference string set is represented by the graph. Here, SUDs could be used as a measure to control the right amount of variation in the string set universe of created genome graphs.

Another future direction is to use expressiveness as a regularization term in the objective function to construct better genome graphs. To ensure efficiency of genome graphs in storing sequences, we can construct genome graphs that minimize their sizes (Qiu and Kingsford, 2021; Pandey *et al.*, 2021). However, reducing the size of a genome graph may result in graphs that are highly expressive, and the distance between these genome graphs will deviate further from distances between true string sets. Adding a SUD term to the objective may address this problem.

7 Acknowledgements

This work was supported in part by the Gordon and Betty Moore Foundation’s Data-Driven Discovery Initiative [GBMF4554 to C.K.], by the US National Institutes of Health [R01GM122935], the US National Science Foundation [DBI-1937540] and the Carnegie Mellon University School of Computer Science Sansom graduate fellowship for computational cancer research to Y.Q.

Conflict of Interest: C.K. is a co-founder of Ocean Genomics, Inc.

References

- Almodaresi, F., Sarkar, H., Srivastava, A., and Patro, R. (2018). A space and time-efficient index for the compacted colored de Bruijn graph. *Bioinformatics*, **34**(13), i169–i177.
- Bolen, C. R., Rubelt, F., Vander Heiden, J. A., and Davis, M. M. (2017). The repertoire dissimilarity index as a method to compare lymphocyte receptor repertoires. *BMC Bioinformatics*, **18**(1), 1–8.
- Dilthey, A., Cox, C., Iqbal, Z., Nelson, M. R., and McVean, G. (2015). Improved genome inference in the MHC using a population reference graph. *Nature Genetics*, **47**(6), 682–688.
- Ebrahimipour Boroojeny, A., Shrestha, A., Sharifi-Zarchi, A., Gallagher, S. R., Sahinalp, S. C., and Chitsaz, H. (2020). Graph traversal edit distance and extensions. *Journal of Computational Biology*, **27**(3), 317–329.
- Garrison, E., Sirén, J., Novak, A. M., Hickey, G., Eizenga, J. M., Dawson, E. T., Jones, W., Garg, S., Markello, C., Lin, M. F., *et al.* (2018). Variation graph toolkit improves read mapping by representing genetic variation in the reference. *Nature Biotechnology*, **36**(9), 875–879.
- Hatcher, E. L., Zhdanov, S. A., Bao, Y., Blinkova, O., Nawrocki, E. P., Ostapchuk, Y., Schäffer, A. A., and Brister, J. R. (2017). Virus variation resource—improved response to emergent viral outbreaks. *Nucleic Acids Research*, **45**(D1), D482–D490.
- Holley, G. and Melsted, P. (2020). Bifrost: highly parallel construction and indexing of colored and compacted de Bruijn graphs. *Genome Biology*, **21**(1), 1–20.
- Iqbal, Z., Caccamo, M., Turner, I., Flicek, P., and McVean, G. (2012). De novo assembly and genotyping of variants using colored de Bruijn graphs. *Nature Genetics*, **44**(2), 226–232.
- Jain, C., Zhang, H., Gao, Y., and Aluru, S. (2020). On the complexity of sequence-to-graph alignment. *Journal of Computational Biology*, **27**(4), 640–654.
- Kingsford, C., Schatz, M. C., and Pop, M. (2010). Assembly complexity of prokaryotic genomes using short reads. *BMC Bioinformatics*, **11**(1), 21.
- Kusner, M., Sun, Y., Kolkin, N., and Weinberger, K. (2015). From word embeddings to document distances. In *International Conference on Machine Learning*, pages 957–966. PMLR.
- Lee, C., Grasso, C., and Sharlow, M. F. (2002). Multiple sequence alignment using partial order graphs. *Bioinformatics*, **18**(3), 452–464.
- Lee, H. and Kingsford, C. (2018). Kourami: graph-guided assembly for novel human leukocyte antigen allele discovery. *Genome Biology*, **19**(1), 1–16.
- Lefranc, M.-P. and Lefranc, G. (2001). *The immunoglobulin factsbook*. Academic press.
- Levenshtein, V. I. *et al.* (1966). Binary codes capable of correcting deletions, insertions, and reversals. In *Soviet Physics Doklady*, volume 10, pages 707–710. Soviet Union.
- Levina, E. and Bickel, P. (2001). The Earth Mover’s distance is the mallows distance: Some insights from statistics. In *Proceedings Eighth IEEE International Conference on Computer Vision. ICCV 2001*, volume 2, pages 251–256. IEEE.
- Li, H., Feng, X., and Chu, C. (2020). The design and construction of reference pangenome graphs with minigraph. *Genome Biology*, **21**(1), 1–19.
- Mangul, S. and Koslicki, D. (2016). Reference-free comparison of microbial communities via de Bruijn graphs. In *Proceedings of the 7th ACM International Conference on Bioinformatics, Computational Biology, and Health Informatics*, pages 68–77.
- Minkin, I. and Medvedev, P. (2020). Scalable pairwise whole-genome homology mapping of long genomes with BubbZ. *IScience*, **23**(6), 101224.
- Minkin, I., Pham, S., and Medvedev, P. (2017). TwoPaCo: An efficient algorithm to build the compacted de Bruijn graph from many complete genomes. *Bioinformatics*, **33**(24), 4024–4032.
- Morris, L. G., Riaz, N., Desrichard, A., Şenbabaoğlu, Y., Hakimi, A. A., Makarov, V., Reis-Filho, J. S., and Chan, T. A. (2016). Pan-cancer analysis of intratumor heterogeneity as a prognostic determinant of survival. *Oncotarget*, **7**(9), 10051.
- Pandey, P., Gao, Y., and Kingsford, C. (2021). VariantStore: an index for large-scale genomic variant search. *Genome Biology*, **22**(1), 1–25.
- Paten, B., Diekhans, M., Earl, D., John, J. S., Ma, J., Suh, B., and Haussler, D. (2011). Cactus graphs for genome comparisons. *Journal of Computational Biology*, **18**(3), 469–481.
- Paten, B., Novak, A. M., Eizenga, J. M., and Garrison, E. (2017). Genome graphs and the evolution of genome inference. *Genome Research*, **27**(5), 665–676.
- Paten, B., Eizenga, J. M., Rosen, Y. M., Novak, A. M., Garrison, E., and Hickey, G. (2018). Superbubbles, ultrabubbles, and cacti. *Journal of Computational Biology*, **25**(7), 649–663.
- Pevzner, P. A., Tang, H., and Waterman, M. S. (2001). An Eulerian path approach to DNA fragment assembly. *Proceedings of the National Academy of Sciences of USA*, **98**(17), 9748–9753.
- Polevikov, E. and Kolmogorov, M. (2019). Synteny paths for assembly graphs comparison. In *19th International Workshop on Algorithms in Bioinformatics (WABI 2019)*. Schloss Dagstuhl-Leibniz-Zentrum fuer Informatik.
- Qiu, Y. and Kingsford, C. (2021). Constructing small genome graphs via string compression. *Bioinformatics*, **37**(Supplement 1), i205–i213.
- Rautiainen, M. and Marschall, T. (2020). GraphAligner: rapid and versatile sequence-to-graph alignment. *Genome Biology*, **21**(1), 1–28.
- Rubner, Y., Tomasi, C., and Guibas, L. J. (2000). The Earth Mover’s distance as a metric for image retrieval. *International Journal of Computer Vision*, **40**(2), 99–121.
- Sievers, F., Wilm, A., Dineen, D., Gibson, T. J., Karplus, K., Li, W., Lopez, R., McWilliam, H., Remmert, M., Söding, J., *et al.* (2011). Fast, scalable generation of high-quality protein multiple sequence alignments using clustal omega. *Molecular Systems Biology*, **7**(1), 539.
- Sirén, J., Garrison, E., Novak, A. M., Paten, B., and Durbin, R. (2020). Haplotype-aware graph indexes. *Bioinformatics*, **36**(2), 400–407.
- Wasserstein, L. N. *et al.* (1969). Markov processes over denumerable products of spaces describing large systems of automata. *Problems of Information Transmission*, **5**(3), 47–52.
- Zerbino, D. R. and Birney, E. (2008). Velvet: algorithms for de novo short read assembly using de bruijn graphs. *Genome Research*, **18**(5), 821–829.
- Zhao, L., Lee, V. H., Ng, M. K., Yan, H., and Bijlsma, M. F. (2019). Molecular subtyping of cancer: current status and moving toward clinical applications. *Briefings in Bioinformatics*, **20**(2), 572–584.