



HashAlign: Hash-Based Alignment of Multiple Graphs

Mark Heimann^(✉), Wei Lee, Shengjie Pan, Kuan-Yu Chen, and Danai Koutra

Computer Science and Engineering, University of Michigan, Ann Arbor, USA
{mheimann, weile, jessepan, kyuchen, dkoutra}@umich.edu

Abstract. Fusing or aligning two or more networks is a fundamental building block of many graph mining tasks (e.g., recommendation systems, link prediction, collective analysis of networks). Most past work has focused on formulating *pairwise* graph alignment as an optimization problem with varying constraints and relaxations. In this paper, we study the problem of *multiple* graph alignment (collectively aligning multiple graphs at once) and propose HASHALIGN, an efficient and intuitive hash-based framework for network alignment that leverages structural properties and other node and edge attributes (if available) simultaneously. We introduce a new construction of LSH families, as well as robust node and graph features that are tailored for this task. Our method quickly aligns multiple graphs while avoiding the all-pairwise-comparison problem by expressing all alignments in terms of a chosen ‘center’ graph. Our extensive experiments on synthetic and real networks show that, on average, HASHALIGN is 2× faster and 10 to 20% more accurate than the baselines in *pairwise* alignment, and 2× faster while 50% more accurate in *multiple* graph alignment.

1 Introduction

Much of the data that is generated daily naturally form graphs, such as interactions between users in social media, communication via email or phone calls, question answering in forums, interactions between proteins, and more. Additionally, graphs may be inferred from non-network data [17]. For joint analysis, it is often desirable to fuse multiple graph data sources by finding the corresponding nodes across them. This task, known as graph alignment or matching, is the focus of our work. It is a core graph theoretical problem that has attracted significant interest, both in academia and industry, due to its numerous applications: identifying users in social networks [19], matching similar documents in lingual matching [4], brain graph alignment in neuroscience, protein-protein alignment [4, 5], chemical compound comparison, and more.

In many applications, the goal is to align *multiple* (more than two) networks at once. Most existing methods get as input two networks, so they handle multiple network alignment by expensively computing all pairwise alignments. In

M. Heimann and W. Lee—These authors contributed equally to this work.

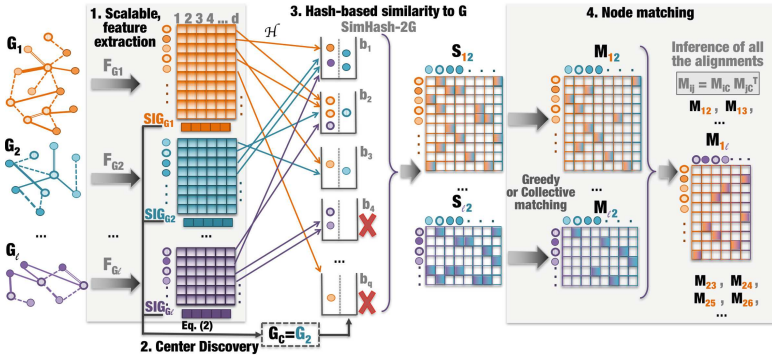


Fig. 1. Overview of proposed approach: HASHALIGN with input l undirected, weighted, attributed graphs (node/edge attributes are denoted with different shades/lines). The framework consists of four parts: (1) scalable, fast, robust, node-ID-invariant feature extraction per graph; (2) ‘center’ graph discovery, which is $G_C = G_2$ in this example; (3) efficient, hash-based similarity computation, S_{iC} , between each graph G_i and G_C (buckets with red crosses do not contribute any pairwise similarity computations, and thus help with efficiency); and (4) node matching computation to find at most one matching per node in M_{ij} .

this paper, we seek to devise an efficient method that collectively aligns multiple networks and can readily adapt to the existence or not of other node/edge information in addition to the graph topology, without increasing its complexity.

Problem 1 (Multiple Graph Alignment with Side Information). **Given** l graphs, $G_1(\mathcal{V}_1, \mathcal{E}_1), \dots, G_l(\mathcal{V}_l, \mathcal{E}_l)$, where $\mathcal{V}_i$ and $\mathcal{E}_i$ are the node and edge sets of graph G_i , respectively, *with or without node/edge attributes*, we seek to **find** the correspondence between their nodes efficiently, so that the input graphs are as close to each other as possible.

To solve this problem, we propose HASHALIGN, an unsupervised method that is based on three key ideas: (i) inferring the similarity between nodes in different graphs based on structural properties and node/edge attributes; (ii) leveraging Locality Sensitive Hashing (LSH) [6] to minimize the number of pairwise node comparisons; (iii) choosing a ‘center’ graph out of l input graphs to which to align all the others, thereby avoiding solving $\binom{l}{2}$ pairwise graph alignment problems (instead solving $l - 1$ alignments and quickly inferring the others by applying simple transformations in the form of sparse matrix multiplications). Figure 1 contains a pictorial overview. Our main contributions are:

- **Flexible Framework.** We propose an efficient and accurate hashing-based family of algorithms, HASHALIGN, which solves the multiple network alignment problem. Our method is general and can readily incorporate any available node and edge attributes. HASHALIGN can be used as a standalone

alignment method or provide its solution to initialize optimization problems for pairwise alignment (e.g., [4]).

- **Methods.** As part of our framework, we propose problem-specific choices of node and graph features, and introduce a new, robust construction of hash families.
- **Experiments.** We conduct extensive experiments on synthetic and real data, which show that HASHALIGN is 2–10× faster than the baselines that tackle either the multiple or pairwise alignment problem, while being equally or up to 50% more accurate.

For reproducibility, the code is available at <https://github.com/GemsLab/HashAlign.git>. Additional supplementary material is provided at <https://markheimann.github.io/papers/HashAlign-PAKDD18-full.pdf>.

2 Related Work

We review work that is relevant to our problem space and choices of techniques:

Graph Alignment. Scalable methods for pairwise graph alignment include a distributed, belief-propagation-based method for protein alignment [5], a message-passing algorithm for aligning sparse networks when some [4] or all [18] possible matchings are considered, alignment of bipartite networks [13, 14], and attributed graph alignment [20]. *Multiple* network alignment, however, poses a further scalability challenge. For instance, the recent optimization-based formulation of [16] solves a bipartite matching problem in $O(n^3)$ time using the Hungarian algorithm. Zhang and Yu [19] introduce the notion of transitivity between graphs to align social networks more scalably with some partial node matchings (anchor links) known *a priori*. Our method HASHALIGN preserves this notion of transitivity for any type of network and requires no anchor links.

Locality-Sensitive Hashing. This technique for efficient similarity search has been used to accelerate the well-known k -nearest neighbor algorithm, often offering theoretical and practical improvements even over sophisticated data structures such as k - d trees [3]. It has also found use in matching problems in other domains, such as ontology matching in information retrieval [8]. In our proposed method, we leverage LSH to efficiently find nodes that are similar. For networks, [12] uses MinHash to find sets of similar nodes in a *single* attributed graph by relying on the adjacency matrix as features, but this is not applicable to the graph alignment setting. Thus, we introduce node-ID invariant representations and adapt LSH to find similarities *across* networks. Our contribution is orthogonal to prior works: a framework for network alignment, HASHALIGN, in which we propose design choices geared toward our specific domain.

3 Proposed Formulation: Two-Graph Alignment

In this section, we first introduce the alignment problem for two graphs. We then describe our proposed approach, and in the next section we extend it to

Table 1. Symbols and definitions. We use bold capital letters for matrices, bold lowercase letters for vectors and normal lowercase letters for scalars.

Symbols	Definitions
$G_p = (\mathcal{V}_p, \mathcal{E}_p)$	Graph p with vertex set $\mathcal{V}_p$ and edge set $\mathcal{E}_p$
$ \mathcal{V}_p = n_p, \mathcal{E}_p = m_p$	Number of nodes and edges in graph G_p , resp.
$\mathbf{s}_{G_p}(v)$	$1 \times d_s$ vector of the structural invariants (e.g., PageRank) for node $v \in G_p$
$\mathbf{a}_n G_p(v), \mathbf{a}_e G_p(v)$	$1 \times d_{an}$ vector of node/edge attributes for node $v \in G_p$
$\mathbf{A}_{nG_p}, \mathbf{A}_{eG_p}$	The stacked node/edge attr. matrices of size $n_p \times d_{an}$ and $n_p \times n_p \times d_{ae}$
$d = d_s + d_{an} + d_{ae}$	Total number of (structural, node, and edge) features
$\mathbf{f}_{G_p}(v), \mathbf{F}_{G_p}$	$1 \times d$ all-feature vec. for node $v \in G_p$ and the resp. stacked $n_p \times d$ mat.
SIG_{G_p}	$1 \times 5d$ graph ‘signature’ vector representing graph G_p
$d(G_i, G_j)$	Distance between graphs G_i and G_j
$\mathbf{S}_{ij}$	Sparse $n_i \times n_j$ similarity matrix between graph G_i and G_j
$\mathbf{M}_{ij}$	$n_i \times n_j$ alignment between graph G_i and G_j
b_i	Bucket i (hashing)
Z	Number of bands (hashing)

the multiple graph alignment problem. Table 1 summarizes the main notations used in our analysis.

3.1 Definition: Relaxed Two-Graph Alignment Problem

The typical graph alignment problem aims to find a one-to-one matching between the nodes of two input graphs. This problem is important, but in many applications it suffices to solve a relaxed version of it: finding a *small set of nodes* that are *likely* to correspond to a given node. Thus, we relax the original alignment problem as follows:

Problem 2 (Relaxed two-graph alignment). Given two graphs, $G_1(\mathcal{V}_1, \mathcal{E}_1)$ and $G_2(\mathcal{V}_2, \mathcal{E}_2)$, which may be (un)directed, (un)weighted and attributed / plain, **we seek to efficiently find a sparse**, weighted bipartite graph $G_S = (\mathcal{V}_1 \cup \mathcal{V}_2, \mathcal{E}_S)$ with edges representing potential matching pairs and being weighted by the likelihood of the match:

$$\forall \text{ potential match } (u, v), u \in \mathcal{V}_1, v \in \mathcal{V}_2, \exists e \in \mathcal{E}_S : w_e = \text{sim}(u, v)$$

and $|\mathcal{E}_S| < \alpha \cdot \max\{n_1, n_2\}$ where $\alpha \in \mathbb{Z}$ ($\alpha > 1$) controls the density of G_S .

To make sure that nodes are matched only to a few of their closest counterparts, the main requirement in Problem 2 is that G_S is sparse, i.e., $|\mathcal{E}_S| \ll n_1 \times n_2$. Most graph alignment methods find 1-1 matchings between the vertex sets [4, 7], and a few approaches relax the requirements of the typical optimization problem to find probabilistic matchings [14], but each method targets a different type of graph (e.g., unipartite, undirected) and most of them rely only on the network structure. In this work we propose a different, intuitive similarity-based approach that encompasses all these settings, leveraging a suitably rich node representation to achieve superior accuracy.

A naive similarity-based method is to: (i) compute all the pairwise similarities between the nodes in G_1 and G_2 , and (ii) keep only the edges with similarities greater than a user-specified threshold. Although this approach results in a sparse graph G_S , it has several drawbacks. First, it is computationally expensive, since it computes *all* $n_1 \times n_2$ pairs of similarities and later applies the threshold for edge filtering. Second, the threshold is arbitrary and affects the potential node matchings significantly. Third, it is not clear how to choose the ‘right’ node representation for similarity computations. Our proposed approach uses hashing to overcome all these issues.

3.2 Node Representation: Handling Node and Edge Attributes

Our framework, HASHALIGN, requires a vector representation of each node. We want these to be comparable across graphs and also leverage node/edge attributes seamlessly (Fig. 2).

We propose to represent each node u with a vector $\mathbf{f}(u)$ of structural features and node/edge attributes (if available). The benefit of this representation is that it can be adjusted to the type of graphs and available information *without any* changes in the problem formulation. Furthermore, it is *node-ID invariant* and can thus be meaningfully compared across graphs. This is not true of representation learning methods like DeepWalk and node2vec, which sample context nodes by their IDs with random walks [9] and thus are not applicable to our multi-network setting [10]. Specifically, in Step 1 of our framework (Fig. 1), we concatenate d_s structural features, d_{an} node attributes, and d_{ae} edge attributes:

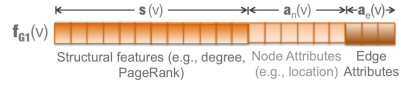


Fig. 2. Proposed feature-based, node-ID invariant representation of vertex v .

- **Structural features** $\mathbf{s} \in \mathbb{R}^{1 \times d_s}$. Examples include the so-called local features (e.g., degree variants) and egonet features. The egonet of node u is defined as the induced subgraph of u and its neighbors, and structural features specific to the egonet include its number of edges, its degree, and more. In addition to these features, we also consider features that combine locality with globality, such as PageRank and various types of centrality. We choose specific structural features that are most robust to noise (Sect. 5).
- **Node attributes** $\mathbf{a}_n \in \mathbb{R}^{1 \times d_{an}}$. If a graph contains node attributes, the node feature vectors $\mathbf{f}$ can be extended to include those. Numerical features can be simply concatenated with the structural features, while categorical attributes can be incorporated by using 1-hot encoding and concatenated to the previously formed feature vector.
- **Edge attributes** $\mathbf{a}_e \in \mathbb{R}^{1 \times d_{ae}}$. We propose converting numerical edge features to node attributes by applying an aggregate function $\xi : \mathcal{E}^{deg_u} \rightarrow \mathbb{R}$ (where the domain is the set of edges incident to $u \in \mathcal{V}$ and deg_u is its degree). Examples for $\xi()$ include sum, average, standard deviation, etc. For

categorical features, we propose to encode the distribution of values per feature. For example, if a feature has q possible values, then q entries with their frequencies will be concatenated with the previous features.

3.3 Proposed Hashing-Based Computation of Potential Matchings

Now we have, for each node u in graph G_p , a real-valued vector $\mathbf{f}_{G_p}(u) \in \mathbb{R}^d$ constructed as described in Sect. 3.2. We propose to use Locality Sensitive Hashing (LSH) [6] to find a small number of potential matchings between nodes across graphs (i.e., nodes with high similarity) scalably, without computing all pairs of $n_1 \times n_2$ similarities. In a nutshell, given a similarity function, LSH reduces the dimensionality of high-dimensional data while preserving their local similarities; that is, it efficiently maps similar data points (in our case, nodes) to the same buckets with high probability. Our proposed hashing approach for alignment takes as input the $\mathbf{F}_{G_1} \in \mathbb{R}^{n_1 \times d}$ feature matrix (with row-wise node representations) for G_1 and $\mathbf{F}_{G_2} \in \mathbb{R}^{n_2 \times d}$ for G_2 , and hashes them row-wise using an LSH family $\mathcal{H}$.

Definition 1 (LSH-2G). *Given $\mathcal{V}_1$ and $\mathcal{V}_2$, the nodes in graph G_1 and G_2 respectively, along with a similarity function $\phi : \mathbb{R}^d \times \mathbb{R}^d \rightarrow [0, 1]$, $\mathcal{H}$ is an LSH-2G family of hash functions such that the probability of two nodes $u, v \in \mathcal{V}_1$ hashing to the same bucket is equal to their similarity, and additionally the probability of two nodes $u \in \mathcal{V}_1$ and $v \in \mathcal{V}_2$ hashing to the same bucket is equal to their similarity: $\Pr[h(u) = h(v)] = \phi(\mathbf{f}_{G_1}(u), \mathbf{f}_{G_2}(v))$.*

We propose an LSH-2G family based on the standard measure of cosine similarity (with Euclidean distance in the supplementary material.) We introduce SimHash-2G, a modified version of SimHash [3] that is based on LSH-2G described above. SimHash-2G chooses K randomly generated column vectors $\{\mathbf{r}_1, \dots, \mathbf{r}_K\} \in \mathbb{R}^d$ that follow the standard Gaussian distribution (i.e., K random hyperplanes). The LSH-2G family consists of K hash functions: $h_k(u) = \text{sign}(\mathbf{f}_{G_p}(u) \cdot \mathbf{r}_k)$. Each of these projects node u on either side of the random hyperplane $\mathbf{r}_k$ (positive or negative sign). For random hyperplane k , the probability of two nodes $u \in \mathcal{V}_1$ and $v \in \mathcal{V}_2$ being mapped to the same bucket is $\Pr[h_k(u) = h_k(v)] = 1 - \frac{\theta_{uv}}{\pi}$, where $\theta_{uv} = \cos^{-1} \frac{\mathbf{f}_{G_1}(u) \cdot \mathbf{f}_{G_2}(v)}{\|\mathbf{f}_{G_1}(u)\|_2 \|\mathbf{f}_{G_2}(v)\|_2}$. The angle $1 - \frac{\theta_{uv}}{\pi}$ captures the proximity of u and v . SimHash-2G computes only the similarity for pairs of nodes according to our proposed SKD-construction (see below).

If a hash function $h_i \in \mathcal{H}$ maps two nodes to the same bucket, that indicates that they *could* be similar, but there is some probability of error. The technique of amplification creates a new LSH family $\mathcal{G}$ with hash function g defined over the functions in $\mathcal{H} = \{h_1, h_2, \dots, h_K\}$, in order to reduce that probability of error. A standard technique is AND-construction where $g(u) = g(v) \implies \forall i \ h_i(u) = h_i(v)$.

However, the AND-construction is too strict and may lead to many false negatives when finding node matchings. To ameliorate that we use the banding

technique: (i) we split each feature vector into Z equal bands, and (ii) per band z , we apply a corresponding LSH-2G family $\mathcal{H}_z$ using AND-construction. In each band, a node can fall into only one bucket, and thus collides with nodes in that same bucket (potential matchings). To handle the observed *skewed* distribution of nodes to buckets and guarantee that each node will have some potential matchings, we introduce the notion of ‘importance’ of a node collision within a band and propose the SKD-construction.

Definition 2 (Importance σ_{tot} of node collision). *Given two nodes $u \in \mathcal{V}_1$ and $v \in \mathcal{V}_2$, and an LSH-2G family $\mathcal{H} = \{h_1, \dots, h_K\}$ s.t. $\forall j \ h_j(u) = h_j(v)$ (i.e., both nodes are mapped to bucket $b_{\mathcal{H}}$), we define the importance of their collision based on $\mathcal{H}$ as the inverse of the size of the corresponding bucket: $\sigma_{\mathcal{H}}(u, v) = \frac{1}{|b_{\mathcal{H}}|}$. The total importance score of a node pair collision over all bands and their corresponding LSH families $\mathcal{H}' = \{\mathcal{H}_1, \dots, \mathcal{H}_Z\}$ is defined as: $\sigma_{tot}(u, v) = \sum_{\mathcal{H} \in \mathcal{H}'} \sigma_{\mathcal{H}}(u, v) \cdot \mathbb{1}_{h_j(u)=h_j(v), \forall h_j \in \mathcal{H}}$.*

Intuitively, the importance of a collision is higher if a few nodes are mapped to a bucket, as the bucket has higher discriminative power. The notion of importance tackles the skewness that we observe in the mapped nodes in the graph alignment setting. Based on this definition, we propose the SKD-construction (where SKD stands for SKewedD).

Definition 3 (SKD-construction). *Given $u \in \mathcal{V}_1$ and $v \in \mathcal{V}_2$, and LSH-2G families $\mathcal{H}' = \{\mathcal{H}_1, \dots, \mathcal{H}_Z\}$, a new family $\mathcal{G}$ with hash function g is based on SKD-construction:*

$$g(u) = g(v) \implies \sigma_{tot}(u, v) \in TOP_{\alpha}(u),$$

where $TOP_{\alpha}(u)$ is the set of top- α total importance scores $\sigma_{tot}(u, v')$ for $v' \in \mathcal{V}_2$, and α is the small factor that controls the density of G_S in Problem 2.

Intuitively, SKD-construction computes the pairwise similarities of nodes that collide often (but not always, like AND-construction) and have important collisions that manage to distinguish the nodes (i.e., it penalizes functions that lead to skewed results).

3.4 From Similarities to Matchings

As shown in Step 3 of Fig. 1, the hashing approach that we introduced returns a small number of high similarities between the nodes of graphs G_1 and G_2 , giving us an $n_1 \times n_2$ sparse matrix $\mathbf{S}$ with node similarities. Here we provide ways to use the similarity information in $\mathbf{S}$ to find the node matchings or correspondences $\mathbf{M} \in \mathbb{Z}^{n_1 \times n_2}$:

- **Greedy matching:** Assuming that the higher the similarity score, the more likely two nodes are to match [14, 20], we can greedily make independent decisions for the best match of each node in G_1 through a function $\chi : \mathcal{V}_1 \rightarrow \mathcal{V}_2$

s.t. $\chi(u) = \operatorname{argmax}_v \{\mathbf{S}_{uv}\}$. Since nodes are matched independently, this is very efficient and parallelizable, but may match more than one node in graph G_1 to the same node in G_2 . It is a preferred method for very large networks or networks of different sizes, and also when multiple potential matchings are desired. In the latter case, it can be trivially extended by updating the function $\chi()$ to return more top potential matchings (instead of only the best one).

- **Collective matching:** An alternative is to leverage existing approaches that find 1-to-1 matchings collectively, given a similarity matrix $\mathbf{S}$. In Sect. 5 we consider scalable options for doing so and study their trade-offs.

4 HASHALIGN: Multiple Graph Alignment

In this section, we extend our HASHALIGN framework to multiple graph alignment, extending the formal definition of the relaxed 2-graph alignment problem.

Problem 3 (Relaxed multiple graph alignment). Given a set of graphs, $\mathcal{G} = \{G_1(\mathcal{V}_1, \mathcal{E}_1), \dots, G_l(\mathcal{V}_l, \mathcal{E}_l)\}$, which may be (un)directed, (un)weighted and attributed / plain, we seek to efficiently **find** a sparse, weighted bipartite graph $G_{Sij} = (\mathcal{V}_i \cup \mathcal{V}_j, \mathcal{E}_{Sij})$ for each pair of graphs $\langle G_i, G_j \rangle$, s.t. $\mathcal{E}_{Sij}$ has the potential matching pairs between their vertex sets and the weights describe how likely the nodes are to match.

Efficient Computation. The key insight to reduce computation is to use one of the l graphs as the ‘baseline’ graph G_C and align the remaining $l-1$ graphs with it in parallel. This approach avoids computing $O(l^2)$ pairwise graph alignments, instead leading to $l-1$ matching matrices $\mathbf{M}_{2C}, \dots, \mathbf{M}_{lC}$ (w.l.o.g. we choose graph $G_C = G_1$ in our notation, but we will explain next the choice of G_C). Inspired by the idea of transitivity [19], which requires node matching consistency between pairs of graphs (Fig. 3), we efficiently infer the remaining matching matrices $\mathbf{M}_{ij}$ (where $i, j \neq C$) via sparse matrix multiplications (Step 4 in Fig. 1): $\mathbf{M}_{ij} = \mathbf{M}_{iC} \cdot \mathbf{M}_{jC}^T$.

Choice of G_C . To reduce the induced alignment errors and their propagation to the inferred matchings, we propose the ‘center’ graph (i.e., the graph in $\mathcal{G}$ with the minimum total distance from the remaining graphs) as the baseline graph G_C (Step 2 in Fig. 1):

$$\operatorname{argmin}_C \sum_j d(G_C, G_j) = \sum_j \|\operatorname{SIG}_{G_C} - \operatorname{SIG}_{G_j}\|_2,$$

where $d(G_C, G_j)$ is the distance between G_C and G_j , and SIG is a graph ‘signature’, which we create by applying an aggregate feature function $\xi()$ over a

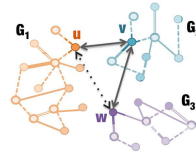


Fig. 3. Node matching consistency: If $u = v$ and $v = w$, then u should match to w (by transitivity).

Algorithm 1. HASHALIGN

Input: (1) $\mathcal{G}=\{G_1, G_2, \dots, G_l\}$; (2) [OPT] Per graph i , node/edge attr.
 $\mathbf{A}_{n_{G_i} \times d_{a_n}}^{n \times d_{a_n}} / \mathbf{A}_{e_{G_i} \times d_{a_e}}^{n \times n \times d_{a_e}}$

Output: A set of matching matrices $\{\mathbf{M}_{ij}\}$ for $i, j \in \{1, \dots, l\}$

```

1: /* STEPS 1&2: NODE REPRESENTATION AND CENTER DISCOVERY */
2: For  $G \in \mathcal{G}$  do
3:    $\mathbf{F}_G = \text{extractFeatures}(G, \mathbf{A}_G^{n \times d_{a_n}}, \mathbf{A}_G^{n \times n \times d_{a_e}})$  ▷ Sec. 3.2
4:  $G_C = \text{findCenter}(\xi(\mathbf{F}_{G1}), \xi(\mathbf{F}_{G2}), \dots, \xi(\mathbf{F}_{Gl}))$  ▷ Eq. (4) & aggregate function
    $\xi() = \text{SIG}$ 
5: /* STEP 3: HASH-BASED SIMILARITY (assuming  $q$  buckets in total) */
6:  $\{b_1, \dots, b_q\} = \text{SimHash-2G}(\mathbf{F}_{G1}, \dots, \mathbf{F}_{Gl})$  ▷ Sec. 3.3 (or EDHash-2G in Appendix B)
7:  $\{\mathbf{S}_{1C}, \mathbf{S}_{2C}, \dots, \mathbf{S}_{lC}\} = \text{computeSparseSimilarities}(b_1, \dots, b_q)$  ▷ SKD-construction
8: /* STEP 4: NODE MATCHING */
9:  $\{\mathbf{M}_{1C}, \mathbf{M}_{2C}, \dots, \mathbf{M}_{lC}\} = \text{GREEDY or COLLECTIVE}(\mathbf{S}_{1C}, \mathbf{S}_{2C}, \dots, \mathbf{S}_{lC})$  ▷ Sec. 3.4
10: For  $i, j \in \{1, \dots, l\}$  do
11:    $\mathbf{M}_{ij} = \mathbf{M}_{iC} \times \mathbf{M}_{jC}^T$  ▷ Sec. 4

```

graph’s nodes. In our work, we use the mean, median, standard deviation, skewness, and kurtosis of each of the d features, giving us a $5d$ -dimensional vector (shown in Step 1 of Fig. 1). Intuitively and empirically, the center graph being as close as possible the other graphs can make the center-based alignments more precise.

Hash-Based Similarity. After hashing all the feature-based node vectors of all the graphs in $\mathcal{G}$ as described in Sect. 3.3, we compute the similarity scores for possibly matching pairs of nodes according to the SKD-construction. We only compute the similarity between nodes in the center graph (right hand-side in the buckets in step 2 of Fig. 1) and nodes in the *peripheral*, or non-center, graphs (left hand-side).

Putting Everything Together. We propose HASHALIGN, a fast, hash-based, multiple graph alignment approach, which is described at a high level in Algorithm 1 (and pictorially in Fig. 1, where $Z = 1$ for simplicity.) It consists of four main steps: (i) node representation, (ii) ‘center’ graph identification, (iii) hash-based similarity, and (iv) node matching. In line 7 of Algorithm 1, SimHash-2G is applied to l graphs in parallel.

Computational Complexity of HASHALIGN. Our framework makes two main substitutions for computational savings. First, it replaces full pairwise similarity computations that are quadratic in the number of nodes with hashing in only $O(K \cdot n_p \cdot d)$ time for graph G_p with n_p nodes, if we use K hash functions on d -dimensional feature vectors. Second, it replaces all $\binom{l}{2}$ pairwise network alignments with only $l - 1$ pairwise network alignments to a center graph (chosen in $O(l^2 \cdot d)$ time), inferring the remainder with sparse matrix multiplications. More details are given in the supplementary material.

5 Experimental Analysis

In this section, we seek to answer the following questions: (1) How robust is our framework compared to baselines for different levels of noise in the graphs (both in the structure and node/edge attributes)? (2) How could HASHALIGN help existing alignment methods perform better and how could these help our method? (3) How do our methods scale when aligning multiple graphs collectively? We answer these questions on three datasets, described in Table 2. We also include additional experiments, such as a *sensitivity analysis* of HASHALIGN to different parameters, in the supplementary material.

Baselines. We consider 3 baseline methods commonly used in the literature: **NetAlign** [4], **Final** [20], and **IsoRank** [18]. We compare their performance against our method, HASHALIGN, where we infer alignments greedily from the hashing-based node similarities. The baselines accept a matrix $\mathbf{L}$ representing prior alignment information between the nodes of the original graphs. By default, we provide a thresholded similarity matrix based on the node attributes to assure good performance based on the attribute information, even for the baselines that are not formulated for it (**NetAlign** and **IsoRank**). We also consider two variants of HASHALIGN, namely HASHALIGN-NA and HASHALIGN-FN, which respectively use **NetAlign** and **Final** to infer alignments from the node similarities as the final step of HASHALIGN (Sect. 3.4).

Data. We evaluate our proposed algorithms on three datasets along with the synthetic data that we generated from them (via permutations and added noise, as in [14, 20]). Formally, given a graph G_1 with adjacency matrix $\mathbf{A}$, we create a noisy graph G_2 with matrix $\mathbf{B} = \mathbf{PAP}^\top$ (i.e., a permutation of itself), where $\mathbf{P}$ is a randomly generated permutation matrix (i.e., with one nonzero entry per row/column). Synthetic noise is applied to both graph structure and labels throughout our experiments to simulate real-world scenarios where the graphs are matchable but different. The noise level p indicates that with probability p , Gaussian noise with $std = 1$ is added to an edge weight; a binary edge label is flipped; or a categorical node/edge label value is changed.

Table 2. Description of real datasets.

Datasets	# Nodes	# Edges	Graph type	Labels	Description
Connectome [1]	941	9,622	Undirected	-	fMRI-inferred graphs
E-mail [2]	1,133	5,451	Undirected	5	Email communications
DBLP [20]	42,252	210,320	Undirected	1	Coauthorship network

Evaluation Metric. Following the literature, we compute the alignment accuracy as $\frac{\# \text{ correct matchings}}{\# \text{ total matchings}}$, where the total number of matchings between G_i and G_j is equal to the minimum number of nodes between the two graphs, $\min\{n_i, n_j\}$.

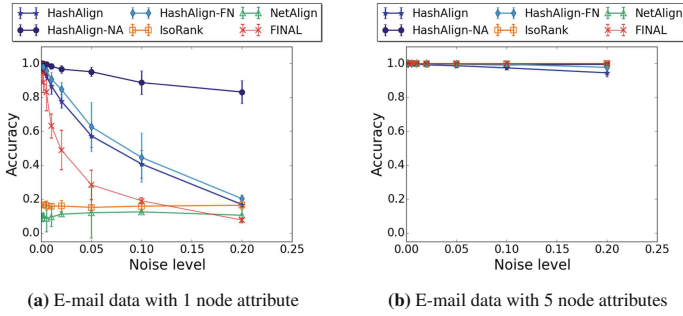


Fig. 4. E-mail dataset (experimental results on other datasets are similar): Effectiveness w.r.t. noise on both attributes and the graph structure. Methods based on the HASHALIGN framework achieve highest accuracy, particularly with limited node attribute information.

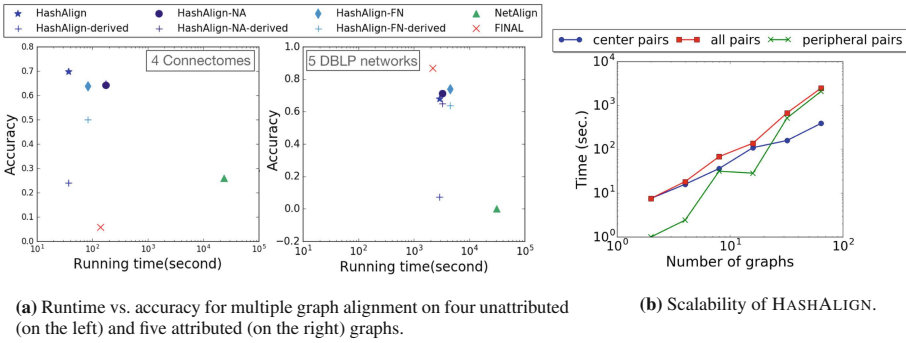


Fig. 5. (a) HASHALIGN has stable efficiency across different kinds of networks. **Final** is fast, but its accuracy is subject to whether node/edge labels exist. (b) HASHALIGN scales linearly in terms of alignment with the center graph.

Experimental Setup. We used the following structural attributes: degree, betweenness centrality, PageRank, egonet degree, average neighbor degree, and egonet connectivity. We chose attributes that are robust to noise (graphs to be aligned are often seen as noisy permutations of each other). More details are given in the supplementary material.

To test the ability of HASHALIGN to incorporate different kinds of features, we also generate synthetic node/edge attributes, if none are available in the real data. For each noise level p , we generate 3 pairs of graphs and report the average accuracy (along with a 95% confidence interval). HASHALIGN is implemented in Python2.7, and the structural feature extraction is based on SNAP [15]. We ran the experiments on Intel(R) Xeon(R) CPU E5 @ 3.50 GHz and 256 GB RAM.

Two-Graph Alignment. We aligned pairs of graphs on all the datasets (with G_1 the real graph and G_2 its noisy permutation at noise level p , generated as described above) and got consistent results. Only the result from the E-mail

network is shown for brevity. With only 1 binary attribute (Fig. 4a), **NetAlign** and **IsoRank** perform poorly because the similarity matrix $\mathbf{L}$ built using just 1 attribute is not informative enough. However, these methods work significantly better in our framework, and the gap between HASHALIGN variants and others grows as noise levels increase. In Fig. 4b, HASHALIGN-NA achieves perfect results in the presence of 5 node attributes, though all methods perform essentially perfectly with abundant node attribute information.

Multiple Graph Alignment. We evaluate HASHALIGN against other methods for multiple graph matching on two datasets: five connectome networks [1] without any labels, and four DBLP co-author networks extracted from the whole DBLP dataset following the settings in [20] with one categorical label (the most frequent conference that an author attends.) Both experiments are conducted with $p = 2\%$ noise.

Figure 5a shows how different methods perform in terms of efficiency and accuracy. When there is no label information to help guide the alignment process (i.e., in the case of connectomes), HASHALIGN achieves best accuracy with short running time for peripheral-center graph pairs alignment, followed by HASHALIGN-FN and HASHALIGN-NA. As for the DBLP networks, since the label with 29 *distinct* values is very discriminative, **Final** can achieve very good efficiency, while HASHALIGN and its variants also have comparable performance. However without node labels, **Final** matches less than 10% of all node pairs, which can be boosted to over 60% if we feed it the hash-based similarity matrices of HASHALIGN. Pairwise graph alignment is the most computationally expensive for large numbers of graphs (see Fig. 5b), but for fewer graphs of the sizes in Fig. 5a, computing all pairwise alignments yields the highest accuracy, and is thus our recommendation if computational resources are not an issue. However, center graph alignment (the ‘*derived*’ versions of HASHALIGN variants) often still outperforms the baselines, and in some cases matches the accuracy of the full pairwise comparisons (e.g., HASHALIGN-NA on the connectome data.)

These results clearly show that HASHALIGN leads to significant improvement over existing methods with regard to both accuracy and runtime. In summary, we see that on average, HASHALIGN (including its variants) are $2\times$ faster and 10 to 20% more accurate than the baselines in pairwise alignment, and $2\times$ faster while up to 50% more accurate in multiple graph alignment. However, these existing methods may have their place within our framework (see Step 4 of Fig. 1), where they may be used to accurately infer alignments from the hashing-based node similarities.

We also verify that the proposed method scales as the number of graphs grows by generating up to 64 synthetic graphs from the aforementioned connectome network with $p = 0.02$ noise, $Z = 2$ and $K = 40$. As shown in Fig. 5b, HASHALIGN’s runtime scales linearly in terms of alignment with the center graph. The runtime for peripheral graph alignments (i.e., w/o the center graph) using sparse matrix multiplication scales subquadratically, as the slope indicates. We omitted the runtime for feature extraction as it is linear on the number of graphs, and does not contribute much to the runtime.

6 Conclusions

We study the problem of multiple graph alignment and propose HASHALIGN, an intuitive, fast and effective similarity-based approach that readily handles any type of input graph. Our method adapts LSH to graph alignment, with a new construction technique and an appropriate node-ID-invariant node representation for this task. Leveraging the rule of matching transitivity, it scales up to many graphs while avoiding solving the expensive alignment task for each pair of graphs separately. Our experiments on real data (incl. sensitivity analysis in the supplementary material) show that HASHALIGN can stand alone as a multi-network alignment tool or be combined with existing methods that require a small set of possible matchings as input. In most cases, it is more accurate, more robust to noise, and/or faster than the baselines. Our work suggests that hashing is a promising direction for scaling up network alignment. Future work could include extending HASHALIGN to use learned node representations specifically designed for multi-network problems, as in the very recent work of [11]. Here one challenge would be devising suitable graph signatures for efficient multiple graph alignment.

Acknowledgements. This material is based upon work supported in part by the National Science Foundation under Grant No. IIS 1743088, and the University of Michigan. Any opinions, findings, and conclusions or recommendations expressed in this material are those of the author(s) and do not necessarily reflect the views of the National Science Foundation or other funding parties. The U.S. Government is authorized to reproduce and distribute reprints for Government purposes notwithstanding any copyright notation here on.

References

1. COBRE (2012). http://fcon_1000.projects.nitrc.org/indi/retro/cobre.html
2. Konect: Koblenz network collection (2016). <http://konect.uni-koblenz.de/networks/>
3. Andoni, A., Indyk, P.: Near-optimal hashing algorithms for approximate nearest neighbor in high dimensions. In: FOCS. IEEE (2006)
4. Bayati, M., Gleich, D.F., Saberi, A., Wang, Y.: Message-passing algorithms for sparse network alignment. ACM TKDD **7**(1), 3:1–3:31 (2013)
5. Bradde, S., Braunstein, A., Mahmoudi, H., Tria, F., Weigt, M., Zecchina, R.: Aligning graphs and finding substructures by a cavity approach. Europhys. Lett. **89**, 37009 (2010)
6. Datar, M., Immorlica, N., Indyk, P., Mirrokni, V.S.: Locality-sensitive hashing scheme based on p-stable distributions. In: SCG, pp. 253–262. ACM (2004)
7. Ding, C.H.Q., Li, T., Jordan, M.I.: Nonnegative matrix factorization for combinatorial optimization: spectral clustering, graph matching, and clique finding. In: ICDM (2008)
8. Duan, S., Fokoue, A., Hassanzadeh, O., Kementsietsidis, A., Srinivas, K., Ward, M.J.: Instance-based matching of large ontologies using locality-sensitive hashing. In: Cudré-Mauroux, P., Heflin, J., Sirin, E., Tudorache, T., Euzenat, J., Hauswirth,

- M., Parreira, J.X., Hendler, J., Schreiber, G., Bernstein, A., Blomqvist, E. (eds.) ISWC 2012. LNCS, vol. 7649, pp. 49–64. Springer, Heidelberg (2012). https://doi.org/10.1007/978-3-642-35176-1_4
9. Goyal, P., Ferrara, E.: Graph embedding techniques, applications, and performance: a survey. arXiv preprint [arXiv:1705.02801](https://arxiv.org/abs/1705.02801) (2017)
10. Heimann, M., Koutra, D.: On generalizing neural node embedding methods to multi-network problems. In: KDD MLG Workshop (2017)
11. Heimann, M., Shen, H., Koutra, D.: Node representation learning for multiple networks: The case of graph alignment. arXiv preprint [arXiv:1802.06257](https://arxiv.org/abs/1802.06257) (2018)
12. Khan, K.U., Nawaz, W., Lee, Y.K.: Set-based unified approach for attributed graph summarization. In: IEEE BDCC, December 2014. <https://doi.org/10.1109/BDCloud.2014.108>
13. Koutra, D., Faloutsos, C.: Individual and collective graph mining: principles, algorithms, and applications. *Synth. Lect. Data Min. Knowl. Discov.* **9**(2), 1–206 (2017)
14. Koutra, D., Tong, H., Lubensky, D.: Big-align: fast bipartite graph alignment. In: ICDM. IEEE (2013)
15. Leskovec, J., Sosič, R.: SNAP: a general-purpose network analysis and graph-mining library. *ACM TIST* **8**(1), 1 (2016)
16. Malmi, E., Chawla, S., Gionis, A.: Lagrangian relaxations for multiple network alignment. *Data Mining Knowl. Discov.* **31**, 1–28 (2017)
17. Safavi, T., Sripada, C., Koutra, D.: Scalable hashing-based network discovery. In: ICDM. IEEE (2017)
18. Singh, R., Xu, J., Berger, B.: Global alignment of multiple protein interaction networks with application to functional orthology detection. *PNAS* **105**(35), 12763–12768 (2008)
19. Zhang, J., Yu, P.S.: Multiple anonymized social networks alignment. In: ICDM. IEEE (2015)
20. Zhang, S., Tong, H.: Final: Fast attributed network alignment. In: KDD. ACM (2016)