

Scalable Hierarchical Clustering with Tree Grafting

Nicholas Monath*
nmonath@cs.umass.edu
CICS, UMass Amherst

Ari Kobren*
akobren@cs.umass.edu
CICS, UMass Amherst

Akshay Krishnamurthy
akshay@cs.umass.edu
Microsoft Research

Michael R. Glass
mrglass@us.ibm.com
IBM

Andrew McCallum
mccallum@cs.umass.edu
CICS, UMass Amherst

ABSTRACT

We introduce GRINCH, a new algorithm for large-scale, non-greedy hierarchical clustering with general linkage functions that compute arbitrary similarity between two point sets. The key components of GRINCH are its rotate and graft subroutines that efficiently reconfigure the hierarchy as new points arrive, supporting discovery of clusters with complex structure. GRINCH is motivated by a new notion of separability for clustering with linkage functions: we prove that when the linkage function is consistent with a ground-truth clustering, GRINCH is guaranteed to produce a cluster tree containing the ground-truth, independent of data arrival order. Our empirical results on benchmark and author coreference datasets (with standard and learned linkage functions) show that GRINCH is more accurate than other scalable methods, and orders of magnitude faster than hierarchical agglomerative clustering.

CCS CONCEPTS

• **Computing methodologies** → **Cluster analysis**; *Online learning settings*.

KEYWORDS

Clustering, Hierarchical Clustering

ACM Reference Format:

Nicholas Monath, Ari Kobren, Akshay Krishnamurthy, Michael R. Glass, and Andrew McCallum. 2019. Scalable Hierarchical Clustering with Tree Grafting. In *The 25th ACM SIGKDD Conference on Knowledge Discovery and Data Mining (KDD '19)*, August 4–8, 2019, Anchorage, AK, USA. ACM, New York, NY, USA, 11 pages. <https://doi.org/10.1145/3292500.3330929>

1 INTRODUCTION

Best-first, bottom-up, hierarchical agglomerative clustering (HAC) is one of the most widely-used clustering algorithms, proving effective for a wide variety of applications such as analyzing gene expression data [10], community detection in social networks [4],

and scientific author disambiguation [8]. One capability that contributes significantly to HAC’s prevalence is that it can be used to construct a clustering according to any cluster-level scoring function, also known as a *linkage function* [22, 24]. This is crucial for applications such as entity resolution, in which the quality of a cluster is typically a learned function of a group of data points [8, 28, 31].

While effective, HAC requires $O(n^2 \log n)$ computation for general linkage functions, making it infeasible to run on datasets of even moderate size. One option for circumventing this computational problem is to use an online or mini-batch variant of the algorithm. However, both HAC variants make irrecoverable, greedy merges and are thus sensitive to data arrival order. Non-greedy, *incremental algorithms* provide a more robust alternative to their online counterparts [21, 34]. Like online approaches, incremental methods consume data points, one at a time, but when new data arrives, incremental algorithms can revisit previous clustering decisions. However, current incremental algorithms fail in two ways: they are only capable of reconsidering clustering decisions at a *local* level and they do not support arbitrary linkage functions [21, 34].

In this paper we introduce GRINCH, a hierarchical, incremental (non-greedy) clustering algorithm that can cluster with any linkage function. GRINCH builds a cluster tree over the incoming data points, one at a time, attempting to keep similar data points near one another in the tree. Robustness to suboptimal data arrival order is achieved by employing both *local and global* tree rearrangements. Local rearrangements are performed using a rotate subroutine, which recursively swaps a child with its aunt. Global rearrangements are performed via a graft subroutine, in which GRINCH may steal a subtree from one part of the hierarchy and merge it with another similar, but distant, subtree. Grafting is a key for both our theoretical and empirical results and supports the discovery of clusters that exhibit (single or sparse) linked structures—an important feature of clustering algorithms used in practice [11].

Theoretically, we define a notion of *model-based separation* that characterizes the relationship between a linkage function and a dataset. For generality, we adopt a graph-theoretic formalism, where data points correspond to vertices of an unknown graph whose connected components form a ground truth clustering. Model-based separation suggests that the linkage function value is high for two item sets if the induced subgraph is connected (see Subsection 2.1). We prove that under this condition, the ground-truth clusters are a tree-consistent partition of the hierarchy built by GRINCH.

In experiments, we show that GRINCH is efficient and builds trees with higher dendrogram purity than other clustering algorithms on large scale datasets. The experiments are performed with a common and important linkage function—average linkage—as well

*The first two authors contributed equally.

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than the author(s) must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.
KDD '19, August 4–8, 2019, Anchorage, AK, USA

© 2019 Copyright held by the owner/author(s). Publication rights licensed to ACM.
ACM ISBN 978-1-4503-6201-6/19/08...\$15.00
<https://doi.org/10.1145/3292500.3330929>

as a linkage function that measures the cosine similarity between two cluster centroid representations. We also perform experiments on two author coreference datasets using learned linkage functions, and demonstrate that GRINCH is more efficient and accurate than the baselines. Our experiments reveal that GRINCH outperforms competitors that only make local tree rearrangements, highlighting the power of the graft subroutine and the robustness of GRINCH.

2 LINKAGE FUNCTIONS FOR CLUSTERING

Clustering is the problem of constructing a partition $C = \{C_1, \dots, C_K\}$ of a dataset $X = \{x_i\}_{i=1}^N$, such that $\bigcup_{C \in C} C = X$ and $\forall C, C' \in C, C \cap C' = \emptyset$. The partition is known as a *clustering* of X .

Most algorithms construct clusterings using pairwise similarities among data points. But, pairwise similarities cannot capture many complex relationships, e.g., data points x_1 and x_2 are similar when clustered with data point x_3 , but are otherwise dissimilar. A natural generalization that can capture these types of relationships are similarities defined over sets of data points, which we refer to as linkage functions. Formally, a linkage function is $f : 2^X \times 2^X \rightarrow \mathbb{R}$.

Clustering with linkage functions is ubiquitous, especially in HAC (from which the name linkage function is derived). In HAC, many popular linkage functions such as single-, complete- and average-linkage are computed from pairwise distance functions. More complex, set-wise linkage functions are used in applications such as image segmentation, noun-phrase coreference and entity resolution; in the latter two domains, these functions are often learned [6, 14, 22, 32, 33]. A unique capability of HAC is that it can easily support an arbitrary linkage function. This flexibility is essential to combat the ill-posed nature of clustering.

2.1 Model-based Separation

Our goal is to design an algorithm that, like HAC, can support arbitrary linkage functions, but is dramatically faster. In developing clustering algorithms, it is often useful to consider various assumptions about the *separability* of the underlying data. For example, in the pairwise setting one of the strongest data assumptions is known as *strict separation* [1]. This assumption holds that any data point in ground-truth cluster C_i is more similar to every other data point in C_i than any data point from a different ground-truth cluster, C_j . Popular instantiations of HAC (e.g., single-, average- and complete-linkage) probably succeed under strict separation, providing some theoretical motivation for these algorithms.

We introduce a notion of *model-based separation* for clustering with a linkage function. Since linkage functions may operate on data of any type, we formalize the definition in terms of a graph, where the data points correspond to vertices.

DEFINITION 1 (MODEL-BASED SEPARATION). Let $G = (X, E)$ be a graph. Let $f : 2^X \times 2^X \rightarrow \mathbb{R}$ be a linkage function that computes the similarity of two groups of vertices and let $\phi : 2^X \times 2^X \rightarrow \{0, 1\}$ be a function that returns 1 if the union of its arguments is a connected subgraph of G . Then f separates G if

$$\forall s_0, s_1, s_2 \subseteq X, \quad \phi(s_0, s_1) > \phi(s_0, s_2) \Rightarrow f(s_0, s_1) > f(s_0, s_2)$$

In words, for a linkage function f to separate a graph G , take any two sets of vertices, s_0 and s_1 , such that $s_0 \cup s_1$ is connected in G ,

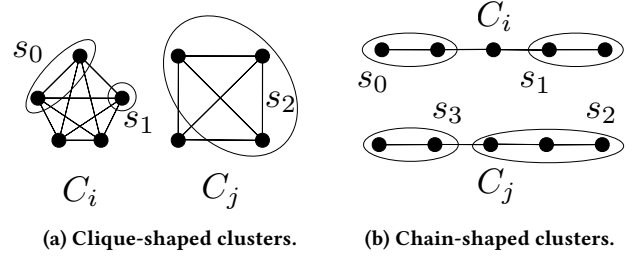


Figure 1: Model-based separation. Figure 1a shows two clique-shaped clusters with data points as vertices in a graph. If f separates the graph then $f(s_0, s_1) > \max[f(s_0, s_2), f(s_1, s_2)]$ because s_0 and s_1 form a connected subgraph. In Figure 1b, even if f separates the graph, it is possible for $f(s_0, s_1) < f(s_1, s_2)$. However, $f(s_1, s_2) < f(s_2, s_3)$.

i.e., $\phi(s_0, s_1) = 1$. Then, for any set s_2 such that $\phi(s_0, s_2) = 0$, the score of f on input (s_0, s_1) must be greater than on input (s_0, s_2) .

Model-based separation offers a non-standard view of clustering. Specifically, the data points of a dataset are treated as vertices in a graph with latent edges. The ground-truth clusters are the connected components of the graph and the goal of clustering is to discover these components using a linkage function.

We provide the following two examples to help build intuition about model-based separation. The examples are used throughout the remainder of our discussion.

EXAMPLE 1 (CLIQUE). Consider a graph $G = (X, E)$ in which each connected component is a clique. Then if f separates G , every vertex in a connected component, C_i , is more similar to all other vertices in C_i than any vertex in connected component C_j , where similarity is defined by f . Thus, clique-structured connected components exactly capture strict separation.

EXAMPLE 2 (CHAIN). Consider a graph $G = (X, E)$ in which each connected component is chain-structured. According to Definition 1, two vertices that are part of the same chain but do not share an edge may be dissimilar under f even if f separates G . However, any two segments of the chain connected by an edge are similar under f .

A visual illustration of both clique and chain style clusters is depicted in Figure 1. As we will see, chain structured connected components pose a challenge to existing incremental algorithms, something we resolve with GRINCH (Section 3).

2.2 Cluster Trees

In most clustering problems, the appropriate number of clusters is unknown a priori. HAC addresses this uncertainty by building a cluster tree over data points.

DEFINITION 2 (CLUSTER TREE [23]). A binary *cluster tree* $\mathcal{T}$ on a dataset $X = \{x_i\}_{i=1}^N$ is a collection of subsets such that $C_0 = \{x_i\}_{i=1}^N \in \mathcal{T}$ and for each $C_i, C_j \in \mathcal{T}$ either $C_i \subset C_j$, $C_j \subset C_i$ or $C_i \cap C_j = \emptyset$. For any $C \in \mathcal{T}$, if $\exists C' \in \mathcal{T}$ with $C' \subset C$, then there exists two $C_L, C_R \in \mathcal{T}$ that partition C .

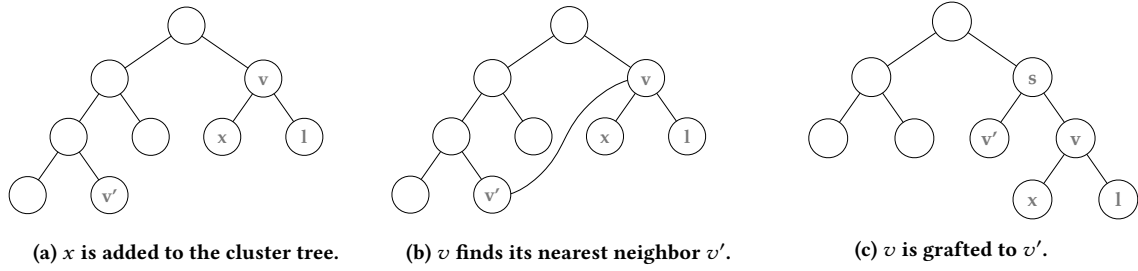


Figure 2: The graft subroutine. Dotted lines denote new nodes and mergers. Before x is added, l and v' are in disjoint subtrees despite belonging to the same ground-truth cluster. The addition of x creates the subtree with root v and initiates the graft.

Given a cluster tree, $\mathcal{T}$, any set of disjoint subtrees whose leaves cover X represents a valid clustering and is referred to as a *tree consistent partition* [16]. Thus, cluster trees compactly encode multiple alternative clusterings, allowing for a clustering to be selected as a post-processing step. Another advantage of using cluster trees is that they often facilitate efficient search and naturally group similar data points near one another in the hierarchy.

We relate model-based separation, cluster trees and HAC in the following fact:

FACT 1. *Let f be a linkage function that separates G . Then running HAC under f returns a cluster tree, $\mathcal{T}$, such that the connected components of G are a tree-consistent partition of $\mathcal{T}$.*

To see why, notice that in each iteration of HAC, the highest scoring pair of remaining subtrees is merged. Since f separates G , a merger resulting in a subtree that corresponds to a connected subgraph of G has higher score than any merger resulting in a disconnected subgraph of G . Even though HAC can construct a cluster tree that contains the ground-truth clustering as a tree-consistent partition, the algorithm costs $O(n^2 \log n)$ for general linkage functions and does not scale to large datasets. We will verify this claim empirically in our experiments (Section 4).

3 ROTATIONS, GRAFTING AND GRINCH

In this section, we derive an efficient, incremental algorithm called GRINCH that can be used to construct clusterings under any linkage function. Like HAC, the backbone of GRINCH is a cluster tree. We begin the discussion by analyzing a greedy, incremental variant of HAC and when it fails. Then, we introduce two subroutines, rotate and graft, that can be used to enhance robustness. Finally, we present our algorithm, GRINCH.

3.1 Online HAC and Rotations

An efficient alternative to HAC is its online variant that merges each incoming data point with its nearest neighbor seen so far (ONLINE). For now, let us consider the setting in which a nearest neighbor is found using a linkage function, f . Let f separate a graph G and let ground-truth clusters be cliques in G (i.e., the data is strictly separated). Even in this simple case, ONLINE may construct a cluster tree in which the ground-truth clustering is not a tree consistent partition. To see why, consider a stream in which the first two data points, x_1 and x_2 , are of the same ground-truth cluster and the third data point, x_3 is of a different cluster. Assume, without loss of

generality, that ONLINE adds x_3 as a sibling of x_1 . Then the ground-truth clustering is not a tree consistent partition of the resulting tree (and all subsequent trees).

To recover from such mistakes, local tree rearrangements may be applied. Previous work uses *rotations*, which swap a child and its aunt in the tree, to correct local errors induced by unfavorable arrival order [21]. While originally designed to be used with pairwise distances, the condition under which rotations should be applied can be extended to linkage functions:

$$f(v, s(v)) < f(v, \text{aunt}(v)) \quad (1)$$

where the functions $s(\cdot)$ and $\text{aunt}(\cdot)$ return the sibling and aunt of their input, respectively. In words, if a node $v \in \mathcal{T}$ achieves a higher score under f with its aunt than with its sibling, then the aunt and sibling should be swapped. Now, let us revisit the example above. Since x_1 and x_2 are both vertices in the same clique in G , they are connected by an edge. Then, by model-based separation, $f(x_1, x_2) > f(x_1, x_3)$, so a rotation will be applied, producing a tree that contains the ground-truth clustering.

Unfortunately, the ONLINE algorithm, augmented with the ability to perform rotations (ROTATE), cannot always recover the connected components of a graph that is separated by f . In particular, ROTATE cannot reliably recover chains (Example 2). By virtue of being a local operation, rotations can only be used to provably recover connected components that are clique-structure.

3.2 Subtree Grafting

We introduce a non-local tree rearrangement called a *graft*, which facilitates the discovery of chain-structured connected components. At a high level, the graft procedure with respect to a node $v \in \mathcal{T}$ searches $\mathcal{T}$ for a node v' that is both similar to v and dissimilar from its current sibling, $s(v')$. If such a subtree is found, v' is disconnected from its parent and made a sibling of v . A visual illustration of a successful graft is depicted in Figure 2.

In detail, a graft searches the leaves of $\mathcal{T}$ for the nearest neighbor leaf of v called l . Then it checks whether the following holds:

$$f(v, l) > \max[f(v, s(v)), f(l, s(l))] \quad (2)$$

i.e., v and l prefer each other to their current siblings according to f . If the condition succeeds, merge v and l . If the condition fails because l prefers its sibling to v , retest the condition at v and l 's parent, $\text{par}(l)$; if the condition fails because v prefers its sibling to l , then retest the condition at $\text{par}(v)$ and l . Continue to

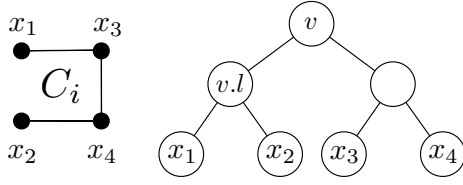


Figure 3: Poorly structured tree. Even though v 's leaves form a connected subgraph in the graph on the left of the Figure (i.e., they all belong to cluster C_i), $v.l$'s descendant leaves, x_1 and x_2 , are disconnected. An attempt to graft either x_1 or x_2 from a node whose descendants are not in C_i may succeed.

check recursively until the condition succeeds or until the first time two nodes, v_1 and v_2 , are reached such that one is the ancestor of the other. Pseudocode for the graft subroutine can be found in Algorithm 1. In the algorithm, `par` returns the parent of a node, `lca` returns the lowest common ancestors of its arguments and `makeSib` merges its arguments and returns their new parent. `NN` performs a nearest neighbor search and `constrNN` performs a nearest neighbor search that excludes its second argument from the result.

Algorithm 1 `graft( $v, \mathcal{T}, f$ )`

```

 $l \leftarrow \text{constrNN}(v, \text{lvs}(v), f, \mathcal{T})$ 
 $v' \leftarrow \text{lca}(v, l); st \leftarrow v$ 
while  $v \neq v' \wedge l \neq v' \wedge s(v) \neq l$  do
  if  $f(v, l) > \max[f(v, s(v)), f(l, s(l))]$  then
     $z \leftarrow s(v); v \leftarrow \text{merge}(v, l)$ 
     $\text{restruct}(z, \text{lca}(z, v), f)$ 
  break
  if  $f(v, l) < f(l, s(l))$  then  $l \leftarrow \text{par}(l)$ 
  if  $f(v, l) < f(v, s(v))$  then  $v \leftarrow \text{par}(v)$ 
if  $v = st$  then output:  $v'$  else output:  $v$ 
```

3.3 Tree Restructuring

While the `graft` subroutine facilitates discovery of chain-structured clusters, poorly structured trees are susceptible to having the `graft` subroutine disconnect previously discovered ground-truth clusters. As an example, consider Figure 3, in which $\text{lvs}(v)$ form the connected subgraph C_i (i.e., they all belong to the same ground-truth cluster). Consider v 's left child, $v.l$, and its descendants, which form a *disconnected* subgraph. An attempt to graft either descendant, x_1 or x_2 , may succeed, even when initiated from a node (not depicted) whose descendants are not connected to C_i . After such a graft, $\mathcal{T}$ cannot contain a tree-consistent partition that matches the ground-truth clustering.

Notice that a subtree can defend against spurious grafts by ensuring that each of its descendant subtrees is connected. For example, in Figure 3, if x_2 and x_3 were swapped, then each descendant subtree of v would be connected. Moreover, after such a swap, grafts from nodes whose descendants were not part of C_i would necessarily fail (assuming that f separates the graph).

During tree construction, the only step that can result in a connected subtree with disconnected descendants is the `graft` subroutine (a rigorous proof is included in the supplement). We introduce

Algorithm 2 `restruct( $z, r, f$ )`

```

while  $z \neq r$  do
   $as \leftarrow \{s(a) \text{ for } a \in \text{ancs}(z) \setminus \text{ancs}(r)\}$ 
   $m \leftarrow \text{argmax}_{a \in as} f(z, a)$ 
  if  $f(z, s(z)) < f(z, m)$  then  $\text{swap}(s(z), m)$ 
   $z \leftarrow \text{par}(z)$ 
```

Algorithm 3 `Insert( $x_i, \mathcal{T}, f$ )`

```

 $l \leftarrow \text{NN}(x_i, f, \mathcal{T}); t \leftarrow \text{makeSib}(x_i, l)$ 
while  $f(x_i, s(x_i)) < f(\text{aunt}(x_i), s(x_i))$  do
   $\text{rotate}(x_i, \text{aunt}(x_i))$ 
 $p \leftarrow \text{par}(x_i)$ 
while  $p \neq \text{null}$  do
   $p \leftarrow \text{graft}(p, \mathcal{T}, f)$ 
```

the `restruct` (*restructure*) subroutine, which is performed after a successful graft, and reorganizes a subtree with the intent of making each of its descendants connected. Let v' be a node that was just grafted, v be the previous sibling of v' (i.e., before the graft) and let $r = \text{lca}(v, v')$ be the current least common ancestor of v and v' . `restruct` is initiated from v . First, the siblings of the ancestors of v (until r) are collected. Then, we find the node in the collection most similar to v . If that node is more similar to v than v 's current sibling (according to f), the two are swapped. The intuition here is that if a graft left v and its new sibling disconnected, then the swap serves as a mechanism to restore the connectedness of v 's parent. Such swaps are attempted from the ancestors of v until r . Pseudocode appears in Algorithm 2.

3.4 Grinch

Using the `rotate`, `graft` and `restruct` tree rearrangement routines, we derive a new algorithm called **GRINCH**, which stands for: **G**rafting and **R**otation-based **IN**cremental **H**ierarchical clustering. The steps of the algorithm are as follows: when a new data point, x_i , arrives, find x_i 's nearest neighbor, l , among the leaves of $\mathcal{T}$. Add x_i to $\mathcal{T}$ as a sibling of l . Then, apply the `rotate` subroutine while Equation 1 is true. Finally, attempt to graft recursively from each ancestor of x_i . Each time a graft is successful, restructure the tree to group similar items together. Pseudocode for **GRINCH** can be found in Algorithm 3.

THEOREM 1. Let $\mathcal{X} = \{x_i\}_{i=1}^N$ be a dataset with ground-truth clustering $\mathcal{C}^* = \{C_1, \dots, C_k\}$. Let f separate a graph G on vertices $\mathcal{X}$ and let each cluster $C \in \mathcal{C}^*$ be a connected component in G . Then **GRINCH** recovers a cluster tree such that $\mathcal{C}^*$ is a tree consistent partition of $\mathcal{T}$ regardless of the input order.

The proof of Theorem 1 can be found in the appendix.

4 EXPERIMENTS

We experiment with **GRINCH** to assess its scalability and accuracy. We begin by demonstrating that **GRINCH** outperforms other incremental clustering algorithms on a synthetic dataset. Observing that

some of the steps of GRINCH are underutilized, we present four approximations of GRINCH’s algorithmic components. We apply each approximation in turn and show that together they dramatically improve GRINCH’s scalability without compromising its clustering quality. Then, we compare the approximate variant of GRINCH to state-of-the-art large scale hierarchical clustering methods. To showcase the flexibility of GRINCH, we also provide experimental results in entity resolution, where the linkage function is learned. Finally, we provide analysis of the graft subroutine—GRINCH’s distinguishing feature—and perform experiments to demonstrate the algorithm’s robustness on adversarial data orderings.

Dendrogram Purity. Before beginning, we briefly review *dendrogram purity*, a preferred method of holistically evaluating hierarchical clusterings [5, 16, 21]. Dendrogram purity is computed as follows: Let $C^* = \{C_1, \dots, C_k\}$ be the ground-truth clustering of a dataset $\mathcal{X}$, and let $\mathcal{P}^* = \{(x, x') | x, x' \in \mathcal{X}, C^*(x) = C^*(x')\}$ be the set of all data point pairs that belong to the same ground-truth clusters. Then the dendrogram purity (DP) of a cluster tree, $\mathcal{T}$ is:

$$DP(\mathcal{T}) = \frac{1}{|\mathcal{P}^*|} \sum_{(x, x') \in \mathcal{P}^*} \text{pur}(\text{lvs}(\text{lca}(x, x')), C^*(x))$$

where $\text{lca}(x, x')$ returns the least common ancestor of x and x' in $\mathcal{T}$, $\text{lvs}(\cdot)$ returns the descendant leaves of its argument, and $\text{pur}(\cdot, C^*(x))$ takes a collection of leaves and computes the fraction that belong to ground-truth cluster $C^*(x)$.

An implementation of GRINCH and code to reproduce these experiments is available ¹.

4.1 Synthetic Data Experiment

In our first experiment, we compare GRINCH to other incremental hierarchical clustering algorithms on a synthetic dataset in order to begin to understand GRINCH’s empirical performance characteristics in a controlled manner. The data is generated so that it satisfies model-based separation with respect to cosine similarity. In particular, the dataset contains 2,500 10,000-dimensional binary vectors that belong to 100 clusters, with 25 points per cluster. Points in cluster k have bits $100(k-1)$ to $100k-1$ set randomly to 1 with probability 0.1. All other bits are set to 0. This way, across cluster points have cosine similarity 0 and within cluster points can have either 0 or non-zero cosine similarity. In other words, two points, x_1 and x_2 , in the same cluster can appear to be dissimilar and end up in distant regions of the tree. The representation of each internal node in the GRINCH tree is the sum of the vectors of its descendant leaves. We compute the cosine similarity between two nodes v and v' as the cosine similarity between their aggregated vectors (we refer to this as *cosine linkage* in the following sections). We compare GRINCH, ROTATE and ONLINE.

The experimental results reveal that GRINCH achieves perfect dendrogram purity (1.0), which is expected given GRINCH’s correctness guarantee. ROTATE achieves a dendrogram purity of 0.872 while ONLINE achieves 0.854. ROTATE and ONLINE do not construct trees of perfect purity because of their inability to globally rearrange a cluster hierarchy.

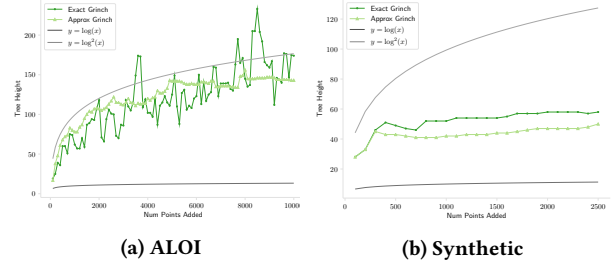


Figure 4: The height of trees built with exact and approximate (capping, single elimination, and single NN search) GRINCH are shown for a subset of ALOI and the Synthetic data described in Section 4.1. We notice that trees are fairly balanced and not chain structured. The height is close to (if not less than) to $\log^2 N$.

4.2 Scalability and Approximations

4.2.1 Worst-case running time analysis. We measure the worst-case running time of the algorithm in terms of the number of linkage function calls. The insert method consists of four components: nearest neighbor search, rotations, grafting, and restructuring. Under particular assumptions (N -point metric space with expansion constant $c > 2$), exact nearest neighbor search is possible in $O(c^{12} \log(N))$ time per query where N is the number of data points [3]. Otherwise, exhaustive $O(N)$ time is required unless approximate nearest neighbor search methods are used. We will use $O(T)$ to refer to the time spent on nearest neighbor search. Let H be the height of the GRINCH tree, i.e., the length of the longest root to leaf path. There can be at most $O(H)$ rotations applied in a call to insert. In the extreme, grafts could be recursively applied moving all of the $N-1$ leaves (those other than the new point’s sibling) one at a time in separate grafts. Since each graft involves a nearest neighbor search, there would be in the worst case $O(NT)$ computations from grafts in insert. There can be at most $O(H)$ swaps made in a given restruct call, and each of these swaps involves considering $O(H)$ ancestors, resulting in $O(H^2)$ time overall. This results in $O(NT + H^2)$ time per data point. While there is no guarantee on the height of trees built with GRINCH, we find in practice that they are close to $\log^2 N$ as shown in Figure 4. If the trees had such a height, the running time would further be reduced to: $O(TN + \log^4(N))$ per data point.

4.2.2 Approximate GRINCH. Some of the algorithmic steps of GRINCH, which are required to prove its correctness, are seldom invoked in practice. For example, and perhaps expectedly, a graft is unlikely to succeed between two nodes close to the root of the tree. Therefore, we introduce a handful of approximations designed to have little effect on the quality of the clusterings constructed by GRINCH, but also designed to make the algorithm significantly faster in practice.

(1) **Capping.** The recursive subroutines *rotate*, *graft*, and *restruct*, improve performance, but are also computationally expensive to check, and often are not applied. Moreover, we notice that tree rearrangements that occur close to the root do not have a significant, instantaneous effect on dendrogram purity. Therefore, we introduce *rotation*, *graft* and *restructure caps*, which

¹<https://github.com/iesl/grinch>

prohibit rotations, grafts and restructures from occurring above a height, h .

- (2) **Single Elimination Mode.** The graft subroutine generally improves GRINCH's clustering performance, and is essential in attaining perfect purity on the synthetic dataset, but we find that graft attempts are rejected many more times than they are accepted. However, at times, we observe that a sequence of recursive grafts are accepted when initiated close to the leaves. Therefore, to limit the number of attempted grafts while retaining these graft sequences, we introduce *single elimination mode*. In this mode, the recursive grafting procedure terminates after a graft between v and v' fails because both prefer their current siblings to a merge.
- (3) **Single Nearest Neighbor Searching.** GRINCH makes heavy use of nearest neighbor search under the linkage function f . Rather than perform nearest neighbor search anew for each graft, when a data point arrives, we perform a single k -NN search ($k \in [25, 50]$) and only consider these nodes during subsequent grafts (until the next data point arrives).
- (4) **Navigable Small World Graphs.** Instead of performing nearest neighbor computations exactly, we can perform them approximately. To this end, we employ a *navigable small world* nearest neighbor graph (NSW)—a data structure inspired by decentralized search in small world networks [19, 20, 30]. To find the nearest neighbor of a data point, x_i , in an NSW, begin at a random node, v . If the similarity between x_i and v is maximal among all neighbors of v , terminate; otherwise, move to the neighbor of v most similar to x_i . To insert a new data point, x_j , find its k nearest neighbors and add edges between those neighbors and a new data point [26]. Thus, NSWs are constructed online. In practice, we simultaneously construct a hierarchical clustering and an NSW over the data points stored in the tree's leaves.

To measure the effects of our approximations on the speed and quality of the resulting algorithm, we conduct the following ablation. We run GRINCH on our synthetically generated dataset as well as a random 5k subset of the ALOI [12] dataset and measure dendrogram purity, time, and the number of calls made to rotate, graft and restruct. We repeat the procedure multiple times, each time adding one of the following approximations, in order: capping, single elimination, single nearest neighbor search and approximate nearest neighbor search. Capping is performed at height 100. We also experiment with removal of the graft and rotate subroutines.

The result of the ablation is contained in Table 1. We observe that, for both datasets, each of the approximations reduces the computational cost of algorithm without effecting the resulting DP. However, once grafts are removed, the DP drops by 3% on ALOI and 12% on the synthetic datasets. When rotate is also removed, DP drops by an additional 6% and 2%, respectively.

Having verified that on a subset of ALOI our approximations improve scalability at little expense in terms of dendrogram purity, in the following experiments we report results for GRINCH in single elimination mode and with the rotation cap set to $h = 100$. We did not tune the cap parameter and used that same setting for all datasets. We mildly tune the number of nearest neighbors connected in NSW construction setting based on the dimensionality of the data, using values of 5 for CovType, 50 for ALOI, and 100 for

the other datasets. This can be tuned based on the nearest neighbor performance of the structure (independent of the clustering algorithm).

4.3 Large Scale Clustering

We compare GRINCH with the following 4 algorithms: **ONLINE** - an online hierarchical clustering algorithm that consumes one data point at a time and places it as a sibling of its nearest neighbor; **ROTATE** - an incremental algorithm that places a data point next to its nearest neighbor and then performs rotations until Equation 1 holds; **MB-HAC** - the mini-batch version of HAC, which keeps a buffer of size b , runs a single step of HAC using the data points in the buffer and then adds the next record to the buffer; **HAC** - best-first, bottom-up hierarchical agglomerative clustering and **PERCH** - a state-of-the-art large scale hierarchical clustering method.

We run each algorithm on 5 large scale clustering datasets: CovType, a dataset of forest covertype, ALOI [12], a 50K subset of the Imagenet ILSVRC12 dataset [27] and the Speaker dataset [13], and a 100K subset of ImageNet containing all 17K classes not just the subset in ILSVRC12. Datasets have 500K, 50K, 100K, 36K, and 100K instances, respectively. We run each HAC variant under two different linkage functions: average linkage and cosine linkage. To compute the cosine similarity between two nodes, v and v' , first, for each node, compute the sum of the vectors contained at their descendant leaves. Then, compute the cosine similarity between the aggregated vectors.

Results are displayed in Table 2, where we record the dendrogram purity averaged over 5 replicates of each algorithm, where for each replicate we randomize the arrival order of the data. The table reveals that GRINCH—under both linkage functions—outperforms the corresponding versions of ROTATE and ONLINE on all datasets except for on the CovType dataset where the methods all seem to perform equally well. This underscores the power of the graft subroutine. GRINCH with approximate nearest neighbor search even outperforms PERCH, which uses exact nearest neighbor search, on ALOI. Recall that, unlike the HAC variants, PERCH employs a specific linkage function. Seeing as the HAC variants outperform PERCH on Speaker suggests that the ability to equip various linkage functions can be advantageous. HAC is best on Speaker, but cannot scale to ALOI.

4.4 Author Coreference

Bibliographic databases, like PubMed, DBLP, and Google Scholar, contain citation records that must be attributed to the corresponding authors. For some records, the attribution process is easy, but for many others, the identities of a publication's authors are ambiguous. For example, DBLP contains hundreds of citations written by different authors named "Wei Wang" that currently cannot be disambiguated². Intuitively, author coreference datasets often exhibit chain like structures because a single citation written by a prolific author (perhaps in a short-lived collaboration) may only be similar to a small number of that author's other citations and dissimilar from the rest.

Following previous work, we train a linkage function to predict the likelihood that a group of citation records were all written by

²<https://dblp.uni-trier.de/pers/hd/w/Wang:Wei>

Approx.	ALOI					Synthetic				
	DP	Time (s)	# Rotate	# Graft	# Restr.	DP	Time (s)	# Rotate	# Graft	# Restr.
GRINCH (No Approx).	0.533	85.371	7107	2435	1088	1.0	160.307	2558	578	203
w/ Cap (100)	0.533	48.452	6495	2157	686	0.993	164.328	2558	578	194
w/ Single Elimn	0.534	39.019	6574	1586	533	0.997	157.622	2523	526	184
w/ Single NN	0.540	22.226	6441	1516	570	0.993	83.014	2517	415	148
w/ no Restruct	0.538	14.292	6477	1634	0	0.993	82.262	2476	426	0
w/ no Graft	0.506	12.748	6747	0	0	0.872	82.055	2259	0	0
w/ no Rotate	0.442	14.793	0	0	0	0.854	80.526	0	0	0

Table 1: Ablation. Each row in the table represents GRINCH with the corresponding approximation applied in addition to all approximations contained in previous rows. The first 4 approximations significantly decreases the computational cost of GRINCH, but do not compromise DP. The ablation is performed for the first 5000 points of ALOI and the Synthetic datasets.

Algorithm	Linkage	CovType	ILSVRC12 (50k)	ALOI	Speaker	ImageNet (100k)
GRINCH	Avg	0.43 ± 0.00	0.557 ± 0.003	0.504 ± 0.002	0.480 ± 0.003	0.065 ± 0.000
GRINCH	CS	0.43 ± 0.00	0.544 ± 0.005	0.499 ± 0.003	0.478 ± 0.003	0.062 ± 0.000
ROTATE	Avg	0.43 ± 0.01	0.545 ± 0.004	0.476 ± 0.004	0.407 ± 0.003	0.063 ± 0.001
ROTATE	CS	0.44 ± 0.01	0.513 ± 0.007	0.472 ± 0.003	0.406 ± 0.003	0.062 ± 0.000
ONLINE	—	0.44 ± 0.01	0.527 ± 0.004	0.435 ± 0.004	0.317 ± 0.002	0.0589
PERCH [21]	—	0.45 ± 0.004	0.53 ± 0.003	0.44 ± 0.004	0.37 ± 0.002	0.065 ± 0.000
PERCH-BC [21]	—	0.45 ± 0.004	0.36 ± 0.005	0.37 ± 0.008	0.09 ± 0.001	0.03 ± 0.00
MB-HAC [21]	Best Reported	0.44 ± 0.005	0.43 ± 0.005	0.30 ± 0.002	0.01 ± 0.002	—
HAC [21]	Average	—	0.54	—	0.55	—

Table 2: Dendrogram Purity results for GRINCH and baseline methods. We compare two linkage functions: approximate average linkage (Avg) and cosine similarity linkage (CS).

the same author [8, 28, 31]. We train our model, that utilizes features such as coauthor names, title, venue, year, etc., by running HAC and, at each step, use the model to predict the precision of merging two groups of records as in [24].

We compare the 5 HAC variants in author coreference on two datasets with labeled author identities: **Rexa** [8] and **PSU-DBLP** [15]. Evaluation is done using the pairwise F1-score of a predicted flat clustering against the ground-truth clustering. To compute pairwise F1-score, each pair of citations that appears in both the same ground-truth and predicted clusters is considered a true positive; each pair of citations that belongs to different ground-truth clusters but the same predicted cluster is considered a false positive. None of the authors represented in the test set, have any publications in the training set.

Figure 3 shows the precision, recall, and pairwise F1-score achieved by each method. The results show that GRINCH outperforms the other scalable methods on both datasets and even outperforms HAC on DBLP. This behavior may stem from overfitting of the learned linkage function, which is exploited by HAC; since GRINCH only approximates HAC, it can be thought of as a form of regularization. Again, GRINCH outperforms ONLINE and ROTATE on both datasets underscoring the importance of the rotate and graft procedures.

4.5 Significance of Grafting

The results above indicate that GRINCH—even when employing a number of approximations—constructs trees with higher dendrogram purity than other scalable methods in a comparable amount of

Algorithm	Rexa			DBLP		
	Pre	Rec	F	Pre	Rec	F
GRINCH	0.808	0.883	0.844	0.809	0.620	0.701
ROTATE	0.864	0.641	0.734	0.876	0.554	0.678
ONLINE	0.850	0.209	0.331	0.827	0.151	0.255
MB-HAC-Med.	0.807	0.881	0.843	0.375	0.631	0.461
MB-HAC-Sm.	0.922	0.333	0.483	0.697	0.151	0.247
HAC	0.805	0.887	0.844	0.741	0.600	0.664

Table 3: Precision, recall and F-Score of various methods on the Rexa and DBLP datasets.

time. Interestingly, GRINCH only differs from the algorithm ROTATE in its use of the graft (and subsequent restrict) subroutine. To better understand the significance of grafting, we compare GRINCH and the algorithm ROTATE on the first 5000 points of ALOI.

Figure 5a shows that dendrogram purity as a function of the number of data points inserted for both GRINCH and ROTATE and the first 5000 points of ALOI. Echoing the results above, by 1000 points, GRINCH dominates ROTATE.

Figure 5b shows the instantaneous and cumulative change in dendrogram purity due to grafts made by GRINCH. That is, for the i th data point, x_i , we record the dendrogram purity after x_i is inserted and rotations are performed (i.e., what would be executed by ROTATE). Then, we perform grafting (if appropriate) and record the dendrogram purity *after* all recursive grafts have been completed. The difference between the dendrogram purity after

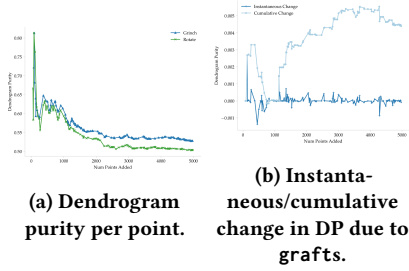


Figure 5: Figure 5a shows the dendrogram purity of two trees, one built by GRINCH and the other built by rotate, on the first 5000 points of ALOI. The dendrogram purity of the tree built GRINCH is greater than that of the tree built by rotate. Figure 5b plots the instantaneous and cumulative change in dendrogram purity due to grafts. While GRINCH achieves 3% larger dendrogram purity than rotate

grafting and before grafting (but after rotations) is the instantaneous change in dendrogram purity due to grafts; the sum of instantaneous changes is the cumulative change.

Note the y -axis of Figure 5b, which reveals that even the most instantaneously significant grafts only lead to minute changes in dendrogram purity (of about 0.001). Moreover, after 5000 points, the cumulative change in dendrogram purity due to grafts is less than 0.005—hardly accounting for the difference in dendrogram purity between the tree built by GRINCH and the tree built by ROTATE (of 0.03). We conclude from these measurements that the increase in performance due to the graft subroutine is related to the rearrangement of small numbers of points. These rearrangements do not immediately have significant impact on dendrogram purity, but they do have significant long-term affects. To make this hypothesis more concrete, consider the case in which two dissimilar data points from the same cluster are split between two distant regions of the tree early on in clustering. The points are never merged (via a graft) and so each point draws a significant portion of the cluster’s other data points to its location in the tree. This has dire consequences with respect to dendrogram purity. If a graft is performed early on to correct the split, an adverse scenario like this can be averted.

4.6 Robustness

For completeness, we perform an experiment used in previous work to test an incremental clustering algorithm’s robustness to data point arrival order [21]. In the experiment, a dataset is ordered in two specific ways. **Round-Robin** – Randomly determine an ordering of ground-truth clusters. Then, construct a data point arrival order such that the i th data point is a member of cluster $i \bmod K$, where K is the number of clusters and $\bmod$ returns the remainder when its first argument is divided by its second. **Sorted** – Randomly determine an ordering of ground-truth clusters. All points of cluster C_i arrive before any point of cluster C_{i+1} arrives. As in previous work, we perform a robustness experiments with the ALOI dataset.

Table 4 shows that GRINCH achieves higher dendrogram purity than both PERCH and mini-batch HAC (with 2 different batch sizes)

on data ordered using the Round Robin ordering scheme. Under this arrival order, MB-HAC performs poorly showing its lack of robustness. When the data is in Sorted order—which makes for easier clustering for MB-HAC—GRINCH outperforms PERCH and is competitive with MB-HAC.

Method	Round.	Sort.
GRINCH	0.503	0.457
PERCH	0.446	0.351
MB-HAC (5K)	0.299	0.464
MB-HAC (2K)	0.171	0.451

Table 4: DP for adversarial arrival orders (ALOI).

5 RELATED WORK

The family of online and incremental clustering methods is diverse, however all algorithms in this family optimize for specific linkage functions. PERCH, from which the rotate procedure is inspired, performs rearrangements to satisfy a condition similar to complete-linkage [21]. BIRCH is another top-down hierarchical clustering algorithm that attempts to minimize a k -center style cost at each node in the tree [34]. BIRCH also includes a non-greedy reassignment step but has been shown to produce low quality trees in practice. Liberty et al [25] propose a flat clustering algorithm that optimizes k -means cost. Since their algorithm runs in the online setting, after a data point arrives and is assigned to a cluster, it may never be reassigned. While not incremental, some work focuses on designing highly scalable algorithms for specific linkage functions. Particular attention is paid to single-linkage because of its connection to the minimum spanning tree problem. For example, recent work develops massively parallel algorithms for single-linkage [2].

When clustering with linkage functions, probabilistic approaches can provide an alternative to HAC. For example, split-merge Markov Chain Monte Carlo (MCMC) methods perform clustering by randomly splitting and merging clusters according to a proposal function [18]. An algorithm similar to split-merge MCMC has even been used for author coreference [31]. This algorithm employs a custom linkage function on structured records and works by maintaining a forest—each tree corresponding to a cluster—and randomly proposing mergers and splits of various branches. Unlike GRINCH, this algorithm relies on sampling to escape local minima. As the number of items grows, the likelihood of sampling a merge or split that will be accepted decreases rapidly.

Our work is partially inspired by complex linkage functions that are used for clustering. One example is Bayesian hierarchical clustering (BHC)—a recursive, probabilistic, hierarchical model for data [16]. Fitting BHC models is performed by running HAC with BHC as the linkage function. Because HAC is inefficient, randomized approaches for fitting BHC have also been proposed, but each of these methods still runs HAC as a subroutine on small, randomly selected subsets of data [17]. HAC-style algorithms are also used to do probabilistic, hierarchical community detection and alongside learned models for entity resolution [4, 24].

Model-based separation is related to recently proposed definitions of *perfect hierarchical clustering structure* [7, 29], in which pairwise similarities between data points lead to a tree that can be

discovered by HAC that has minimal cost. The costs used in these works are variants of Dasgupta’s cost [9]. Perfect hierarchical clustering structures are a special case of model-based separation, in which single-, average-, or complete-linkage is used. Model-based separation is strictly more general, allowing for linkage functions that compute the similarity of two point sets arbitrarily, rather than as a function of pairwise data point similarities.

6 CONCLUSION

This paper introduces GRINCH, an incremental algorithm for hierarchical clustering under any linkage function. The algorithm relies on three subroutines, rotate, graft, and restruct, that help it to discover complex cluster structure regardless of data arrival order. We introduce model-based separation for clustering with linkage functions and prove that GRINCH always returns a tree with perfect dendrogram purity when running in the separated setting. We describe an efficient implementation of GRINCH and present an empirical evaluation demonstrating that GRINCH is more accurate than other baseline approaches and more scalable than HAC. We believe that GRINCH is an asset for large clustering problems in which the data points engage in complicated relationships and clusters are best modeled by learned linkage function.

ACKNOWLEDGMENTS

This work was supported in part by the Center for Data Science and the Center for Intelligent Information Retrieval, in part by the National Science Foundation under Grant No. NSF-1763618, in part by the Chan Zuckerberg Initiative under the project “Scientific Knowledge Base Construction,” in part by International Business Machines Corporation Cognitive Horizons Network agreement number W1668553, and in part using high performance computing equipment obtained under a grant from the Collaborative R&D Fund managed by the Massachusetts Technology Collaborative. Any opinions, findings and conclusions or recommendations expressed in this material are those of the authors and do not necessarily reflect those of the sponsor.

REFERENCES

- [1] M.-F. Balcan, A. Blum, and S. Vempala. 2008. A discriminative framework for clustering via similarity functions. *STOC*.
- [2] M. Bateni, S. Behnezhad, M. Derakhshan, M. Hajiaghayi, R. Kiveris, S. Lattanzi, and V. Mirrokni. 2017. Affinity Clustering: Hierarchical Clustering at Scale. *NeurIPS*.
- [3] A. Beygelzimer, S. Kakade, and J. Langford. 2006. Cover trees for nearest neighbor. *ICML*.
- [4] C. Blundell and Y. W. Teh. 2013. Bayesian hierarchical community discovery. *NeurIPS*.
- [5] C. Blundell, Y. W. Teh, and K. A. Heller. 2010. Bayesian Rose Trees. *UAI*.
- [6] K. Clark and C. D. Manning. 2016. Improving Coreference Resolution by Learning Entity-Level Distributed Representations. *ACL*.
- [7] V. Cohen-Addad, V. Kanade, F. Mallmann-Trenn, and C. Mathieu. 2018. Hierarchical clustering: Objective functions and algorithms. *SODA*.
- [8] A. Culotta, P. Kanani, R. Hall, M. Wick, and A. McCallum. 2007. Author disambiguation using error-driven machine learning with a ranking loss function. *Workshop on Information Integration on the Web*.
- [9] S. Dasgupta. 2016. A cost function for similarity-based hierarchical clustering. *STOC*.
- [10] M. B. Eisen, P. T. Spellman, P. O. Brown, and D. Botstein. 1998. Cluster analysis and display of genome-wide expression patterns. *PNAS*.
- [11] M. Ester, H. Kriegel, J. Sander, X. Xu, et al. 1996. A density-based algorithm for discovering clusters in large spatial databases with noise. *KDD*.
- [12] J. Geusebroek, G. J. Burghouts, and A. W.M. Smeulders. 2005. The Amsterdam library of object images. *IJCV*.
- [13] C. S. Greenberg, D. Bansé, G. R. Doddington, D. Garcia-Romero, J. J. Godfrey, T. Kinnunen, A. F. Martin, A. McCree, M. Przybocki, and D. A. Reynolds. 2014. The NIST 2014 speaker recognition i-vector machine learning challenge. *Odyssey*.
- [14] A. Haghighi and D. Klein. 2010. Coreference resolution in a modular, entity-centered model. *NAACL-HLT*.
- [15] H. Han, H. Zha, and C. L. Giles. 2005. Name disambiguation in author citations using a k-way spectral clustering method. *JCDL*.
- [16] K. Heller and Z. Ghahramani. 2005. Bayesian hierarchical clustering. *ICML*.
- [17] K. Heller and Z. Ghahramani. 2005. Randomized algorithms for fast Bayesian hierarchical clustering.
- [18] S. Jain and R. M. Neal. 2004. A split-merge Markov chain Monte Carlo procedure for the Dirichlet process mixture model. *Journal of computational and Graphical Statistics*.
- [19] J. Kleinberg. 2000. The small-world phenomenon: An algorithmic perspective. *STOC*.
- [20] J. Kleinberg. 2006. Complex networks and decentralized search algorithms. *ICM*.
- [21] A. Kobren, N. Monath, A. Krishnamurthy, and A. McCallum. 2017. A Hierarchical Algorithm for Extreme Clustering. *KDD*.
- [22] P. Kohli, P. H. S. Torr, and L. Ladicky. 2009. Robust higher order potentials for enforcing label consistency. *IJCV*.
- [23] A. Krishnamurthy, S. Balakrishnan, M. Xu, and A. Singh. 2012. Efficient active algorithms for hierarchical clustering. *ICML*.
- [24] H. Lee, M. Recasens, A. Chang, M. Surdeanu, and D. Jurafsky. 2012. Joint entity and event coreference resolution across documents. *EMNLP/CoNLL*.
- [25] E. Liberty, R. Sriharsha, and M. Sviridenko. 2016. ALENEX. *Workshop on Algorithm Engineering and Experiments*.
- [26] Y. Malkov, A. Ponomarenko, A. Logvinov, and V. Krylov. 2014. Approximate nearest neighbor algorithm based on navigable small world graphs. *Information Systems*.
- [27] O. Russakovsky, J. Deng, H. Su, J. Krause, S. Satheesh, S. Ma, Z. Huang, A. Karpathy, A. Khosla, M. Bernstein, et al. 2015. Imagenet large scale visual recognition challenge. *IJCV*.
- [28] S. Singh, A. Subramanya, F. Pereira, and A. McCallum. 2011. Large-scale cross-document coreference using distributed inference and hierarchical models. *ACL*.
- [29] D. Wang and Y. Wang. 2018. An Improved Cost Function for Hierarchical Cluster Trees. *arXiv preprint*.
- [30] D. J. Watts and S. H. Strogatz. 1998. Collective dynamics of ‘small-world’ networks. *Nature*.
- [31] M. Wick, S. Singh, and A. McCallum. 2012. A discriminative hierarchical model for fast coreference at large scale. *ACL*.
- [32] S. Wiseman, A. M. Rush, and S. M. Shieber. 2016. Learning global features for coreference resolution. *NAACL-HLT*.
- [33] L. Zhang, M. Song, Z. Liu, X. Liu, J. Bu, and C. Chen. 2013. Probabilistic graphlet cut: Exploiting spatial structure cue for weakly supervised image segmentation. *CVPR*.
- [34] T. Zhang, R. Ramakrishnan, and M. Livny. 1996. BIRCH: an efficient data clustering method for very large databases. *SIGMOD*.

A APPENDIX

A.1 Proof of Theorem 1

Define the following properties:

DEFINITION 3 (STRONG CONNECTIVITY). Let $G = (X, E)$ be a graph and let $\mathcal{T}[v]$ be a tree rooted at a node v with leaves, $\text{lhs}(v) = X' \subseteq X$. v is **connected** if X' is a connected subgraph of G . v is **strongly connected** if every descendant of v is connected. v is a **maximal strongly connected node** if v is strongly connected and $\text{par}(v)$ is not strongly connected. Finally, the tree $\mathcal{T}$ satisfies **strong connectivity** if all connected nodes in $\mathcal{T}$ are strongly connected.

DEFINITION 4 (COMPLETENESS). Let $G = (X, E)$ be a graph and let $\mathcal{T}[v]$ be a tree rooted at a node v with leaves, $\text{lhs}(v) = X' \subseteq X$. Then v is **complete** if X' is a connected component in G . The tree $\mathcal{T}$ satisfies **completeness** if the set of connected components of G are (the leaves of) a tree consistent partition of $\mathcal{T}$.

According to Theorem 1, GRINCH always constructs a tree that satisfies completeness. To prove the theorem, we will show that after the addition of each new data point, the resulting tree satisfies strong connectivity and completeness. We analyze various subroutines of GRINCH and demonstrate how they preserve strong connectivity, completeness or both. In the proceeding lemmas and proofs, let $G = (X, E)$ be a graph and let f be a model that separates G .

LEMMA 1 (ROTATION LEMMA). Let $\mathcal{T}$ be a tree with $\text{lhs}(\mathcal{T}) = X$, and let x be a new data point to be added to $\mathcal{T}$. Then all nodes that were strongly connected before the addition of x are strongly connected after the addition of x , i.e., rotations preserve strong connectivity.

Note: while rotations preserve strong connectivity, they do not guarantee completeness. Therefore rotations are insufficient for proving Theorem 1.

PROOF. Let v be a maximal strongly connected node in $\mathcal{T}$ and assume that x is added as a leaf of v (rotations have not yet been applied). Consider two cases: (1) there exists an edge between x and some leaf in $\text{lhs}(v)$, and (2) there does not exist an edge between x and any leaf in $\text{lhs}(v)$.

Case 1: Let $L \subseteq \text{lhs}(v)$ be the set of v 's descendant leaves to which x is connected. Then x is initially added as a sibling of its nearest neighbor leaf, x' , and $x' \in L$ because f separates G . $\text{par}(x)$ is strongly connected because there exists an edge between x and x' .

The addition of x does not disconnect v or any strongly connected descendant of v . To see why, consider the siblings of the ancestors of x' before the addition of x . Any such sibling that was connected to x' , is, after the addition of x , also connected to $\text{par}(x)$ and thus remains strongly connected. Nodes that are not ancestors of x cannot be disconnected and thus, before rotations, strong connectivity is preserved.

Now consider subsequent rotations. By the logic above, x and its sibling, $x' = s(x)$, are connected. If a rotation succeeds then x and $\text{aunt}(x)$ are swapped. So long as $\text{aunt}(x)$ and $s(x)$ form a connected subgraph in G , i.e., $\phi(s(x), \text{aunt}(x)) = \phi(x, s(x)) = 1$, then the rotation preserves strong connectivity.

The only way for a rotation to disrupt strong connectivity is if x and $\text{aunt}(x)$ are swapped, and $s(x)$ and $\text{aunt}(x)$ do not form a connected subgraph in G , i.e., $\phi(x, s(x)) > \phi(s(x), \text{aunt}(x))$. But, because f separates G , $\phi(x, s(x)) > \phi(s(x), \text{aunt}(x)) \implies f(x, s(x)) > f(s(x), \text{aunt}(x))$ and so, in this case, a rotation will not be performed and the procedure terminates.

Case 2: If there does not exist an edge between x and any leaf in $\text{lhs}(v)$, then after x is made a sibling of some leaf $x'' \in \text{lhs}(v)$, v is no longer strongly connected and so strong connectivity has not been preserved. Since v was strongly connected before the addition of x , there exists an edge between $\text{lhs}(s(x))$ and $\text{lhs}(\text{aunt}(x))$. Since f separates G , $f(x, s(x)) < f(s(x), \text{aunt}(x))$, which triggers the rotate subroutine. Rotations proceed with respect to x at least until x is no longer a descendant of v , and thus, v remains strongly connected. Strongly connected nodes that are not descendants of v are unaffected by the rotations and so strong connectivity is preserved. $\square$

LEMMA 2 (GRAFTING LEMMA 1). Let $\mathcal{T}$ satisfy strong connectivity and completeness. Let v be a node in $\mathcal{T}$ such that v is either a maximal strongly connected node or not strongly connected. Then a graft operation initiated from v preserves strong connectivity and completeness.

PROOF. Let $\mathcal{T}$ be strongly connected and complete. Since $\text{lhs}(v)$ is not a strict subset of any connected component in G , there does not exist a non-empty subset s in $\text{lhs}(\mathcal{T}) \setminus \text{lhs}(v)$ such that $s \cup \text{lhs}(v)$ is a connected subgraph in G . For any node v' that is strongly connected but not maximal, there must be an edge connecting $\text{lhs}(v')$ and $\text{lhs}(s(v'))$ and $s(v')$ must be strongly connected, so $f(v', s(v')) > f(v, v')$. Therefore, an attempt to make any such v' the sibling of v fails.

If v'' is a maximal strongly connected node, an attempt to make v'' the sibling of v may succeed but this does not disconnect any strongly connected subtrees in $\mathcal{T}$. The same is true if v'' is not strongly connected. $\square$

LEMMA 3 (GRAFTING LEMMA 2). Let $\mathcal{T}$ be a tree such that $\text{lhs}(\mathcal{T}) = X$ and let $\mathcal{T}$ satisfy strong connectivity. Let v be strongly connected and let $\text{lhs}(v)$ be a strict subset of the vertices in some connected component, C , in G . Then, a graft initiated from v returns a node v' such that v' is strongly connected and $\text{lhs}(v) \subset \text{lhs}(v')$.

PROOF. Since $\text{lhs}(v)$ are a strict subset of the vertices in the connected component, C , there exists a non-empty subset s in $\text{lhs}(\mathcal{T}) \setminus \text{lhs}(v)$ such that $s \cup \text{lhs}(v)$ constitute the vertices in C . Let ℓ maximize $f(v, \ell)$ over all $\text{lhs}(\mathcal{T}) \setminus \text{lhs}(v)$. By the fact that $\text{lhs}(v)$ is a strict subset of a connected component, there must exist an edge between $\text{lhs}(v)$ and ℓ . Note that ℓ is the leaf found when the constrained nearest neighbor search from v is initiated in the first line of graft (Algorithm 1).

If $f(v, \ell) < f(\ell, s(\ell))$, then there must exist an edge between ℓ and a node in $\text{lhs}(s(\ell))$ and so $\text{par}(\ell)$ is strongly connected. If $f(v, \ell) < f(v, s(v))$, then there must exist an edge between a node in $\text{lhs}(v)$ and a node in $\text{lhs}(s(v))$ and so $\text{par}(v)$ is strongly connected. In both of these cases, we do not merge v with ℓ , but instead attempt another merge between two strongly connected nodes, either: $\text{par}(v)$ with ℓ , v with $\text{par}(\ell)$, or $\text{par}(v)$ with $\text{par}(\ell)$.

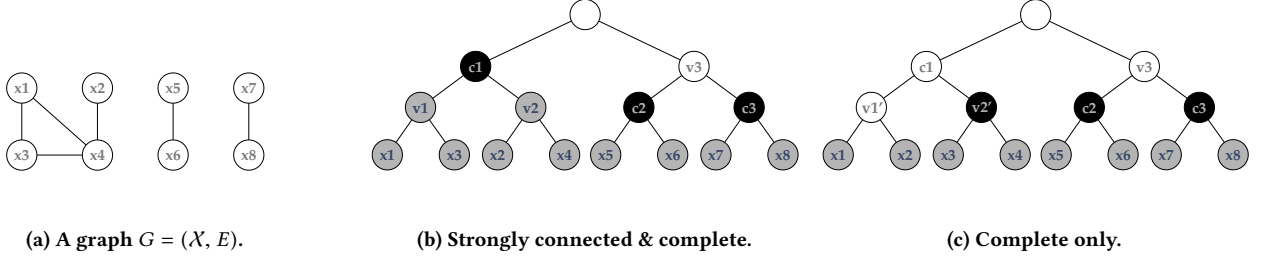


Figure 6: A graph G with 3 connected components (Figure 6a). In Figure 6b and Figure 6c, black-bordered nodes are strongly connected, thick-black-border nodes are maximal, gray bordered nodes are connected (but not strongly) and nodes with dashed borders are disconnected. The tree in Figure 6b satisfies strong connectivity and completeness. The tree in Figure 6c does not satisfy strong connectivity because v_1 is disconnected.

As before, the two nodes we are attempting to merge also have an edge between them.

Let v_1 and v_2 be two nodes involved in a merge and let $v_1 \in \text{ancs}(v)$ and $v_2 \in \text{ancs}(\ell)$. If at some point

$$f(v_1, v_2) > \max[f(v_1, s(v_1)), f(v_2, s(v_2))]$$

then v_2 is made a sibling of v_1 and the new parent of v_1 is returned. Since v_1 and v_2 are strongly connected and there exists an edge between $\text{lvs}(v_1)$ and $\text{lvs}(v_2)$, $\text{par}(v_1)$, which is created by the merge, is strongly connected, and the lemma holds.

If a merge is never performed, the recursion stops when $v_1 = v_2 = \text{lca}(v, \ell)$. In this case, the lca , which we return, is already strongly connected and, by definition, its leaves are a superset of $\text{lvs}(v)$. $\square$

LEMMA 4 (RESTRUCTURING LEMMA). *Let $v \in \mathcal{T}$ be strongly connected. Let $a \in \text{ancs}(v)$ be the deepest connected ancestor of v such that: a is not strongly connected, and all siblings of the nodes on the path from v to a are strongly connected. Then restruct on inputs v and a restructures $\mathcal{T}[a]$ so that a satisfies strong connectivity.*

PROOF. Let z be the deepest ancestor of v that is strongly connected with parent $\text{par}(z)$ that is disconnected. Since $\text{par}(z)$ is disconnected (but by assumption both z and $s(z)$ are connected), there are no edges between $\text{lvs}(z)$ and $\text{lvs}(s(z))$.

Let a' be a child of a and without loss of generality, $a' \notin \text{ancs}(z)$. Since a is the deepest connected ancestor of z , there must exist an edge between $\text{lvs}(z)$ and $\text{lvs}(a')$.

When computing the argmax of $f(z, \cdot)$ in the `restruct` method, a node, z' , that is connected to z will be returned and then swapped with $s(z)$. The new parent of z is strongly connected because z and z' are both strongly connected and there exists an edge between $\text{lvs}(z)$ and $\text{lvs}(z')$. Any subsequent swap attempted from a disconnected node with a connected ancestor succeeds and produces a new parent that is strongly connected.

Since a is connected and a swap among the descendants of a do not change $\text{lvs}(a)$, swapping preserves the connectedness of a . Therefore, swaps proceed until the node a is reached at which point a must be strongly connected.

Note that a swap attempt between a strongly connected node and a node to which it is not connected fails, because f separates G . A swap attempt between a connected node and a node to which it is connected succeeds and produces a new parent that is strongly connected. $\square$

We now prove Theorem 1.

PROOF. We show by induction that if GRINCH is used to build a tree, $\mathcal{T}$, over vertices, $\mathcal{X}$, then the connected components of G are a tree consistent partition in $\mathcal{T}$. Furthermore, $\mathcal{T}$ satisfies strong connectivity.

Clearly, the theorem holds for the base case: a tree with a single node.

Let $\mathcal{X} = \text{lvs}(\mathcal{T})$. Assume the inductive hypothesis: that $\mathcal{T}$ satisfies completeness and strong connectivity. Now vertex x arrives.

If there does not exist an edge between x and any other vertex in $\mathcal{X}$, then after rotations, $\mathcal{T}'$ satisfies completeness. Since $\forall a \in \text{ancs}(x)$, $\text{lvs}(a)$ is not a strict subset of any connected component in G , by Grafting Lemma 1, subsequent graft attempts from the ancestors of x preserve strong connectivity and completeness and so the theorem holds.

Assume that x is connected to some set of leaves $s \subseteq \text{lvs}(\mathcal{T})$. Since $\mathcal{T}$ satisfies strong connectivity, by the Rotation Lemma, after x is added and rotations terminate, $\mathcal{T}'$ satisfies strong connectivity. Note that $\mathcal{T}'$ may not satisfy completeness if, before the arrival of x , the leaves in s formed at least 2 distinct connected subgraphs in G .

After rotations, a series of graft attempts are performed. Consider the first graft initiated at $\text{par}(x)$. By Grafting Lemma 2, the attempt returns a strongly connected ancestor of x whose leaves are a strict superset of $\text{lvs}(x)$. If a merge is performed that moves a node v and makes it a sibling of v' , then strong connectivity may be violated. However, notice that the only nodes that can be disconnected by such a merge are the node that, prior to the merge, were ancestors of v and also descendants of $a = \text{lca}(v, v')$.

After the merge, a is restructured, and by the Restructuring Lemma, the resulting tree satisfies strong connectivity. Subsequent calls to `graft` proceed from a . Notice that each invocation of `graft` returns a new strongly connected node with a strictly larger number of descendant leaves, until the resulting tree satisfies completeness. Therefore, successive grafting followed by restructuring eventually returns a node whose leaves are a connected component of G . Ultimately, after rotations and grafting, $\mathcal{T}'$ must satisfy completeness and strong connectivity. $\square$