

Scalable Column Concept Determination for Web Tables Using Large Knowledge Bases

Dong Deng^y Yu Jiang^y Guoliang Li^y Jian Li^z Cong Yu[#]

^yDepartment of Computer Science, Tsinghua University, Beijing, China.

^zInstitute for Interdisciplinary Information Sciences, Tsinghua University, Beijing, China.

[#] Google Research, New York, USA

dd11@mails.thu.edu.cn, sunlight07@126.com, f.iguoliang, lijian83_g@tsinghua.edu.cn, congyu@google.com

ABSTRACT

Tabular data on the Web has become a rich source of structured data that is useful for ordinary users to explore. Due to its potential, tables on the Web have recently attracted a number of studies [6, 18] with the goals of understanding the semantics of those Web tables and providing effective search and exploration mechanisms over them. An important part of table understanding and search is column concept determination, i.e., identifying the most appropriate concepts associated with the columns of the tables. The problem becomes especially challenging with the availability of increasingly rich knowledge bases that contain hundreds of millions of entities [10, 31].

In this paper, we focus on an important instantiation of the column concept determination problem, namely, the concepts of a column are determined by fuzzy matching its cell values to the entities within a large knowledge base. We provide an efficient and scalable MapReduce-based solution that is scalable to both the number of tables and the size of the knowledge base and propose two novel techniques: knowledge concept aggregation and knowledge entity partition. We prove that both the problem of finding the optimal aggregation strategy and that of finding the optimal

Table 1: Example table from the Web (<http://celebrity.psyphil.com/actor/Tom-Hanks/filmography>).

1994	Forrest Gump	Robert Zemeckis	Tom Hanks, Robin Wright, Sally Field, ...
1998	Saving Private Ryan	Steven Spielberg	Tom Hanks, Tom Sizemore, Matt Damon, ...
1999	The Green Mile	Frank Darabont	Tom Hanks, David Morse, Bonnie Hunt, ...
1995	Toy Story	John Lasseter	Tom Hanks, Tim Allen, Wallace Shawn, ...
2002	Catch Me If You Can	Steven Spielberg	Tom Hanks, Christopher Walken, Leonardo DiCaprio, ...
2006	The Da Vinci Code	Ron Howard	Tom Hanks, Ian McKellen, Jean Reno, ...

Figure 1: Example knowledge base based on Wikipedia.

to denote the maximum matching, called fuzzy overlap of T and C , and $|V_C \cap E_T|$ to denote its value.

We can extend exact-matching functions by replacing $V_C \cap E_T$ with $V_C \tilde{\cap} E_T$ to define fuzzy-matching functions. Next we take fuzzy Jaccard similarity as an example.

$$\text{Fuzzy Jaccard Similarity: } \text{SIM}(C, T) = \frac{|V_C \tilde{\cap} E_T|}{|V_C| + |E_T| - |V_C \tilde{\cap} E_T|}.$$

For example, consider a column with two cell values $\{\text{"Shark Night 3D"}, \text{"Da Vinci Code"}\}$ and a type with two entities $\{\text{"Shark Night"}, \text{"The Da Vinci Code"}\}$. Suppose we use the Jaccard similarity. "Shark Night 3D" and "Shark Night" are fuzzy matching with similarity 0.67. "Da Vinci Code" and "The Da Vinci Code" are fuzzy matching with similarity 0.75. Thus the fuzzy overlap size is $0.67 + 0.75 = 1.42$ and the fuzzy Jaccard similarity is $\frac{1.42}{2+2-1.42} = 0.55$.

Scalability Requirement: Since there are hundreds of millions of columns and millions of entities, existing machine learning based methods [18, 27] are neither efficient nor scalable for such large datasets and it calls for scalable methods. To this end, we extend existing MapReduce-based similarity-join methods and develop effective techniques to support efficient and scalable column concept determination, while maintaining the annotation quality.

3. RELATED WORKS

MapReduce: We adopt the MapReduce framework to address the scalability challenge. MapReduce [8] is a dominant framework to support data-intensive parallel computation in shared-nothing clusters. In MapReduce, data is represented as $\langle \text{key}, \text{value} \rangle$ pairs and stored in a distributed file system (DFS) across different cluster nodes. MapReduce contains two core functions: **Map** and **Reduce**. The map phase loads and processes different partitions of the input data in parallel and outputs $\langle \text{key}, \text{value} \rangle$ pairs. The pairs are then hash-partitioned, sorted by the key, and sent across the cluster in a shuffle phase. The pairs that share the same key are merged in a sorted order and sent to the same reduce node, which then processes the pairs and writes new $\langle \text{key}, \text{value} \rangle$ pairs to files on the DFS.

Knowledge Bases: DBPedia [3] and Freebase [5] are similar projects that extract information from Wikipedia to cre-

ate an extensive structured knowledge base and supplement those data with knowledge that are collaboratively created by thousands of data-loving users. Freebase contained about 2 thousand types and 40 million entities. Probase [31] is a universal, probabilistic taxonomy which is harnessed from billions of Web pages using "isA" and "such as" patterns. Probase includes more than 2.7 million types and 22 million entities. Finally, Yago [25] is a large semantic knowledge base which is derived from Wikipedia and WordNet, and it contains 0.3 million types and 9 million entities.

Annotating Web Tables: Google's WebTables project [6] is the first to start exploring table data on the Web, which has been followed up by a number of other studies [9, 18, 27, 34]. Google has since launched a Table Search project [35] at <http://research.google.com/tables>, where external users can explore the table data. More recently, Google has also started exploring stitching similar tables on the same page [19] to provide better utility from the table data.

At the core of managing Web table data is the annotation problem, i.e., finding meanings for the cells and columns in the table. Existing studies used machine learning techniques to annotate tables [18] that are neither as scalable nor as efficient as the approaches we present in this paper. Other relevant studies (although those do not address the column concept identification problem) include: Yakout et. al. [34] studied how to augment entities with attribute values and discovering attributes using Web tables. Pimplikar et. al. [21] studied how to answer table queries based on column keywords. Quercini et. al. [22] utilized search engines to discover and annotate entities in Web tables. Hignette [15] annotated tables with ontology concepts based on relevance degrees. Wang et. al. [30] utilized the interconnections of Probase to understand Web tables.

Extracting Table Data into RDF Data: There are many studies on extracting table data into RDF data [12, 14]. Guo et. al. [12] used a schema mapping technique to extract and integrate data from tables into RDF. Han et. al. [14] studied how to translate spreadsheet data into RDF. Assem et. al. [26] studied how to annotate and covert quantitative data tables into RDF. Different from these methods, our work focuses on making fuzzy matching based annotation

techniques more scalable and efficient. How to apply our technique to the annotation techniques proposed by those studies is a challenge that we intend to pursue in the future.

Similarity Join: There are many recent studies on similarity join, which, given two sets of objects, find all similar object pairs based on a given similarity function and a threshold [32, 23, 33, 11, 2, 28]. Vernica et al. [28] and Metwally et al. [20] proposed a MapReduce framework to support similarity joins. However they required a similarity threshold and their techniques cannot be extended to support our ranking problem. Kim et al. [17] studied parallel algorithm for top- k similarity join with Euclidean distance constraint using MapReduce while we focus on finding top- k types for each column. These methods do not allow fuzzy matching between entities and cell values.

4. SIMILARITY-JOIN FRAMEWORK

In this section, we extend MapReduce-based similarity joins to support scalable column concept determination. We first take the overlap similarity as an example in Section 4.1 and then extend it to support other functions in Section 4.2.

4.1 Framework for Overlap Similarity

If we use a type to annotate a column, the type should share at least one common entity/cell value with the column. Thus for each column, we first find the types that share common entities with the column and then rank these types to return top- k types. To implement this idea in MapReduce, we introduce a two-stage similarity-join based framework. In the first stage, for each column, we identify the types which share at least one common entity/cell value with the column. In the second stage, we count the number of their shared entities/cell values and retrieve top k types with the largest similarity. Algorithm 1 illustrates the pseudo-code.

Stage 1: For each column, find the list of types that share at least one entity/cell value with the column.

Map: $\langle C, \{v_1, v_2, \dots\} \rangle \rightarrow \langle v_i, C \rangle$. $\langle T, \{e_1, e_2, \dots\} \rangle \rightarrow \langle e_i, T \rangle$.

It takes as input a sequence of columns (each column includes a set of cell values) from Web tables and a set of types (each type includes a set of entities) from a knowledge base. For each column C , it outputs key-value pairs $\langle v_i, C \rangle$ where v_i is any cell value of C . For each type T , it outputs key-value pairs $\langle e_i, T \rangle$ where e_i is any entity of T .

Reduce: $\langle v_i = e_i, list(C/T) \rangle \rightarrow \langle C, list(T) \rangle$.

It takes as input a set of key-value pairs $\langle v_i = e_i, list(C/T) \rangle$, where $v_i = e_i$ is a cell value/an entity, and $list(C/T)$ is a list of columns/types that contain the cell value/entity. The **Reducer** separates $list(C/T)$ into a list of columns, denoted by $list(C)$ and a list of types, denoted by $list(T)$. (We use a flag to differentiate C/T to be a column or a type.) For each column C in $list(C)$, it outputs a key-value pair $\langle C, list(T) \rangle$.

Combine: We can use **Combine** to do a local reduce. As it is not a focus of this paper, we will not discuss the details.

Stage 2: Compute top- k types of every column.

Map: $\langle C, list(T) \rangle \rightarrow \langle C, list(T) \rangle$.

It takes as input key-value pairs $\langle C, list(T) \rangle$, where $list(T)$ is a list of types that share an entity/cell value with C . It outputs the same key-value pairs as input.

Reduce: $\langle C, list(list(T)) \rangle \rightarrow \langle C, topk(T) \rangle$.

It takes as input key-value pairs $\langle C, list(list(T)) \rangle$, where $list(list(T))$ is a list of lists of types that share common enti-

Algorithm 1: Similarity-Join Based Framework

```

// the first stage
1 Map( $\langle C/T, \{v_1, v_2, \dots\} / \{e_1, e_2, \dots\} \rangle$ )
2   foreach column  $\langle C, \{v_1, v_2, \dots\} \rangle$  do output( $\langle v_i, C \rangle$ );
3   foreach type  $\langle T, \{e_1, e_2, \dots\} \rangle$  do output( $\langle e_i, T \rangle$ );
4 Reduce ( $\langle v_i = e_i, list(C/T) \rangle$ )
5   foreach  $C/T$  in  $list(C/T)$  do
6     if  $C/T$  is a type then add  $T$  to  $list(T)$ ;
7     else if  $C/T$  is a column then add  $C$  to  $list(C)$ ;
8   foreach  $C$  in  $list(C)$  do output( $\langle C, list(T) \rangle$ );
// the second stage
9 Map ( $\langle C, list(T) \rangle$ )
10  output( $\langle C, list(T) \rangle$ );
11 Reduce ( $\langle C, list(list(T)) \rangle$ )
12  Initialize a hash map  $\mathcal{H}$ ;
13  foreach  $T$  in  $list(list(T))$  do
14    if  $T \in \mathcal{H}$  then  $\mathcal{H}(T) = \mathcal{H}(T) + 1$ ;
15    else  $\mathcal{H}(T) = 1$ ;
16  topk( $T$ ) = top- $k$  types with largest  $\mathcal{H}(T)$ ;
17  output( $\langle C, topk(T) \rangle$ );

```

ties/cell values with C . We use a hash based method to compute the overlap similarity of C and a type T in $list(list(T))$. We first initialize a hash table. Then for each type in $list(list(T))$, if it is not in the hash table, we add it into the hash table and set the occurrence number as 1; otherwise we increase the number by 1. After scanning all types, the occurrence number of a type is exactly the overlap similarity of the type with the column. Based on the number, we rank the types and add the top- k types with the largest overlap similarity into $topk(T)$. Finally it outputs key-value pairs $\langle C, topk(T) \rangle$.

EXAMPLE 1. Figure 2 shows four columns and four types. Figure 3 illustrates a running example that uses our framework to label each column. In the first stage, the **Map** step reads the column file and the type file. For each column (type), it outputs $\langle \text{cell value, column} \rangle (\langle \text{entity, type} \rangle)$ pairs. For C_1 , it outputs six key-value pairs $\langle v_1, C_1 \rangle, \dots, \langle v_6, C_1 \rangle$. For T_2 , it outputs two key-value pairs $\langle e_1, T_2 \rangle, \langle e_2, T_2 \rangle$. In the **Reduce** step, it reads the key-value pairs grouped by the key (cell value or entity). Consider $\langle v_1 = e_1, \{C_1, C_2, T_1, T_2\} \rangle$. T_1, T_2 share an entity e_1 with C_1, C_2 . It outputs $\langle C_1, \{T_1, T_2\} \rangle$ and $\langle C_2, \{T_1, T_2\} \rangle$. In the second stage, the **Map** step outputs the same key-value pairs as the input. The **Reduce** step computes top- k types for each column by counting the occurrence number of each type in the key-value pair. Consider $\langle C_1, \{\{T_1, T_2\}, \{T_1, T_2\}, \{T_1, T_3\}, \{T_1, T_3\}\} \rangle$. T_1 appears four times, and T_2 and T_3 appear twice. Thus the top-1 type for C_1 is T_1 . Similarly the top-1 type for C_2 (C_3) is T_2 (T_3). Notice that we cannot label C_4 since it does not share any entity with any type.

4.2 Extension to Other Functions

Exact-matching Similarity Functions: Different from the overlap similarity, besides $|V_C \cap E_T|$, other exact-matching functions, e.g., weighted Jaccard, also rely on $|V_C|$, $|E_T|$, and the weight of each entity/cell value. To support these functions, we add these parameters into the key-value pairs, and in the **Reduce** step of the second stage, we compute the sim-

Column C ₁		Type T ₁ (/film)	
v ₁	Forrest Gump	e ₁	Forrest Gump
v ₂	Catch Me If You Can	e ₂	Catch Me If You Can
v ₃	Saving Private Ryan	e ₃	Saving Private Ryan
v ₄	Schindler's List	e ₄	Schindler's List
v ₅	Shark Night 3D	e ₅	Shark Night
v ₆	Da Vinci Code	e ₆	The Da Vinci Code
v ₇	The Terminal	e ₇	The Green Mile
v ₈	Full Metal Jacket	e ₈	Heartbreak Ridge

Column C ₂		Type T ₂ (/film/drama)	
v ₁	Forrest Gump	e ₁	Forrest Gump
v ₂	Catch Me If You Can	e ₂	Catch Me If You Can
v ₇	The Terminal	e ₇	The Green Mile

Column C ₃		Type T ₃ (/film/war)	
v ₃	Saving Private Ryan	e ₃	Saving Private Ryan
v ₄	Schindler's List	e ₄	Schindler's List
v ₈	Full Metal Jacket	e ₈	Heartbreak Ridge

Column C ₄		Type T ₄ (/film/thriller)	
v ₅	Shark Night 3D	e ₅	Shark Night
v ₆	Da Vinci Code	e ₆	The Da Vinci Code

Figure 2: Example columns and types.

ilarity based on these parameters. Thus we can extend this framework to support other exact-matching functions.

Fuzzy-matching Similarity Functions: To support fuzzy-matching functions, we slightly modify the framework. For each cell value (or entity), we generate its signatures. A cell value and an entity is similar only if they share at least one common *signature*. There are many signature strategies, e.g., q -gram [11] and prefix signature [28]. We can use any of them in our method. Next we modify the basic framework to support fuzzy-matching functions based on the signatures.

Stage 1: For each column, find the list of types that share at least one common *signature* with the column.

Map: $\langle C, \{v_1, \dots\} \rangle \rightarrow \langle \text{sig}, \langle C, v_i \rangle \rangle$. $\langle T, \{e_1, \dots\} \rangle \rightarrow \langle \text{sig}, \langle T, e_j \rangle \rangle$.

For each column C , for each cell value v_i , it computes the signatures and for each signature sig , generates key-value pairs, $\langle \text{sig}, \langle C, v_i \rangle \rangle$. For each type T , for each entity e_j , it computes the signatures and for each signature sig , generates key-value pairs $\langle \text{sig}, \langle T, e_j \rangle \rangle$.

Reduce: $\langle \text{sig}, \text{list}(\langle C, v_i \rangle / \langle T, e_j \rangle) \rangle \rightarrow \langle C, \text{list}(\langle T, v_i, e_j \rangle) \rangle$.

It separates the list into the column list and the type list. For each column, for each type, if the column has a cell value that shares a common signature with an entity of the type, it adds the type into a type list $\text{list}(\langle T, v_i, e_j \rangle)$. Finally it outputs key-value pairs $\langle C, \text{list}(\langle T, v_i, e_j \rangle) \rangle$.

Stage 2: Compute top- k types of every column.

Map: $\langle C, \text{list}(\langle T, v_i, e_j \rangle) \rangle \rightarrow \langle C, \text{list}(\langle T, v_i, e_j \rangle) \rangle$.

It outputs the same key-value pairs as the input.

Reduce: $\langle C, \text{list}(\text{list}(\langle T, v_i, e_j \rangle)) \rangle \rightarrow \langle C, \text{topk}(T) \rangle$.

For each column, it gets a list of lists of types. Then it computes the fuzzy-matching similarity using existing algorithms [29]. Based on the fuzzy-matching similarity, we rank the types and identify the top- k types with the largest fuzzy-matching similarity to annotate the column.

EXAMPLE 2. Consider column C_4 and type T_4 . Suppose the signatures of $v_5 = \text{"Shark Night 3D"}$ and $e_5 = \text{"Shark Night"}$ are $\text{sig}_1 = \text{"Shark"}$. The signatures of $v_6 = \text{"Da Vinci Code"}$ and $e_6 = \text{"The Da Vinci Code"}$ are $\text{sig}_2 = \text{"Vinci"}$. In the *Map* step of the first stage, we generate pairs $\langle \text{sig}_1, \langle C_4, v_5 \rangle \rangle$ and $\langle \text{sig}_2, \langle C_4, v_6 \rangle \rangle$ for C_4 and $\langle \text{sig}_1, \langle T_4, v_5 \rangle \rangle$ and $\langle \text{sig}_2, \langle T_4, v_6 \rangle \rangle$ for T_4 . In the *Reduce* step, we generate $\langle C_4, \langle T_4, v_5, e_5 \rangle \rangle$ and

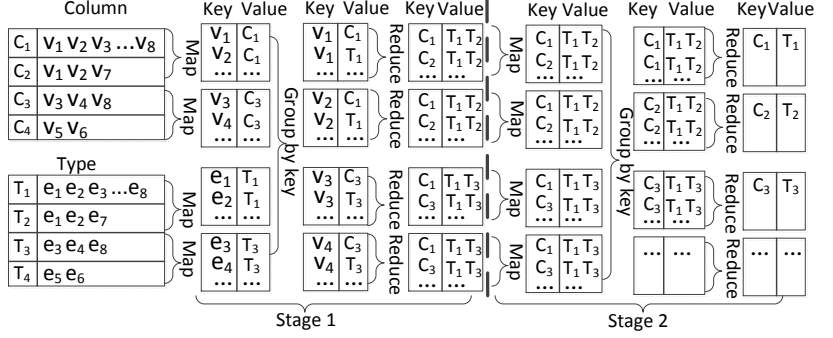


Figure 3: Running example of our framework ($v_1 = e_1$, $v_2 = e_2$, $v_3 = e_3$, $v_4 = e_4$. $v_5 \sim e_5$, $v_6 \sim e_6$ ($\sim$: similar)).

$\langle C_4, \langle T_4, v_6, e_6 \rangle \rangle$. In the *Reduce* step of the second stage, we compute the fuzzy similarity. The top-1 type of C_4 is T_4 . Thus the fuzzy matching functions can annotate C_4 by tolerating inconsistencies between entities and cell values.

4.3 Observations

Optimization Goal: To improve the performance, our optimization goal is to reduce the number of key-value pairs. Although we can use the prefix-filtering technique [28], it has the following problems. First, it depends on a given threshold and cannot support our top- k problem. Second, it involves other expensive overhead[28], e.g., pre-defining a global order and selecting high-quality prefixes based on the global order. To address these problems, we propose effective pruning technique based on the following observations.

Observation 1: We can aggregate the types into different groups. For each entity, to generate the key-value pairs, we use a group to replace the types in the group. Since each group contains multiple types, this method can reduce the number of key-value pairs. In addition, a group may be used by multiple entities, and we can share the computations for the types in the same group. We will discuss how to utilize groups to label columns in Section 5.

Observation 2: We can partition the cell values into different partitions. For each column, we utilize the partitions to generate its key-value pairs. Since each partition contains multiple cell values, this method can reduce the number of key-value pairs. Moreover, if the cell values of a column fall into one partition, we can utilize the partition to directly annotate the column and prune other unnecessary partitions. We will discuss the details in Section 6.

In the following sections, we focus on improving the efficiency and scalability. For simplicity, we take overlap similarity as an example, and our techniques can be easily applied to other similarity functions.

5. KNOWLEDGE TYPE AGGREGATION

For each entity, we want to reduce the number of key-value pairs produced for the entity. Let $\mathcal{L}(e)$ denote the list of types that contain entity e . We aggregate types in $\mathcal{L}(e)$ and generate several disjoint groups $\mathcal{G}_1(e)$, $\mathcal{G}_2(e)$, ..., $\mathcal{G}_l(e)$. Instead of producing key-value pairs $\langle e, T \rangle$ for every type $T \in$

$L(e)$, we generate key-value pairs $he;G(e)i$ for every group $G(e)$ for $1 \leq i \leq l$. That is we use group $G(e)$ to replace (multiple) types in the group. Obviously this method can reduce the number of key-value pairs and can improve the performance. There are several research challenges in this type aggregation based method. The first one is how to utilize the groups to compute top- k types for each column. We formally introduce an aggregation framework to address this issue in Section 5.1. The second challenge is to quantify different aggregation strategies. We propose a cost based model to analyze the group quality in Section 5.2. The third one is to efficiently generate high-quality groups. We devise efficient algorithms to address this challenge in Section 5.3.

5.1 Aggregation Framework

Given a knowledge base, for each entity e , we aggregate the types in its inverted list $L(e)$ and partition them into l disjoint groups $G_1(e), G_2(e), \dots, G_l(e)$, such that

- (1) Subset: $G(e)$ is a subset of $L(e)$ for $1 \leq i \leq l$;
- (2) Completeness: $\bigcup_{1 \leq i \leq l} G(e) = L(e)$; and
- (3) Disjoint: $G(e) \cap G_j(e) = \emptyset$ for $i \neq j$.

Let $G_T = \{G_1, G_2, \dots, G_g\}$ denote the set of distinct groups for all entities. Let H_T^G denote a hash table from the distinct groups to lists of types, i.e., each group G is associated with a list of types contained in G , $H_T^G(G)$.

For example, consider the types and entities in Figure 2. $L(e_1) = \{T_1, T_2\}$, $L(e_2) = \{T_1, T_2\}$, $L(e_3) = \{T_1, T_3\}$, $L(e_4) = \{T_1, T_3\}$, $L(e_5) = \{T_1, T_4\}$, $L(e_6) = \{T_1, T_4\}$, $L(e_7) = \{T_1, T_2\}$, $L(e_8) = \{T_1, T_3\}$. The basic framework will generate 16 key-value pairs. Suppose we aggregate them into three groups, $G_1 = \{T_1, T_2\}$, $G_2 = \{T_1, T_3\}$, $G_3 = \{T_1, T_4\}$. We generate 8 key-value pairs $he_1;G_1i$, $he_2;G_1i$, $he_3;G_2i$, $he_4;G_2i$, $he_5;G_3i$, $he_6;G_3i$, $he_7;G_1i$, and $he_8;G_2i$.

Suppose we can get the groups of each entity, i.e., $he;fG_1(e), G_2(e), \dots, G_l(e)g$. We will discuss how to generate the high-quality groups in Section 5.3. Next we discuss how to use the groups to compute top- k types of each column. We still employ a two-stage MapReduce framework. Algorithm 2 illustrates the pseudo-code.

Stage 1: For each column, find the list of types that share at least one entity/cell value with the column.

Map $hC;f v_1; v_2; \dots; g; i \mapsto h v_i; C. he;fG_1(e); \dots; G_l(e)g; i \mapsto he;G(e)i$.

The input and output for the columns are the same as those of the basic framework. For the types, the input is a set of key-value pairs $he;fG_1(e); G_2(e); \dots; G_l(e)gi$. For each entity, it outputs key-value pairs $he;G(e)i$ for $1 \leq i \leq l$.

Reduce $h v = e; list(C=G)i \mapsto h C; list(G)i$.

Different from the basic framework, for each column, it outputs key-value pairs $hC; list(G)i$, where G is a group which has a type sharing a common entity with C .

Stage 2: Compute top- k types of every column.

Map $hC; list(G)i \mapsto h C; list(G)i$.

Similarly it outputs key-value pairs $hC; list(G)i$.

Reduce $hC; list(list(G))i \mapsto h C; topk(T)i$.

For each column C , it gets a list of lists of groups. We still use a hash based method to compute the overlap similarity of C and a type $T \in G \in list(list(G))$. We first initialize a hash table. Then for each group G in $list(list(G))$, we get the list of types in the group, i.e., $H_T^G(G)$. For each type in $H_T^G(G)$, we compute its occurrence number similar as the basic framework. Finally we identify top- k types with the

Algorithm 2: Aggregation-based Algorithm

```
// the first stage
1 Map hC=e;f v1; v2; ::; g; i ↦ hC=fG1(e); G2(e); ::; g; i
2   foreach column hC;f v1; v2; ::; g; i do output( h vi; C);
3   foreach entity he;fG1(e); G2(e); ::; g; i do
4     output( he; G(e)i);
4 Reduce( h v = e; list(C=G)i
5   foreach C=T in list(C=G) do
6     if C=T is a group then add G to list(G);
7     else if
```


(a) Variable gadget $\mathcal{X}^g$ (b) Clause gadget $\mathcal{C}^g$

Figure 14: Variable and Clause Gadgets. Each vertex represents a distinct type and each hyperedge represents an entity.

Figure 15: An example of reducing 3-CNF to finding aggregation strategy problem.

PROOF. First, given any aggregation strategy, we can easily check whether the cost is no larger than t in polynomial time.

Next we prove the decision problem is NPC by a reduction from the 3-CNF-SAT problem. In particular, given any instance ϕ of a 3-CNF-SAT problem, we construct an instance of our decision problem $\langle \mathcal{H}, t \rangle$ such that the 3-CNF formula ϕ is satisfiable if and only if there exists an aggregation strategy of $\mathcal{H}$ with cost no larger than t . In the reduction we let $t = 16n + 24k$ where n and k are respectively the numbers of variables and clauses in ϕ . The construct of $\mathcal{H}$ is as follows (an example can be found in Figure 15). For each variable x_i in ϕ , we create a variable gadget $\mathcal{X}_i^g$ which consists of three

entities $e_1^i = \{T_1^i, T_2^i, T_3^i, T_4^i, T_5^i, T_8^i\}$, $e_2^i = \{T_2^i, T_3^i, T_6^i, T_8^i\}$ and $e_3^i = \{T_4^i, T_5^i, T_7^i, T_8^i\}$. Recall that the tip T_6^i is labeled with the literal x_i and tip T_7^i is labeled with literal $\bar{x}_i$. For each clause C_j in ϕ , we create a clause gadget $\mathcal{C}_j^g$ which consists of seven entities $e_4^j = \{T_9^j, T_{10}^j, T_{11}^j\}$, $e_5^j = \{T_{11}^j, T_{13}^j, T_{14}^j\}$, $e_6^j = \{T_{10}^j, T_{12}^j, T_{13}^j\}$, $e_7^j = \{T_{10}^j, T_{11}^j, T_{13}^j\}$, $e_8^j = \{T_{10}^j, T_{11}^j\}$, $e_9^j = \{T_{11}^j, T_{12}^j\}$ and $e_{10}^j = \{T_{13}^j, T_{14}^j\}$. Recall that the tips T_{11}^j , T_{12}^j and T_{13}^j are labeled with the three literals of clause C_j . The variable gadgets and clause gadgets are connected as follows. We identify together each tip in $\mathcal{C}_j^g$ with the tip in the variable gadget that has the same label (i.e., they corresponds to the same literal). This finishes the construction of $\mathcal{H}$.

The reduction creates 3 entities and 8 types for each variable x_i and 7 entities and 6 types for each clause C_j . We can perform the reduction algorithm which totally creates $3n + 7k$ entities and $8n + 6k$ types in polynomial time as it produces each entity or type in the constant time.

We next show that the 3-CNF formula ϕ is satisfiable if and only if there exists an aggregation strategy for $\mathcal{H}$ with cost no larger than $t = 16n + 24k$.

First, suppose there is a truth assignment that satisfies ϕ . We prove there must exist an aggregation strategy with cost no larger than $t = 16n + 24k$ as follows. For each $1 \leq i \leq n$, if $x_i = \text{True}$, we aggregate variable gadget $\mathcal{X}_i^g$ using aggregation strategy $\mathcal{S}_{V_1}$ with cost 16 and $G_x^i = \langle T_6^i \rangle$. If $x_i = \text{False}$, we aggregate variable gadget $\mathcal{X}_i^g$ using aggregation strategy $\mathcal{S}_{V_2}$ with cost 16 and $G_x^i = \langle T_7^i \rangle$. By using this aggregation strategy to aggregate every variable gadget $\mathcal{X}_i^g (1 \leq i \leq n)$, the total cost is $16n$. For each $1 \leq j \leq k$, as ϕ is satisfied, C_j contains at least one literal $l_m (1 \leq m \leq 3)$ whose value is **True**. Suppose l_m is the first literal in C_j whose value is evaluated to be **True** in the truth assignment. We aggregate clause gadget $\mathcal{C}_j^g$ using aggregation strategy $\mathcal{S}_{L_m}$ with cost 24 and $G_C^j = \langle T_{l_m}^j \rangle$. Moreover, we can see that the group G_C^j also exists in the aggregation strategy chosen for variable gadget $\mathcal{X}_i^g$ since l_m is **True** (more specifically, $G_C^j = G_x^i$). Hence, the additional cost $c_j = 0$ (the cost of $|G_C^j|$ has been already counted in the variable clause). Using this strategy, the total cost to aggregate all $\mathcal{C}_j^g$ is $24k$ and the additional cost is 0. Based on the discussion above, the overall cost to aggregate entity set $\mathcal{H}$ is $16n + 24k$.

Next suppose that there is an aggregation strategy of $\mathcal{H}$ with cost no larger than $16n + 24k$. We prove there must exist an assignment that satisfies ϕ . For any variable gadget $\mathcal{X}_i^g (1 \leq i \leq n)$, the minimum aggregation cost is 16. Thus the minimum cost to aggregate all variable gadgets is $16n$. For any clause gadget $\mathcal{C}_j^g (1 \leq j \leq k)$, the minimum aggregation cost is 24. Thus the minimum cost to aggregate all clause gadgets is $24k$. Thus to achieve an aggregation cost of $\mathcal{H}$ no larger than $16n + 24k$, we need to apply the aggregation strategy with minimum cost to all the variable gadgets and clause gadgets and all the additional costs $c_j (1 \leq j \leq k)$ should be 0. According to Property 1 and Property 2, there must exist a group G_x^i in each variable gadget $\mathcal{X}_i^g$ which only contains one of the two tips labeled with x_i and $\bar{x}_i$ and there must exist a group G_C^j in each clause gadgets $\mathcal{C}_j^g$ which only contains one of the three tips labeled with the literals. To make all the additional cost $c_j = 0$, each group G_C^j should also exist in its corresponding variable gadgets $\mathcal{X}_i^g$ which is exactly G_x^i . We set $x_i = \text{True}$ if $G_x^i = G_C^j$ contains only the tip labeled with x_i . Otherwise we set $x_i = \text{False}$ if $G_x^i = G_C^j$ contains only the tip labeled with $\bar{x}_i$. By this assignment, any clause $C_j (1 \leq j \leq k)$ is satisfied as C_j contains a literal x_i when $x_i = \text{True}$ (as G_C^j contains the tip labeled with x_i) and contains literal $\bar{x}_i$ when $x_i = \text{False}$ (as G_C^j contains the tip labeled with $\bar{x}_i$). Note that this assignment will not set $x_i = 1$ and $x_i = 0$ simultaneously. This is because based on Property 1, any aggregation of a variable gadget of cost 16 contains exactly one singleton group which contain exactly one tip. Thus the Theorem 4 is proved. $\square$