

Towards Universal Dense Blocking for Entity Resolution

Tianshu Wang
Institute of Software, CAS
HIAS, UCAS
China
tianshu2020@iscas.ac.cn

Hongyu Lin
Institute of Software, CAS
China
hongyu@iscas.ac.cn

Xianpei Han
Institute of Software, CAS
China
xianpei@iscas.ac.cn

Xiaoyang Chen
University of CAS
China
chenxiaoyang19@mails.ucas.ac.cn

Boxi Cao
Institute of Software, CAS
China
boxi2020@iscas.ac.cn

Le Sun
Institute of Software, CAS
China
sunle@iscas.ac.cn

ABSTRACT

Blocking is a critical step in entity resolution, and the emergence of neural network-based representation models has led to the development of dense blocking as a promising approach for exploring deep semantics in blocking. However, previous advanced self-supervised dense blocking approaches require domain-specific training on the target domain, which limits the benefits and rapid adaptation of these methods. To address this issue, we propose UNIBLOCKER, a dense blocker that is pre-trained on a domain-independent, easily-obtainable tabular corpus using self-supervised contrastive learning. By conducting domain-independent pre-training, UNIBLOCKER can be adapted to various downstream blocking scenarios without requiring domain-specific fine-tuning. To evaluate the universality of our entity blocker, we also construct a new benchmark covering a wide range of blocking tasks from multiple domains and scenarios. Our experiments show that the proposed UNIBLOCKER, without any domain-specific learning, significantly outperforms previous self- and unsupervised dense blocking methods and is comparable and complementary to the state-of-the-art sparse blocking methods.

CCS CONCEPTS

• **Information systems** → **Entity resolution**; *Deduplication*.

KEYWORDS

Dense Blocking, Entity Resolution, LLM, Pre-Training

1 INTRODUCTION

Entity resolution (ER) [18, 37] aims to identify and merge duplicate records that refer to the same real-world entity, serving as a critical component in data cleaning, knowledge integration, and information management [9, 28, 43]. A key challenge in entity resolution is the quadratic time complexity of pairwise record comparisons. To overcome this, the blocking-and-matching pipeline has become the mainstream solution in ER systems, where the blocking step groups records into small sets using computable features, and the matching step accurately matches records in each set with complex techniques. Blocking is essential for efficiently discarding incorrect options and reducing the time complexity in entity resolution.

The boost of neural network-based representation models has led to the development of *dense blocking*, a promising approach to exploring deep semantics in the blocking step. Specifically, dense

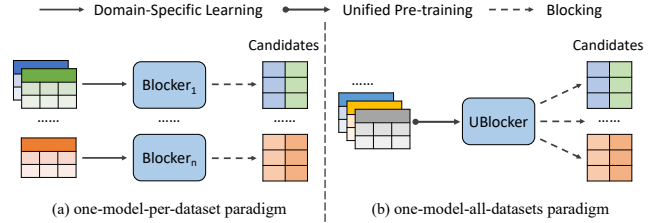


Figure 1: Illustration of *one-model-per-dataset* paradigm and *one-model-all-datasets* paradigm for dense blocking.

blocking utilizes an encoder, a neural network architecture, to convert entity records into dense semantic vectors. The distances between these vectors serve as a similarity measure between pairs of records. By performing nearest neighbor searches for each record, this technique generates a set of potential matching pairs or blocks. Compared to conventional syntactic-based sparse blocking approaches [46], dense blocking is more effective in modeling deep semantic information while maintaining scalability. As a result, dense blocking has become a research hotspot in the field of entity blocking [1, 15, 36, 51, 55, 63, 65].

Unfortunately, recent studies [40, 63] have shown that, in contrast to supervised dense blocking methods [4, 65], self-supervised dense blocking methods [51, 55] do not offer substantial benefits and may be inferior to the direct use of pre-trained word or sentence embeddings. As illustrated in Figure 1(a), we argue that this is because these methods learn *domain-specific* models, which means that for each new blocking task, a unique representation model needs to be trained using the available records. Dense blocking methods in this paradigm demand considerable time and computational resources for each blocking. Learning on domain-specific data can also lead to overfitting or collapse, making the outcomes not worth the effort. However, the potential of self-supervised dense blocking learning from extensive data remains largely unexplored.

In this paper, we aim to extend the methodology of dense blocking from the one-model-per-dataset to a more generalized one-model-all-datasets paradigm, as shown in Figure 1. Motivated by recent advances in large language models [5, 48], we develop a universal dense blocker, UNIBLOCKER, by exploiting the power of cross-domain self-supervised learning to pre-train on large amounts of tabular data. Figure 2 shows the overall pre-training procedure

of UNIBLOCKER. UNIBLOCKER is built upon a Transformer-based architecture and trained with contrastive learning on GitTables [26], a large-scale relational tabular corpus. Specifically, we filter and process the tables from this corpus into records, generate positive and negative pairs using data paraphrasing techniques at different levels, and obtain record representations by passing them through the encoder. The contrastive learning algorithm refines the model’s capability to represent records. Based on the aforementioned process, UNIBLOCKER, pre-trained on a heterogeneous, multi-domain, and multi-topic corpus of 1 million tables, acquires the capability to universally block entity collections across diverse domains and scenarios without further domain-specific learning.

To ascertain the effectiveness and universality of UNIBLOCKER, we gathered and constructed a new benchmark for universal blocking evaluation. The benchmark comprises 19 datasets spread across 9 domains and 4 scenarios, which enable a thorough evaluation of the performance of UNIBLOCKER and other blocking techniques. Unlike previous studies that evaluate methods on partially labeled datasets and typically with identical schema records, this new benchmark emphasizes the ability of a blocking technique to transfer seamlessly across domains and scenarios without requiring extra resources. As a result, the constructed benchmark evaluates entity blockers based on both effectiveness and scalability, considers both homogeneous and heterogeneous entity record collections [43], and includes datasets of different scales from various domains.

Thorough experiments have been carried out on the constructed benchmark. The experimental results highlight the effectiveness and universality of UNIBLOCKER, verifying the feasibility and benefits of universal dense blocking. Specifically, experiments show that UNIBLOCKER, without any additional domain-specific training, outperforms state-of-the-art (SOTA) self- and unsupervised dense blocking systems by almost 10% in mean average precision. Furthermore, UNIBLOCKER is comparable to SOTA sparse blocking methods and is more effective in some complex and dirty scenarios where duplicate records are semantically similar. In addition, UNIBLOCKER demonstrated similar and even better efficiency when blocking large volumes of records. Overall, these experimental results provide strong empirical support for the effectiveness, universality, and scalability of UNIBLOCKER.

Contribution. Our contributions can be summarized as follows:

- We provide new insight into the limitations of existing self-supervised dense blocking methods that follow the one-model-per-dataset paradigm, and present our vision for one-model-all-datasets universal dense blocking.
- We propose a unified contrastive pre-training framework for developing the universal dense blocker, UNIBLOCKER. We design several optimizations to represent structural records better.
- We construct a new benchmark for universal blocking evaluation, which covers a wide range of blocking tasks from multiple domains and scenarios. This also lays the foundation for future research on automatic parameter search or good attribute identification for blocking [40, 47].
- We conduct thorough experiments to demonstrate the effectiveness and universality of UNIBLOCKER by comparing it with the state-of-the-art dense and sparse blocking systems.

2 PRELIMINARIES AND BACKGROUND

In this section, we first introduce the problem definition and dense blocking. Then we discuss the limitations of blocking benchmarks.

2.1 Problem Definition

ER involves identifying records that refer to the same real-world entity, often called matches or duplicates. To avoid the quadratic number of comparisons with complicated entity matching algorithms, the blocking step is introduced in ER systems, which uses basic features to quickly scan all record pairs to achieve a reasonable trade-off between the number of comparisons and the loss of performance. Formally, the blocking step is defined in Definition 1.

DEFINITION 1 (BLOCKING). *Given entity record collections consisting of n records $\mathcal{E} = \{e_i \mid 1 \leq i \leq n\}$, where each record e is composed of attributes and values $e = \{(attr_i, val_i) \mid 1 \leq i \leq k\}$, the blocking step aims to generate candidate pairs $\mathcal{C} = \{(e_i, e_j)\} \subseteq \mathcal{E} \times \mathcal{E}$ to approximate the golden duplicate pairs $\mathcal{M} = \{(e_i, e'_i)\} \subseteq \mathcal{E} \times \mathcal{E}$ under a given recall or comparison budget.*

Efficiency requirements for blocking require blocking algorithms to be scalable in terms of time and resource usage as records increase, and to generate a preferably linear quantity of candidates.

2.2 Dense Blocking

Dense blocking is an emerging technique for identifying potential matches in entity resolution with neural models. Unlike dense entity matching [2, 35] where two records are simultaneously passed to a neural network, dense blocking transforms entity records into dense vectors independently and discovers potential matches through nearest neighbor search. Compared to conventional sparse blocking methods, dense blocking allows for the modeling of deep semantic information, which reduces labor-intensive involvement [20, 30, 40]. The technique has evolved from using context-free word embeddings [52] to contextual pre-training language models [31, 63] and from supervised learning [65] to self-supervised learning [51, 55]. The search for potential matches has also shifted from threshold-based [52, 65] to cardinality-based [29, 51]. Despite the significant progress made by self-supervised blocking methods [51, 55], the one-model-per-dataset paradigm limits their application, making it necessary to develop the one-model-for-all-datasets paradigm.

2.3 Benchmarks for Blocking

Unlike entity matching, the standard benchmarks for evaluating blocking methods are limited. Some studies [65] have utilized non-public datasets, making experiment replication more challenging. Some other works [51] leverage entity matching datasets [35] directly, which may introduce evaluation bias for blocking due to selection bias during dataset construction. The process of building an entity matching dataset without golden labels involves blocking and labeling. Sampling bias occurs when a particular blocking method is used to generate candidates, or when only a subset of candidates are labeled by an unspecified selection. Despite numerous recent works on entity blocking [40, 47, 63], the evaluation datasets used in these studies are limited to tables with identical schema. Therefore, existing benchmarks typically focus only on datasets with biases or in a specific scenario. As a result, previous

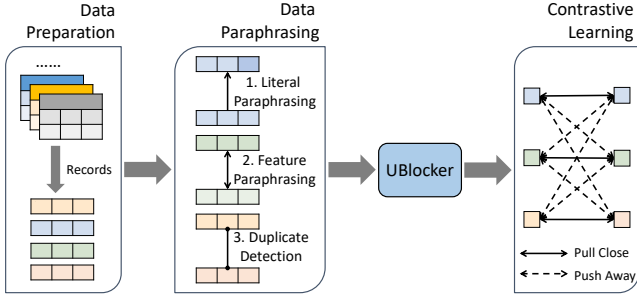


Figure 2: Overview of UNIBLOCKER pre-training procedure.

benchmarks cannot effectively evaluate the universal performance of entity blockers in a general scenario.

3 UNIBLOCKER: A UNIVERSAL DENSE BLOCKER FOR ENTITY BLOCKING

In this section, we propose a contrastive pre-training framework to develop UNIBLOCKER, a universal dense blocker capable of blocking records across domains without domain-specific learning.

3.1 Overview

Figure 2 shows the overall pre-training procedure of UNIBLOCKER, which is a deep Transformer-based neural network trained on a large-scale relational tabular corpus using the contrastive learning algorithm. Specifically, UNIBLOCKER starts by performing column detection and data cleaning to transform tables into sets of records. Subsequently, each record is transformed into a sequence of embeddings, and a Transformer-based encoder is employed to generate a dense semantic vector. To learn an effective record encoder, we leverage data paraphrasing techniques to generate positive and negative record pairs and contrastive learning is conducted by pulling positive pairs closer while pushing negative pairs further apart.

In the following, we will explain the process of acquiring the pre-training data from the tabular corpus. We will then describe the model architecture of UNIBLOCKER and demonstrate the approach by which UNIBLOCKER is trained using self-supervised contrastive learning on tabular records.

3.2 Pre-training Data Conditioning

The divergence between unstructured natural language and structured records (i.e., tuples of attributes and values) makes the direct application of pre-trained natural language models for table-related tasks suboptimal [13]. To fill this gap, in this paper, we propose to address this problem by pre-training UNIBLOCKER on large-scale relational tables to align the unstructured natural language models with structured and semi-structured records.

Specifically, we leverage GitTables [26] as our pre-training corpus because it is the first large table corpus derived from offline relational tables. Despite the initial filtering conducted by the GitTables authors, we find that GitTables still contains a very large number of tables with various data types, including annotated data for machine learning, tabular data of statistics or features, log data, and so on. Since not all tables contain entities suitable for learning

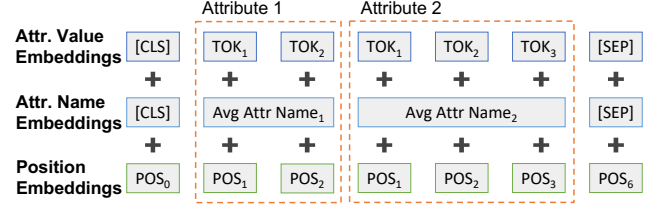


Figure 3: Record input representation. Record embeddings are the sum of token embeddings of attribute value, attribute name embeddings and relative position embeddings.

entity blocking, we used decision trees to learn a heuristic filtering algorithm based on column type detection to remove non-entity tables. Our filtering method heuristically determines the type of each column and then removes non-entity tables based on the results. More details about data conditioning, including column type detection and filtering, can be found in our supplementary materials.

After cleaning the corpus, we transform all tables into sets of entity records. Each row in a table is regarded as a record. These sets of records are then fed into the UNIBLOCKER to learn a universal entity blocker, as illustrated below.

3.3 Encoder Architecture

In this section, we present how we transform a record into its dense representation. Specifically, we propose a new record embedding technique that facilitates the conversion of structured records into a sequence of embeddings, allowing Transformer-based architectures to better model individual records for blocking. To achieve this, we transform the common tiled serialization into a combination of multiple embeddings. A Transformer-based encoder is then adapted to obtain the representation of the records.

Record Embedding. Serialization is a prerequisite for passing entity records to the Transformer architecture. In ER, it is common to serialize each record $e = \{(\text{attr}_i, \text{val}_i)\}_{1 \leq i \leq k}$ by concatenating attribute name and value tuples sequentially [31]. However, we have observed that this serialization technique is effective for matching but less effective for blocking. The reason for this is that blocking does not involve interactions between record pairs. Consequently, the aforementioned method can suffer from performance degradation when attribute values are missing or of limited length and the input sequence would mainly consist of attribute names. Additionally, the Transformer’s use of position embeddings introduces the order of attributes as a significant factor. This has sparked debates about attribute shuffling in various studies [31, 34, 55].

Based on these observations, we propose a record embedding method with bare-bone serialization, as shown in Figure 3. First, we concatenate all non-null attribute values of a record to form the input sequence: $\text{val}_1 \text{ } \text{val}_2 \text{ } \dots \text{ } \text{val}_k$, where $\text{ } \text{ }$ stands for space. We then incorporate attribute name information into token embeddings within a sequence by augmenting the embedding of each token with the average embedding of the corresponding attribute names. Finally, to overcome the problem of misaligned focus on attributes, we propose a per-attribute position embedding method

Table 1: Literal Data Paraphrasing Actions.

Level	Action	Description
Char.	substitute	Simulate keyboard distance and OCR engine error
	insert, substitute, swap, delete	Apply paraphrasing randomly
Token	substitute	Find suitable word from dictionary
	insert, split, crop, swap, delete	Apply paraphrasing randomly
Cell	substitute	Substitute attribute name from dictionary
	crop, delete	Apply paraphrasing randomly

as a replacement for the original position embedding. The per-attribute position embedding is obtained by resetting the position id at the beginning of each attribute value and converting it to embedding sequences. The resulting serialized record embedding includes attribute name and relative position information.

Transformer-based Encoder. In recent years, Transformer-based pre-trained language models (PLMs) have attracted considerable attention and achieved state-of-the-art performance in numerous areas, including ER and DI&P [29, 31, 55], demonstrating the power of this architecture. Despite many practices in table pre-training [13], none of them are trained for record representations and can be applied to blocking directly, so we employ the Transformer-based architecture and build on existing PLMs.

Given a record, we first transform it into its corresponding record embedding with the abovementioned technique. This embedding is then fed into a series of self-attention and feed-forward neural network modules. Our goal in this study is to obtain dense record representations for blocking, so we leverage the encoder-only architecture. After the record undergoes transformation through the Transformer-based encoder, we use the average embeddings of all tokens in the last layer as the final record representation. For more information on the Transformer architecture, please refer to [53].

3.4 Data Paraphrasing for Positive Pair Generation

One of the main challenges in self-supervised learning is to generate positive samples that are optimal for downstream tasks [59]. To address this challenge, we leverage two data paraphrasing techniques, namely literal paraphrasing and feature paraphrasing, from different perspectives to generate positive pairs. We also introduce duplicate detection that identifies potential positives in the original data, preventing performance degradation due to false negatives.

Literal Paraphrasing The presence of redundant, missing, misplaced, misspelled, and noisy records in entity resolution has received widespread attention [19, 38]. Taking advantage of the record embedding design, we can focus on simulating the dirtiness of records in the real scenario. Based on this intuition, we perform literal paraphrasing at three levels: character, token, and attribute. The paraphrasing actions are summarized in Table 1.

Feature Paraphrasing. Our approach differs from Sudowoodo that applies the cutoff operation at the embedding level [55], while we employ dropout for feature paraphrasing across all Transformer layers. Dropout is a technique that prevents overfitting of neural networks during training by randomly deactivating neurons with

a probability p . When deployed, the encoder creates a slightly varied feature representation of the same input. Therefore, we encode identical records with dropout twice to obtain different representations serving as positive samples for training, which has been confirmed to be highly effective by prior research [21].

Duplicate Detection. In addition to manual data paraphrasing, existing “duplicate” records can also serve as positive pairs for training. Two records are considered “duplicate” if they share high similarity and can be transformed into each other by a few steps of literal paraphrasing. To avoid performance degradation due to false negatives, we propose a heuristic method to detect duplicates. This method involves computing the similarities between records using various string similarity measures [12]. For each pair of records in the batch, if any of the computed similarities exceed a predetermined threshold s , we consider them to be positive pairs.

From another perspective, duplicate detection can be seen as weak supervised learning based on sparse similarity measures. Since the blocking phase cannot afford such computationally intensive measures, models can be trained in this way to generate similar representations for similar records across different measures.

3.5 Model Training

In the following, we will describe how we train UNIBLOCKER with the aforementioned techniques. Formally, given a batch of N records donated as $E = \{e_i\}_{1 \leq i \leq N}$, we generate positive samples for each of them by applying literal paraphrasing, resulting in $E' = \{e'_i\}_{1 \leq i \leq N}$. We then apply record embedding to these $2N$ instances and feed them into the encoder with feature paraphrasing, obtaining their representations $Z = \{z_i\}_{1 \leq i \leq N}$ and $Z' = \{z'_i\}_{1 \leq i \leq N}$. For the constructed pairs from these $2N$ records, we consider the pairs generated through data paraphrasing and duplicate detection as positive pairs, while the remaining pairs in the batch are regarded as negative pairs. Finally, since each instance may have multiple positive samples, we adapt Circle Loss [50] as our loss function. Suppose for each sample e and its representation z we have J positives and K negatives, where $J + K = N$. We denote the similarity of their paraphrased version representations to z as $\{s_p^j\}_{1 \leq j \leq J}$ and $\{s_n^k\}_{1 \leq k \leq K}$ respectively. The loss function is described as:

$$\mathcal{L}_{\text{Circle}} = \sum_{i=1}^N \log \left[1 + \sum_{j=1}^J \sum_{k=1}^K \exp \left(\gamma \left(s_n^k - s_p^j + m \right) \right) \right] \quad (1)$$

where γ is a scaling factor acting similar to τ in InfoNCE loss [8] and m is a margin for better similarity separation.

4 BENCHMARKING UNIVERSAL BLOCKING

This section presents a new benchmark and evaluation to overcome the limitations of blocking benchmarks, as discussed in Section 2.3. This aims to assess the universality of blocking methods.

4.1 Overview of Proposed Benchmark

We propose a new benchmark that prioritizes diversity while ensuring accuracy by collecting public datasets with full golden labels. Specifically, we focus on diversity across multiple factors:

Table 2: Datasets of our proposed benchmark. We superscript datasets appeared in previous literature with asterisk.

Dataset	Domain	Size	# Attr.	# Pos
Effectiveness				
census	census	841	5	344
cora	citation	1879	18	64578
notebook	notebook	343	13	2152
notebook2	notebook	1661	1	2814
altosight	product	835	4	4082
altosight2	product	2006	5	4393
fodors-zagats_homo [*]	restaurant	553, 331	6	112
dblp-acm [*]	citation	2616, 2294	4	2224
dblp-scholar [*]	citation	2616, 64263	4	5438
abt-buy_homo [*]	product	1081, 1092	2	1098
walmart-amazon_homo [*]	electronics	2554, 22074	5	1154
fodors-zagats_heter	restaurant	553, 331	4,5	112
imdb-dbpedia [*]	movies	27615, 23182	4,7	22863
amazon-google [*]	product	1363, 3226	4,4	1300
abt-buy_heter	product	1081, 1092	2,3	1098
walmart-amazon_heter	electronics	2554, 22074	16,21	1154
movies	movies	18984	14,31,10	4135
Scalability				
songs	musics	1E+6	7	-
citeseer-dblp	citation	1.8E+6, 2.5E+6	6	-

- **Diverse Scenarios.** The entity resolution scenario includes two aspects, the number of collections and schema heterogeneity. Unlike previous benchmarks, we cover the case where the input is from a single or double collection of records. For input from a single collection, which is often known as dirty ER or deduplication, there is no schema heterogeneity and all records have the same attributes. For input from two collections (i.e., record linkage), we consider two scenarios where the schema is homogeneous or heterogeneous. Although schema-agnostic methods can handle both scenarios, little work has been done on the performance differences between these two scenarios. Additionally, we cover the scenario where entity collections come from multiple sources, which is more challenging due to schema heterogeneity.
- **Diverse Domains.** For each scenario, we aim to include datasets from as many different domains as possible. The inconsistency of duplicate records from different domains places different demands on blocking methods.
- **Diverse Properties.** In addition to the domains, the properties of datasets are also critical factors that can affect the performance of different methods. For instance, a table with a single attribute may be considered textual, whereas a table with many attributes may pose a challenge for finding key attributes. Additionally, the similarity between duplicate and non-duplicate records and the average number of duplicates per record are internal properties of the dataset that may influence method performance.

By incorporating the diversity of these factors, our benchmark provides a comprehensive evaluation of blocking methods. This

also lays the foundation for future research on automatic parameter search or good attribute identification for blocking [40, 47].

Table 2 provides an overview of the datasets included in our benchmark, which comprises a total of 17 datasets from 9 different domains. These datasets include 6 from one source, 10 from two sources and 1 from three sources. For the datasets sourced from two sources, we include 5 homogeneous and 5 heterogeneous datasets. The number of attributes per entity record ranges from 1 to 31, with an average number of duplicates per record ranging from 0.02 to 1. Different “notebook” and “altosight” datasets originate from different years of the SIGMOD Programming Contest. After preprocessing and labeling, most datasets are found to be homogeneous, especially those from DeepMatcher [35]. For the heterogeneous datasets, we include both artificially heterogeneous datasets, constructed by renaming and merging attributes as [54], and naturally heterogeneous datasets where raw data was available. For example, we constructed “fodors-zagats_heter” and “abt-buy_heter” by renaming and merging attributes. In contrast, “walmart-amazon_heter” comprises the original heterogeneous records with more than 15 attributes and labels from “walmart-amazon_homo”. The multi-source “movies” dataset is a merge of three datasets (IMDB-TMDB, IMDB-TVDB, and TMDb-TVDB) obtained from JedaIToolkit [45]. Details of the datasets and processing can be found in the supplementary material.

4.2 Evaluation

The evaluation metrics commonly used to assess the effectiveness of entity blocking methods, which are also used in our benchmark, are pair completeness (PC) and pair quality (PQ). These metrics are widely employed in the literature and provide a means to quantify the performance of entity blocking methods:

- (1) Pair completeness (also known as recall), is the fraction of true matched pairs identified: $PC = |C \cap M|/|M|$
- (2) Pair quality (also called precision), is the fraction of matched pairs among the candidate pairs: $PQ = |C \cap M|/|C|$

Balancing pair completeness and pair quality is a challenging trade-off, as it is difficult to achieve both simultaneously under efficiency constraints. Since the matching phase provides further assess to pair quality and discarded pairs cannot be recovered, the blocking phase typically prioritizes pair completeness. However, the variability in the number of candidate pairs generated by different blocking methods and the complexity of the internal parameters make it difficult to fairly compare the performance of different methods. To address this issue, a configuration optimization evaluation setting [40] has been proposed, which aims to tune all parameters of a given method to achieve a threshold for pair completeness. In this paper, we also follow this setting to set the threshold of pair completeness to 90%. In practice, however, it is unrealistic to search for all parameters in a limited amount of time due to the missing golden labels and the huge parameter space. Therefore, we stick default parameters for all methods to assess their universality across all datasets. In addition, for cardinality-based methods using k-nearest neighbor (kNN) search, we establish an upper bound on k of 100 as a comparison budget and check the candidates in nearest neighbor order. We evaluate the blockers based on whether they can achieve pair completeness greater than 90%. If both blockers

meet this criterion, we compare their maximum pair quality. Otherwise, we compare their pair completeness. This approach allows us to assess the effectiveness of the blockers under different conditions and provide a comprehensive evaluation of their performance.

In addition, we leverage mean average precision (mAP) to evaluate the overall performance of kNN approaches. mAP, which is the area under the precision-recall curve (AUC-PR), is calculated as $\sum_i (R_i - R_{i-1}) P_i$ where R_i and P_i represent PC and PQ at $k = i$.

5 EXPERIMENTS

In the following section, we conduct thorough experiments to determine the effectiveness and universality of UNIBLOCKER, as well as to identify its strengths and weaknesses in comparison to sparse blocking methods. We first evaluate the performance of UNIBLOCKER against both domain-specific and unified sentence-based dense blocking methods. Then we compare UNIBLOCKER with sparse blocking methods. Finally, we conduct ablation studies.

5.1 Experimental Setup

The baselines involved in our experiments include:

- **DeepBlocker** [51]: DeepBlocker is a design space exploration for self-supervised blocking solutions. We use the best overall solution Autoencoder for comparison.
- **Sudowoodo** [55]: Sudowoodo is a contrastive, PLMs-based, self-supervised learning framework used to solve various data integration and preparation tasks, including blocking.
- **STransformer** [49]: STransformer is a sentence embedding framework trained on sentence pairs using siamese networks. It’s the best blocking solution with pre-trained embeddings [63].
- **Blocking Workflows** [45]: Blocking Workflows is the general term for the sparse blocking pipeline, including block building, block cleaning, and comparison cleaning.
- **Sparkly** [47]: Sparkly is a kNN blocker that uses Lucene to perform top-k blocking. The well-known tf/idf measure has shown impressively superior performance in [47].

We implemented UNIBLOCKER using Python 3.10, Pytorch 1.13, and the Transformers [58] library. We chose to use the all-mpnet-base-v2 from STransformer [49] as the initial pre-training weight and set the maximum sequence length to 256. We set the probability of applying literal paraphrasing actions to 0.01, the dropout probability for feature paraphrasing to 0.15, and the similarity threshold s for duplicate detection to 0.85. The scaling factor γ and margin m of the Circle Loss are the default 80 and 0.4, respectively. We used the AdamW optimizer with a learning rate of $1e-5$, a linear learning rate schedule without warm-up steps, and a batch size of 128 for pre-training. The total number of pre-training steps is 10,000.

For all baselines, we followed their proposed implementations and settings. As discussed in Section 4.2, we kept consistent parameters for each method across all datasets to assess universality. For Blocking Workflows, we adhered to the DBW setting [40]. For Sparkly, we employed the default Okapi BM25 scoring function and the 3-gram tokenizer. In terms of dense blocking baselines, we only serialized record values for effectiveness, using FAISS as the indexer. If the dataset contains more than two entity collections, we merge them into a single one to convert it to dirty ER blocking.

5.2 Comparison with Dense Blocking Methods

UNIBLOCKER vs Domain-Specific Methods. To assess the effectiveness and universality of UNIBLOCKER, we first compare UNIBLOCKER with SOTA domain-specific self-supervised dense blocking methods. Experimental results presented in Table 3 indicate that UNIBLOCKER outperforms previous domain-specific dense blocking methods by almost 10% mAP on average. In contrast to the lack of a clear winner between DeepBlocker and Sudowoodo, UNIBLOCKER consistently delivers the best performance among them on almost all datasets, further proving its effectiveness by reducing the number of nearest neighbors required. These results indicate that UNIBLOCKER not only addresses the problem of limited downstream data, but also eliminates the need for domain-specific model training. Pre-training allows deep models to reach their full potential and learn a more general representation of the structured records. In conclusion, our experimental results provide compelling evidence for the advantage of self-supervised dense blocking learning on large-scale open-domain data.

Finding 1. UNIBLOCKER outperforms SOTA domain-specific dense blockers, showing the feasibility of developing universal dense blockers via a unified domain-independent pre-training.

UNIBLOCKER vs Unified Sentence-Based Models. To explore the necessity of pre-training on structured data, we compare UNIBLOCKER with the sentence embedding language model. The performance of STransformer, which served as our initialized model for pre-training, is detailed in Table 3. Our experimental results show that UNIBLOCKER outperforms STransformer by over 10% mAP on average. STransformer performs relatively poorly on some less challenging but more structured datasets, such as “fodors-zagats_homo”, “abt-buy_homo”, and “walmart-amazon_homo”. This highlights the barriers to applying language models directly to blocking due to the divergences between unstructured natural language and structured records. Overall, our results emphasize the necessity of training the universal blocker on structured and semi-structured records.

Finding 2. UNIBLOCKER outperforms sentence-based models, supporting the need to align unstructured language models with structured and semi-structured records by pre-training.

Domain-Specific Fine-tuning for UNIBLOCKER. To investigate the potential benefits of domain-specific fine-tuning for UNIBLOCKER, we further fine-tune UNIBLOCKER in the one-model-per-dataset paradigm. The mAP changes of the fine-tuned models are also presented in Table 3. Comparing the experimental results with and without fine-tuning, we observe that UNIBLOCKER derives limited benefit from fine-tuning except for “imdb-dbpedia”. This finding suggests that the domain-independent pre-training of UNIBLOCKER has resulted in a universal blocking capability that is sufficient to address domain-specific blocking challenges without the need for further fine-tuning. Consequently, this finding reaffirms the effectiveness of UNIBLOCKER in solving blocking problems across multiple domains.

Finding 3. UNIBLOCKER benefits little from domain-specific self-supervised fine-tuning, indicating that UNIBLOCKER can learn sufficient knowledge from domain-independent pre-training.

Table 3: Overall results of different dense blocking methods. As discussed in Section 4.2, if both methods exceed the preset PC threshold (90%), we evaluate their PQ, otherwise, we compare their PC. K is the number of nearest neighbors needed to reach the reported PC and ΔmAP is the change after domain-specific fine-tuning. The best results are highlighted in bold.

Datasets	Domain-Specific Learning						Unified Pre-training					
	DeepBlocker			Sudowoodo			STransformer			UniBLOCKER		
	mAP	PC / PQ	K	mAP	PC / PQ	K	mAP	PC / PQ	K	mAP ΔmAP	PC / PQ	K
census	9.83	90.41 / 1.44	33	8.73	90.12 / 1.43	39	12.32	91.86 / 7.03	9	16.04 -0.05	90.99 / 15.14	5
cora	41.84	66.49 / 29.48	100	27.29	46.47 / 25.58	100	50.80	77.64 / 37.54	100	51.21 $+0.03$	77.19 / 37.39	100
notebook	27.20	90.29 / 7.67	97	29.69	90.20 / 10.22	84	28.49	90.01 / 10.50	84	43.36 $+0.83$	90.38 / 16.90	52
notebook2	40.54	90.05 / 2.69	62	33.48	86.57 / 2.25	100	34.39	90.12 / 2.28	86	41.55 $+0.01$	90.37 / 3.97	52
altosight	8.94	52.96 / 3.13	100	24.77	85.82 / 6.33	100	26.12	90.00 / 7.72	85	38.83 -0.20	90.32 / 13.74	51
altosight2	12.56	46.93 / 1.18	100	16.59	53.83 / 1.86	100	22.99	81.72 / 2.44	100	29.04 -0.02	89.80 / 2.81	100
fodors-zagats_homo	60.51	100.00 / 21.01	1	59.90	99.11 / 20.83	1	53.42	93.75 / 9.85	2	60.51 $+0.00$	100.00 / 21.01	1
dblp-acm	90.29	97.39 / 82.80	1	90.02	97.21 / 82.65	1	88.21	95.28 / 81.00	1	91.88 -0.06	99.19 / 84.33	1
dblp-scholar	67.85	90.35 / 20.52	9	57.40	90.01 / 2.16	85	69.85	91.49 / 23.38	8	73.12 $+0.31$	91.55 / 31.19	6
abt-buy_homo	43.55	90.06 / 1.02	90	69.79	90.52 / 15.31	6	42.83	90.52 / 7.07	13	79.80 $+0.36$	93.16 / 31.51	3
walmart-amazon_homo	39.37	90.64 / 2.73	15	44.13	90.29 / 4.08	10	30.73	90.03 / 1.51	27	53.22 -0.44	93.33 / 14.06	3
fodors-zagats_heter	60.51	100.00 / 21.01	1	56.07	92.86 / 19.51	1	20.69	90.18 / 0.29	65	59.35 $+0.00$	98.21 / 20.64	1
imdb-dbpedia	0.00	0.00 / 0.00	100	8.42	43.63 / 0.36	100	8.07	46.06 / 0.38	100	23.93 $+2.19$	59.53 / 0.49	100
amazon-google	43.98	90.15 / 1.12	77	52.29	90.23 / 7.17	12	46.20	90.46 / 6.64	13	62.31 -0.17	90.38 / 17.24	5
abt-buy_heter	45.18	90.06 / 1.43	64	71.73	90.79 / 18.43	5	47.17	91.70 / 9.31	10	80.89 -0.01	93.16 / 31.51	3
walmart-amazon_heter	25.88	84.84 / 0.38	100	28.57	90.12 / 0.67	61	32.47	90.12 / 1.94	21	30.18 -0.65	90.03 / 0.75	54
movies	7.53	89.84 / 0.25	100	2.62	76.71 / 0.26	100	3.06	68.54 / 0.19	100	8.07 $+0.02$	90.04 / 2.20	13
Mean	36.80	80.03 / 11.64	61.76	40.09	82.62 / 12.89	53.24	36.34	85.85 / 12.30	48.47	49.60 $+0.13$	89.86 / 20.29	32.35

Table 4: Overall results of sparse, dense blocking methods and their ensemble. The best results are emphasized in bold and the better ensemble results are underlined. Backslash (\) indicates runtime errors.

Datasets	Sparse				Dense			Ensemble		
	Blocking Workflows	Sparkly3			UniBLOCKER			UniBLOCKER + Sparkly3		
		PC / PQ	mAP	PC / PQ	K	mAP	PC / PQ	K	mAP	PC / PQ
census	45.93 / 11.10	25.55	95.06 / 21.95	4	16.04	90.99 / 15.14	5	18.11	95.93 / 17.66	4
cora	97.51 / 65.32	49.60	74.06 / 37.94	100	51.21	77.19 / 37.39	100	<u>51.79</u>	<u>82.39 / 28.91</u>	100
notebook	87.45 / 12.35	40.93	90.33 / 14.83	61	43.36	90.38 / 16.90	52	<u>45.04</u>	90.06 / 16.79	39
notebook2	\	41.47	90.12 / 3.70	55	41.55	90.37 / 3.97	52	<u>41.96</u>	90.05 / 2.82	42
altosight	50.22 / 15.30	46.04	90.69 / 26.67	27	38.83	90.32 / 13.74	51	44.00	90.54 / 21.89	24
altosight2	\	28.50	90.03 / 3.51	85	29.04	89.80 / 2.81	100	28.46	<u>90.19 / 3.65</u>	49
fodors-zagats_homo	100.00 / 4.73	60.51	100.00 / 21.01	1	60.51	100.00 / 21.01	1	56.41	100.00 / 12.81	1
dblp-acm	100.00 / 4.18	91.56	98.74 / 83.94	1	91.88	99.19 / 84.33	1	87.81	99.69 / 75.87	1
dblp-scholar	100.00 / 0.10	74.27	92.16 / 31.40	6	73.12	91.55 / 31.19	6	71.60	92.24 / 26.17	5
abt-buy_homo	89.00 / 14.90	80.58	92.80 / 31.39	3	79.80	93.16 / 31.51	3	80.35	<u>94.26 / 36.38</u>	2
walmart-amazon_homo	99.77 / 0.27	56.93	90.47 / 20.44	2	53.22	93.33 / 14.06	3	54.92	94.37 / 14.26	2
fodors-zagats_heter	100.00 / 3.64	59.52	98.21 / 20.64	1	59.35	98.21 / 20.64	1	56.31	100.00 / 12.61	1
imdb-dbpedia	52.91 / 5.46	37.34	69.96 / 0.58	100	23.93	59.53 / 0.49	100	34.35	<u>71.43 / 0.31</u>	100
amazon-google	94.87 / 3.90	61.60	91.62 / 17.48	5	62.31	90.38 / 17.24	5	<u>64.05</u>	<u>92.54 / 20.30</u>	3
abt-buy_heter	90.05 / 14.80	80.85	90.06 / 45.70	2	80.89	93.16 / 31.51	3	81.24	94.99 / 36.38	2
walmart-amazon_heter	99.73 / 0.01	33.85	80.68 / 0.36	100	30.18	90.03 / 0.75	54	<u>38.18</u>	<u>90.03 / 1.11</u>	20
movies	87.17 / 0.27	14.48	90.37 / 2.35	12	8.07	90.04 / 2.20	13	12.00	<u>90.45 / 3.90</u>	5
Mean	86.31 / 10.42	51.97	89.73 / 22.58	33.24	49.60	89.86 / 20.29	32.35	50.98	<u>91.72 / 19.52</u>	23.53

5.3 Comparison with Sparse Blocking Methods

Effectiveness Evaluation. Initially, we compare the effectiveness of UniBLOCKER and SOTA sparse blocking methods. Our experimental results, shown in Table 4, reveal that Sparkly surpasses UniBLOCKER by about 2% mAP on average. However, it’s worth noting that UniBLOCKER is comparable to Sparkly on simple homogeneous datasets, and blocking better on some more difficult datasets, such as “walmart-amazon_heter” and “notebook”. In addition, the ensemble of UniBLOCKER and Sparkly can improve PC by up to 5% on

dataset “cora”. Lastly, self-supervised dense retrieval models have outperformed BM25 as more high-quality data become available and model parameters expand. This suggests that the performance of UniBLOCKER will continue to improve over time, showing great promise for universal self-supervised dense blocking.

Finding 4. UniBLOCKER is comparable to SOTA sparse blocking methods and they complement each other to a certain extent.

Efficiency Analysis. Furthermore, we perform efficiency analysis on scalability datasets from our proposed benchmark. To this end, we downsample the entity collections from 10^2 to 10^6 and compare

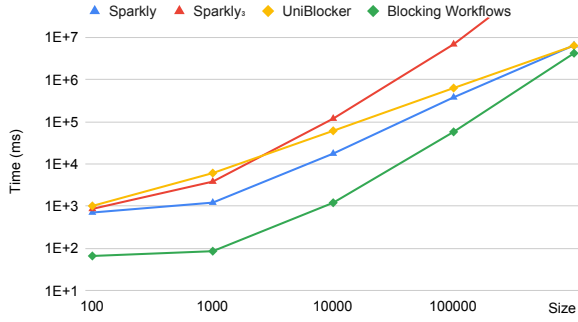


Figure 4: Average blocking times for varying approaches on scalability datasets with different record sizes.

Table 5: Average performance in ablation studies. UNIBLOCKER “w/o OD” means pre-training with sufficient out-domain data instead of open-domain GitTables. “RE”, “LP”, “FP”, and “DD” stand for record embedding, literal paraphrasing, feature paraphrasing, and duplicate detection, respectively.

Methods	mAP	PC	PQ	k
UNIBLOCKER w/o OD	46.48	89.37	18.00	35.76
UNIBLOCKER w/o RE	48.31	89.29	19.72	33.82
UNIBLOCKER w/o LP	49.09	89.40	20.23	34.71
UNIBLOCKER w/o FP	48.64	89.49	19.75	35.94
UNIBLOCKER w/o DD	48.31	89.06	19.65	36.12
UNIBLOCKER	49.60	89.86	20.29	32.35

the blocking time of different approaches. As shown in Figure 4, sparse blocking methods are efficient for small numbers of records, but the difference decreases as the number of records increases. We can see that UNIBLOCKER has the slowest time growth curve because it generates fixed-length dense representations for records. However, for Sparkly with a 3-gram tokenizer, the number of tokens and the number of record pairs with the same tokens grow non-linearly, posing challenges for index storage and nearest neighbor search. Therefore, q-gram technology improves the effectiveness but sacrifices some scalability at larger scales.

Finding 5. Sparse blocking methods are more efficient than UNIBLOCKER for small numbers of records, but the difference decreases as the number of records increases.

5.4 Ablation Study

We now perform ablation studies on UNIBLOCKER to validate the importance of specific components, as depicted in Table 5.

Effectiveness of Pre-training Data. To explore the effect of pre-training data diversity on UNIBLOCKER, we replace the training data extracted from open-domain GitTables with out-domain tables sourced from scalability datasets in our benchmark. This helps to establish whether the pre-training data serves solely to bridge the gap between natural language and structured records, or if it also contributes to the universality of UNIBLOCKER. Experimental results in Table 5 show that UNIBLOCKER achieves a significant performance

gain over the model trained on out-domain data, suggesting that open-domain data pre-training is crucial for its universality.

Effectiveness of Record Embedding. After replacing the proposed record embedding (RE) with the previous serialization way, a performance gap of 1.29% on mAP is observed, as presented in Table 5. We believe that the limitations of the prior serialization way have been concealed, as it was mainly used for entity matching where cross-record interactions occur within the same model. For blocking, however, only one record is passed to the model at a time. The inclusion of extra tokens and attribute names can impact the final representation of records, particularly when attribute values are short or highly missing. The proposed Record Embedding, therefore, can help Transformer-based architectures to better represent a single record and be more suitable for blocking.

Effectiveness of Data Paraphrasing. We then experiment with the effect of different data paraphrasing approaches by removing them in turn. As shown in Table 5 again, we find that both feature paraphrasing (FP) and lexical paraphrasing (LP) contribute to the performance of UNIBLOCKER, among which LP is more pronounced because this type of paraphrasing is more refined, diverse, and stable, without the chance of introducing false positives. Moreover, due to the complexity of the real data, the application of duplicate detection (DD) achieves a 1.29% mAP improvement.

Finding 7. Pre-training on open-domain data and the techniques applied in UNIBLOCKER are effective and necessary.

6 RELATED WORK

6.1 Blocking

Blocking, a crucial step in addressing the inherently quadratic complexity of entity resolution, has been approached in numerous ways [46]. Blocking Workflows [45], which originated in the late 1960s [18], involve block building [18, 41, 65], block cleaning [39, 42], and comparison cleaning [42, 44] to group records into blocks using signatures derived from records. However, challenges arise in precisely controlling candidate pair numbers, tuning pipeline [20, 30], and automating the tuning [40]. Sparse join [23, 47] operates by identifying syntactically similar pairs in collections using either threshold-based [7] or cardinality-based [3] conditions within a filter-verification framework. While effective, they sometimes struggle with loose conditions, necessitating approximation techniques [22, 33, 64] for scalability. Dense blocking represents records as dense vectors and performs similarity searches to find potential pairs [15]. Despite improving record representation through the use of pre-trained language models [31] or via self-supervised [51, 55] and active learning [29], performance is constrained by the divergences between natural language and structured records and the absence of record-specific pre-trained models.

6.2 Table Pre-training

Recent developments in table pre-training, motivated by the success of pre-training paradigms, have significantly impacted various downstream tasks [13]. Early stages used pre-training word embeddings on relational or web datasets [6, 24], which later evolved to Transformer-based language models. In the area of natural language processing, the focus is on tasks integrating tabular and textual data, such as table question answering [25, 61], fact verification [16, 32],

and table summarization [60], which typically take whole tables as input. In terms of data integration and preparation, the emphasis is on table-centric tasks such as table interpretation, cleaning, and argumentation, with most studies [11, 27, 56] pre-training models on web tables and supervised fine-tuning on specific tasks. The advent of data lakes has sparked interest in table union search, with contrastive pre-training frameworks [10, 14, 17] taking table columns as input. In the field of machine learning, there have been efforts [57, 62] to compare the performance of pre-trained deep models and traditional methods in classification and regression tasks on tabular tables comprising categorical and continuous features. Notably, none of these works pre-trained models for entity record representations and can be directly applied to blocking.

7 CONCLUSION

The main focus of this paper is to explore the potential of developing a universal dense blocker that can be applied across diverse scenarios and domains through large-scale self-supervised learning. To achieve this, we propose a unified contrastive pre-training framework that exploits the power of cross-domain pre-training for developing the universal dense blocker, UNIBLOCKER. The pre-training is performed via self-supervised contrastive learning on large-scale relational tables. We also introduce a new record embedding technique that allows Transformer-based architectures to better model individual records. In addition, we propose three data paraphrasing techniques to generate positives for training. To assess the effectiveness and universality of UNIBLOCKER, we create a new benchmark for universal blocking evaluation that covers a wide range of blocking tasks from multiple domains and scenarios. Experimental results demonstrate the effectiveness and universality of UNIBLOCKER, holding great promise for universal dense blocking.

REFERENCES

- [1] Mario Almagro, David Jiménez, Diego Ortego, Emilio Almazán, and Eva Martínez. 2022. Block-SCL: Blocking Matters for Supervised Contrastive Learning in Product Matching. *arXiv:2207.02008* (July 2022). *arXiv:2207.02008* [cs] <http://arxiv.org/abs/2207.02008>
- [2] Nils Barlaug and Jon Atle Gulla. 2021. Neural Networks for Entity Matching: A Survey. *ACM Trans. Knowl. Discov. Data* 15, 3 (June 2021), 1–37. <https://doi.org/10.1145/3442200>
- [3] Christian Böhm and Florian Krebs. 2002. High Performance Data Mining Using the Nearest Neighbor Join. In *Proceedings of the 2002 IEEE International Conference on Data Mining (ICDM 2002), 9–12 December 2002, Maebashi City, Japan*. IEEE Computer Society, 43–50. <https://doi.org/10.1109/ICDM.2002.1183884>
- [4] Alexander Brinkmann, Roece Shraga, and Christian Bizer. 2023. SC-Block: Supervised Contrastive Blocking within Entity Resolution Pipelines. <https://doi.org/10.48550/arXiv.2303.03132> *arXiv:2303.03132* [cs]
- [5] Tom B. Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, Sandhini Agarwal, Ariel Herbert-Voss, Gretchen Krueger, Tom Henighan, Rewon Child, Aditya Ramesh, Daniel M. Ziegler, Jeffrey Wu, Clemens Winter, Christopher Hesse, Mark Chen, Eric Sigler, Mateusz Litwin, Scott Gray, Benjamin Chess, Jack Clark, Christopher Berner, Sam McCandlish, Alec Radford, Ilya Sutskever, and Dario Amodei. 2020. Language Models are Few-Shot Learners. In *Advances in Neural Information Processing Systems 33: Annual Conference on Neural Information Processing Systems 2020, NeurIPS 2020, December 6–12, 2020, virtual*. Hugo Larochelle, Marc'Aurelio Ranzato, Raia Hadsell, Maria-Florina Balcan, and Hsuan-Tien Lin (Eds.). <https://proceedings.neurips.cc/paper/2020/hash/1457c0d6bfc4967418bfb8ac142f64a-Abstract.html>
- [6] Riccardo Cappuzzo, Paolo Papotti, and Saravanan Thirumuruganathan. 2020. Creating Embeddings of Heterogeneous Relational Datasets for Data Integration Tasks. In *Proc. 2020 ACM SIGMOD Int. Conf. Manag. Data*. ACM, Portland OR USA, 1335–1349. <https://doi.org/10.1145/3318464.3389742>
- [7] Surajit Chaudhuri, Venkatesh Ganti, and Raghav Kaushik. 2006. A Primitive Operator for Similarity Joins in Data Cleaning. In *Proceedings of the 22nd International Conference on Data Engineering, ICDE 2006, 3–8 April 2006, Atlanta, GA, USA*. IEEE Computer Society, 5. <https://doi.org/10.1109/ICDE.2006.9>
- [8] Ting Chen, Simon Kornblith, Mohammad Norouzi, and Geoffrey E. Hinton. 2020. A Simple Framework for Contrastive Learning of Visual Representations. In *Proceedings of the 37th International Conference on Machine Learning, ICML 2020, 13–18 July 2020, Virtual Event (Proceedings of Machine Learning Research, Vol. 119)*. PMLR, 1597–1607. <http://proceedings.mlr.press/v119/chen20j.html>
- [9] Vassilis Christophides, Vasilis Efthymiou, Themis Palpanas, George Papadakis, and Kostas Stefanidis. 2021. An Overview of End-to-End Entity Resolution for Big Data. *ACM Comput. Surv.* 53, 6 (Nov. 2021), 1–42. <https://doi.org/10.1145/3418896>
- [10] Tianji Cong, Fatemeh Nargesian, and H. V. Jagadish. 2023. Pylon: Semantic Table Union Search in Data Lakes. *CoRR abs/2301.04901* (2023). <https://doi.org/10.48550/arXiv.2301.04901> *arXiv:2301.04901*
- [11] Xiang Deng, Huan Sun, Alyssa Lees, You Wu, and Cong Yu. 2020. TURL: Table Understanding through Representation Learning. *Proc. VLDB Endow.* 14, 3 (Nov. 2020), 307–319. <https://doi.org/10.14778/3430915.3430921>
- [12] AnHai Doan, Alon Y. Halevy, and Zachary G. Ives. 2012. *Principles of Data Integration*. Morgan Kaufmann. <http://research.cs.wisc.edu/dibook/>
- [13] Haoyu Dong, Zhoujun Cheng, Xinyi He, Mengyu Zhou, Anda Zhou, Fan Zhou, Ao Liu, Shi Han, and Dongmei Zhang. 2022. Table Pre-Training: A Survey on Model Architectures, Pre-Training Objectives, and Downstream Tasks. In *Proc. Thirty-First Int. Jt. Conf. Artif. Intell. International Joint Conferences on Artificial Intelligence Organization*, Vienna, Austria, 5426–5435. <https://doi.org/10.24963/ijcai.2022/761>
- [14] Yuyang Dong, Chuan Xiao, Takuma Nozawa, Masafumi Enomoto, and Masafumi Oyamada. 2022. DeepJoin: Joinable Table Discovery with Pre-trained Language Models. *CoRR abs/2212.07588* (2022). <https://doi.org/10.48550/arXiv.2212.07588> *arXiv:2212.07588*
- [15] Muhammad Ebraheem, Saravanan Thirumuruganathan, Shafiq Joty, Mourad Ouzzani, and Nan Tang. 2018. Distributed Representations of Tuples for Entity Resolution. *Proc. VLDB Endow.* 11, 11 (July 2018), 1454–1467. <https://doi.org/10.14778/3236187.3236198>
- [16] Julian Eisenschlos, Maharshi Gor, Thomas Müller, and William W. Cohen. 2021. MATE: Multi-view Attention for Table Transformer Efficiency. In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing, EMNLP 2021, Virtual Event / Punta Cana, Dominican Republic, 7–11 November, 2021*. Association for Computational Linguistics, 7606–7619. <https://doi.org/10.18653/v1/2021.emnlp-main.600>
- [17] Grace Fan, Jin Wang, Yuliang Li, Dan Zhang, and René J. Miller. 2022. Semantics-aware Dataset Discovery from Data Lakes with Contextualized Column-based Representation Learning. *CoRR abs/2210.01922* (2022). <https://doi.org/10.48550/arXiv.2210.01922> *arXiv:2210.01922*
- [18] Ivan P. Fellegi and Alan B. Sunter. 1969. A Theory for Record Linkage. *J. Amer. Statist. Assoc.* 64, 328 (Dec. 1969), 1183–1210. <https://doi.org/10.1080/01621459.1969.10501049>
- [19] Cheng Fu, Xianpei Han, Jiaming He, and Le Sun. 2020. Hierarchical Matching Network for Heterogeneous Entity Resolution. In *Proc. Twenty-Ninth Int. Jt. Conf. Artif. Intell. International Joint Conferences on Artificial Intelligence Organization*, Yokohama, Japan, 3665–3671. <https://doi.org/10.24963/ijcai.2020/507>
- [20] Sainyam Galhotra, Donatella Firmani, Barna Saha, and Divesh Srivastava. 2021. BEER: Blocking for Effective Entity Resolution. In *Proc. 2021 Int. Conf. Manag. Data*. ACM, Virtual Event China, 2711–2715. <https://doi.org/10.1145/3448016.3452747>
- [21] Tianyu Gao, Xingcheng Yao, and Danqi Chen. 2021. SimCSE: Simple Contrastive Learning of Sentence Embeddings. In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing, EMNLP 2021, Virtual Event / Punta Cana, Dominican Republic, 7–11 November, 2021*. Association for Computational Linguistics, 6894–6910. <https://doi.org/10.18653/v1/2021.emnlp-main.552>
- [22] Aristides Gionis, Piotr Indyk, and Rajeev Motwani. 1999. Similarity Search in High Dimensions via Hashing. In *VLDB '99, Proceedings of 25th International Conference on Very Large Data Bases, September 7–10, 1999, Edinburgh, Scotland, UK*. Morgan Kaufmann, 518–529. <http://www.vldb.org/conf/1999/P49.pdf>
- [23] Luis Gravano, Panagiotis G. Ipeirotis, H. V. Jagadish, Nick Koudas, S. Muthukrishnan, and Divesh Srivastava. 2001. Approximate String Joins in a Database (Almost) for Free. In *VLDB 2001, Proceedings of 27th International Conference on Very Large Data Bases, September 11–14, 2001, Roma, Italy*. Morgan Kaufmann, 491–500. <http://www.vldb.org/conf/2001/P491.pdf>
- [24] Michael Günther, Maik Thiele, Julius Gonsior, and Wolfgang Lehner. 2021. Pre-Trained Web Table Embeddings for Table Discovery. In *aIDM '21: Fourth Workshop in Exploiting AI Techniques for Data Management, Virtual Event, China, 25 June, 2021*. ACM, 24–31. <https://doi.org/10.1145/3464509.3464892>
- [25] Jonathan Herzig, Paweł Krzysztof Nowak, Thomas Müller, Francesco Piccinno, and Julian Martin Eisenschlos. 2020. TaPas: Weakly Supervised Table Parsing via Pre-training. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics, ACL 2020, Online, July 5–10, 2020*. Association for Computational Linguistics, 4320–4333. <https://doi.org/10.18653/v1/2020.acl>

- main.398
- [26] Madelon Hulsebos, Çağatay Demiralp, and Paul Groth. 2022. GitTables: A Large-Scale Corpus of Relational Tables. arXiv:2106.07258 (April 2022). arXiv:2106.07258 [cs] <http://arxiv.org/abs/2106.07258>
 - [27] Hiroshi Iida, Dung Thai, Varun Manjunatha, and Mohit Iyyer. 2021. TAB-BIE: Pretrained Representations of Tabular Data. In *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, NAACL-HLT 2021, Online, June 6-11, 2021*. Association for Computational Linguistics, 3446–3456. <https://doi.org/10.18653/v1/2021.naacl-main.270>
 - [28] Ihab F. Ilyas and Xu Chu. 2019. *Data Cleaning*. ACM, New York, NY, USA. <https://doi.org/10.1145/3310205>
 - [29] Arjit Jain, Sunita Sarawagi, and Prithviraj Sen. 2021. Deep Indexed Active Learning for Matching Heterogeneous Entity Representations. *Proc. VLDB Endow.* 15, 1 (Sept. 2021), 31–45. <https://doi.org/10.14778/3485450.3485455>
 - [30] Han Li, Pradap Konda, Paul Suganthan G C, Anhui Doan, Benjamin Snyder, Youngchoon Park, Ganesh Krishnan, Rohit Deep, and Vijay Raghavendra. 2018. MatchCatcher: A Debugger for Blocking in Entity Matching. *OpenProceedings.org*. <https://doi.org/10.5441/002/EDBT.2018.18>
 - [31] Yuliang Li, Jinfeng Li, Yoshihiko Suhara, Anhui Doan, and Wang-Chiew Tan. 2020. Deep Entity Matching with Pre-Trained Language Models. *Proc. VLDB Endow.* 14, 1 (Sept. 2020), 50–60. <https://doi.org/10.14778/3421424.3421431>
 - [32] Qian Liu, Bei Chen, Jiaqi Guo, Morteza Ziyadi, Zeqi Lin, Weizhu Chen, and Jian-Guang Lou. 2022. TAPEx: Table Pre-training via Learning a Neural SQL Executor. In *The Tenth International Conference on Learning Representations, ICLR 2022, Virtual Event, April 25-29, 2022*. OpenReview.net. <https://openreview.net/forum?id=O50443AsCP>
 - [33] Yury A. Malkov and Dmitry A. Yashunin. 2020. Efficient and Robust Approximate Nearest Neighbor Search Using Hierarchical Navigable Small World Graphs. *IEEE Trans. Pattern Anal. Mach. Intell.* 42, 4 (2020), 824–836. <https://doi.org/10.1109/TPAMI.2018.2889473>
 - [34] Zhengjie Miao, Yuliang Li, and Xiaolan Wang. 2021. Rotom: A Meta-Learned Data Augmentation Framework for Entity Matching, Data Cleaning, Text Classification, and Beyond. In *Proc. 2021 Int. Conf. Manag. Data*. ACM, Virtual Event China, 1303–1316. <https://doi.org/10.1145/3448016.3457258>
 - [35] Sidharth Mudgal, Han Li, Theodoros Rekatsinas, Anhui Doan, Youngchoon Park, Ganesh Krishnan, Rohit Deep, Esteban Arcaute, and Vijay Raghavendra. 2018. Deep Learning for Entity Matching: A Design Space Exploration. In *Proc. 2018 Int. Conf. Manag. Data*. ACM, Houston TX USA, 19–34. <https://doi.org/10.1145/3183713.3196926>
 - [36] John Bosco Mugeni and Toshiyuki Amagasa. 2022. A Graph-Based Blocking Approach for Entity Matching Using Pre-Trained Contextual Embedding Models. In *Proc. 37th ACM SIGAPP Symp. Appl. Comput.* ACM, Virtual Event, 357–364. <https://doi.org/10.1145/3477314.3507689>
 - [37] Felix Naumann and Melanie Herschel. 2010. *An Introduction to Duplicate Detection*. Morgan & Claypool Publishers. <https://doi.org/10.2200/S00262ED1V01Y201003DTM003>
 - [38] Hao Nie, Xianpei Han, Ben He, Le Sun, Bo Chen, Wei Zhang, Suhui Wu, and Hao Kong. 2019. Deep Sequence-to-Sequence Entity Matching for Heterogeneous Entity Resolution. In *Proc. 28th ACM Int. Conf. Inf. Knowl. Manag.* ACM, Beijing China, 629–638. <https://doi.org/10.1145/3357384.3358018>
 - [39] George Papadakis, George Alexiou, George Papastefanatos, and Georgia Koutrika. 2015. Schema-agnostic vs Schema-based Configurations for Blocking Methods on Homogeneous Data. *Proc. VLDB Endow.* 9, 4 (2015), 312–323. <https://doi.org/10.14778/2856318.2856326>
 - [40] George Papadakis, Marco Fisichella, Franziska Schoger, George Mandilaras, Nikolaus Augsten, and Wolfgang Nejdl. 2022. How to Reduce the Search Space of Entity Resolution: With Blocking or Nearest Neighbor Search? arXiv:2202.12521 (Oct. 2022). arXiv:2202.12521 [cs] <http://arxiv.org/abs/2202.12521>
 - [41] George Papadakis, Ekaterini Ioannou, Claudia Niederée, Themis Palpanas, and Wolfgang Nejdl. 2011. Eliminating the redundancy in blocking-based entity resolution methods. In *Proceedings of the 2011 Joint International Conference on Digital Libraries, JCDL 2011, Ottawa, ON, Canada, June 13-17, 2011*. ACM, 85–94. <https://doi.org/10.1145/1998076.1998093>
 - [42] George Papadakis, Ekaterini Ioannou, Themis Palpanas, Claudia Niederée, and Wolfgang Nejdl. 2013. A Blocking Framework for Entity Resolution in Highly Heterogeneous Information Spaces. *IEEE Trans. Knowl. Data Eng.* 25, 12 (2013), 2665–2682. <https://doi.org/10.1109/TKDE.2012.150>
 - [43] George Papadakis, Ekaterini Ioannou, Emmanouil Thanos, and Themis Palpanas. 2021. *The Four Generations of Entity Resolution*. Springer International Publishing, Cham. <https://doi.org/10.1007/978-3-031-01878-7>
 - [44] George Papadakis, Georgia Koutrika, Themis Palpanas, and Wolfgang Nejdl. 2014. Meta-Blocking: Taking Entity Resolution to the Next Level. *IEEE Trans. Knowl. Data Eng.* 26, 8 (2014), 1946–1960. <https://doi.org/10.1109/TKDE.2013.54>
 - [45] George Papadakis, George Mandilaras, Luca Gagliardi, Giovanni Simonini, Emmanouil Thanos, George Giannakopoulos, Sonia Bergamaschi, Themis Palpanas, and Manolis Koubarakis. 2020. Three-Dimensional Entity Resolution with JedAI. *Information Systems* 93 (Nov. 2020), 101565. <https://doi.org/10.1016/j.is.2020.101565>
 - [46] George Papadakis, Dimitrios Skoutas, Emmanouil Thanos, and Themis Palpanas. 2021. Blocking and Filtering Techniques for Entity Resolution: A Survey. *ACM Comput. Surv.* 53, 2 (March 2021), 1–42. <https://doi.org/10.1145/3377455>
 - [47] Derek Paulsen, Yash Govind, and Anhui Doan. 2023. Sparkly: A Simple yet Surprisingly Strong TF/IDF Blocker for Entity Matching. *Proc. VLDB Endow.* 16, 6 (April 2023), 1507–1519. <https://doi.org/10.14778/3583140.3583163>
 - [48] Colin Raffel, Noam Shazeer, Adam Roberts, Katherine Lee, Sharan Narang, Michael Matena, Yanqi Zhou, Wei Li, and Peter J. Liu. 2020. Exploring the Limits of Transfer Learning with a Unified Text-to-Text Transformer. *J. Mach. Learn. Res.* 21 (2020), 140:1–140:67. <http://jmlr.org/papers/v21/20-074.html>
 - [49] Nils Reimers and Iryna Gurevych. 2019. Sentence-BERT: Sentence Embeddings using Siamese BERT-Networks. In *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing, EMNLP-IJCNLP 2019, Hong Kong, China, November 3-7, 2019*. Association for Computational Linguistics, 3980–3990. <https://doi.org/10.18653/v1/D19-1410>
 - [50] Yifan Sun, Changmao Cheng, Yuhang Zhang, Chi Zhang, Liang Zheng, Zhongdao Wang, and Yichen Wei. 2020. Circle Loss: A Unified Perspective of Pair Similarity Optimization. In *2020 IEEE/CVF Conference on Computer Vision and Pattern Recognition, CVPR 2020, Seattle, WA, USA, June 13-19, 2020*. Computer Vision Foundation / IEEE, 6397–6406. <https://doi.org/10.1109/CVPR42600.2020.00643>
 - [51] Saravanan Thirumuruganathan, Han Li, Nan Tang, Mourad Ouzzani, Yash Govind, Derek Paulsen, Glenn Fung, and Anhui Doan. 2021. Deep Learning for Blocking in Entity Matching: A Design Space Exploration. *Proc. VLDB Endow.* 14, 11 (July 2021), 2459–2472. <https://doi.org/10.14778/3476249.3476294>
 - [52] Saravanan Thirumuruganathan, Shameem A. Puthiya Parambath, Mourad Ouzzani, Nan Tang, and Shafiq Joty. 2018. Reuse and Adaptation for Entity Resolution through Transfer Learning. arXiv:1809.11084 (Sept. 2018). arXiv:1809.11084 [cs, stat] <http://arxiv.org/abs/1809.11084>
 - [53] Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N. Gomez, Lukasz Kaiser, and Illia Polosukhin. 2017. Attention is All you Need. In *Advances in Neural Information Processing Systems 30: Annual Conference on Neural Information Processing Systems 2017, December 4-9, 2017, Long Beach, CA, USA*. 5998–6008. <https://proceedings.neurips.cc/paper/2017/hash/3f5ee243547dee91fbd053c1c4a845aa-Abstract.html>
 - [54] Jin Wang, Yuliang Li, and Wataru Hirota. 2021. Machamp: A Generalized Entity Matching Benchmark. In *Proc. 30th ACM Int. Conf. Inf. Knowl. Manag.* ACM, Virtual Event Queensland Australia, 4633–4642. <https://doi.org/10.1145/3459637.3482008>
 - [55] Runhui Wang, Yuliang Li, and Jin Wang. 2022. Sudowoodo: Contrastive Self-Supervised Learning for Multi-Purpose Data Integration and Preparation. arXiv:2207.04122 (Oct. 2022). arXiv:2207.04122 [cs] <http://arxiv.org/abs/2207.04122>
 - [56] Zhiruo Wang, Haoyu Dong, Ran Jia, Jia Li, Zhiyi Fu, Shi Han, and Dongmei Zhang. 2021. TUTA: Tree-based Transformers for Generally Structured Table Pre-training. In *KDD '21: The 27th ACM SIGKDD Conference on Knowledge Discovery and Data Mining, Virtual Event, Singapore, August 14-18, 2021*. ACM, 1780–1790. <https://doi.org/10.1145/3447548.3467434>
 - [57] Zifeng Wang and Jimeng Sun. 2022. TransTab: Learning Transferable Tabular Transformers Across Tables. *CoRR* abs/2205.09328 (2022). <https://doi.org/10.48550/arXiv.2205.09328>
 - [58] Thomas Wolf, Lysandre Debut, Victor Sanh, Julien Chaumond, Clement Delangue, Anthony Moi, Pierric Cistac, Tim Rault, Rémi Louf, Morgan Funtowicz, Joe Davison, Sam Shleifer, Patrick von Platen, Clara Ma, Yacine Jernite, Julien Plu, Canwen Xu, Teven Le Scao, Sylvain Gugger, Mariama Drame, Quentin Lhoest, and Alexander M. Rush. 2020. Transformers: State-of-the-Art Natural Language Processing. In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing: System Demonstrations*. Association for Computational Linguistics, Online, 38–45. <https://www.aclweb.org/anthology/2020.emnlp-demos.6>
 - [59] Tete Xiao, Xiaolong Wang, Alexei A. Efros, and Trevor Darrell. 2021. What Should Not Be Contrastive in Contrastive Learning. In *9th International Conference on Learning Representations, ICLR 2021, Virtual Event, Austria, May 3-7, 2021*. OpenReview.net. <https://openreview.net/forum?id=CZ8Y3NzuVzO>
 - [60] Xinyu Xing and Xiaojun Wan. 2021. Structure-Aware Pre-Training for Table-to-Text Generation. In *Findings of the Association for Computational Linguistics: ACL/IJCNLP 2021, Online Event, August 1-6, 2021 (Findings of ACL, Vol. ACL/IJCNLP 2021)*. Association for Computational Linguistics, 2273–2278. <https://doi.org/10.18653/v1/2021.findings-acl.200>
 - [61] Pengcheng Yin, Graham Neubig, Wen-tau Yih, and Sebastian Riedel. 2020. TabERT: Pretraining for Joint Understanding of Textual and Tabular Data. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics, ACL 2020, Online, July 5-10, 2020*. Association for Computational Linguistics, 8413–8426. <https://doi.org/10.18653/v1/2020.acl-main.745>
 - [62] Jinsung Yoon, Yao Zhang, James Jordon, and Mihaela van der Schaar. 2020. VIME: Extending the Success of Self- and Semi-supervised Learning to Tabular Domain. In *Advances in Neural Information Processing Systems 33: Annual*

- Conference on Neural Information Processing Systems 2020, NeurIPS 2020, December 6-12, 2020, virtual.* <https://proceedings.neurips.cc/paper/2020/hash/7d97667a3e056acab9aaf653807b4a03-Abstract.html>
- [63] Alexandros Zeakis, George Papadakis, Dimitrios Skoutas, and Manolis Koubarakis. 2023. Pre-Trained Embeddings for Entity Resolution: An Experimental Analysis. *Proc. VLDB Endow.* 16, 9 (July 2023), 2225–2238. <https://doi.org/10.14778/3598581.3598594>
- [64] Jiaqi Zhai, Yin Lou, and Johannes Gehrke. 2011. ATLAS: a probabilistic algorithm for high dimensional similarity search. In *Proceedings of the ACM SIGMOD International Conference on Management of Data, SIGMOD 2011, Athens, Greece, June 12-16, 2011*. ACM, 997–1008. <https://doi.org/10.1145/1989323.1989428>
- [65] Wei Zhang, Hao Wei, Bunyamin Sisman, Xin Luna Dong, Christos Faloutsos, and David Page. 2020. AutoBlock: A Hands-off Blocking Framework for Entity Matching. In *Proc. 13th Int. Conf. Web Search Data Min.* ACM, Houston TX USA, 744–752. <https://doi.org/10.1145/3336191.3371813>

Algorithm 1: Heuristic Rule for Column Type Detection.

Input: A table column *col* with its name *col.name* and list of values *col.values*.

Output: One of the following predefined column types: {identifier, numeric, url, date, category, entity, text}

```

1 if number of unique values in col.values is 100% of its length
  then
2   return identifier
3 if col.values are numeric or boolean then
4   return numeric
5 if col.values are url then
6   return url
7 if col.values are datetime then
8   return date
9 value_lens ← number of token of col.values
10 if variance of value_lens is 0 then
11   return category
12 else
13   if average of value_len less than 10 then
14     return entity
15   else
16     return text

```

A DETAILS OF PRE-TRAINING DATA CONDITIONING

This section further introduces the details of pre-training data conditioning:

Column Type Detection. The column type detection algorithm is shown in Algorithm 1, originally used for tabular corpus analysis. Specifically, we define seven column types, namely identifier, numeric, URL, date, category, entity, and text. To begin, we determine whether a given column is an identifier by checking whether the column value selectivity is 100%. Following this, we assess whether the column values are numeric (either boolean or float). For columns representing URL or date information, we utilize regular expressions and the Python package `dateutil`¹ to match patterns. We segregate the remaining columns into either category or entity and text types, depending on whether the length variance of column values is equal to zero. We differentiate entity from text by examining if the number of tokens is less than 10.

Table filtering. After the column type detection, we apply additional filtering to exclude tables that do not contain suitable entity records. We randomly select 100 tables for manual labeling and train a decision tree using features like the number of each column type and the type of the first column. To make the filtering algorithm more robust, we refine the decision tree manually to create the final filtering algorithm. Specifically, we filter out tables containing statistics where more than half of the columns are comprised of numeric, URL, or date information. We also exclude tables where

Algorithm 2: Nearest Neighbor Blocker

Input: list of entity record collections *collections*, record vectorizer *vectorizer*, nearest neighbor indexer *indexer*, and nearest neighbor number *k*

Output: list of candidate pair sets arranged in nearest neighbor ranking *candidates*

```

1 if number of collections > 2 then
2   merge collections into one
3 source ← vectorize(collections[0], vectorizer)
4 target ← vectorize(collections[-1], vectorizer)
5 index ← buildIndex(target, indexer)
6 indices ← [] // empty list
7 for batch_queries in chunk(source) do
8   batch_indices ← batchSearch(batch_queries, index, k)
9   extend batch_indices to the end of indices
  // convert indices to list of candidate pair sets
  // arranged in nearest neighbor ranking candidates
10 candidates ← []
11 for i in range(k) do
12   candidate_pairs ← {} // empty set
13   for j in range(len(source)) do
14     candidate_pair ← (j, indices[j][i])
15     if candidate_pair did not appear before then
16       insert candidate_pair into candidate_pairs
17   append candidate_pairs to the end of candidates
18 return candidates

```

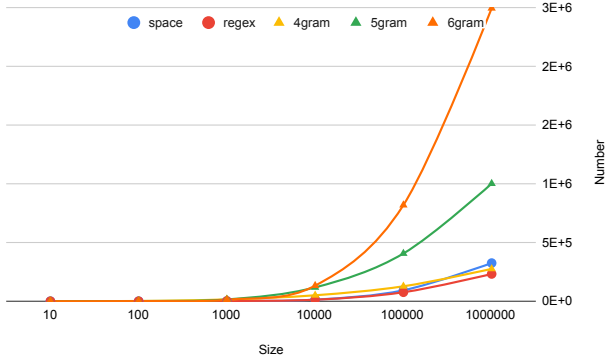
the first column is a date, which is typically associated with log data. Finally, we filter tables that are not written in English.

B BENCHMARK DISCUSSION

Our proposed benchmark value can be further validated by experimental results. On the one hand, Table 4 shows that existing blocking methods can effectively address the homogeneous and simple datasets that previous studies mainly focused on in Table 2. There is still room for improvement on some of our newly introduced datasets, providing new challenges and opportunities for future research. On the other hand, our proposed benchmark is the first to evaluate the general performance of different methods in different scenarios using the same parameters on all datasets, laying the foundation for future automatic, data-driven sparse blocking tuning approaches or good attribute identification for blocking. Finally, by comparing the performance differences of UNIBLOCKER on artificially and naturally heterogeneous datasets with their corresponding homogeneous datasets, we can see that the mainstream approach of artificially constructing dirty data does not pose much of a challenge to deep models. How to better perform entity resolution directly on real highly heterogeneous data (such as “walmart-amazon_heter”) is also a future research direction.

¹<https://github.com/dateutil/dateutil>

Figure 5: Average number of tokens using different tokenizers with increasing entity record size on two scalability datasets.

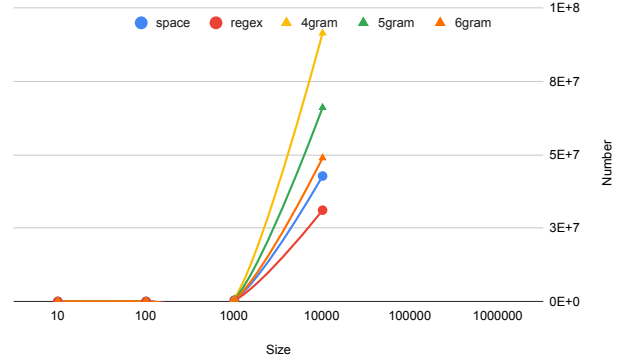


C NEAREST NEIGHBOR BLOCKER

This section presents how to use UNIBLOCKER for blocking any number of entity record collections. It is worth mentioning that our general implementation of the nearest neighbor blocker can apply different sparse or dense vectorizers and indexers for a thorough and fair comparison of different methods. As shown in Algorithm 2, the nearest neighbor blocker operates on a list of entity collections, using a record vectorizer and a nearest neighbor indexer to identify k nearest neighbors for each record. In cases with more than two collections, they are unified into a single one. The first and last collections are vectorized as the *source* and *target* respectively, with the latter being indexed. The algorithm iteratively searches for k nearest neighbors of source vectors and stores these indices. Subsequently, candidate pairs are formulated, each constituted of a source index and its corresponding nearest neighbor index from the target. Each candidate pair is prioritized and deduplicated according to its nearest neighbor ranking, to facilitate progressive entity matching or evaluation. This sequence of candidate pair sets forms the final output of the algorithm.

Besides offering a more precise assessment for blocking, our suggested benchmark value can be further corroborated by experimental results. On one hand, Table 4 demonstrates that current sparse and dense blocking methods effectively tackle homogeneous and simple datasets primarily studied in Table 2. However, there is scope for improvement on some newly introduced datasets, which present new challenges and opportunities for future research. On the other hand, our suggested benchmark is the first to evaluate general performance across various scenarios using the same parameters for all datasets, setting the groundwork for future data-driven sparse blocking tuning approaches [40] or more versatile dense blocking model evaluations. Lastly, by comparing UNIBLOCKER’s performance differences on artificially and naturally heterogeneous datasets with their homogeneous counterparts, we observe that artificial dirty data construction does not significantly challenge deep models. A future research direction includes improving entity resolution on highly heterogeneous real data (such as "walmart-amazon_heter").

Figure 6: Average number of overlapping records (i.e. records sharing at least one common token) using different tokenizers with increasing entity record size on two scalability datasets.



D IMPACT OF TOKENIZERS ON TOKENS AND OVERLAPPING RECORDS.

Although q-gram is an effective fuzzing matching technique, it sacrifices blocking scalability. In this section, we present a concise overview of the impact of varying tokenizers on the number of tokens and overlapping records (i.e. records sharing at least one common token). As shown in Figure 5, the q-gram tokenizers produce far more tokens than the regular tokenizers do. The ideal upper bound on the number of tokens produced by the q-gram tokenizer is the m^q , where m is the number of characters. With the increase of q , the growth rate of token number will accelerate rapidly. Exponentially growing tokens pose a serious challenge to the size and speed of building index, which is what the filter-verification framework needs. Furthermore, though a smaller q produces fewer tokens, as shown in Figure 6, this produces more overlapping records. Sub-quadratic record pairs with shared tokens may also pose a challenge for finding similar records. As a result, dense blocking with fixed dimension vectors not only models semantic information but is also more scalable than q-gram sparse blocking methods.

E IMPACT OF ANN SEARCH ON EFFECTIVENESS

Table 6 shows the overall performances of sparse and dense blocking methods with the same approximate nearest neighbor algorithm (ANN) and configuration. In this simple setting, we can find that the average mAP of the sparse blocking method decreases by 3.09 compared to the precision search, but the dense blocking method only decreases by 1. Therefore, sparse join may receive more influence from ANN search than dense join.

Table 6: Overall performance of sparse, dense kNN join with approximate nearest neighbor search.

Datasets	Sparse kNN					Dense kNN				
	AP	PC	PQ	F1	k	AP	PC	PQ	F1	k
census	69.89	90.34	15.28	26.14	6	73.17	90.88	23.06	36.78	4
cora	66.57	90.25	9.16	16.63	10	73.21	90.06	22.85	36.45	4
notebook	39.08	90.15	14.96	25.67	43	36.59	90.13	11.19	19.90	59
notebook2	24.45	82.92	2.48	4.81	100	26.22	83.36	2.55	4.95	100
altosight	49.12	90.08	1.95	3.82	44	48.03	90.15	3.74	7.18	23
altosight2	11.89	93.60	9.24	16.82	7	20.42	95.06	15.47	26.61	5
fodors-zagats_homo	47.37	71.46	34.73	46.74	100	48.01	72.04	34.44	46.60	100
dblp-acm	89.04	96.63	82.15	88.80	1	90.94	98.29	83.56	90.33	1
dblp-scholar	67.23	90.03	26.29	40.69	7	68.44	90.99	23.25	37.03	8
abt-buy_homo	55.74	92.86	19.51	32.25	1	58.79	97.32	20.45	33.80	1
walmart-amazon_homo	57.98	96.43	20.26	33.49	1	59.95	99.11	20.83	34.42	1
fodors-zagats_heter	12.83	30.69	0.28	0.56	100	6.86	26.28	0.22	0.43	100
imdb-dbpeda	3.11	75.28	0.20	0.40	100	4.80	85.37	0.25	0.50	100
amazon-google	31.65	90.06	11.15	19.84	70	34.85	90.06	8.50	15.53	98
abt-buy_heter	40.44	90.12	2.61	5.07	72	38.28	90.16	3.10	5.99	65
walmart-amazon_heter	29.76	85.18	0.38	0.77	100	25.99	83.19	0.38	0.75	100
movies	48.76	90.47	4.09	7.82	10	47.81	90.12	6.79	12.63	6
Average	43.82	85.09	14.98	21.78	45.41	44.84	86.03	16.51	24.11	45.59