

Deep Entity Matching with Pre-Trained Language Models

Yuliang Li, Jinfeng Li,
Yoshihiko Suhara

Megagon Labs
{yuliang,jinfeng,yoshi}@megagon.ai

AnHai Doan
University of Wisconsin Madison
anhai@cs.wisc.edu

Wang-Chiew Tan
Megagon Labs
wangchiew@megagon.ai

ABSTRACT

We present DITTO, a novel entity matching system based on pre-trained Transformer-based language models. We fine-tune and cast EM as a sequence-pair classification problem to leverage such models with a simple architecture. Our experiments show that a straightforward application of language models such as BERT, DistilBERT, or RoBERTa pre-trained on large text corpora already significantly improves the matching quality and outperforms previous state-of-the-art (SOTA), by up to 29% of F1 score on benchmark datasets. We also developed three optimization techniques to further improve DITTO’s matching capability. DITTO allows domain knowledge to be injected by highlighting important pieces of input information that may be of interest when making matching decisions. DITTO also summarizes strings that are too long so that only the essential information is retained and used for EM. Finally, DITTO adapts a SOTA technique on data augmentation for text to EM to augment the training data with (difficult) examples. This way, DITTO is forced to learn “harder” to improve the model’s matching capability. The optimizations we developed further boost the performance of DITTO by up to 9.8%. Perhaps more surprisingly, we establish that DITTO can achieve the previous SOTA results with at most half the number of labeled data. Finally, we demonstrate DITTO’s effectiveness on a real-world large-scale EM task. On matching two company datasets consisting of 789K and 412K records, DITTO achieves a high F1 score of 96.5%.

PVLDB Reference Format:

Yuliang Li, Jinfeng Li, Yoshihiko Suhara, AnHai Doan, and Wang-Chiew Tan. Deep Entity Matching with Pre-Trained Language Models. PVLDB, 14(1): 50 - 60, 2021.
doi:10.14778/3421424.3421431

PVLDB Artifact Availability:

The source code, data, and/or other artifacts have been made available at <https://github.com/megagonlabs/ditto>.

1 INTRODUCTION

Entity Matching (EM) refers to the problem of determining whether two data entries refer to the same real-world entity. Consider the two datasets about products in Figure 1. The goal is to determine the set of pairs of data entries, one entry from each table so that each pair of entries refer to the same product.

If the datasets are large, it can be expensive to determine the pairs of matching entries. For this reason, EM is typically accompanied by a pre-processing step, called *blocking*, to prune pairs of entries that are unlikely matches to reduce the number of candidate pairs to consider. As we will illustrate, correctly *matching* the candidate pairs requires substantial language understanding and domain-specific knowledge. Hence, entity matching remains a challenging task even for the most advanced EM solutions.

We present DITTO, a novel EM solution based on pre-trained Transformer-based language models (or *pre-trained language models* in short). We cast EM as a sequence-pair classification problem to leverage such models, which have been shown to generate highly contextualized embeddings that capture better language understanding compared to traditional word embeddings. DITTO further improves its matching capability through three optimizations: (1) It allows domain knowledge to be added by highlighting important pieces of the input that may be useful for matching decisions. (2) It summarizes long strings so that only the most essential information is retained and used for EM. (3) It augments training data with (difficult) examples, which challenges DITTO to learn “harder” and also reduces the amount of training data required. Figure 2 depicts DITTO in the overall architecture of a complete EM workflow.

There are 9 candidate pairs of entries to consider for matching in total in Figure 1. The blocking heuristic that matching entries must have one word in common in the **title** will reduce the number of pairs to only 3: the first entry on the left with the first entry on the right and so on. Perhaps more surprisingly, even though the 3 pairs are highly similar and look like matches, only the first and last pair of entries are true matches. Our system, DITTO, is able to discern the nuances in the 3 pairs to make the correct conclusion for every pair while some state-of-the-art systems are unable to do so.

The example illustrates the power of language understanding given by DITTO’s pre-trained language model. It understands that *instant immersion spanish deluxe 2.0* is the same as *instant immers spanish dlux 2* in the context of software products even though they are syntactically different. Furthermore, one can explicitly emphasize that certain parts of a value are more useful for deciding matching decisions. For books, the domain knowledge that the grade level or edition is important for matching books can be made explicit to DITTO, simply by placing tags around the grade/edition values. Hence, for the second candidate pair, even though the titles are highly similar (i.e., they overlap in many words), DITTO is able to focus on the grade/edition information when making the matching decision. The third candidate pair shows the power of language understanding for the opposite situation. Even though the entries look dissimilar DITTO is able to attend to the right parts of a value (i.e., the *manf./modelno* under different attributes) and also understand the semantics of the model number to make the right decision.

This work is licensed under the Creative Commons BY-NC-ND 4.0 International License. Visit <https://creativecommons.org/licenses/by-nc-nd/4.0/> to view a copy of this license. For any use beyond those covered by this license, obtain permission by emailing info@vldb.org. Copyright is held by the owner/author(s). Publication rights licensed to the VLDB Endowment.
Proceedings of the VLDB Endowment, Vol. 14, No. 1 ISSN 2150-8097.
doi:10.14778/3421424.3421431

title	manf./modelno	price		title	manf./modelno	price
instant immersion spanish deluxe 2.0	topics entertainment	49.99	✓	instant immers spanish dlux 2	NULL	36.11
adventure workshop 4th-6th grade 7th edition	encore software	19.99	✗	encore inc adventure workshop 4th-6th grade 8th edition	NULL	17.1
sharp printing calculator	sharp el1192bl	37.63	✓	new-sharp shr-el1192bl two-color printing calculator 12-digit lcd black red	NULL	56.0

Figure 1: Entity Matching: determine the matching entries from two datasets.

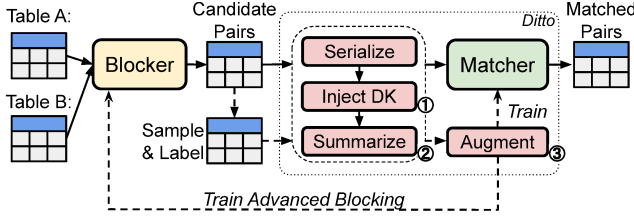


Figure 2: An EM system architecture with DITTO as the matcher. In addition to the training data, the user of DITTO can specify (1) a method for injecting domain knowledge (DK), (2) a summarization module for keeping the essential information, and (3) a data augmentation (DA) operator to strengthen the training set.

Contributions In summary, the following are our contributions:

- We present DITTO, a novel EM solution based on pre-trained language models (LMs) such as BERT. We fine-tune and cast EM as a sequence-pair classification problem to leverage such models with a simple architecture. To our knowledge, DITTO is one of the first EM solutions that leverage pre-trained Transformer-based LMs¹ to provide deeper language understanding for EM.
- We also developed three optimization techniques to further improve DITTO’s matching capability through injecting domain knowledge, summarizing long strings, and augmenting training data with (difficult) examples. The first two techniques help DITTO focus on the right information for making matching decisions. The last technique, data augmentation, is adapted from [30] for EM to help DITTO learn “harder” to understand the data invariance properties that may exist but are beyond the provided labeled examples and also, reduce the amount of training data required.
- We evaluated the effectiveness of DITTO on three benchmark datasets: the Entity Resolution benchmark [25], the Magellan dataset [24], and the WDC product matching dataset [37] of various sizes and domains. Our experimental results show that DITTO consistently outperforms the previous SOTA EM solutions in all datasets and by up to 31% in F1 scores. Furthermore, DITTO consistently performs better on dirty data and is more label efficient: it achieves the same or higher previous SOTA accuracy using less than half the labeled data.
- We applied DITTO to a real-world large-scale matching task on two company datasets, containing 789K and 412K entries respectively. To deploy an end-to-to EM pipeline efficiently, we developed an advanced blocking technique to help reduce the number of pairs to consider for DITTO. DITTO obtains high accuracy, 96.5% F1 on a holdout dataset. The blocking phase also helped speed up the end-to-end EM deployment significantly, by up to 3.8 times, compared to naive blocking techniques.

¹There is a concurrent work [6] which applies a similar idea.

- Finally, we open-source DITTO at <https://github.com/megagonlabs/ditto>.

Outline Section 2 overviews DITTO and pre-trained LMs. Section 3 describes how we optimize DITTO with domain knowledge, summarization, and data augmentation. Our experimental results are described in Section 4 and the case study is presented in Section 5. We discuss related work in Section 6 and conclude in Section 7. A full version of the paper is available at [27].

2 BACKGROUND AND ARCHITECTURE

We present the main concepts behind EM and provide some background on pre-trained LMs before we describe how we fine-tune the LMs on EM datasets to train EM models. We also present a simple method for reducing EM to a sequence-pair classification problem so that pre-trained LMs can be used for solving the EM problem.

Notations DITTO’s EM pipeline takes as input two collections D and D' of data entries (e.g., rows of relational tables, XML documents, JSON files, text paragraphs) and outputs a set $M \subseteq D \times D'$ of pairs where each pair $(e, e') \in M$ is thought to represent the same real-world entity (e.g., person, company, laptop, etc.). A data entry e is a set of key-value pairs $e = \{(\text{attr}_i, \text{val}_i)\}_{1 \leq i \leq k}$ where attr_i is the attribute name and val_i is the attribute’s value represented as text. Note that our definition of data entries is general enough to capture both structured and semi-structured data such as JSON files.

As described earlier, an end-to-end EM system consists of a *blocker* and a *matcher*. The goal of the blocking phase is to quickly identify a small subset of $D \times D'$ of candidate pairs of high recall (i.e., a high proportion of actual matching pairs are that subset). The goal of a matcher (i.e., DITTO) is to accurately predict, given a pair of entries, whether they refer to the same real-world entity.

2.1 Pre-trained language models

Unlike prior learning-based EM solutions that rely on word embeddings and customized RNN architectures to train the matching model (See Section 6 for a detailed summary), DITTO trains the matching models by fine-tuning pre-trained LMs in a simpler architecture.

Pre-trained LMs such as BERT [13] and GPT-2 [39] have demonstrated good performance on a wide range of NLP tasks. They are typically deep neural networks with multiple Transformer layers [49], typically 12 or 24 layers, pre-trained on large text corpora such as Wikipedia articles in an unsupervised manner. During pre-training, the model is self-trained to perform auxiliary tasks such as missing token and next-sentence prediction. Studies [9, 48] have shown that the shallow layers capture lexical meaning while the deeper layers capture syntactic and semantic meanings of the input sequence after pre-training.

A specific strength of pre-trained LMs is that it learns the semantics of words better than conventional word embedding techniques such as word2vec, GloVe, or FastText. This is largely because the Transformer architecture calculates token embeddings from all the tokens in the input sequence and thus, the embeddings it generates are *highly-contextualized* and captures the semantic and contextual understanding of the words. Consequently, such embeddings can capture polysemy, i.e., discern that the same word may have different meanings in different phrases. For example, the word *Sharp* has different meanings in “*Sharp resolution*” versus “*Sharp TV*”. Pre-trained LMs will embed “*Sharp*” differently depending on the context while traditional word embedding techniques such as FastText always produce the same vector independent of the context. Such models can also understand the opposite, i.e., that different words may have the same meaning. For example, the words *immersion* and *immers* (respectively, (*deluxe*, *dlux*) and (2.0, 2)) are likely the same given their respective contexts. Thus, such language understanding capability of pre-trained LMs can improve the EM performance.

2.2 Fine-tuning pre-trained language models

A pre-trained LM can be fine-tuned with task-specific training data so that it becomes better at performing that task. Here, we fine-tune a pre-trained LM for the EM task with a labeled training dataset consisting of positive and negative pairs of matching and non-matching entries as follows:

- (1) Add task-specific layers after the final layer of the LM. For EM, we add a simple fully connected layer and a softmax output layer for binary classification.
- (2) Initialize the modified network with parameters from the pre-trained LM.
- (3) Train the modified network on the training set until it converges.

The result is a model fine-tuned for the EM task. See Appendix A of the full version [27] for the model architecture. In DITTO, we fine-tune the popular 12-layer BERT model [13], RoBERTa [28], and a 6-layer smaller but faster variant DistilBERT [43]. However, our proposed techniques are independent of the choice of pre-trained LMs and DITTO can potentially perform even better with larger pre-trained LMs. The pair of data entries is serialized (see next section) as input to the LM and the output is a match or no-match decision. DITTO’s architecture is much simpler when compared to many state-of-the-art EM solutions today [14, 32]. Even though the bulk of the “work” is simply off-loaded to pre-trained LMs, we show that this simple scheme works surprisingly well in our experiments.

2.3 Serializing the data entries for Ditto

Since LMs take token sequences (i.e., text) as input, a key challenge is to convert the candidate pairs into token sequences so that they can be meaningfully ingested by DITTO.

DITTO serializes data entries as follows: for each data entry $e = \{(\text{attr}_i, \text{val}_i)\}_{1 \leq i \leq k}$, we let

$\text{serialize}(e) ::= [\text{COL}] \text{attr}_1 [\text{VAL}] \text{val}_1 \dots [\text{COL}] \text{attr}_k [\text{VAL}] \text{val}_k$,

where [COL] and [VAL] are special tokens for indicating the start of attribute names and values respectively. For example, the first entry of the second table is serialized as:

[COL] title [VAL] instant immers spanish dlux 2 [COL] manf./modelno [VAL] NULL [COL] price [VAL] 36.11

To serialize a candidate pair (e, e') , we let

$\text{serialize}(e, e') ::= [\text{CLS}] \text{serialize}(e) [\text{SEP}] \text{serialize}(e') [\text{SEP}]$,

where [SEP] is the special token separating the two sequences and [CLS] is the special token necessary for BERT to encode the sequence pair into a 768-dimensional vector which will be fed into the fully connected layers for classification.

Other serialization schemes There are different ways to serialize data entries so that LMs can treat the input as a sequence classification problem. For example, one can also omit the special tokens “[COL]” and/or “[VAL]”, or exclude attribute names attr_i during serialization. We found that including the special tokens to retain the structure of the input does not hurt the performance in general and excluding the attribute names tend to help only when the attribute names do not contain useful information (e.g., names such as $\text{attr}_1, \text{attr}_2, \dots$) or when the entries contain only one column. A more rigorous study on this matter is left for future work.

Heterogeneous schemas As shown, the serialization method of DITTO does not require data entries to adhere to the same schema. It also does not require that the attributes of data entries to be matched prior to executing the matcher, which is a sharp contrast to other EM systems such as DeepER [14] or DeepMatcher² [32]. Furthermore, DITTO can also ingest and match hierarchically structured data entries by serializing nested attribute-value pairs with special start and end tokens (much like Lisp or XML-style parentheses structure).

3 OPTIMIZATIONS IN DITTO

As we will describe in Section 4, the basic version of DITTO, which leverages only the pre-trained LM, is already outperforming the SOTA on average. Here, we describe three further optimization techniques that will facilitate and challenge DITTO to learn “harder”, and consequently make better matching decisions.

3.1 Leveraging Domain Knowledge

Our first optimization allows domain knowledge to be injected into DITTO through *pre-processing* the input sequences (i.e., serialized data entries) to emphasize what pieces of information are potentially important. This follows the intuition that when human workers make a matching/non-matching decision on two data entries, they typically look for spans of text that contain key information before making the final decision. Even though we can also train deep learning EM solutions to learn such knowledge, we will require a significant amount of training data to do so. As we will describe, this pre-processing step on the input sequences is lightweight and yet can yield significant improvements. Our experiment results show that with less than 5% of additional training time, we can improve the model’s performance by up to 8%.

There are two main types of domain knowledge that we can provide DITTO.

²In DeepMatcher, the requirement that both entries have the same schema can be removed by treating the values in all columns as one value under one attribute.

Span Typing The type of a span of tokens is one kind of domain knowledge that can be provided to DITTO. Product id, street number, publisher are examples of span types. Span types help DITTO avoid mismatches. With span types, for example, DITTO is likelier to avoid matching a street number with a year or a product id.

Table 1 summarizes the main span types that human workers would focus on when matching three types of entities in our benchmark datasets.

Table 1: Main span types for matching entities in our benchmark datasets.

Entity Type	Types of Important Spans
Publications, Movies, Music	Persons (e.g., Authors), Year, Publisher
Organizations, Employers	Last 4-digit of phone, Street number
Products	Product ID, Brand, Configurations (num.)

The developer specifies a recognizer to type spans of tokens from attribute values. The recognizer takes a text string v as input and returns a list $\text{recognizer}(v) = \{(s_i, t_i, \text{type}_i)\}_{i \geq 1}$ of start/end positions of the span in v and the corresponding type of the span. DITTO’s current implementation leverages an open-source Named-Entity Recognition (NER) model [46] to identify known types such as persons, dates, or organizations and use regular expressions to identify specific types such as product IDs, last 4 digits of phone numbers, etc.

After the types are recognized, the original text v is replaced by a new text where special tokens are inserted to reflect the types of the spans. For example, a phone number “(866) 246-6453” may be replaced with “(866) 246 - [LAST] 6453 [/LAST]” where [LAST]/[/LAST] indicates the start/end of the last 4 digits and additional spaces are also added because of tokenization. In our implementation, when we are sure that the span type has only one token or the NER model is inaccurate in determining the end position, we drop the end indicator and keep only the start indicator token.

Intuitively, these newly added special tokens are additional signals to the self-attention mechanism that already exists in pre-trained LMs, such as BERT. If two spans have the same type, then DITTO picks up the signal that they are likelier to be the same and hence, they are aligned together for matching. In the above example,

“..246- [LAST] 6453 [/LAST] .. [SEP] .. [LAST] 0000 [/LAST]..”

when the model sees two encoded sequences with the [LAST] special tokens, it is likely to take the hint to align “6453” with “0000” without relying on other patterns elsewhere in the sequence that may be harder to learn.

Span Normalization The second kind of domain knowledge that can be passed to DITTO rewrites syntactically different but equivalent spans into the same string. This way, they will have identical embeddings and it becomes easier for DITTO to detect that the two spans are identical. For example, we can enforce that “VLDB journal” and “VLDBJ” are the same by writing them as VLDBJ. Similarly, we can enforce the general knowledge that “5 %” vs. “5.00 %” are equal by writing them as “5.0%”.

The developer specifies a set of rewriting rules to rewrite spans. The specification consists of a function that first identifies the spans of interest before it replaces them with the rewritten spans. DITTO contains a number of rewriting rules for numbers, including

rules that round all floating point numbers to 2 decimal places and dropping all commas from integers (e.g., “2,020” $\rightarrow$ “2020”). For abbreviations, we allow the developers to specify a dictionary of synonym pairs to normalize all synonym spans to be the same.

3.2 Summarizing long entries

When the value is an extremely long string, it becomes harder for the LM to understand what to pay attention to when matching. In addition, one limiting factor of Transformer-based pre-trained LMs is that there is a limit on the sequence length of the input. For example, the input to BERT can have at most 512 sub-word tokens. It is thus important to summarize the serialized entries down to the maximum allowed length while retaining the key information. A common practice is to truncate the sequences so that they fit within the maximum length. However, the truncation strategy does not work well for EM in general because the important information for matching is usually not at the beginning of the sequences.

There are many ways to perform summarization [31, 40, 42]. In DITTO’s current implementation, we use a TF-IDF-based summarization technique that *retains non-stopword tokens with the high TF-IDF scores*. We ignore the start and end tags generated by span typing in this process and use the list of stop words from scikit-learn library [35]. By doing so, DITTO feeds only the most informative tokens to the LM. We found that this technique works well in practice. Our experiment results show that it improves the F1 score of DITTO on a text-heavy dataset from 41% to over 93% and we plan to add more summarization techniques to DITTO’s library in the future.

3.3 Augmenting training data

We describe how we apply data augmentation to augment the training data for entity matching.

Data augmentation (DA) is a commonly used technique in computer vision for generating additional training data from existing examples by simple transformation such as cropping, flipping, rotation, padding, etc. The DA operators not only add more training data, but the augmented data also allows to model to learn to make predictions invariant of these transformations.

Similarly, DA can add training data that will help EM models learn “harder”. Although labeled examples for EM are arguably not hard to obtain, invariance properties are very important to help make the solution more robust to dirty data, such as missing values (NULLs), values that are placed under the wrong attributes or missing some tokens.

Next, we introduce a set of DA operators for EM that will help train more robust models.

Augmentation operators for EM The proposed DA operators are summarized in Table 2. If s is a serialized pair of data entries with a match or no-match label l , then an augmented example is a pair (s', l) , where s' is obtained by applying an operator o on s and s' has the same label l as before.

The operators are divided into 3 categories. The first category consists of span-level operators, such as `span_del` and `span_shuffle`. These two operators are used in NLP tasks [30, 55] and shown to be effective for text classification. For `span_del`, we randomly delete from s a span of tokens of length at most 4 without special tokens

Table 2: Data augmentation operators in DITTO. The operators are 3 different levels: span-level, attribute-level, and entry-level. All samplings are done uniformly at random.

Operator	Explanation
span_del	Delete a randomly sampled span of tokens
span_shuffle	Randomly sample a span and shuffle the tokens' order
attr_del	Delete a randomly chosen attribute and its value
attr_shuffle	Randomly shuffle the orders of all attributes
entry_swap	Swap the order of the two data entries e and e'

(e.g., [SEP], [COL], [VAL]). For span_shuffle, we sample a span of length at most 4 and randomly shuffle the order of its tokens.

These two operators are motivated by the observation that making a match/no-match decision can sometimes be “too easy” when the candidate pair of data entries contain multiple spans of text supporting the decision. For example, suppose our negative examples for matching company data in the existing training data is similar to what is shown below.

[CLS] . . . [VAL] Google LLC . . . [VAL] (866) 246-6453 [SEP] . . .
[VAL] Alphabet inc . . . [VAL] (650) 253-0000 [SEP]

The model may learn to predict “no-match” based on the phone number alone, which is insufficient in general. On the other hand, by corrupting parts of the input sequence (e.g., dropping phone numbers), DA forces the model to learn beyond that, by leveraging the remaining signals, such as the company name, to predict “no-match”.

The second category of operators is attribute-level operators: attr_del and attr_shuffle. The operator attr_del randomly deletes an attribute (both name and value) and attr_shuffle randomly shuffles the order of the attributes of both data entries. The motivation for attr_del is similar to span_del and span_shuffle but it gets rid of an attribute entirely. The attr_shuffle operator allows the model to learn the property that the matching decision should be independent of the ordering of attributes in the sequence.

The last operator, entry_swap, swaps the order of the pair (e, e') with probability 1/2. This teaches the model to make symmetric decisions (i.e., $F(e, e') = F(e', e)$) and helps double the size of the training set if both input tables are from the same data source.

MixDA: interpolating the augmented data Unlike DA operators for images which almost always preserve the image labels, the operators for EM can distort the input sequence so much that the label becomes incorrect. For example, the attr_del operator may drop the company name entirely and the remaining attributes may contain no useful signals to distinguish the two entries.

To address this issue, DITTO applies MixDA, a recently proposed data augmentation technique for NLP tasks [30] illustrated in Figure 3. Instead of using the augmented example directly, MixDA computes a convex interpolation of the original example with the augmented examples. Hence, the interpolated example is somewhere in between, i.e., it is a “partial” augmentation of the original example and this interpolated example is expected to be less distorted than the augmented one.

The idea of interpolating two examples is originally proposed for computer vision tasks [61]. For EM or text data, since we cannot directly interpolate sequences, MixDA interpolates their representations by the language model instead. In practice, augmentation

with MixDA slows the training time because the LM is called twice. However, the prediction time is not affected since the DA operators are only applied to training data. Formally, given an operator o (e.g., span deletion) and an original example s , to apply o on s with MixDA (as Figure 3 illustrates),

- (1) Randomly sample λ from a Beta distribution $\lambda \sim \text{Beta}(\alpha, \alpha)$ with a hyper-parameter $\alpha \in [0, 1]$ (e.g., 0.8 in our experiments);
- (2) Denote by $\text{LM}(s)$ the LM representation of a sequence s . Let

$$\text{LM}(s'') = \lambda \cdot \text{LM}(s) + (1 - \lambda) \cdot \text{LM}(\text{augment}(s, o)).$$

Namely, $\text{LM}(s'')$ is the convex interpolation between the LM outputs of s and the augmented $s' = \text{augment}(s, o)$;

- (3) Train the model by feeding $\text{LM}(s'')$ to the rest of the network and back-propagate. Back-propagation updates both the LM and linear layer’s parameters.

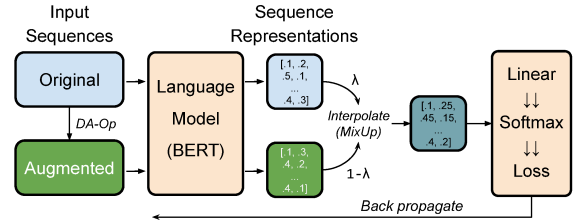


Figure 3: Data augmentation with MixDA.

4 EXPERIMENTS

We present the experiment results on benchmark datasets for EM: the ER Benchmark datasets [25], the Magellan datasets [24] and the WDC product data corpus [37]. DITTO achieves new SOTA results on all these datasets and outperforms the previous best results by up to 31% in F1 score. The results show that DITTO is more robust to dirty data and performs well when the training set is small. DITTO is also more label-efficient as it achieves the previous SOTA results using only 1/2 or less of the training data across multiple subsets of the WDC corpus. Our ablation analysis shows that (1) using pre-trained LMs contributes to over 50% of DITTO’s performance gain and (2) all 3 optimizations, domain knowledge (DK), summarization (SU) and data augmentation (DA), are effective. For example, SU improves the performance on a text-heavy dataset by 52%, DK leads to 1.2% average improvement on the ER-Magellan datasets and DA improves on the WDC datasets by 2.53% on average. In addition, we show in Appendix B in the full version [27] that although DITTO leverages deeper neural nets, its training and prediction time is comparable to the SOTA EM systems.

4.1 Benchmark datasets

We experimented with all the 13 publicly available datasets used for evaluating DeepMatcher [32]. These datasets are from the ER Benchmark datasets [25] and the Magellan data repository [12]. We summarize the datasets in Table 3 and refer to them as ER-Magellan. These datasets are for training and evaluating matching models for various domains including products, publications, and businesses. Each dataset consists of candidate pairs from two structured tables of entity records of the same schema. The pairs are sampled from the results of blocking and manually labeled. The positive rate (i.e., the ratio of matched pairs) ranges from 9.4% (Walmart-Amazon) to 25% (Company). The number of attributes ranges from 1 to 8.

Table 3: The 13 datasets divided into 4 categories of domains. The datasets marked with † are text-heavy (Textual). Each dataset with * has an additional dirty version to test the models’ robustness against noisy data.

Datasets	Domains
Amazon-Google, Walmart-Amazon*	software / electronics
Abt-Buy†, Beer	product
DBLP-ACM*, DBLP-Scholar*, iTunes-Amazon*	citation / music
Company†, Fodors-Zagats	company / restaurant

Among the datasets, the Abt-Buy and Company datasets are text-heavy meaning that at least one attributes contain long text. Also, following [32], we use the dirty version of the DBLP-ACM, DBLP-Scholar, iTunes-Amazon, and Walmart-Amazon datasets to measure the robustness of the models against noise. These datasets are generated from the clean version by randomly emptying attributes and appending their values to another randomly selected attribute.

Each dataset is split into the training, validation, and test sets using the ratio of 3:1:1. The same split of the datasets is also used in the evaluation of other EM solutions [17, 23, 32]. We list the size of each dataset in Table 5.

The WDC product data corpus [37] contains 26 million product offers and descriptions collected from e-commerce websites [54]. The goal is to find product offer pairs that refer to the same product. To evaluate the accuracy of product matchers, the dataset provides 4,400 manually created golden labels of offer pairs from 4 categories: computers, cameras, watches, and shoes. Each category has a fixed number of 300 positive and 800 negative pairs. For training, the dataset provides for each category pairs that share the same product ID such as GTINs or MPNs mined from the product’s webpage. The negative examples are created by selecting pairs that have high textual similarity but different IDs. These labels are further reduced to different sizes to test the models’ label efficiency. We summarize the different subsets in Table 4. We refer to these subsets as the WDC datasets.

Table 4: Different subsets of the WDC product data corpus. Each subset (except Test) is split into a training set and a validation set with a ratio of 4:1 according to the dataset provider [37]. The last column shows the positive rate (%POS) of each category in the xLarge set. The positive rate on the test set is 27.27% for all the categories.

Categories	Test	Small	Medium	Large	xLarge	%POS
Computers	1,100	2,834	8,094	33,359	68,461	14.15%
Cameras	1,100	1,886	5,255	20,036	42,277	16.98%
Watches	1,100	2,255	6,413	27,027	61,569	15.05%
Shoes	1,100	2,063	5,805	22,989	42,429	9.76%
All	4,400	9,038	25,567	103,411	214,736	14.10%

Each entry in this dataset has 4 attributes: title, description, brand, and specTable. Following the setting in [37] for DeepMatcher, we allow DITTO to use any subsets of attributes to determine the best combination. We found in our experiments that DITTO achieves the best performance when it uses only the title attribute. We provide further justification of this choice in Appendix F of the full version.

4.2 Implementation and experimental setup

We implemented DITTO in PyTorch [34] and the Transformers library [56]. We currently support 4 pre-trained models: DistilBERT [43], BERT [13], RoBERTa [28], and XLNet [59]. We use the base uncased variant of each model in all our experiments. We further apply the half-precision floating-point (fp16) optimization to accelerate the training and prediction speed. In all the experiments, we fix the max sequence length to be 256 and the learning rate to be 3e-5 with a linearly decreasing learning rate schedule. The batch size is 32 if MixDA is used and 64 otherwise. The training process runs a fixed number of epochs (10, 15, or 40 depending on the dataset size) and returns the checkpoint with the highest F1 score on the validation set. We conducted all experiments on a p3.xlarge AWS EC2 machine with 4 V100 GPUs (1 GPU per run).

Compared methods. We compare DITTO with the SOTA EM solution DeepMatcher. We also consider other baseline methods including Magellan [24], DeepER [14], and follow-up works of DeepMatcher [17, 23]. We also compare with variants of DITTO without the data augmentation (DA) and/or domain knowledge (DK) optimization to evaluate the effectiveness of each component. We summarize these methods below. We report the average F1 of 5 repeated runs in all the settings.

- **DeepMatcher:** DeepMatcher [32] is the SOTA matching solution. Compared to DITTO, DeepMatcher customizes the RNN architecture to aggregate the attribute values, then compares/aligns the aggregated representations of the attributes. DeepMatcher leverages FastText [5] to train the word embeddings. When reporting DeepMatcher’s F1 scores, we use the numbers in [32] for the ER-Magellan datasets and numbers in [37] for the WDC datasets. We also reproduced those results using the open-sourced implementation.
- **DeepMatcher+:** Follow-up work [23] slightly outperforms DeepMatcher in the DBLP-ACM dataset and [17] achieves better F1 in the Walmart-Amazon and Amazon-Google datasets. According to [32], the Magellan system ([24], based on classical ML models) outperforms DeepMatcher in the Beer and iTunes-Amazon datasets. We also implemented and ran DeepER [14], which is another RNN-based EM solution. We denote by DeepMatcher+ (or simply DM+) the best F1 scores among DeepMatcher and these works aforementioned. We summarize in Appendix C of [27] the implementation details and performance of each method.
- **DITTO:** This is the full version of our system with all 3 optimizations, domain knowledge (DK), TF-IDF summarization (SU), and data augmentation (DA) turned on. See the details below.
- **DITTO(DA):** This version only turns on the DA (with MixDA) and SU but does not have the DK optimization. We apply one of the span-level or attribute-level DA operators listed in Table 2 with the entry_swap operator. We compare the different combinations and report the best one. Following [30], we apply MixDA with the interpolation parameter λ sampled from a Beta distribution Beta(0.8, 0.8).
- **DITTO(DK):** With only the DK and SU optimizations on, this version of DITTO is expected to have lower F1 scores but train much faster. We apply the span-typing to datasets of each domain according to Table 1 and apply the span-normalization on the number spans.

- **Baseline:** This base form of DITTO corresponds simply to fine-tuning a pre-trained LM on the EM task. We did not apply any optimizations on the baseline. For each ER-Magellan dataset, we tune the LM for the baseline and found that RoBERTa generally achieves the best performance. Thus, we use RoBERTa in the other 3 DITTO variants (DITTO, DITTO(DA), and DITTO(DK)) by default across all datasets. The Company dataset is the only exception, where we found that the BERT model performs the best. For the WDC benchmark, since the training sets are large, we use DistilBERT across all settings for faster training.

There is a concurrent work [6], which also applies pre-trained LM to the entity matching problem. The proposed method is similar to the baseline method above, but due to the difference in the evaluation methods ([6] reports the best epoch on the test set, instead of the validation set), the reported results in [6] is not directly comparable. We summarize in Appendix E of [27] the difference between DITTO and [6] and explain why the reported results are different.

4.3 Main results

Table 5 shows the results of the ER-Magellan datasets. Overall, DITTO (with optimizations) achieves significantly higher F1 scores than the SOTA results (DM+). DITTO without optimizations (i.e., the baseline) achieves comparable results with DM+. DITTO outperforms DM+ in all 13 cases and by up to 31% (Dirty, Walmart-Amazon) while the baseline outperforms DM+ in 12/13 cases except for the Company dataset with long text.

In addition, we found that DITTO is better at datasets with small training sets. Particularly, the average improvement on the 7 smallest datasets is 15.6% vs. 1.48% on average on the rest of datasets. DITTO is also more robust against data noise than DM+. In the 4 dirty datasets, the performance degradation of DITTO is only 0.57 on average while the performance of DM+ degrades by 8.21. These two properties make DITTO more attractive in practical EM settings.

Moreover, in Appendix D of [27], we show an evaluation of DITTO’s label efficiency on 5 of the ER-Magellan medium-size datasets. In 4/5 cases, when trained on less than 20% of the original training data, DITTO is able to achieve close or even better performance than DM+ when the full training sets are in use.

DITTO also achieves promising results on the WDC datasets (Table 6). DITTO achieves the highest F1 score of 94.08 when using all the 215k training data, outperforming the previous best result by 3.92. Similar to what we found in the ER-Magellan datasets, the improvements are higher on settings with fewer training examples (to the right of Table 6). The results also show that DITTO is more *label efficient* than DeepMatcher. For example, when using only 1/2 of the data (Large), DITTO already outperforms DeepMatcher with all the training data (xLarge) by 2.89 in All. When using only 1/8 of the data (Medium), the performance is within 1% close to DeepMatcher’s F1 when 1/2 of the data (Large) is in use. The only exception is the shoes category. This may be caused by the large gap of the positive label ratios between the training set and the test set (9.76% vs. 27.27% according to Table 4).

4.4 Ablation study

Next, we analyze the effectiveness of each component (i.e., LM, SU, DK, and DA) by comparing DITTO with its variants without these optimizations. The results are shown in Table 5 and Figure 4.

Table 5: F1 scores on the ER-Magellan EM datasets. The numbers of DeepMatcher+ (DM+) are the highest available found in [17, 23, 32] or re-produced by us.

Datasets	DM+	DITTO	DITTO (DA)	DITTO (DK)	Baseline	Size
Structured						
Amazon-Google	70.7	75.58 (+4.88)	75.08	74.67	74.10	11,460
Beer	78.8	94.37 (+15.57)	87.21	90.46	84.59	450
DBLP-ACM	98.45	98.99 (+0.54)	99.17	99.10	98.96	12,363
DBLP-Google	94.7	95.6 (+0.9)	95.73	95.80	95.84	28,707
Fodors-Zagats	100	100.00 (+0.0)	100.00	100.00	98.14	946
iTunes-Amazon	91.2	97.06 (+5.86)	97.40	97.80	92.28	539
Walmart-Amazon	73.6	86.76 (+13.16)	85.50	83.73	85.81	10,242
Dirty						
DBLP-ACM	98.1	99.03 (+0.93)	98.94	99.08	98.92	12,363
DBLP-Google	93.8	95.75 (+1.95)	95.47	95.57	95.44	28,707
iTunes-Amazon	79.4	95.65 (+16.25)	95.29	94.48	92.92	539
Walmart-Amazon	53.8	85.69 (+31.89)	85.49	80.67	82.56	10,242
Textual						
Abt-Buy	62.8	89.33 (+26.53)	89.79	81.69	88.85	9,575
Company	92.7	93.85 (+1.15)	93.69	93.15	41.00	112,632

Table 6: F1 scores on the WDC product matching datasets. The numbers for DeepMatcher (DM) are taken from [37].

Size Methods	xLarge (1/1)		Large (1/2)		Medium (1/8)		Small (1/20)	
	DM	Ditto	DM	Ditto	DM	Ditto	DM	Ditto
Computers	90.80	95.45	89.55	91.70	77.82	88.62	70.55	80.76
		+4.65		+2.15		+10.80		+10.21
Cameras	89.21	93.78	87.19	91.23	76.53	88.09	68.59	80.89
		+4.57		+4.04		+11.56		+12.30
Watches	93.45	96.53	91.28	95.69	79.31	91.12	66.32	85.12
		+3.08		+4.41		+11.81		+18.80
Shoes	92.61	90.11	90.39	88.07	79.48	82.66	73.86	75.89
		-2.50		-2.32		+3.18		+2.03
All	90.16	94.08	89.24	93.05	79.94	88.61	76.34	84.36
		+3.92		+3.81		+8.67		+8.02

The use of a pre-trained LM contributes to a large portion of the performance gain. In the ER-Magellan datasets (excluding Company), the average improvement of the baseline compared to DeepMatcher+ is 7.75, which accounts for 78.5% of the improvement of the full DITTO (9.87). While DeepMatcher+ and the baseline DITTO (essentially fine-tuning a LM) are comparable on the Structured datasets, the baseline performs much better on all the Dirty datasets and the Abt-Buy dataset. This confirms our intuition that the language understanding capability is a key advantage of DITTO over existing EM solutions. The Company dataset is a special case because the length of the company articles (3,123 words on average) is much greater than the max sequence length of 256. The SU optimization increases the F1 score of this dataset from 41% to over 93%. In the WDC datasets, across the 20 settings, LM contributes to 3.41 F1 improvement on average, which explains 55.3% of improvement of the full DITTO (6.16).

The DK optimization is more effective on the ER-Magellan datasets. Compared to the baseline, the improvement of DITTO(DK) is 1.08 on average and is up to 5.88 on the Beer dataset while the improvement is only 0.22 on average on the WDC datasets. We inspected

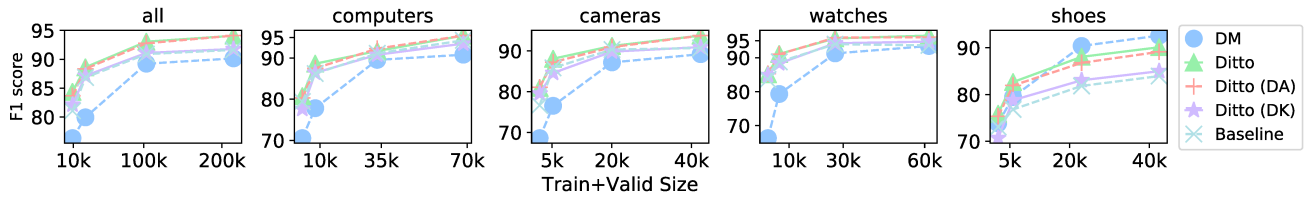


Figure 4: F1 scores on the WDC datasets of different versions of DITTO. DM: DeepMatcher.

the span-typing output and found that only 66.2% of entry pairs have spans of the same type. This is caused by the current NER module not extracting product-related spans with the correct types. We expect DK to be more effective if we use an NER model trained on the product domain.

DA is effective on both datasets and more significantly on the WDC datasets. The average F1 score of the full DITTO improves upon DITTO(DK) (without DA) by 1.39 and 2.53 respectively in the two datasets. In the WDC datasets, we found that the span_del operator always performs the best while the best operators are diverse in the ER-Magellan datasets. We list the best operator for each dataset in Table 7. We note that there is a large space of tuning these operators (e.g., the MixDA interpolation parameter, maximal span length, etc.) and new operators to further improve the performance.

Table 7: Datasets that each DA operator achieves the best performance. The suffix (S)/(D) and (Both) denote the clean/dirty version of the dataset or both of them. All operators are applied with the entry_swap operator.

Operator	Datasets
span_shuffle	DBLP-ACM (Both), DBLP-Google (Both), Abt-Buy
span_del	Walmart-Amazon(D), Company, all of WDC
attr_del	Beer, iTunes-Amazon(S), Walmart-Amazon(S)
attr_shuffle	Fodors-Zagats, iTunes-Amazon(D)

5 CASE STUDY: EMPLOYER MATCHING

We present a case of applying DITTO to a real-world EM task. An online recruiting platform would like to join its internal employer records with newly collected public records to enable downstream aggregation tasks. Given two tables *A* and *B* (internal and public) of employer records, the goal is to find, for each record in table *B*, a record in table *A* that represents the same employer. Both tables have 6 attributes: name, addr, city, state, zipcode, and phone. Our goal is to find matches with both high precision and recall.

Basic blocking. Our first challenge is size of the datasets. Table 8 shows that both tables are of nontrivial sizes even after deduplication. The first blocking method we designed is to only match companies with *the same zipcode*. However, since 60% of records in Table *A* do not have the zipcode attribute and some large employers have multiple sites, we use a second blocking method that returns for each record in Table *B* the top-20 most similar records in *A* ranked by the TF-IDF cosine similarity of name and addr attributes. We use the union of these two methods as our blocker, which produces 10 million candidate pairs.

Data labeling. We labeled 10,000 pairs sampled from the results of each blocking method (20,000 labels in total). We sampled pairs

Table 8: Sizes of the two employer datasets to be matched.

	TableA		TableB		#Candidates
	original	deduplicated	original	deduplicated	Basic blocking
Size	789,409	788,094	412,418	62,511	10,652,249

of high similarity with higher probability to increase the difficulty of the dataset to train more robust models. The positive rate of all the labeled pairs is 39%. We split the labeled pairs into training, validation, and test sets by the ratio of 3:1:1.

Applying DITTO. The user of DITTO does not need to extensively tune the hyperparameters but only needs to specify the domain knowledge and choose a data augmentation operator. We observe that the street number and the phone number are both useful signals for matching. Thus, we implemented a simple recognizer that tags the first number string in the addr attribute and the last 4 digits of the phone attribute. Since we would like the trained model to be robust against the large number of missing values, we choose the attr_del operator for data augmentation.

We plot the model’s performance in Figure 5. DITTO achieves the highest F1 score of 96.53 when using all the training data. DITTO outperforms DeepMatcher (DM) in F1 and trains faster (even when using MixDA) than DeepMatcher across different training set sizes.

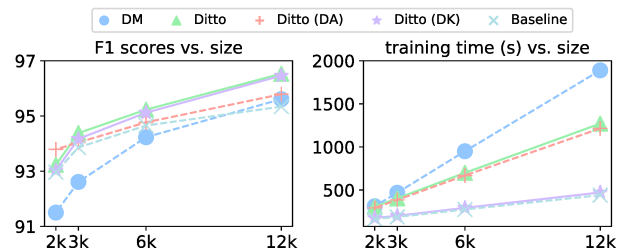


Figure 5: F1 and training time for the employer matching models.

Advanced blocking. Optionally, before applying the trained model to all the candidate pairs, we can use the labeled data to improve the basic blocking method. We leverage Sentence-BERT [41], a variant of the BERT model that trains sentence embeddings for sentence similarity search. The trained model generates a high-dimensional (e.g., 768 for BERT) vector for each record. Although this model has a relatively low F1 (only 92%) thus cannot replace DITTO, we can use it with vector similarity search to quickly find record pairs that are likely to match. We can greatly reduce the matching time by only testing those pairs of high cosine similarity. We list the running time for each module in Table 9. With this technique, the overall EM process is accelerated by 3.8x (1.69 hours vs. 6.49 hours with/without advanced blocking).

Table 9: Running time for blocking and matching with DITTO. Advanced blocking consists of two steps: computing the representation of each record with Sentence-BERT [41] (Encoding) and similarity search by blocked matrix multiplication [1] (Search). With advanced blocking, we only match each record with the top-10 most similar records according to the model.

	Basic Blocking	Encoding (GPU)	Search (CPU)	Matching	
				(top-10)	(ALL)
Time (s)	537.26	2,229.26	1,981.97	1,339.36	22,823.43

6 RELATED WORK AND DISCUSSION

EM solutions have tackled the blocking problem [2, 8, 16, 33, 52] and the matching problem with rules [11, 15, 45, 51], crowdsourcing [18, 22, 50], or machine learning [4, 10, 18, 24, 44].

Recently, EM solutions used deep learning and achieved promising results [14, 17, 23, 32, 62]. DeepER [14] trains EM models based on the LSTM [21] neural network architecture with word embeddings such as GloVe [36]. DeepER also proposed a blocking technique to represent each entry by the LSTM’s output. Our advanced blocking technique based on Sentence-BERT [41], described in Section 5, is inspired by this. Auto-EM [62] improves deep learning-based EM models by pre-training the EM model on an auxiliary task of entity type detection. DITTO also leverages transfer learning by fine-tuning pre-trained LMs, which are more powerful in language understanding. We did not compare DITTO with Auto-EM in experiments because the entity types required by Auto-EM are not available in our benchmarks. However, we expect that pre-training DITTO with EM-specific data/tasks can improve the performance of DITTO further and is part of our future work. DeepMatcher introduced a design space for applying deep learning to EM. Following their template architecture, one can think of DITTO as replacing both the attribute embedding and similarity representation components in the architecture with a single pre-trained LM such as BERT, thus providing a much simpler overall architecture.

All systems, Auto-EM, DeepER, DeepMatcher, and DITTO formulate matching as a binary classification problem. The first three take a pair of data entries of the same arity as input and aligns the attributes before passing them to the system for matching. On the other hand, DITTO serializes both data entries as one input with structural tags intact. This way, data entries of different schemas can be uniformly ingested, including hierarchically formatted data such as those in JSON. Our serialization scheme is not only applicable to DITTO, but also to other systems such as DeepMatcher. In fact, we serialized data entries to DeepMatcher under one attribute using our scheme and observed that DeepMatcher improved by as much as 5.2% on some datasets.

A concurrent work [6] also applies pre-trained LMs to the entity matching problem and achieves good performance. While the proposed method in [6] is similar to the baseline version of DITTO, DITTO can be further optimized using domain knowledge, data augmentation, and summarization. We also present a comprehensive experiment analysis on more EM benchmarks using a more standard evaluation method. We provide a detailed comparison between DITTO and [6] in Appendix E of our full version [27].

External knowledge is known to be effective in improving neural network models in NLP tasks [7, 47, 53, 58]. To incorporate

domain knowledge, DITTO modularizes the way domain knowledge is incorporated by allowing users to specify and customize rules for preprocessing input entries. Data augmentation (DA) has been extensively studied in computer vision and has recently received more attention in NLP [30, 55, 57]. We designed a set of DA operators suitable for EM and apply them with MixDA [30], a recently proposed DA strategy based on convex interpolation. To our knowledge, this is the first time data augmentation has been applied to EM.

Active learning is a recent trend in EM to train high-quality matching models with limited labeling resources [19, 23, 29, 38]. Under the active learning framework, the developer interactively labels a small set of examples to improve the model while the updated model is used to sample new examples for the next labeling step. Although active learning’s goal of improving label efficiency aligns with data augmentation in DITTO, they are different solutions, which can be used together; active learning requires human interaction in each iteration, whereas data augmentation does not. According to [29], one needs to adjust the model size and/or the training process such that the response time becomes acceptable for user interaction in active learning. Thus, we consider applying it to DITTO is not straightforward because of the relatively long fine-tuning time of the DITTO. We leave this aspect to future development of DITTO.

Discussion. Like other deep learning-based EM solutions, DITTO requires a non-trivial amount of labeled training examples (e.g., the case study requires 6k examples to achieve 95% F1) and DITTO’s DA and DK optimizations help reduce the labeling requirement to some extent. Currently, the LMs that we have tested in DITTO are pre-trained on general English text corpora and thus might not capture well EM tasks with a lot of numeric data and/or specific domains such as the scientific domain. For domain-specific tasks, a potential solution is to leverage specialized LMs such as SciBERT [3] or BioBERT [26] trained on scientific and biology corpus respectively. For numeric data, a good candidate solution would be a hybrid neural network similar to [20, 60] that combines the numeric features with the textual features.

7 CONCLUSION

We present DITTO, an EM system based on fine-tuning pre-trained Transformer-based language models. DITTO uses a simple architecture to leverage pre-trained LMs and is further optimized by injecting domain knowledge, text summarization, and data augmentation. Our results show that it outperforms existing EM solutions on all three benchmark datasets with significantly less training data. DITTO’s good performance can be attributed to the improved language understanding capability mainly through pre-trained LMs, the more accurate text alignment guided by the injected knowledge, and the data invariance properties learned from the augmented data. We plan to further explore our design choices for injecting domain knowledge, text summarization, and data augmentation. In addition, we plan to extend DITTO to other data integration tasks beyond EM, such as entity type detection and schema matching with the ultimate goal of building a BERT-like model for tables.

REFERENCES

- [1] Firas Abuzaid, Geet Sethi, Peter Bailis, and Matei Zaharia. 2019. To Index or Not to Index: Optimizing Exact Maximum Inner Product Search. In *Proc. ICDE '19*. IEEE, 1250–1261.
- [2] Linkage Rohan Baxter, Rohan Baxter, Peter Christen, et al. 2003. A comparison of fast blocking methods for record. (2003).
- [3] Iz Beltagy, Kyle Lo, and Arman Cohan. 2019. SciBERT: A pretrained language model for scientific text. *arXiv preprint arXiv:1903.10676* (2019).
- [4] Mikhail Bilenko and Raymond J Mooney. 2003. Adaptive duplicate detection using learnable string similarity measures. In *Proc. KDD '03*. 39–48.
- [5] Piotr Bojanowski, Edouard Grave, Armand Joulin, and Tomas Mikolov. 2017. Enriching word vectors with subword information. *TACL* 5 (2017), 135–146.
- [6] Ursin Brunner and Kurt Stockinger. 2020. Entity matching with transformer architectures-a step forward in data integration. In *EDBT*.
- [7] Qian Chen, Xiaodan Zhu, Zhen-Hua Ling, Diana Inkpen, and Si Wei. 2018. Neural natural language inference models enhanced with external knowledge. In *Proc. ACL '18*. 2406–2417.
- [8] Peter Christen. 2011. A survey of indexing techniques for scalable record linkage and deduplication. *TKDE* 24, 9 (2011), 1537–1555.
- [9] Kevin Clark, Urvashi Khandelwal, Omer Levy, and Christopher D. Manning. 2019. What Does BERT Look at? An Analysis of BERT’s Attention. In *Proc. BlackBoxNLP '19*. 276–286.
- [10] William W Cohen and Jacob Richman. 2002. Learning to match and cluster large high-dimensional data sets for data integration. In *Proc. KDD '02*. 475–480.
- [11] Nilesh Dalvi, Vibhor Rastogi, Anirban Dasgupta, Anish Das Sarma, and Tamas Sarlos. 2013. Optimal hashing schemes for entity matching. In *Proc. WWW '13*. 295–306.
- [12] Sanjib Das, AnHai Doan, Paul Suganthan G. C., Chaitanya Gokhale, Pradap Konda, Yash Govind, and Derek Paulsen. [n.d.]. The Magellan Data Repository. <https://sites.google.com/site/anhaidgroup/projects/data>.
- [13] Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2019. BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding. In *Proc. NAACL-HLT '19*. 4171–4186.
- [14] Muhammad Ebraheem, Saravanan Thirumuruganathan, Shafiq Joty, Mourad Ouzzani, and Nan Tang. 2018. Distributed representations of tuples for entity resolution. *PVLDB* 11, 11 (2018), 1454–1467.
- [15] Ahmed Elmagarmid, Ihab F Ilyas, Mourad Ouzzani, Jorge-Arnulfo Quiané-Ruiz, Nan Tang, and Si Yin. 2014. NADEEF/ER: generic and interactive entity resolution. In *Proc. SIGMOD '14*. 1071–1074.
- [16] Jeffrey Fisher, Peter Christen, Qing Wang, and Erhard Rahm. 2015. A clustering-based framework to control block sizes for entity resolution. In *Proc. KDD '15*. 279–288.
- [17] Cheng Fu, Xianpei Han, Le Sun, Bo Chen, Wei Zhang, Suhui Wu, and Hao Kong. 2019. End-to-end multi-perspective matching for entity resolution. In *Proc. IJCAI '19*. AAAI Press, 4961–4967.
- [18] Chaitanya Gokhale, Sanjib Das, AnHai Doan, Jeffrey F Naughton, Narasimhan Rampalli, Jude Shavlik, and Xiaojin Zhu. 2014. Corleone: Hands-off crowdsourcing for entity matching. In *Proc. SIGMOD '14*. 601–612.
- [19] Sairam Gurajada, Lucian Popa, Kun Qian, and Prithviraj Sen. 2019. Learning-Based Methods with Human-in-the-Loop for Entity Resolution. In *CIKM*. 2969–2970.
- [20] Jonathan Herzig, Paweł Krzysztof Nowak, Thomas Müller, Francesco Piccinno, and Julian Martin Eisenschlos. 2020. TAPAS: Weakly Supervised Table Parsing via Pre-training. *arXiv preprint arXiv:2004.02349* (2020).
- [21] Sepp Hochreiter and Jürgen Schmidhuber. 1997. Long short-term memory. *Neural computation* 9, 8 (1997), 1735–1780.
- [22] Adam Marcus Eugene Wu David Karger and Samuel Madden Robert Miller. 2011. Human-powered Sorts and Joins. *PVLDB* 5, 1 (2011).
- [23] Jungo Kasai, Kun Qian, Sairam Gurajada, Yunyao Li, and Lucian Popa. 2019. Low-resource Deep Entity Resolution with Transfer and Active Learning. In *Proc. ACL '19*. 5851–5861.
- [24] Pradap Konda, Sanjib Das, Paul Suganthan G. C., AnHai Doan, Adel Ardalan, Jeffrey R. Ballard, Han Li, Fatemah Panahi, Haojun Zhang, Jeffrey F. Naughton, Shishir Prasad, Ganesh Krishnan, Rohit Deep, and Vijay Raghavendra. 2016. Magellan: Toward Building Entity Matching Management Systems. *PVLDB* 9, 12 (2016), 1197–1208.
- [25] Hanna Köpcke, Andreas Thor, and Erhard Rahm. 2010. Evaluation of entity resolution approaches on real-world match problems. *PVLDB* 3, 1-2 (2010), 484–493.
- [26] Jinhyuk Lee, Wonjin Yoon, Sungdong Kim, Donghyeon Kim, Sunkyu Kim, Chan Ho So, and Jaewoo Kang. 2020. BioBERT: a pre-trained biomedical language representation model for biomedical text mining. *Bioinformatics* 36, 4 (2020), 1234–1240.
- [27] Yuliang Li, Jinfeng Li, Yoshihiko Suhara, AnHai Doan, and Wang-Chiew Tan. 2020. Deep Entity Matching with Pre-Trained Language Models. *arXiv preprint arXiv:2004.00584* (2020).
- [28] Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Mandar Joshi, Danqi Chen, Omer Levy, Mike Lewis, Luke Zettlemoyer, and Veselin Stoyanov. 2019. RoBERTa: A robustly optimized bert pretraining approach. *arXiv preprint arXiv:1907.11692* (2019).
- [29] Venkata Vamsikrishna Meduri, Lucian Popa, Prithviraj Sen, and Mohamed Sarwat. 2020. A Comprehensive Benchmark Framework for Active Learning Methods in Entity Matching. In *SIGMOD*. 1133–1147.
- [30] Zhengjie Miao, Yuliang Li, Xiaolan Wang, and Wang-Chiew Tan. 2020. Snippet: Semi-supervised Opinion Mining with Augmented Data. In *Proc. WWW '20*.
- [31] Rada Mihalcea and Paul Tarau. 2004. TextRank: Bringing order into text. In *Proc. EMNLP '04*. 404–411.
- [32] Sidharth Mudgal, Han Li, Theodoros Rekatsinas, AnHai Doan, Youngchoon Park, Ganesh Krishnan, Rohit Deep, Esteban Arcaute, and Vijay Raghavendra. 2018. Deep learning for entity matching: A design space exploration. In *Proc. SIGMOD '18*. 19–34.
- [33] George Papadakis, Dimitrios Skoutas, Emmanouil Thanos, and Themis Palpanas. 2019. Blocking and Filtering Techniques for Entity Resolution: A Survey. *arXiv preprint arXiv:1905.06167* (2019).
- [34] Adam Paszke, Sam Gross, Francisco Massa, Adam Lerer, James Bradbury, Gregory Chanan, Trevor Killeen, Zeming Lin, Natalia Gimelshein, Luca Antiga, et al. 2019. PyTorch: An imperative style, high-performance deep learning library. In *Proc. NeurIPS '19*. 8024–8035.
- [35] Fabian Pedregosa, Gaël Varoquaux, Alexandre Gramfort, Vincent Michel, Bertrand Thirion, Olivier Grisel, Mathieu Blondel, Peter Prettenhofer, Ron Weiss, Vincent Dubourg, et al. 2011. Scikit-learn: Machine learning in Python. *the Journal of machine Learning research* 12 (2011), 2825–2830.
- [36] Jeffrey Pennington, Richard Socher, and Christopher D Manning. 2014. Glove: Global vectors for word representation. In *Proc. EMNLP '14*. 1532–1543.
- [37] Anna Primpeli, Ralph Peeters, and Christian Bizer. 2019. The WDC training dataset and gold standard for large-scale product matching. In *Companion Proc. WWW '19*. 381–386.
- [38] Kun Qian, Lucian Popa, and Prithviraj Sen. 2017. Active learning for large-scale entity resolution. In *CIKM*. 1379–1388.
- [39] Alec Radford, Jeff Wu, Rewon Child, David Luan, Dario Amodei, and Ilya Sutskever. 2019. Language Models are Unsupervised Multitask Learners. (2019).
- [40] Alec Radford, Jeffrey Wu, Rewon Child, David Luan, Dario Amodei, and Ilya Sutskever. 2019. Language models are unsupervised multitask learners. *OpenAI Blog* 1, 8 (2019), 9.
- [41] Nils Reimers and Iryna Gurevych. 2019. Sentence-BERT: Sentence embeddings using Siamese BERT-networks. In *Proc. EMNLP-IJCNLP '19*. 3982–3992.
- [42] Alexander M Rush, Sumit Chopra, and Jason Weston. 2015. A neural attention model for abstractive sentence summarization. In *Proc. EMNLP '15*.
- [43] Victor Sanh, Lysandre Debut, Julien Chaumond, and Thomas Wolf. 2019. DistilBERT, a distilled version of BERT: smaller, faster, cheaper and lighter. In *Proc. EMNLP '19*.
- [44] Sunita Sarawagi and Anuradha Bhamidipaty. 2002. Interactive deduplication using active learning. In *Proc. KDD '02*. 269–278.
- [45] Rohit Singh, Venkata Vamsikrishna Meduri, Ahmed Elmagarmid, Samuel Madden, Paolo Papotti, Jorge-Arnulfo Quiané-Ruiz, Armando Solar-Lezama, and Nan Tang. 2017. Synthesizing entity matching rules by examples. *PVLDB* 11, 2 (2017), 189–202.
- [46] Spacy. [n.d.]. <https://spacy.io/api/entityrecognizer>.
- [47] Yu Sun, Shuohuan Wang, Yukun Li, Shikun Feng, Xuyi Chen, Han Zhang, Xin Tian, Danxiang Zhu, Hao Tian, and Hua Wu. 2019. ERNIE: Enhanced representation through knowledge integration. *arXiv preprint arXiv:1904.09223* (2019).
- [48] Ian Tenney, Dipanjan Das, and Ellie Pavlick. 2019. BERT Rediscovered the Classical NLP Pipeline. In *Proc. ACL '19*. 4593–4601.
- [49] Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. 2017. Attention is all you need. In *Proc. NIPS '17*. 5998–6008.
- [50] Jiannan Wang, Tim Kraska, Michael J Franklin, and Jianhua Feng. 2012. CrowdER: crowdsourcing entity resolution. *PVLDB* 5, 11 (2012), 1483–1494.
- [51] Jiannan Wang, Guoliang Li, Jeffrey Xu Yu, and Jianhua Feng. 2011. Entity matching: How similar is similar. *PVLDB* 4, 10 (2011), 622–633.
- [52] Qing Wang, Mingyuan Cui, and Huizhi Liang. 2015. Semantic-aware blocking for entity resolution. *TKDE* 28, 1 (2015), 166–180.
- [53] Xiang Wang, Xiangnan He, Yixin Cao, Meng Liu, and Tat-Seng Chua. 2019. KGAT: Knowledge Graph Attention Network for Recommendation. In *Proc. KDD '19*. 950–958.
- [54] WDC Product Data Corpus. [n.d.]. <http://webdatacommons.org/largescaleproductcorpus/v2>.
- [55] Jason Wei and Kai Zou. 2019. EDA: Easy Data Augmentation Techniques for Boosting Performance on Text Classification Tasks. In *Proc. EMNLP-IJCNLP '19*. 6382–6388.
- [56] Thomas Wolf, Lysandre Debut, Victor Sanh, Julien Chaumond, Clement Delangue, Anthony Moi, Pierric Cistac, Tim Rault, Rémi Louf, Morgan Funtowicz, et al. 2019. HuggingFace’s Transformers: State-of-the-art Natural Language Processing. *arXiv preprint arXiv:1910.03771* (2019).

- [57] Qizhe Xie, Zihang Dai, Eduard Hovy, Minh-Thang Luong, and Quoc V Le. 2019. Unsupervised data augmentation. *arXiv preprint arXiv:1904.12848* (2019).
- [58] Bishan Yang and Tom Mitchell. 2017. Leveraging Knowledge Bases in LSTMs for Improving Machine Reading. In *Proc. ACL '17*. 1436–1446.
- [59] Zhilin Yang, Zihang Dai, Yiming Yang, Jaime Carbonell, Russ R Salakhutdinov, and Quoc V Le. 2019. XLNet: Generalized autoregressive pretraining for language understanding. In *Proc. NeurIPS '19*. 5754–5764.
- [60] Pengcheng Yin, Graham Neubig, Wen-tau Yih, and Sebastian Riedel. 2020. TaBERT: Pretraining for Joint Understanding of Textual and Tabular Data. *arXiv preprint arXiv:2005.08314* (2020).
- [61] Hongyi Zhang, Moustapha Cisse, Yann N Dauphin, and David Lopez-Paz. 2018. mixup: Beyond empirical risk minimization. In *Proc. ICLR '18*.
- [62] Chen Zhao and Yeye He. 2019. Auto-EM: End-to-end Fuzzy Entity-Matching using Pre-trained Deep Models and Transfer Learning. In *Proc. WWW '19*. 2413–2424.