



Universiteit
Leiden
The Netherlands

Can BERT dig it? Named entity recognition for information retrieval in the archaeology domain

Brandsen, A.; Verberne, S.; Lambers, K; Wansleebeben, M.

Citation

Brandsen, A., Verberne, S., Lambers, K., & Wansleebeben, M. (2022). Can BERT dig it? Named entity recognition for information retrieval in the archaeology domain. *Journal On Computing And Cultural Heritage*, 15(3), 1-18. doi:10.1145/3497842

Version: Publisher's Version

License: [Creative Commons CC BY 4.0 license](#)

Downloaded from: <https://hdl.handle.net/1887/3464621>

Note: To cite this publication please use the final published version (if applicable).

Can BERT Dig It? Named Entity Recognition for Information Retrieval in the Archaeology Domain

ALEX BRANDSEN, SUZAN VERBERNE, KARSTEN LAMBERS, and
MILCO WANSLEEBEN, Leiden University

The amount of archaeological literature is growing rapidly. Until recently, these data were only accessible through metadata search. We implemented a text retrieval engine for a large archaeological text collection (~658 million words). In archaeological IR, domain-specific entities such as locations, time periods and artefacts play a central role. This motivated the development of a named entity recognition (NER) model to annotate the full collection with archaeological named entities. In this article, we present ArcheoBERTje, a BERT (Bidirectional Encoder Representations from Transformers) model pre-trained on Dutch archaeological texts. We compare the model's quality and output on an NER task to a generic multilingual model and a generic Dutch model. We also investigate ensemble methods for combining multiple BERT models, and combining the best BERT model with a domain thesaurus using conditional random fields. We find that ArcheoBERTje outperforms both the multilingual and Dutch model significantly with a smaller standard deviation between runs, reaching an average F1 score of 0.735. The model also outperforms ensemble methods combining the three models. Combining ArcheoBERTje predictions and explicit domain knowledge from the thesaurus did not increase the F1 score. We quantitatively and qualitatively analyse the differences between the vocabulary and output of the BERT models on the full collection and provide some valuable insights in the effect of fine-tuning for specific domains. Our results indicate that for a highly specific text domain such as archaeology, further pre-training on domain-specific data increases the model's quality on NER by a much larger margin than shown for other domains in the literature, and that domain-specific pre-training makes the addition of domain knowledge from a thesaurus unnecessary.

CCS Concepts: • **Information systems** → **Structured text search**; • **Computing methodologies** → **Information extraction**; **Supervised learning by classification**; **Neural networks**; **Ensemble methods**; • **Applied computing** → **Arts and humanities**;

Additional Key Words and Phrases: Archaeology, language modelling, BERT

ACM Reference format:

Alex Brandsen, Suzan Verberne, Karsten Lambers, and Milco Wansleeben. 2022. Can BERT Dig It? Named Entity Recognition for Information Retrieval in the Archaeology Domain. *J. Comput. Cult. Herit.* 15, 3, Article 51 (September 2022), 18 pages. <https://doi.org/10.1145/3497842>

1 INTRODUCTION

Like in other domains, archaeologists produce large amounts of text about their research. Besides research leading to scholarly output, commercial archaeology companies survey and excavate areas before developers build

Authors' addresses: A. Brandsen (corresponding author), K. Lambers, and M. Wansleeben, Faculty of Archaeology, Leiden University, Einsteinweg 2, Leiden, The Netherlands, 2333 CC; emails: {a.brandsen, k.lambers, m.wansleeben}@arch.leidenuniv.nl; S. Verberne, Leiden Institute for Advanced Computer Science, Leiden University, Niels Bohrweg 1, Leiden, The Netherlands, 2333 CA; email: s.verberne@liacs.leidenuniv.nl.



This work is licensed under a Creative Commons Attribution International 4.0 License.

© 2022 Copyright held by the owner/author(s).

1556-4673/2022/09-ART51

<https://doi.org/10.1145/3497842>

there and might destroy the archaeological remains. For each of these investigations, a report is written and stored in a repository. In the Netherlands, more than 4,000 of these documents are produced every year [42], with the total currently estimated at 70,000. These documents are used to some extent by both academic and commercial archaeologists to do further research.

Currently, this so-called grey literature is underused, as the available search tools only offer metadata search, making searching through these reports time consuming and inaccurate [24]. A strong need for better search tools has been well documented in prior work [6, 24, 43, 53], as the information in the full text of the reports can be of great value. Archaeological information needs are often recall-oriented list questions, consisting of a combination of What, Where and When aspects, e.g., “Find all cremations from the Early Middle Ages in the Netherlands” [6]. These are difficult to satisfy as the previously available search interfaces only offer search on the title, a short description, and sometimes information about the dating and type of archaeology encountered (stored in metadata fields), but the latter two are often missing or incorrectly assigned. Archaeologists want to search in more detail, and are often interested in the so-called by-catch: a single find unlike the rest of an excavation. For example, on an excavation yielding mainly Bronze Age material, a single Medieval cremation most likely will not be mentioned in the metadata, making it difficult to retrieve without manually searching through all PDFs.

To address these needs, we implemented a text retrieval engine for a large collection of archaeological reports in the Netherlands. The retrieval collection contains an export (obtained in 2017) of every PDF file in the DANS repository¹ with the label ‘Archaeology’. This totals more than 60,000 documents and 658 million tokens.

A full text search would alleviate a lot of the current challenges archaeologists face in their search of information, but as Habermehl [24] mentions, even in the relatively structured metadata, both synonymy and polysemy are a challenge, which is likely to be even worse in the free text in the body of the documents:

- Synonymy is a challenge because it leads to a lower recall: as there are numerous ways to write concepts relevant to archaeology, a search for one of these variants will not return the others. Specifically time periods have many synonyms. For example, the ‘Early Middle Ages’ can also be expressed as the ‘Early Medieval Period’, or ‘Merovingian Period’, or as dates that fall within the period, such as ‘600 CE’ and ‘1400 BP’.
- Polysemy, however, causes precision to be lower because one word can have multiple meanings, causing irrelevant meanings to appear in the search results. A good archaeological example is *Swifterbant*, which is a location, a type of pottery, an excavation event, and a time period. This problem of polysemy causes query ambiguity, as a full-text search engine does not know which meaning the user is looking for in their query, and then also does not know which meaning to retrieve from the corpus.

Automatic query expansion is often used to combat problems with synonymy, either by using thesauri or embeddings to add synonyms and similar terms to a query and increase the recall [11, 47]. Unfortunately in the case of time periods, this is difficult, as some time periods span thousands or millions of years, and adding each year with multiple variations (AD, BC, CE, BCE, BP) would result in an extremely large query. Polysemy is usually addressed in web search engines by diversifying search results or query suggestions [10, 46]: for each possible meaning of the ambiguous query, at least one relevant result is shown. For our specific domain, this is not possible because we do not have the large amount of user traffic that generic web search engines have, to be able to learn the different relevant results for any query term.

Instead, we opt for **named entity recognition (NER)** to automatically detect archaeological entities in the corpus, and then allow archaeologists to find these using an entity-based query interface, combined with a full text search. The entity search attempts to solve the polysemy problem, as the user specifies—in the structured query interface—which meaning of a word they are looking for, e.g., the Location² *Swifterbant*. In this case,

¹<https://easy.dans.knaw.nl/ui/home>.

²Entity types will be capitalised from here on for clarity.

only documents where the Location entity *Swifterbant* has been detected will be returned. Although this helps the user specify their query, it also means that entities that have not been correctly identified will not be returned; in other words, errors in the NER output might propagate to retrieval errors. Therefore, to give the user freedom in the query form that best suits their information need, we combine entity search with full-text search.

We have previously published a prototype of our search engine online. The search engine uses Elasticsearch [22] to index the full text, and in the prototype, entities were automatically labelled with a baseline NER model based on **conditional random fields (CRF)**. The resulting entity-based full-text search was experienced as positive by a focus group of archaeologists [6].

However, the baseline NER model offers room for improvement. As prior work on archaeological NER indicated, CRF with common token-, context- and thesaurus-based features leads to relatively low F1 scores, around 0.50 to 0.70 [6, 7]. In the past couple of years, transfer learning, and specifically BERT (Bidirectional Encoder Representations from Transformers) models [19], have been used successfully to get **state-of-the-art (SotA)** results for NER. On general domain benchmarks, the SotA methods yield impressive F1 scores of up to 0.943 [60]. However, in other domains and languages, the performance of NER systems is generally lower [32].

BERT has not been applied to the archaeology domain yet in any language, and we believe this domain could benefit from context-dependent embeddings due to the polysemy presented previously. Two generic Dutch BERT models have been released [17, 18] which can help our research. Prior work on language- and domain-specific BERT models reports mixed results on the effect of pre-training on language- and domain-specific data (see Section 2.4). In this article, we investigate whether BERT can improve NER in the Dutch archaeology domain, and to what extent further pre-training on domain-specific texts improves the quality of the model. We compare Google’s multilingual model [19], the Dutch BERTje model [17], and our own ArcheoBERTje model that we further pre-trained on Dutch excavation reports. We do not compare the Dutch RobBERT model, as it has a different training procedure and longer training times. We analyse the differences between the three models, and we experiment with ensembles to combine multiple models and a domain-specific thesaurus. As there is unfortunately no test collection with relevance assessments available for the Dutch archaeology domain, we do not evaluate the performance of the information retrieval, only the performance of the NER.

We address the following research question:

- (1) To what extent does further pre-training a BERT model with domain-specific training data improve the model’s quality in our highly specific domain?
- (2) Can a domain-specific BERT model be improved by adding domain knowledge from a thesaurus in a CRF ensemble model?
- (3) What errors are made by the models, and what are the differences in predicted entities between the three models?

The contributions of our article are threefold. First, we propose entity-driven full-text search in which the professional user enters a structured query, and documents are filtered for the occurrence of the query entities detected by our new domain-specific BERT model. Second, we show that for a highly specific domain such as archaeology, further pre-training on domain-specific data increases the model’s quality on NER by a much larger margin than shown for other domains in the literature. Third, we show that the domain-specific BERT model outperforms ensemble methods combining different BERT models, and also outperforms a CRF-based ensemble of BERT with explicit domain knowledge from the archaeological thesaurus.

We make our modified training dataset, the pre-trained ArcheoBERTje model, and the fine-tuned ArcheoBERTje model for NER publicly available [5].³

³<https://doi.org/10.5281/zenodo.4739063>, also available via the HuggingFace library for ease of use: <https://huggingface.co/alexbrandsen>.

2 RELATED WORK

In this section, we first summarise different approaches to NER (knowledge-driven and data-driven), followed by a discussion of related work on NER for document retrieval, on IR and NER in the archaeological domain, and we summarise the prior work on domain-specific BERT models.

2.1 Knowledge-Driven and Data-Driven NER

Early NER systems were knowledge based, and relied on thesauri and handcrafted rules to detect entities [41]. These methods are limited by the coverage of the thesaurus. Therefore, data-driven methods have become more popular, typically approaching NER as a supervised machine learning problem.

A highly effective machine learning method is CRF [31], which has become a common baseline for NER. Since 2011, word embeddings have become increasingly important as representations in NER. Especially Word2vec [34] has been used extensively for NER [44, 45]. These embedding-based methods typically feed the embeddings to CRF and/or Bi-LSTM algorithms to make NER predictions.

A big shift in NLP was introduced by Devlin et al. [19], who presented their BERT architecture in 2019. BERT and other contextual embedding architectures are currently achieving SotA results with transfer learning for a large range of NLP tasks, including NER. Two major differences with previous embedding models are (1) that BERT embeddings are contextual, meaning that the same token can have a different embedding based on context, and (2) that it handles out-of-vocabulary words effectively, by dividing tokens into sub-tokens it does have in vocabulary, using the WordPiece [19] or SentencePiece [29] tokeniser.

Recent results indicate that ensemble methods that combining generic and domain-specific BERT models [14], combining BERT with dictionary features [33], or adding a CRF on top of BERT [48] can improve NER quality. In this work, we investigate whether addition of information from a thesaurus can improve NER in a highly specific domain.

2.2 NER for Document Retrieval

In the context of document retrieval, NER can play a role in better ranking or filtering documents based on entities in the query. Guo et al. [23] were the first to address the task of recognising named entities in queries. They found that, despite queries in web search being short, 70% of the queries contained a named entity. They classify the entities according to a pre-defined taxonomy using a weakly supervised topic modelling approach on the query data. Cowan et al. [16] also address NER in queries, but for the travel domain. They use CRF on the queries for extracting the relevant entities.

More recently, the relevance of NER on queries has been emphasised for the e-commerce domain. Wen et al. [57] and Cheng et al. [13] both implement end-to-end query analysis methods for e-commerce search; the extracted queries are then used to filter the retrieved products.

As opposed to the prior work, we do not focus on query analysis but on document analysis; our expert users prefer the use of structured queries, which makes query analysis unnecessary (see Section 4.4). Our documents, however, are long and unstructured (as opposed to the products in e-commerce search), making NER on the document side necessary for matching structured queries to the relevant documents.

2.3 IR and NER in Archaeology

As argued by Richards et al. [43], archaeology has great potential for thesaurus-based IR and NER, as it has a relatively well-controlled vocabulary and there are thesauri of archaeological concepts available in multiple languages. However, unlike some other fields, archaeology terminology partly consists of common words, like ‘pit’, ‘well’ and ‘post’. In addition, words can be archaeological entities or not, depending on the context in which they are used (past or present). For example, the word ‘road’ is not archaeologically relevant in the snippet “pit next to the main road”, but is part of an archaeological entity in the snippet “a Roman road from 34 CE”.

Archaeology has started experimenting with IR relatively recently. The focus of the prior work is on Information or Knowledge Extraction, mainly for automatically generating document metadata. An early study by Amrani et al. [3] aimed specifically at extracting information for archaeology professionals in a knowledge-based approach. A more data-driven approach using machine learning to detect Time Period entities was investigated in the OpenBoek project [38, 39], but since then most studies have been knowledge driven [9, 27, 55, 56].

More recently, Talboom [50] experimented with embeddings in a Bi-LSTM model to recognise zooarchaeological entities (species and specific bones). A notable exception to the Information Extraction research we often see in archaeology is the work by Gibbs and Colley [21], who created a full-text search engine on a small Australian corpus (roughly 1,000 documents) combined with facets based on manually entered metadata.

So far, NLP in the archaeology domain has not benefitted from BERT-based models. We believe it is a good candidate domain for BERT, as the polysemy mentioned in the introduction and the present/past distinction mentioned previously should be easier to detect with the context-dependent embeddings that BERT produces.

2.4 Language- and Domain-Specific BERT Models

The original BERT paper [19] not only presented an English BERT model but also a multilingual model (multi-BERT) trained on data in 104 languages. This model is often used when no single-language model is available [25, 28, 35]. Research by Wu and Dredze [59] shows that multiBERT achieved higher accuracy on NER and other NLP tasks than monolingual models trained with comparable amounts of data. Moon et al. [35] also showed that fine-tuning multiBERT on a mixed language NER dataset provided better results than fine-tuning on individual languages.

However, recent work has shown that for some languages, multiBERT is outperformed by language-specific BERT models [37]. For NER, this has been shown for Finnish [54], Dutch [17], German [12] and Russian [30], among other languages.

For specific domains, it has been shown that further pre-training the English BERT-base model on large amounts of text from that domain increases the quality of the model on multiple tasks, although sometimes by a small margin. BioBERT in the biomedical domain shows an increase in F1 for NER of only 0.62% points [32]. SciBERT, trained on a large amount of scientific texts from different domains, shows an increase in F1 for NER of 2% to 5% points, indicating that domain pre-training is useful for NER [4]. They also show that training BERT from scratch with a domain-specific vocabulary does not increase F1 substantially compared to fine-tuning an existing BERT model with an existing generic vocabulary, gaining only 0.6% points.

When we look at research done on non-English in a specialised domain like our study, there is little prior work. A study in the Russian cyber-security domain shows that the Russian model (RuBERT) outperformed multiBERT, and further pre-training RuBERT with domain-specific documents yielded the highest F1 [51]. In the Spanish biomedical domain, Akhtyamova [2] shows a similar result, although their NER BERT model is trained for 30 epochs, possibly leading to overfitting.

To the best of our knowledge, we are the first to address domain-specific NER for Dutch, as well as the first to automatically label a large archaeological document collection with our domain-specific BERT model for the purpose of professional search.

3 DATA

The unlabelled dataset we use for further pre-training the Dutch BERTje model to ArcheoBERTje consists of more than 60,000 documents and 658 million tokens across 16.6 million sentences, around 2 GB of data. The documents mainly consist of survey/excavation reports but also include other documents such as research plans, appendices, maps and data descriptions.

The labelled training data we use for NER we created previously [7], and consists of 15 documents that have been annotated by archaeology students. Although 15 reports is a relatively low number, these are longer than

Table 1. Descriptions and Examples for Each Entity Type

Entity	Description	Examples
Artefact	An archaeological object found in the ground.	Axe, pot, stake, arrow head, coin
Time Period	A defined (archaeological) period in time.	Middle Ages, Neolithic, 500 BC, 4000 BP
Location	A placename or (part of) an address.	Amsterdam, Steenstraat 1, Lutjebroek
Context	An anthropogenic, definable part of a stratigraphy. Something that can contain Artefacts	Rubbish pit, burial mound, stake hole
Material	The material an Artefact is made of.	Bronze, wood, flint, glass
Species	A species' name (in Latin or Dutch)	Cow, <i>Corvus Corax</i> , oak

Note: Examples are translated from Dutch. Adapted from Brandsen et al. [7], p. 4574.

average documents, totalling 1,343 pages (average 89 pages per document), containing roughly 440,000 tokens and almost 43,000 annotated entities across six categories: Artefacts, Time Periods, Locations, Contexts, Materials and Species (Table 1). The inter-annotator agreement reported is 95% (average pairwise F1 score), so it is of relatively high quality [7]. The data is stored in the BIO annotation schema and is available for download.⁴

The dataset has been split into five folds of three documents each. All methods are evaluated using this fivefold split.

3.1 Pre-processing

For cross validation, we divided the 15 annotated documents across five folds so that each fold has a roughly equal number of tokens. The exact fold split and training data can be found on in the Zenodo repository.

We found that in the dataset, sentences often exceed the maximum sequence length of 512 WordPiece tokens. This is not because sentences actually have more than 512 words, but partly because tables and OCRed maps and images create very long 'sentences' that are not cut up by the sentence detection algorithm. The other cause is that words that are uncommon outside of archaeology are cut up into many sub-tokens by the WordPiece tokeniser, as they do not exist in the vocabulary (also see Section 6.2).

Since sentences longer than 512 tokens will be trimmed, some of the input tokens will not get a prediction. To counteract this, we wrote a pre-processing script that attempts to break at a punctuation mark (',', ';' or ';') between the 60th and 90th token, and if there are none, it inserts a line break after the 90th token. This shortened the sentences sufficiently to have almost no instances where the sentence was longer than 512 WordPiece tokens. Only 136 tokens in the entire dataset fell outside the 512 limit and received no prediction (default "O" label). These tokens only contained two entities, so the effect on the performance metrics will be negligible.

4 METHODS

4.1 Baselines

As the first baseline, we use the method we published previously [7], where we trained a CRF model using common word shape features (e.g., occurrence of uppercase letters, numbers), part-of-speech tags (e.g., noun, verb) and an archaeological thesaurus in a five-word window, and performed hyperparameter optimisation. We used the same features, leading to a micro F1 score of 0.62. This is relatively low when comparing the score to NER in other domains, where F1 scores between 0.8 and 0.9 are common [2, 32].

The second baseline is the standard NER pipeline of spaCy 2.0, with default parameters (architecture: TransitionBasedParser.v2, random seed, max_steps: 20,000, Adam.v1 optimiser with learn_rate of 0.001). This method uses pre-existing Dutch word embeddings (nl_core_news_lg) with a deep convolutional neural network with residual connections, and a transition-based approach as the classifier [26].

⁴Zenodo repository: <http://doi.org/10.5281/zenodo.3544544>.

4.2 Fine-Tuning BERT for Dutch Archaeology and NER

Model training for evaluation. To train ArcheoBERTje, we started with the Dutch BERTje model [17] and further pre-trained the model with our complete unlabelled archaeological collection, split into a 90/10 train and validation set.⁵ We used the same configuration as BERTje, with a batch size of 4. We decided not to train a model from scratch as previous research showed only a minimal increase in quality compared to further pre-training [4] an existing model, and because our corpus is relatively small and would probably not be enough to train an effective model.

To fine-tune the BERT models for the NER task, we used the labelled data and fivefold cross validation as described in Section 3.⁶ For model comparison and to investigate the stability of each model with different random seeds, we trained all three models 10 times per fold, each time using a different seed (1, 2, 4, 8, 16, 32, 64, 128, 254, 512) and report averages over all runs and folds (50 runs in total per BERT model).

Model for full collection labelling. To create the best possible model for inference on the entire corpus, we performed a grid search across hyperparameters as suggested by Devlin et al. [19]. We optimised the hyperparameters with fold 2 as test set, fold 1 as development set and the other folds as training set, as this combination had the median F1 score across all models and folds. The grid search yielded the following optimal parameters for our data: two training epochs, 5×10^{-5} learning rate and 0.1 weight decay. We then fine-tuned the inference model on all labelled data with these hyperparameters. This way we maximise the amount of training data available for training the model that we use to label the full collection.

4.3 Ensemble Methods

As far as we are aware, we are the first to combine a multilingual model, a language-specific model and a domain-specific model into one ensemble method. We evaluate the following ensemble methods (one run over five folds per ensemble):

- Majority voting on the predictions of multiBERT, BERTje and ArcheoBERTje;
- CRF which uses the prediction labels of the three models as features;
- CRF which uses the prediction labels of ArcheoBERTje only;
- CRF which uses the prediction labels of the three models as features, combined with the baseline features;
- CRF which uses the prediction labels of ArcheoBERTje only, combined with the baseline features; and
- CRF which uses the embeddings produced by ArcheoBERTje as features.

The preceding ‘baseline features’ are those adopted from prior work (see Section 4.1) and include word shape, part-of-speech tags and thesaurus features. We optimised the hyperparameters of each CRF ensemble with gradient descent using the L-BFGS method, optimising c1 and c2 (the coefficients for L1 and L2 regularisation). The optimisation was run separately for each fold. All CRF ensembles use a five-token window, taking into account the features from the two tokens before and after the current token.

The thesaurus we use in our CRF baseline and ensembles is the ABR (*Archeologisch Basisregister*) [7, 8], a thesaurus containing time periods (e.g., Bronze Age), artefacts (e.g., axe) and materials (e.g., flint). A token is assigned the binary feature ‘occurs in period/artefact/material list’ if it is part of an n-gram that occurs in the thesaurus. So the token ‘Bronze’ would only be assigned a positive value for the feature if the token ‘age’ follows it.

4.4 Entity-Driven Document Search

Indexing. Before we index the documents, we first run the inference NER model on each page to detect the entities. We then store the entities and full text in a JSON file for each document, together with the relevant metadata (authors, DOI, coordinates, document type, etc.) retrieved from the DANS repository via an API.

⁵We used HuggingFace’s [58] language modelling script version 3.0.2.

⁶We used HuggingFace’s token classification script version 3.0.2.

Search through 60,000 excavation reports from the [DANS archive](#).

Query:

Use an asterisk (*) as wildcard, so "axe" also finds "axes".

Time Period

Optional: specify a starting and end year to search for a particular period.

Start year:

End year:

Concept search

Not getting the right results? Try searching on concepts:

Artefact:

Context:

Species:

Fig. 1. Query interface showing query for “Artefact: urn AND Context: cremation AND startdate < –2000 AND enddate > –800 AND fulltext: upside down”. Interface and query translated to English for the readers’ convenience.

To tackle the synonymy problem for time periods (see Section 1), we use a custom script that translates all extracted Time Period entities to year ranges. It uses regular expressions to convert dates (e.g., ‘100 BCE’, ‘start of the 9th century’) and an extended and customised version of the PeriodO time period gazetteer [40] to translate Time Periods (e.g., ‘Bronze Age’, ‘Medieval period’). These date ranges are added to the JSON and can be used to filter results by allowing users to specify a date range in their query. These JSON files are then sent to an instance of ElasticSearch running on a webserver, which indexes them. At the moment, the retrieval unit is a page, so for any query the terms/entities must occur together on a page. We are aware this is not optimal, as search terms might be split across pages. As such, in future work we will index per document section by using a section detection algorithm.

Query interface and analysis. Our search engine has a faceted search interface in which metadata filters are combined with entity fields and full-text search [52]. We have included facets for document type and subject (metadata fields). In addition, as requested by our target group, we added geographical search via a map functionality, which allows users to draw a rectangle or polygon to search only in a certain region.

At query time, the user can specify if they are looking for a specific entity type and/or specify a date range in which they are interested. The entities and date range are used to filter the result set and can be combined with a standard full text search. This allows for relatively complex queries such as “Artefact: urn AND Context: cremation AND startdate < –2000 AND enddate > –800 AND fulltext: upside down”. This example is a real request entered by an archaeologist, who was looking for upside down urns in the Bronze Age in or around cremations. Users do not need to use complex query syntax but can instead define their query by filling in the relevant fields in the graphical user interface, as shown in Figure 1.

Document ranking. Most archaeological information needs are recall-oriented tasks: the users want a complete list and do not mind having irrelevant results in the (top of) the result set [6]. As the focus of our work is on entity-driven search, we opt for the default ElasticSearch ranking model, consisting of TF-IDF and the field-length norm (the shorter the field, the higher the relevance) [20]. The only field included for ranking is the page text content; other fields are only used for filtering.

Note that we do not evaluate the ranking, because there is no test collection available yet for Dutch archaeological document retrieval. Therefore, the scope of this article is limited to the NER and the evaluation thereof.

5 RESULTS

5.1 Model Stability and Quality

Table 2 shows the micro average precision, recall and F1 score for the three BERT models, compared to the CRF and spaCy baselines. We find that the multilingual BERT model does not outperform the baselines, but the more specialised BERTje and ArcheoBERTje models do, with ArcheoBERTje achieving the highest F1 score.

Table 2. Micro Average Precision, Recall and F1 Score at Token Level (B and I Labels), over 10 Runs with Different Seeds, for Each of the Five Folds (50 Runs Total)

Model	Precision	Recall	F1 (Std.)	Fails
CRF Baseline	0.785	0.526	0.630	n/a
spaCy Baseline	0.717	0.602	0.654	n/a
multiBERT	0.623	0.550	0.583 (0.015)	4
BERTje	0.718	0.682	0.699 (0.005)	0
ArcheoBERTje	0.743	0.729	0.735 (0.004)	0

Note: Standard deviation of F1 over the 10 runs is added in brackets for the BERT models. Standard deviation of precision and recall lies between 0.006 and 0.020. The 'Fails' column indicates the number of times the model failed to learn (F1 = 0).

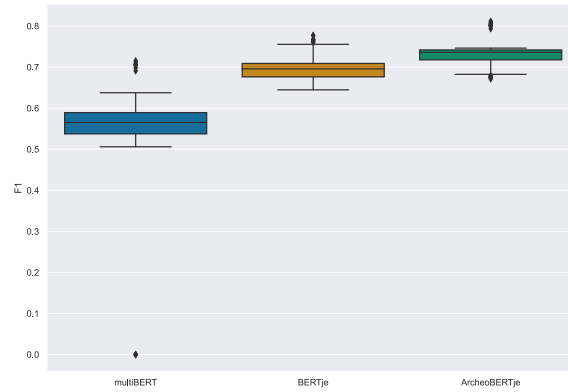


Fig. 2. Distribution of F1 scores over 10 runs with different seeds, for each of the five folds (50 runs per model). The zero scores for multiBERT are runs where the model failed to learn.

We also show the average standard deviation over 10 runs with different seeds for five folds. The standard deviation between runs is very low, between 0.015 and 0.004. The recent work by Tikhomirov et al. [51] reports a standard deviation of 0.015 to 0.008, similar to our results. When comparing the predicted labels of each of the models in a pairwise manner, the differences are significant according to McNemar's test (χ^2 between 650 and 4,276, $p < 0.00001$).

Figure 2 shows the distribution of F1 scores over the 50 runs per model in a boxplot. Here we again see that the standard deviation is low, and that ArcheoBERTje consistently outperforms the other two models. The F1 scores of 0 for multiBERT are outliers, and we assume these are caused by the ADAM optimiser getting stuck in a local minimum where the loss does not decrease. In this local minimum, predicting the majority class (O) seems to yield the highest accuracy, but of course O labels are not taken into account when calculating an F1 score for NER, so we get a score of zero. This can be solved by changing the learning rate, but this would not change the overall view that BERTje and ArcheoBERTje outperform multiBERT, so we did not investigate further on fixing this for multiBERT.

The low standard deviations for ArcheoBERTje indicate that further pre-training with domain-specific data not only increases the model quality on average but also makes the model more stable, reducing the chance of getting a sub-optimal model in a run.

Another way to compare the models is by looking at differences between the errors made. In Table 3, we report the top 10 most frequent error combinations for the three models. Here we can see that quite often, BERTje and

Table 3. The 10 Most Frequent Error Combinations between the Three Models for Which at Least One Model Has the Correct Prediction

Freq.	True	multiBERT	BERTje	ArcheoBERTje
1,137	B-LOC	O	B-LOC	B-LOC
1,122	B-ART	O	B-ART	B-ART
1,015	O	B-ART	O	O
575	B-SPE	O	B-SPE	B-SPE
561	O	B-LOC	O	O
466	B-PER	O	B-PER	B-PER
429	O	O	B-ART	B-ART
425	I-PER	O	I-PER	I-PER
402	B-ART	O	O	B-ART
373	O	I-PER	O	O

Note: Errors are marked in red.

Table 4. Micro F1 Score, Precision and Recall for the Six Ensemble Methods, for One Run over Five Folds

Ensemble	Precision	Recall	F1
ArcheoBERTje (50 runs avg)	0.743	0.729	0.735
ArcheoBERTje (optimised production model)	0.784	0.731	0.757
Majority Voting	0.784	0.695	0.737
CRF with 3 BERT model prediction labels as features	0.786	0.683	0.731
CRF with only production ArcheoBERTje predictions as features	0.786	0.717	0.750
CRF with 3 BERT model prediction labels + baseline features	0.795	0.644	0.712
CRF with production ArcheoBERTje prediction labels + baseline features	0.793	0.649	0.714
CRF with only production ArcheoBERTje embeddings as features	0.767	0.604	0.676

Note: ArcheoBERTje results averaged over 50 runs and the optimised production model are added for comparison. The ArcheoBERTje predictions used as features for CRF are from the production model. The baseline features are the word- and context-based features used for CRF in prior work.

ArcheoBERTje have similar predictions (whether correct or not), whereas multiBERT predicted a different label. We see that multiBERT often misses Locations (LOC), Artefacts (ART) and Species (SPE), and sometimes predicts entities that are not there. The first error combination where ArcheoBERTje outperforms BERTje is number 9, having correctly predicted B-ARTs while the other two models do not. In Sections 5.3 and 6.1, we further analyse the output and errors made by the ArcheoBERTje model to provide insight into the model's behaviour.

5.2 Ensembles

Table 4 shows the results of the ensemble methods.⁷ The highest F1 (0.757) is obtained by the optimised production ArcheoBERTje model.

The highest precision is obtained by the CRF ensemble with the baseline features combined with the predicted labels from all three models. The highest recall is achieved by ArcheoBERTje solo.⁸ Using a CRF with BERT embeddings as features instead of the default BERT classifier (softmax) does not increase performance. Given

⁷As the standard deviation between multiple runs is low, combining multiple runs of the same model in an ensemble model is very unlikely to increase the F1 score, at the expense of a vastly increased computing time and cost. Hence, we do not apply this approach.

⁸For general domain Portuguese NER, Souza et al. [48] show the same pattern: Portuguese BERT has the highest recall, whereas combining BERT with CRF yields the highest precision and F1.

Table 5. Overview of Entities Detected in the Entire Corpus, Showing Total and Unique Counts, Plus the Top Five for Each Entity (Translated from Dutch Where Relevant)

Entity	Total	Unique	Top Five
Artefacts	2,520,492	53,675	pottery, charcoal, flint, bone, brick
Contexts	1,602,124	21,319	pit, ditch, posthole, well, house
Materials	457,031	6,146	wooden, flint, wood, metal, bronze
Locations	3,488,698	147,077	nederland, ', groningen, noord - brabant, gelderland
Species	928,437	34,540	cow, hazel, sheep, goat, pig
Time Periods	4,698,323	98,445	roman period, iron age, 150–210, late medieval, modern
Total	13,695,105	361,202	

the recall-oriented nature of professional search tasks like ours, we prioritise recall over precision for the NER labelling, and use ArcheoBERTje for labelling the full collection.

5.3 Analysis of the Retrieval Collection

After labelling the full retrieval collection with ArcheoBERTje, we analyse the extracted entities. Table 5 shows for each entity type the total frequency and the amount of unique entities. We also show the top 5 entities extracted for each type (translated from Dutch to English).

As we already mentioned in the introduction, archaeologists are interested in the What, Where and When of excavations. And so we see that Artefacts, Locations and Time Periods are the most common entities:

- For Artefacts, we see that pottery and flint are common, which we expected, but apparently also charcoal, which we did not expect, but could be explained by the use of carbon dating, which often uses charcoal as a sample.
- In the Locations category, we see that the second most common entity is an apostrophe ('). Although this is clearly not a location, luckily it will not affect retrieval, as it is not something users would search for, and Elasticsearch does not include apostrophes in its index, so it would not match any documents. We speculate that ArcheoBERTje mislabels apostrophes as locations because of the occurrence of apostrophes in some Dutch place names (e.g., 's Hertogenbosch).
- For Time Periods, the only unexpected entry in the top 5 is “150–210”. When we investigated this further, we found this is actually a soil grain size used in coring reports, which have been incorrectly labelled as a time period by ArcheoBERTje. In addition, 150–210 μm is the grain size for medium course sand, apparently the most common grain size in the Netherlands. When we look further down the Time Period top 100, we also see other common grain sizes: 210–300, 105–150 and 105–210. This is an issue when searching for archaeology between 105 and 300 CE, as these irrelevant coring reports will also be returned. We believe that these errors are made because these numbers come from tables, and as such do not have any sentence context, making them difficult to predict correctly. The most likely way to fix this is by making a post-processing correction on the extracted entities. This is something we will improve in the next version of our NER method.

The grain sizes are also clearly visible in Figure 3, in which we have plotted the frequency of years found in entities in the corpus. The figure shows a number of plateaus, indicating the use of time periods instead of single dates—that is, the last plateau is the Late Middle Ages ending in 1500 CE. These plateaus are not completely flat as single dates and subperiods can cause spikes and smaller sub-plateaus.

The thin spike just after the year 0 can probably be attributed to misclassified entities—that is, the ‘10’ in ‘10-02-2006’ being labelled by ArcheoBERTje as a Time Period and translated to 10 CE. Other than this, we see

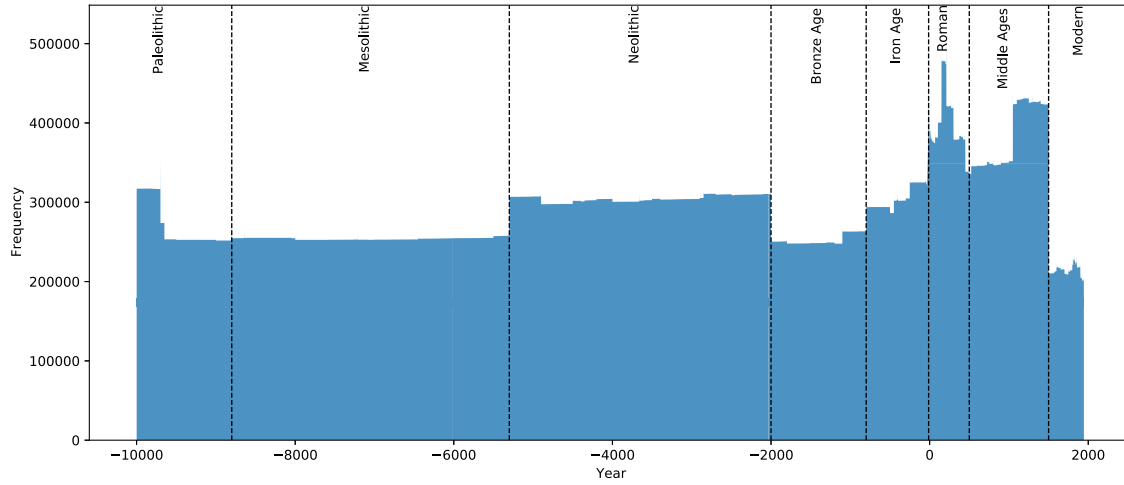


Fig. 3. Graph showing for each year in each detected time period, how often it occurs in our dataset, labelled by ArcheoBERTje. For clarity, years before 10,000 BCE are not included. Major time periods are denoted with dashed lines.

O	391053	1598	351	145	0	700	665	728	17	575	228	368	122
B-ART	1721	6492	265	183	1	38	0	136	0	15	0	136	0
I-ART	717	363	1086	22	5	3	6	1	1	0	2	26	4
B-MAT	197	372	60	559	5	0	0	2	0	10	0	19	1
I-MAT	23	4	32	9	10	0	0	0	0	2	0	1	0
B-PER	1086	14	4	0	0	6997	232	4	1	17	0	5	0
I-PER	1169	2	1	0	0	267	5917	0	0	0	0	0	1
B-CON	1627	182	3	0	0	46	0	3415	24	5	0	0	0
I-CON	93	1	3	0	0	0	0	45	24	0	2	0	0
B-LOC	788	10	0	2	0	21	1	3	0	3545	61	5	0
I-LOC	446	3	0	0	0	0	0	0	0	98	638	0	0
B-SPE	389	168	7	31	0	8	0	8	0	0	0	2186	44
I-SPE	124	8	53	1	4	0	0	0	0	0	0	40	542
	O	B-ART	I-ART	B-MAT	I-MAT	B-PER	I-PER	B-CON	I-CON	B-LOC	I-LOC	B-SPE	I-SPE

Fig. 4. Confusion matrix between true labels and ArcheoBERTje predictions.

a big plateau in the middle (5300–2000 BCE), which represents the Neolithic. This indicates that a large amount of data is available describing this period in the Stone Age.

6 DISCUSSION

6.1 Error Analysis

Figure 4 shows the confusion matrix between labels predicted by ArcheoBERTje and the true labels. The diagonal line and the first row and column are typical for NER. The diagonal shows the true positives, the top row is

Table 6. ArcheoBERTje Precision, Recall and F1 Score for Each Label

	Precision	Recall	F1
B-ART (Artefacts)	0.704	0.722	0.713
I-ART	0.582	0.486	0.530
B-CON (Contexts)	0.787	0.644	0.708
I-CON	0.358	0.143	0.204
B-MAT (Materials)	0.587	0.456	0.514
I-MAT	0.400	0.123	0.189
B-LOC (Locations)	0.831	0.799	0.815
I-LOC	0.685	0.538	0.603
B-SPE (Species)	0.785	0.769	0.777
I-SPE	0.759	0.702	0.729
B-PER (Time Periods)	0.866	0.837	0.851
I-PER	0.867	0.804	0.835
Macro Average	0.684	0.585	0.622
Micro Average	0.784	0.731	0.757

where the model predicted an entity where there is not one, and the first column is where the model predicted O where there should be an entity. We also see the I / B label confusion quite clearly, mainly for Time Periods and Locations, where the model predicts an I instead of a B, or the other way around.

A more interesting error is the confusion between Materials and Artefacts. This is caused by words like “flint”, which can be both an Artefact (“a piece of flint”) or a Material (“a flint axe”). In Dutch, “pottery” has the same issue. Even archaeologists struggle with distinguishing between the two [7], so it is unsurprising that ArcheoBERTje finds this difficult as well. As there is a lot of ambiguity in this entity category, perhaps merging the two categories into one entity type would increase the overall performance. We have seen in previous research that archaeologists will also confuse the two categories when creating queries, so having them both in one search field might not even cause any problems at search time.

Table 6 shows the evaluation per entity type. In general, the I labels are more difficult to predict, and Materials are more difficult than the other entities. In fact, Materials are currently not included in the search engine, as archaeologists find it difficult to differentiate between Materials and Artefacts in their queries, so this will not affect retrieval quality. When we remove Materials from the overall micro F1 score calculation, we get an increase of only around 0.01, as there are only a small number in our training data, around 3,000.

When we look at some of the errors made by ArcheoBERTje in more depth, we find some interesting patterns. For example, for missing B-ART labels, many errors are adjectives that were assigned the O label, such as for “big axe” or “complete pot”, the adjectives are labelled O, and axe / pot are labelled B-ART. This error is not surprising, as most archaeologists would probably find it difficult to define these entities as well. In addition, users are more likely to only search for the base artefact and not include an adjective, so they would search for “pot” not “complete pot”. In a pilot study evaluating our archaeological search engine, we analysed users’ search behaviour and found that of the 148 issued queries, none included an adjective.⁹

For Time Periods, we again see that adjectives are missed from the start of an entity, but also prepositions. Some examples include “from”, “between” and “start of”. Also we find that connecting words between Time Periods are missed, such as “and”, “or” and “±” (used to denote the standard deviation of a carbon dating). Although this does cause some noise, missing adjectives/prepositions or connecting words are not a considerable issue if the main period has been detected. For instance, for “start of 10th century”, if we miss “start of”, this means the year

⁹Extension and publication of this user study is part of our future work.

range is 900 to 1000 CE instead of 900 to 925 CE. Again, as archaeologists care more about recall than precision, this should not hinder their search.

The predicted Context¹⁰ entities also have some interesting anomalies. In particular, we analysed the top 10 most misclassified tokens and found that these are all words that can denote contemporary objects (and thus not a Context) or actual (pre-)historical Contexts. An example is *put*, which can mean a trench dug by archaeologists, or a water well found in an excavation, and both instances of *put* can contain an artefact, leading to similar contexts around these words. Other examples are “house”, “church”, “ditch”, “mine” and “settlement”. It seems that even with the context-dependent embeddings BERT produces, these ambiguous words are still a challenge. Perhaps future language models are more refined and might be able to distinguish between these types of ambiguous terms.

A special case is the word *poel* (pond). We see that this token is always labelled as O, whereas it is in fact a Context. When we checked the sentences this word occurs in, we see they are all very typical of Contexts, i.e., “we found pottery in the pond”, which is similar to sentence structures of other Contexts that are classified correctly. The only possible explanation we can find is that the word *poel* only occurs in one of the documents, so when this document is in the test set, the word does not occur at all in the train or dev set. This confirms the importance of creating train-test splits on the document level, to avoid overfitting. At the same time, this might be an issue that could be potentially alleviated by increasing the size of the training data.

More generally speaking, we see that the BERT models make impossible B and I predictions, i.e., an I label without a B label for the previous token. Unlike CRF, which learns the probabilities of two labels occurring after one another, BERT sees every token as an individual classification task without taking into account the predicted label of the previous token. This might explain why the CRF model with ArcheoBERTje labels as features (see Table 4) outperforms ArcheoBERTje on precision, as it corrects some of these mistakes. Perhaps another approach to correct this is a rule-based postprocessing step that checks the validity of I labels following B labels, and corrects impossible combinations.

During the annotation process, we used a test document of a hundred sentences (1,962 tokens) to calculate the inter-annotator agreement [7]. We added ArcheoBERTje predictions to this data, to see if ArcheoBERTje predictions are more often wrong when humans also have disagreement, indicating that the model mimics human confusion. We disregard tokens where everyone (including ArcheoBERTje) predicts an O label, leaving 292 tokens. In 57.5% of these tokens, all annotators and ArcheoBERTje predict the same label. In 31.5% of tokens, there is some disagreement between annotators, but ArcheoBERTje predicts the same label as the majority, and in 4.4% of tokens, ArcheoBERTje predicts a label different from the majority. In 6.5% of tokens, ArcheoBERTje predicts one label, whereas annotators all predict the same different label. This is only a small sample, but the preceding suggests that BERT models are decently equipped to learn from the majority where there is inter-annotator disagreement.

6.2 Tokenisation Issues

The vocabulary of a BERT model is determined by the collection used for pre-training. The WordPiece tokeniser optimises the set of (sub-word) tokens to maximise the coverage of the collection’s vocabulary. The same tokenisation is applied to the input sentences at inference. An example is shown below, where we compare tokenisation with the multiBERT and BERTje vocabularies. We see that target entities (“*Swifterbant*”, “*aardewerkscherven*” and “*Midden Neolithicum*”) are split up into three or more sub-tokens by the multiBERT and BERTje tokenisers.

Original sentence:

“*In put twee werden 3 Swifterbant aardewerkscherven aangetroffen uit het Midden Neolithicum.*” (“In trench two, 3 Swifterbant pottery shards from the Middle Neolithic were found.”)

¹⁰For clarity, Contexts are defined as an anthropogenic structures or objects that can contain Artefacts, such as rubbish pits, burials, houses, and so on.

multiBERT tokenisation (23 tokens):

In put twee werden 3 Swift ##er ##bant aarde ##werks ##cher ##ven aan ##get ##roffen uit het Midden Neo ##lit ##hic ##um.

BERTje tokenisation (20 tokens), also used for ArcheoBERTje:

In put twee werden 3 Swift ##er ##ban ##t aardewerk ##scher ##ven aangetroffen uit het Midden Neo ##lith ##icum.

As an additional analysis, we trained a SentencePiece tokeniser on our archaeological collection, with the same vocabulary size as the BERTje model (30k).

Archaeology tokenisation (14 tokens):

In put twee werden 3 Swifterbant aardewerk ##scherven aangetroffen uit het Midden Neolithicum .

The examples show that a more specific pre-training corpus would lead to more complete domain words. However, our collection is small for such from-scratch pre-training and the experiments in the sciBERT paper have shown that even a much larger pre-training collection only gives a +0.6% point F1 increase compared to further pre-training the generic model [4].

Understandably, the problem of input sequences longer than 512 tokens was occurring more often with the multilingual model, as the vocabulary (with fixed size) is not solely Dutch. This means that many less common Dutch words are not in the vocabulary, and are cut into many sub-tokens by the WordPiece tokeniser. This effect is aggravated by the Dutch language having a lot of compound words and a much longer average word length (4.8 in English [36] vs. 8 in Dutch [15]).

For our experiments comparing the different BERT models, it was sufficient to split up long sentences in the training and test data as a data pre-processing step. However, for the inference described in Section 5.3, we did not pre-process the text, and as such, entities found in long sentences after 512 SentencePiece tokens will have been assigned the incorrect “O” label, skewing the results. In future research, we will implement an automatic sentence splitting module, similar to the one implemented in FLAIR [1].

7 CONCLUSION

In this article, we have evaluated BERT models for NER in the Dutch archaeological domain, with the purpose of improving our archaeological search engine. We implemented the search engine for a large archaeological text collection, with a structured query interface that allows the specification of entity types. The document collection is automatically annotated with archaeological named entities such as Location, Time Period and Artefact.

In response to our research questions, first we found that fine-tuning a BERT model with domain-specific training data improves the model’s quality by a large margin for the archaeological domain, larger than in related work addressing domain-specific BERT models. We achieve an average F1 of 0.735 after hyperparameter optimisation, and very small standard deviations over runs with different random seeds.

Second, the domain-specific BERT model was superior in F1 and recall than an ensemble combining multiple BERT models, and could not be further improved by adding domain knowledge from a thesaurus in a CRF ensemble model. This indicates that after pre-training and fine-tuning on a domain-specific collection, the BERT model already covers the relevant information from the domain thesaurus. We did find a higher *precision* when we combined all three BERT models in a CRF model and added domain knowledge. However, as almost all information needs in archaeology are recall oriented, and combining models is computationally expensive and environmentally taxing [49], we opt for the ArcheoBERTje model for labelling the full retrieval collection.

Third, our error analysis shows that there is confusion between the Artefact and Material entities, similar to what humans experienced in the annotation process. For Artefacts and Time Periods, a common error is missing the adjective or preposition in an entity. The detection of Time Periods is a bit noisy, with other non-year

numbers erroneously labelled as time ranges. Context entities such as “house” and “ditch” are difficult for the models to distinguish from non-entity words. Creating train-test splits on the document level is important to avoid overfitting, as the consistently misclassified Context *poel* shows, which only occurs in one document. An analysis of tokenisation by each of the models indicates that the multiBERT model is hampered by the rough tokenisation, splitting many relevant terms in sub-words.

In the near future, we will evaluate the entity-driven search engine with users, both in a controlled experiment and in natural search contexts. We will also investigate entity-based query suggestion. Once entities are mapped to a thesaurus or embedded in a semantic space, this allows for query improvement by suggesting parent or sibling entities in the thesaurus or nearest-neighbours in the embedding space.

REFERENCES

- [1] Alan Akbik, Tanja Bergmann, Duncan Blythe, Kashif Rasul, Stefan Schweter, and Roland Vollgraf. 2019. FLAIR: An easy-to-use framework for state-of-the-art NLP. In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics (Demonstrations)*. 54–59. <https://doi.org/10.18653/v1/N19-4010>
- [2] Liliya Akhtyamova. 2020. Named entity recognition in Spanish biomedical literature: Short review and BERT model. In *Proceedings of the 26th Conference of Open Innovations Association (FRUCT’20)*. IEEE, Los Alamitos, CA 1–7. <https://doi.org/10.23919/FRUCT48808.2020.9087359>
- [3] A. Amrani, V. Abajian, and Y. Kodratoff. 2008. A chain of text-mining to extract information in archaeology. In *Proceedings of the Conference on Information and Communication Technologies: From Theory to Applications (ICTTA’08)*. 1–5. <https://doi.org/10.1109/ICTTA.2008.4529905>
- [4] Iz Beltagy, Kyle Lo, and Arman Cohan. 2020. SCIBERT: A pretrained language model for scientific text. In *Proceedings of the Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP’19)*. 3615–3620. <https://doi.org/10.18653/v1/d19-1371>
- [5] A. Brandsen. 2021. ArcheoBERTje—A Dutch BERT model for the archaeology domain. *Zenodo Repository*. Retrieved February 24, 2022 from <https://doi.org/10.5281/zenodo.4739063>
- [6] Alex Brandsen, Karsten Lambers, Suzan Verberne, and Milco Wansleebe. 2019. User requirement solicitation for an information retrieval system applied to Dutch grey literature in the archaeology domain. *Journal of Computer Applications in Archaeology* 2, 1 (3 2019), 21–30. <https://doi.org/10.5334/jcaa.33>
- [7] Alex Brandsen, Suzan Verberne, Milco Wansleebe, and Karsten Lambers. 2020. Creating a dataset for named entity recognition in the archaeology domain. In *Proceedings of the 12th Language Resources and Evaluation Conference*. 4573–4577.
- [8] R. W. Brandt, E. Drenth, M. Montforts, R. H. P. Proos, I. M. Roorda, and R. Wiemer. 1992. *Archeologisch Basisregister*. Technical Report. Rijksdienst voor Cultureel Erfgoed, Amersfoort.
- [9] Kate Byrne and Ewan Klein. 2010. Automatic extraction of archaeological events from text. In *Making History Interactive: Computer Applications and Quantitative Methods in Archaeology 2009*, B. Frischer, J. Crawford, and D. Koller (Eds.). BAR International Series 2079, Oxford, 48–56.
- [10] Gabriele Capannini, Franco Maria Nardini, Raffaele Perego, and Fabrizio Silvestri. 2011. Efficient diversification of web search results. *Proceedings of the VLDB Endowment* 4, 7 (2011), 451–459. <https://doi.org/10.14778/1988776.1988781>
- [11] Claudio Carpineto and Giovanni Romano. 2012. A survey of automatic query expansion in information retrieval. *ACM Computing Surveys* 44, 1 (1 2012), Article 1, 50 pages. <https://doi.org/10.1145/2071389.2071390>
- [12] Branden Chan, Stefan Schweter, and Timo Möller. 2021. German’s next language model. In *Proceedings of the 28th International Conference on Computational Linguistics*. 6788–6796. <https://doi.org/10.18653/v1/2020.coling-main.598>
- [13] Xiang Cheng, Mitchell Bowden, Bhushan Ramesh Bhange, Priyanka Goyal, Thomas Packer, and Faizan Javed. 2020. An end-to-end solution for named entity recognition in eCommerce search. *arXiv:2012.07553* (2020). <http://arxiv.org/abs/2012.07553>.
- [14] Jenny Copara, Nona Naderi, Julien Knafo, Patrick Ruch, and Douglas Teodoro. 2020. Named entity recognition in chemical patents using ensemble of contextual language models. *arXiv:2007.12569* (2020). <http://arxiv.org/abs/2007.12569>.
- [15] Hugo Brandt Corstius. 1981. *Opplerlandse Taal- & Letterkunde*. Querido.
- [16] Brooke Cowan, Sven Zethelius, Brittany Luk, Teodora Baras, Prachi Ukarde, and Daodao Zhang. 2015. Named entity recognition in travel-related search queries. In *Proceedings of the 29th AAAI Conference on Artificial Intelligence*. 3935–3941.
- [17] Wietse de Vries, Andreas van Cranenburgh, Arianna Bisazza, Tommaso Caselli, Gertjan van Noord, and Malvina Nissim. 2019. BERTje: A Dutch BERT Model. *arXiv:1912.09582* (2019). <http://arxiv.org/abs/1912.09582>.
- [18] Pieter Delobelle, Thomas Winters, and Bettina Berendt. 2020. RobBERT: A Dutch RoBERTa-based language model. In *Findings of the Association for Computational Linguistics: EMNLP 2020*. Association for Computational Linguistics, 3255–3265. <https://doi.org/10.18653/v1/2020.findings-emnlp.292>

- [19] Jacob Devlin, Ming Wei Chang, Kenton Lee, and Kristina Toutanova. 2019. BERT: Pre-training of deep bidirectional transformers for language understanding. In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*. 4171–4186. <https://doi.org/10.18653/v1/N19-1423>
- [20] Elasticsearch. 2018. Theory Behind Relevance Scoring. Retrieved February 24, 2022 from <https://www.elastic.co/guide/en/elasticsearch/guide/current/scoring-theory.html>.
- [21] Martin Gibbs and Sarah Colley. 2012. Digital preservation: Online access and historical archaeology ‘grey literature’ from New South Wales, Australia. *Australian Archaeology* 75 (2012), 95–103. <https://doi.org/10.1080/03122417.2012.11681957>
- [22] C. Gormley and Z. Tong. 2015. *Elasticsearch: The Definitive Guide: A Distributed Real-Time Search and Analytics Engine*. O’Reilly Media, Sebastopol.
- [23] Jiafeng Guo, Gu Xu, Xueqi Cheng, and Hang Li. 2009. Named entity recognition in query. In *Proceedings of the 32nd Annual International ACM SIGIR Conference on Research and Development in Information Retrieval (SIGIR’09)*. ACM, New York, NY, 267–274. <https://doi.org/10.1145/1571941.1571989>
- [24] Diederick Habermehl. 2019. *Over Zaaïen En Oogsten, De Kwaliteit En Bruikbaarheid Van Archeologische Rapporten Voor Synthetiserend Onderzoek*. Technical Report. Rijksdienst voor Cultureel Erfgoed, Amersfoort. <https://www.cultureelerfgoed.nl/publicaties/publicaties/2019/01/01/over-zaaien-en-oogsten>.
- [25] Kai Hakala and Sampo Pyysalo. 2019. Biomedical named entity recognition with multilingual BERT. In *Proceedings of the 5th Workshop on BioNLP Open Shared Tasks*. 56–61. <https://doi.org/10.18653/v1/d19-5709>
- [26] M. Honnibal and I. Montani. 2017. spaCy 2: Natural language understanding with Bloom embeddings, convolutional neural networks and incremental parsing. To appear.
- [27] S. Jeffrey, J. Richards, F. Ciravegna, S. Waller, S. Chapman, and Z. Zhang. 2009. The Archaeotools project: Faceted classification and natural language processing in an archaeological context. *Philosophical Transactions: Series A, Mathematical, Physical, and Engineering Sciences* 367, 1897 (June 2009), 2507–2519. <https://doi.org/10.1098/rsta.2009.0038>
- [28] Young Min Kim and Tae Hoon Lee. 2020. Korean clinical entity recognition from diagnosis text using BERT. *BMC Medical Informatics and Decision Making* 20, S7 (Sept. 2020), 242. <https://doi.org/10.1186/s12911-020-01241-8>
- [29] Taku Kudo and John Richardson. 2018. SentencePiece: A simple and language independent subword tokenizer and detokenizer for neural text processing. In *Proceedings of the Conference on Empirical Methods in Natural Language Processing: System Demonstrations (EMNLP’18)*. 66–71. <https://doi.org/10.18653/v1/d18-2012>
- [30] Yuri Kuratov and Mikhail Arkhipov. 2019. Adaptation of deep bidirectional multilingual transformers for Russian language. *arXiv:1905.07213* (2019). <http://arxiv.org/abs/1905.07213>.
- [31] John Lafferty, Andrew McCallum, Fernando C. N. Pereira, and Fernando Pereira. 2001. Conditional random fields: Probabilistic models for segmenting and labeling sequence data. In *Proceedings of the 18th International Conference on Machine Learning*. 282–289.
- [32] Jinhyuk Lee, Wonjin Yoon, Sungdong Kim, Donghyeon Kim, Sunkyu Kim, Chan Ho So, and Jaewoo Kang. 2019. BioBERT: A pre-trained biomedical language representation model for biomedical text mining. *Bioinformatics* 36, 4 (Sept. 2019), 1234–1240. <https://doi.org/10.1093/bioinformatics/btz682>
- [33] Xiangyang Li, Huan Zhang, and Xiao Hua Zhou. 2020. Chinese clinical named entity recognition with variant neural structures based on BERT methods. *Journal of Biomedical Informatics* 107 (July 2020), 103422. <https://doi.org/10.1016/j.jbi.2020.103422>
- [34] Tomas Mikolov, Kai Chen, Greg Corrado, and Jeffrey Dean. 2013. Efficient estimation of word representations in vector space. In *Proceedings of the 1st International Conference on Learning Representations (ICLR’13)—Workshop Track Proceedings*.
- [35] Taesun Moon, Parul Awasthy, Jian Ni, and Radu Florian. 2019. Towards lingua franca named entity recognition with BERT. *arXiv:1912.01389* (2019). <http://arxiv.org/abs/1912.01389>.
- [36] Peter Norvig. 2013. English Letter Frequency Counts: Mayzner Revisited. Retrieved February 24, 2022 from <http://norvig.com/mayzner.html>.
- [37] Debora Nozza, Federico Bianchi, and Dirk Hovy. 2020. What the [MASK]? Making sense of language-specific BERT models. *arXiv:2003.02912* (2020). <http://arxiv.org/abs/2003.02912>.
- [38] Hans Pajmams and Alex Brandsen. 2009. What is in a name: Recognizing monument names from free-text monument descriptions. In *Proceedings of the 18th Annual Belgian-Dutch Conference on Machine Learning (Benelearn)*, M. G. J. van Erp, J. H. Stehouwer, and M. van Zaanen (Eds.). Tilburg Centre for Creative Computing, Tilburg, 2–6.
- [39] H. Pajmams and A. Brandsen. 2010. Searching in archaeological texts: Problems and solutions using an artificial intelligence approach. *PalArch’s Journal of Archaeology of Egypt/Egyptology* 7, 2 (2010), 1–6.
- [40] Adam Rabinowitz, Ryan Shaw, Sarah Buchanan, Patrick Golden, and Eric Kansa. 2016. Making sense of the ways we make sense of the past: The PeriodO project. *Bulletin of the Institute of Classical Studies* 59, 2 (Dec. 2016), 42–55. <https://doi.org/10.1111/j.2041-5370.2016.12037.x>
- [41] Lisa F. Rau. 1991. Extracting company names from text. In *Proceedings of the 7th IEEE Conference on Artificial Intelligence Applications*. IEEE, Los Alamitos, CA, 29–32. <https://doi.org/10.1109/caia.1991.120841>
- [42] RCE. 2017. De Erfgoedmonitor. Retrieved February 24, 2022 from <https://erfgoedmonitor.cultureelerfgoed.nl/mosaic/kerncijfers/>.

- [43] Julian Richards, Douglas Tudhope, and Andreas Vlachidis. 2015. Text mining in archaeology: Extracting information from archaeological reports. In *Mathematics and Archaeology*, Juan A. Barcelo and Igor Bogdanovic (Eds.). CRC Press, Boca Raton, FL, 240–254. <https://doi.org/10.1201/b18530-15>
- [44] Miran Seok, Hye Jeong Song, Chan Young Park, Jong Dae Kim, and Yu Seop Kim. 2016. Named entity recognition using word embedding as a feature. *International Journal of Software Engineering and Its Applications* 10 (2016), 93–104. <https://doi.org/10.14257/ijseia.2016.10.2.08>
- [45] Scharolita Katharina Sienčnik. 2015. Adapting word2vec to named entity recognition. In *Proceedings of the 20th Nordic Conference of Computational Linguistics (NODALIDA'15)*, 239–243.
- [46] Yang Song, Dengyong Zhou, and Li Wei He. 2011. Post-ranking query suggestion by diversifying search results. In *Proceedings of the 34th International ACM SIGIR Conference on Research and Development in Information Retrieval (SIGIR'11)*. ACM, New York, NY, 815–824. <https://doi.org/10.1145/2009916.2010025>
- [47] Andrés Soto, José A. Olivas, and Manuel E. Prieto. 2008. Fuzzy approach of synonymy and polysemy for information retrieval. *Studies in Fuzziness and Soft Computing* 224 (2008), 179–198. https://doi.org/10.1007/978-3-540-76973-6_12
- [48] Fábio Souza, Rodrigo Nogueira, and Roberto Lotufo. 2019. Portuguese named entity recognition using BERT-CRF. *arXiv:1909.10649* (2019). <http://arxiv.org/abs/1909.10649>.
- [49] Emma Strubell, Ananya Ganesh, and Andrew McCallum. 2020. Energy and policy considerations for deep learning in NLP. In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics (ACL'19)*, 3645–3650. <https://doi.org/10.18653/v1/p19-1355>
- [50] Leontien Talboom. 2017. *Improving the Discoverability of Zooarchaeological Data with the Help of Natural Language Processing*. Ph.D. Dissertation. University of York.
- [51] Mikhail Tikhomirov, N. Loukachevitch, Anastasiia Sirotina, and Boris Dobrov. 2020. Using BERT and augmentation in named entity recognition for cybersecurity domain. In *Natural Language Processing and Information Systems*. Lecture Notes in Computer Science, Vol. 12089. Springer, 16–24. https://doi.org/10.1007/978-3-030-51310-8_2
- [52] Daniel Tunkelang. 2009. Faceted search. *Synthesis Lectures on Information Concepts, Retrieval, and Services* 1, 1 (2009), 1–80. <https://doi.org/10.2200/s00190ed1v01y200904icr005>
- [53] Monique Van den Dries. 2016. Is everybody happy? User satisfaction after ten years of quality management in European archaeological heritage management. In *When Valletta Meets Faro, the Reality of European Archaeology in the 21st Century, Proceedings of the International Conference*, P. Florjanowicz (Ed.). Archaeolingua, Lisbon, 126–135.
- [54] Antti Virtanen, Jenna Kanerva, Rami Ilo, Jouni Luoma, Juhani Luotolahti, Tapio Salakoski, Filip Ginter, and Sampo Pyysalo. 2019. Multilingual is not enough: BERT for Finnish. *arXiv:1912.07076* (2019). <http://arxiv.org/abs/1912.07076>.
- [55] Andreas Vlachidis, Ceri Binding, Keith May, and Douglas Tudhope. 2013. Automatic metadata generation in an archaeological digital library: Semantic annotation of grey literature. *Studies in Computational Intelligence* 458 (2013), 187–202. https://doi.org/10.1007/978-3-642-34399-5_10
- [56] A. Vlachidis, D. Tudhope, M. Wansleben, J. Azzopardi, K. Green, L. Xia, and H. Wright. 2017. *D16.4: Final Report on Natural Language Processing*. Technical Report. ARIADNE. http://legacy.ariadne-infrastructure.eu/wp-content/uploads/2019/01/D16.4_Final_Report_on_Natural_Language_Processing_Final.pdf.
- [57] Musen Wen, Deepak Kumar Vasthimal, Alan Lu, Tian Wang, and Aimin Guo. 2019. Building large-scale deep learning system for entity recognition in e-commerce search. In *Proceedings of the 6th IEEE/ACM International Conference on Big Data Computing, Applications, and Technologies (BDCAT'19)*. ACM, New York, NY, 149–154. <https://doi.org/10.1145/3365109.3368765>
- [58] Thomas Wolf, Lysandre Debut, Victor Sanh, Julien Chaumond, Clement Delangue, Anthony Moi, Pierric Cistac, et al. 2020. Transformers: State-of-the-art natural language processing. In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing: System Demonstrations*, 38–45. <https://doi.org/10.18653/v1/2020.emnlp-demos.6>
- [59] Shijie Wu and Mark Dredze. 2020. Are all languages created equal in multilingual BERT? In *Proceedings of the 5th Workshop on Representation Learning for NLP*, 120–130. <https://doi.org/10.18653/v1/2020.repl4nlp-1.16>
- [60] Ikuya Yamada, Akari Asai, Hiroyuki Shindo, Hideaki Takeda, and Yuji Matsumoto. 2020. LUKE: Deep contextualized entity representations with entity-aware self-attention. In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP'20)*, 6442–6454. <https://doi.org/10.18653/v1/2020.emnlp-main.523>

Received June 2021; revised October 2021; accepted November 2021