

Towards Comparable Historical NER: Building a Shared Evaluation Corpus for 18th-Century Historical Texts

Anonymous Submission

Abstract

Named Entity Recognition (NER) is increasingly applied to historical text analysis. However, differences in evaluation materials, metrics, and annotation guidelines across existing NER projects make it difficult to systematically compare different approaches to historical NER. This study addresses this issue by constructing an evaluation corpus through the normalization of four annotated datasets from the long 18th century. We evaluate the performance of the Edinburgh Geoparser, spaCy and BERT on this corpus using five evaluation modes. Results show that even under the most lenient criteria, the highest F1-score remains below 70%, highlighting the challenges of applying existing NER systems to historical texts. Through detailed error analysis, we identify common challenges such as spelling and formatting issues. These findings demonstrate the limitations of NER tools in historical documents. We argue that future work should involve collaboration with historians to ensure that evaluation corpus align with real user needs.

Keywords: named entity recognition, evaluation corpus, historical documents, digital humanities, natural language processing

1 Introduction

With the development of mass digitization in cultural heritage, an increasing number of historical documents have become available in digital formats. In this context, Named Entity Recognition (NER) can be considered a crucial step for extracting structured and meaningful information from digitized historical texts. NER is a subtask of Natural Language Processing (NLP) that involves identifying and classifying entities of common interest from texts, such as persons, places, organizations, etc. It has been widely used in contemporary projects and serves as the basis for many text-mining applications, such as semantic annotation, question answering, ontology population and opinion mining [1]. In recent years, applications in the historical domain have become one of the main developments of NER tasks. Some studies shown that person and place names are the most frequent search elements in various digital libraries [2, 3]. Extracting these key entities by using NER not only helps historians retrieve relevant texts but also benefits downstream tasks such as entity linking and relation extraction. In other words, high-quality NER can greatly support the exploration and study of historical materials.

However, while NER is often regarded as a solved problem in some major contemporary NER evaluation forums with reported success ratios over 95% [4], these results are largely based on contemporary texts, such as newspaper articles, journal articles, blog posts, and other sources [1]. NER applied to historical texts can face more challenges than applied to contemporary well-edited journalism materials. For instance, historical documents are often heterogeneous and may contain OCR errors and spelling variations that have evolved over time, all of which can influence the performance of NER [5]. Moreover, historians' requirements and expectations for named entities and NER may differ from those defined in mainstream NER tasks, and this is rarely discussed in existing research. Given these challenges and continuous advancements in NER technologies, it

Paper under review.

© 2025 by the authors. Licensed under Creative Commons Attribution 4.0 International (CC BY 4.0).

is crucial to investigate the performance of existing NER systems on historical texts and explore ways to adapt them to achieve satisfactory results.

However, it is difficult to systematically compare the shortcomings and advantages of different NER approaches applied to historical materials due to the varying settings adopted in most historical NER projects [6]. More specifically, existing historical NER projects often use different definitions and annotation boundaries for named entities. The measurements they use to report the performance of NER systems are also varying. These variations make it difficult to give comparisons of NER methods across different projects. As a result, it remains difficult to make absolute claims about the performance of NER on historical documents. The incomparability of historical NER approaches makes it difficult to build upon existing work to advance the technologies. To address this issue, we propose the development of a shared evaluation historical corpus to enable fair and systematic comparison of NER systems and drive progress in historical NER, which motivates the study.

This study uses datasets from the long 18th century¹ (1688-1815) as a case study to create an evaluation corpus and establish evaluation standards for assessing the capabilities of NER applied to materials from this period. In the process, we introduce normalised annotations to evaluate the NER performance of three tools representing different approaches: the Edinburgh Geoparser², spaCy³, and a BERT-based NER system⁴, covering methods from rule-based systems to deep learning models. The creation of the corpus makes it possible to employ a set of evaluation measurements to compare the strengths and weaknesses of different NER tools on historical text from this period. Meanwhile, the experience gained from applying NER to these materials could also be transferred to the work on other historical periods.

2 Related Work

2.1 Evaluation Development

Annotated datasets are commonly used to evaluate the performance of systems and models. With the development of NER, various evaluation tasks and datasets in the NLP field have gradually emerged to assess NER performance.

In 1995, the term ‘Named Entity’ was first introduced at the 6th Message Understanding Conference (MUC-6) [7]. The task mainly focused on identifying the names of people, organizations, and geographic locations in texts drawn from the 1993 and 1994 *Wall Street Journal*. In MUC events, system outputs were compared against manually created gold standard annotations, and each extracted object was classified into categories such as **Correct** (exact match), **Partial** (partial match), **Incorrect** (does not match), **Missed** (items present in the gold standard but not identified), and **Spurious** (items incorrectly predicted by the system) [8]. The evaluation framework then calculated performance scores using Precision, Recall and F-measure accordingly. These scores have been widely adapted ever since as the standard way for measuring the NER performance.

Following MUC, system development competitions for languages other than English emerged, such as IREX (Information Retrieval and Extraction Exercise) and CoNLL (Conference on Computational Natural Language Learning), which included not only English but also Japanese and German. The CoNLL-2003 shared task focused on four types of entities: Persons, Locations, Organizations, and Miscellaneous. It included English texts from Reuters Corpus and German texts extracted from the newspaper *Frankfurter Rundschau* [9]. IREX targeted the identification of eight entity types (Persons, Locations, Organizations, Aircrafts, Date, Time, Money and Percent) in a

¹ Typically from around 1688 (the Glorious Revolution) to 1815 (the end of the Napoleonic Wars).

² <https://www.ltg.ed.ac.uk/software/geoparser/>

³ <https://spacy.io/>

⁴ Implemented using the `dslim/bert-base-NER` model (<https://huggingface.co/dslim/bert-base-NER>) from Hugging Face, which is fine-tuned on the CoNLL-2003 English dataset.

Japanese newspaper corpus [10]. In contrast to MUC, they did not include partial matches but only entities that exactly matched the corresponding entities in the data file were counted as correct and contributed to the scores [11]. The CoNLL-2003 dataset remains widely used today, with many NER tools employing it for training or testing models to demonstrate their effectiveness.

The exact match criterion used in CoNLL is not always applicable to certain domains. For example, in biomedical tasks, researchers often value useful information provided by partial matches. To address this, SemEval-2013 (Semantic Evaluation) introduced an extended evaluation regime [12], which included the following four criteria: (1) **Strict**, which requires exact matches of both entity boundaries and types; (2) **Exact**, which allows type misclassification; (3) **Partial**, allowing boundary overlaps and incorrect entity types; and (4) **Type**, which requires correct classification of entity types, regardless of boundaries.

Unlike traditional NER evaluations that rely on straightforward Precision and Recall metrics, the ACE (Automatic Content Extraction) program employed a more complex scoring system. Rather than merely measuring whether entities were correctly identified, it recognizes all mentions of entities and assigns them weighted values [13]. The ACE evaluation calculates scores using a formula in which an entity’s final score is the product of the entity’s inherent value and the sum of the weighted values of all its mentions [14].

In recent years, NER has begun to be applied to the field of history, and historical NER evaluation conferences have also emerged. The HIPE (Identifying Historical People, Places and other Entities) shared task provided an opportunity to evaluate the NER on historical newspapers in French, German, and English [15]. HIPE used both **Strict mode** where correct matches must get both entity type and boundary exactly matched and **Fuzzy mode** where partially matched boundary is allowed.

2.2 System Comparability

With the rise of historical NER, it is important to understand the capabilities of NER systems applied to historical materials. However, most research on historical NER has been conducted only in the context of individual projects, which makes it difficult to compare results. Challenges in system comparability arise from differences in materials, measurements, and annotation guidelines.

Historical texts tend to be highly heterogeneous, including various document types and languages. This heterogeneity leads to great variations between collections, while domain shifts inevitably affect NER performance, involving factors such as transcription quality, multilingualism, and formatting. With these variables, NER methods applied to different materials are difficult to compare. For example, Rodriquez et al. compared the performance of four NER systems on Holocaust survivor testimonies and letters of soldiers [16]. Stanford CRF performed best on the testimonies with a 54% F-score, whereas OpenCalais was the most accurate for the letters with a 36% score. This demonstrates not only performance loss across document types, but also the incomparability of NER applied to different materials.

Although most NER projects adopt standard metrics such as Precision, Recall and F-score, the modes of which entities contribute to the final scores vary. For instance, in CoNLL, only exact matches are considered correct, whereas MUC also includes partial matches in its scoring. As a result, it is difficult to compare NER systems across projects when the evaluation modes differ, and unified measurements are required to enable meaningful comparisons.

In different historical corpora, there is no standardised guideline to aid annotators in making consistent decisions about challenging interpretations that affect the scope and boundaries of annotations. For example, should honorifics like ‘King’ and ‘Queen’ be considered person names? How should entities such as ‘London Bridge’ be annotated? Should ‘Queen’s garden’ be tagged as a person or a place? The different choices make NER faces challenges of different degrees. Determining entity span is another difficulty. For instance, should honorifics following person names

such as ‘Sir Hans Sloane’ be included in person annotation? Identifying the correct span for place names can be even more complex, especially for compound place names, such as ‘cliffs of a high mountain of slego in Ireland’. How to identify references of individual objects and the boundaries between objects and their descriptions consistently become a challenge for researchers.

Ortolja-Baird et al. discussed this issue for encoding Sloane’s data [17]. They consulted several domain experts, and one view suggested that the head noun of the phrase is the best candidate for identifying the object (e.g., ‘corall’ in ‘Red corall growing on a rock with shells’), while another argued that descriptive details like color and size are meaningful and should be included (e.g., ‘Red corall’). Ortolja-Baird noted these interpretive differences depend on data use. Although debates on defining entity boundaries remain, there is a consensus that emphasizes considering context and historical sensitivity. Overall, annotation guidelines vary with background and requirements. Consequently, NER systems assessed in these projects are challenged distinct criteria, which makes it difficult to achieve accurate and systematic comparisons of NER approaches.

3 Method

To address the comparability challenges of historical NER, this paper develops a shared evaluation corpus based on four annotated datasets, including Mary Hamilton Papers, Old Bailey Online, Sloane’s Catalogue and HIPE2020. These historically significant and heterogeneous datasets help reveal the problems faced by NER systems when applied to long 18th-century documents. Moreover, they are already annotated with person and place names using TEI or IOB tags, which helps reduce the cost of corpus construction.

3.1 Datasets

Mary Hamilton (1756-1816), one of the most well-connected women in 18th-century British high society, maintained good relationships with royalty and aristocracy [18]. Her letters provide good opportunities to explore the cultural and social life of the time. Approximately 1,600 digitized letters are available online in XML with TEI tags. After excluding those without relevant annotations, 1,589 were selected for evaluation.

Sir Hans Sloane (1660–1753) was an early modern physician and naturalist, and he left a vast collection that became foundational to British cultural heritage [19]. The Enlightenment Architectures project digitized and TEI-encoded several of Sloane’s manuscript catalogue, and Volume I and Volume V were chosen for evaluation.

Old Bailey Online is a searchable archive of criminal trials held at the Old Bailey criminal court from 1674 to 1913 [20]. Given the corpus size and annotation density, 500 files from trials between 1674 and 1747 were selected.

HIPE2020 is a shared task aimed at evaluating NER performance on historical newspapers [15]. It includes newspapers from 1790 to 1960. According to the research scope, 19 newspapers dated between 1790 and 1840 were selected.

3.2 Annotations

Since all selected data are already annotated for person and place names, manual annotation is unnecessary. However, as discussed previously, the lack of uniform annotation guidelines means annotators from different projects may interpret person and place names inconsistently. For instance, in the Old Bailey dataset, certain nominal mentions such as ‘Prisoner’ and ‘two young men’ are annotated as person entities, whereas other datasets tend to exclude such expressions. This difference stems from variations in text types and research aims. As a record of court proceedings, Old Bailey Online focuses on individuals relevant to trials, including victims, witnesses,

and defendants. In this context, nominal mentions still provide key information about real persons.

The identification of place names is even more complex due to varying document types and subject matter. Mary Hamilton’s letters and HIPE2020 primarily focus on locations with geopolitical borders, such as cities, countries, and Parishes. Specifically, the Mary Hamilton Papers include English locations ranging from towns like ‘Taxal’ and ‘Sandleford’ to broader regions such as ‘Richmond’ and ‘Northamptonshire’, while HIPE2020 favours modern political entities. In contrast, Sloane Catalogue contains place names referring to natural features such as forests, mountains, and rivers. These places tend to be more detailed, e.g., ‘chalk pitts near Greenwich in Kent.’ Old Bailey place names mostly refer to urban buildings and specific sites, such as streets, pubs, and landmarks like ‘White-cross-street’ and ‘High-gate’.

3.3 Normalization

As discussed, the selected historical datasets originate from different projects, each with its own annotation priorities and focus. As a result, the scope and span of annotated entities vary. For example, different corpora give different answers on whether ‘the Duke of Buckingham’ should be tagged as a person name, or whether titles such as ‘Mr.’ in ‘Mr. Bradbourn’ should be included. Our goal is to normalize these annotations to ensure consistency across the evaluation corpus. This process involves standardizing the span and scope of existing annotations by establishing a set of design principles that ensure a unified standard. By doing so, we aim to build a comprehensive and heterogeneous evaluation corpus that serves as a fair benchmark for testing and comparing different NER systems.

Based on the instances of entities in the historical corpora, the scope and span of person and place entities across the different projects are classified and defined in the Table 1 and Table 2:

Person	
Classification	Definition
Names	Proper names of individuals e.g., <i>Ann Petty</i>
Honorifics	Words or expressions implying or expressing high status, politeness, or respect. Including common titles (Mr., Ms. Etc.), professional titles (President, Dr. etc.), nobility titles (King, Duke etc.), religious titles (Bishop, Minister etc.) and military titles (Captain etc.).
Names with honorifics	Honorifics followed by person names e.g., <i>Dr. Campbell</i>
Conjoined entities	Entities that are combined or linked together, typically with a conjunction such as ‘and’, ‘or’ or ‘,’. e.g., <i>Ann Slow alias Ebram, James and William Greffis</i>
Nominal mentions	The use of a noun or noun phrase to mention a person entity without using their proper name. Generally, occupation, status and demonym are used to be the nominal mentions. e.g., <i>Cittizen of London, two young men, Prisoner</i>

Table 1: Span and Scope of Person Entities

The normalization should ensure that person and place annotations in the evaluation corpus, derived from multiple corpora, follow a consistent scope. The normalization rules should be based on the following principles:

- Annotations should have a clear, consistent, and generalizable scope that applies across the included corpora and remains valid if new datasets are added in the future.
- The evaluation corpus should provide a level playing field for assessing various NER systems, helping to reveal their limitations when applied to long 18th-century texts.

To ensure the rules are both universal and effective in exposing the limitations of historical NER, we reviewed annotation guidelines from mainstream NER datasets to get their annotation selections regarding the span and scope of person and place entities. By comparing selections across our historical datasets and four widely used benchmark corpora (MUC, ACE, CoNLL-2003, and OntoNotes), we applied a majority voting approach to derive normalization rules and adjust the

Place	
Classification	Definition
Administrative Locations	Geographic areas that are defined and managed by a government or administrative body e.g., convulsed all <i>Europe</i> , from <i>Virginia</i> , in <i>Kent</i>
Natural Features	Naturally occurring physical formations or areas on the earth’s surface e.g., famed <i>Mount Vesuvius</i> , <i>River Thames</i>
Facilities	Functional, primarily man-made structures, including buildings and similar facilities designed for human habitation e.g., apartments at <i>Louvre</i>
Thoroughfares	Names of streets, squares, roads, highways, addresses, etc. e.g., <i>Oxford Road</i>
Possessive and Descriptive	Constructs that show ownership, possession, or a relationship between things or proper noun is used to describe an extensive entity. e.g., <i>Fairford church</i> , <i>London’s park</i> In this example, ‘Fairford church’ is not a proper noun, while ‘Fariford’ is the descriptive noun of the church (a church in Fairford).
Nested and Embedded Entities	Location entities are embedded in other person or location named entities. e.g., Bay of <i>Naples</i> , duke of <i>Buckingham</i> In this example, ‘Naples’ is a location, and ‘Naples’ is nested in another location entity ‘Bay of Naples’
Locative Designators	Terms used to provide specific information about the geographical context of a place. e.g., <i>River Thames</i> , <i>Cripplegate Parish</i> , <i>Parish of Paddington</i> <i>Parish of..</i> , <i>District of..</i> , <i>City of..</i> are categorized in locative designators rather than locative specifiers (e.g. on the coast of California), because they specifies the nature of the entities as being a parish, a district and a city.
Locative Specifiers	Terms used to provide additional details that specify or clarify the geographical position of a place. e.g., <i>the capital of Corsica</i> , <i>coast of Spain</i>
Locative Modifiers	Terms that add context or specify features of a place name. e.g., <i>North London</i>
Nominal mentions	The use of a noun or noun phrase to mention a location entity without using their proper name. e.g., <i>the Turnpike road</i> , <i>the Park</i>
Compound toponyms	Structures that combine multiple geographical units e.g., <i>Globe Tavern in Fleet-street</i> , <i>Sun Tavern behind the Royal Exchange</i>

Table 2: Span and Scope of Place Entities

existing annotations accordingly. That is, for each category in the Table 1 and Table 2, we adopt the annotation choice used by the majority of the corpora, ensuring resulting rules are generalizable. The detailed majority voting results are provided in Table 4 and Table 5 in Appendix A.

According to the review, all NER corpora classify names, including surnames, given names, and middle names, as person entities, while most exclude conjoined entities and nominal mentions. Among the generic datasets, only ACE and OntoNotes include honorifics, although OntoNotes limits annotation to noble and royal titles, excluding common titles such as ‘Mr.’ or ‘Dr.’ In contrast, honorifics appear more prominently in historical corpora. In historical Europe, they were frequently used to refer to nobility or royalty, often without mentioning specific personal names. For example, ‘Duke of Buckingham’ in the Old Bailey refers to John Smith, and ‘the Queen’ in Mary Hamilton’s letters refers to Queen Charlotte. Excluding such terms could result in losing significant information about these individuals. Therefore, honorifics are included in the normalized evaluation corpus to evaluate NER performance on this feature.

For the majority voting results of place entities, most corpora agree that administrative locations, natural features, and thoroughfares are subtypes of place entities. Regarding facilities, only CoNLL explicitly classifies them as places among the generic corpora. However, facilities are included in the normalized evaluation corpus for two reasons: (1) all historical corpora treat them as places; and (2) although ACE and OntoNotes do not classify them as locations, they define them as a separate entity type ‘Facility,’ which indicates their relevance and importance. For annotation boundaries of place entities, both generic and historical corpora give lower-than-average votes to categories such as ‘Possessive and Descriptive’, ‘Nested and Embedded Entities’, ‘Locative Mod-

ifiers’, and ‘Nominal Mentions’. Therefore, these will be excluded or revised in the normalized corpus. In cases of compound toponyms, historical corpora show partial agreement. Considering their treatment in generic corpora, we decided to normalize such cases into separate location units and include them. Finally, due to ongoing disagreement about locative designators and specifiers in historical corpora, we refer to generic guidelines and choose to retain locative designators while excluding specifiers.

In summary, based on the majority voting results discussed above, existing annotations were normalized to ensure consistency across the evaluation corpus. The normalized evaluation corpus supports a fair comparison of historical NER system performance. For detailed implementation of the normalization and data preprocessing procedures, please refer to the GitHub repository⁵ linked in Appendix B.

4 NER Performance Evaluation

4.1 Tools Selection

Using the normalized corpus, we evaluated the performance of three out-of-the-box NER tools in recognizing persons and places in 18th-century historical texts. We selected Edinburgh Geoparser, spaCy, and BERT as they are representative of different NER approaches, and evaluating BERT also allows us to assess the state-of-the-art advancements in large language models on historical NER tasks.

4.2 Measuring Performance

As discussed in Section 2.1, the performance of typical NER tasks is usually evaluated in terms of Precision, Recall and F1 measures. However, the definition of what counts a correct entity can vary across different tasks. CoNLL adopts a strict evaluation scheme where only entities with exact matches are considered correct, whereas in MUC events, partially correct entities are assigned a 50% weight and contribute to the scores.

To illustrate these differences more intuitively, consider the example of ‘Duke of Buckingham.’ When comparing the recognition results of an NER system with the annotations in the evaluation corpus, several types of outcomes may occur:

- **Complete correctness** occurs when both the entity type and boundary perfectly match the evaluation corpus, such as correctly identifying ‘Duke of Buckingham’ as a Person entity.
- **Boundary errors** arise when the entity type is correct but the boundary is incomplete or excessive, such as tagging ‘Duke’ as a Person but missing ‘of Buckingham.’
- **Type errors** involve incorrect classification regardless of boundary accuracy, for instance, misclassifying ‘Duke of Buckingham’ as a Place rather than a Person, even though ‘Buckingham’ is matched.
- **Missed entities** occur when the system fails to detect entities present in the corpus, omitting ‘Duke of Buckingham’ entirely.
- **Spurious entities** are false positives, where the system incorrectly identifies non-entities as entities, such as tagging ‘Represent’ as a Person.

These categories form the foundation for calculating Precision, Recall, and F-score in Section 4.3, providing a comprehensive assessment of system performance across different types of recognition errors.

⁵ repository link

Given these types of recognition outcomes, we designed five evaluation modes based on the MUC classification of entity results and SemEval’s multiple evaluation models to more accurately score system performance under different criteria. These five evaluation modes reflect varying levels of stringency in determining which entities are eligible to contribute to the final score.

Strict evaluation represents the most rigorous approach, where only entities with perfect type classification and exact boundary matches receive full weight, while all other cases are treated as incorrect with no credit. **Type evaluation** relaxes boundary requirements by giving partial credit (50% weight) to entities with correct classification but overlapping spans, while maintaining strict penalties to errors in type classification, as well as missed and spurious entities.

Partial evaluation introduces the most nuanced scoring system. In addition to fully correct entities and those with correct type but overlapping boundaries, entities with incorrect type but identical or overlapping spans are also considered. These categories are weighted at 100%, 50%, and 25% respectively, recognising that boundary detection remains valuable even when classification fails. When considering partial matches, it is essential to recognize the significance of instances where the NER system assigns incorrect types (misclassifications) but achieves partial or complete boundary matches. These cases still provide valuable insights into the system’s recognition capabilities. For example, in the phrase ‘Queen’s garden’, the evaluation corpus may annotate it as a place entity, whereas the NER system might recognize only ‘Queen’ as a person entity. Our goal is not to attempt to demonstrate the failures of NER systems, but rather to analyse the types of errors and their causes. To achieve this, as shown in equation below, instead of completely discarding misclassified matches, we assign a 25% weight to entities with incorrect classifications that have exact or partial boundary matches. In Equation 1, *COR* denotes the number of completely correct entities, *PAR_type* refers to entities with correct type but partially matching boundaries, and *PAR_weak* represents entities with incorrect type but overlapping or matching boundaries.

$$\begin{aligned}
Precision &= \frac{COR + 0.5 \times PAR_type + 0.25 \times PAR_weak}{related\ annotations\ produced\ by\ NER\ system} \\
Recall &= \frac{COR + 0.5 \times PAR_type + 0.25 \times PAR_weak}{related\ entities\ in\ evaluation\ corpus} \\
F_1 &= \frac{2 \times Precision \times Recall}{Precision + Recall}
\end{aligned} \tag{1}$$

Lenient evaluation mirrors Type mode but awards full weight to entities with both correct type and either exact or overlapping boundaries, emphasising boundary detection over perfect precision. **Ultra-lenient evaluation** builds on the Partial mode framework but assigns full 100% weight to all three categories, providing maximum credit for any form of entity recognition regardless of classification accuracy.

4.3 Results and Discussions

We evaluated the performance of three NER tools: Edinburgh Geoparser, spaCy and BERT by using the normalized evaluation corpus. Table 3 presents the evaluation results of these NER tools across the five evaluation modes. For more details about the testing code, please refer to Appendix B.

The results show that even in the ultra lenient evaluation mode, the highest score remains below 70%, which indicates that NER faces significant challenges when applied to historical documents. In addition to the inherent limitations of the approaches, features specific to historical documents, such as multilingual content, noise introduced during digitisation, and language variation over time, making it difficult for NER systems trained on contemporary data to perform effectively [5]. Overall, the BERT-based model demonstrates the best performance across all evaluation modes, which

	Strict	Type	Partial	Lenient	Ultra-Lenient
Edinburgh Geoparser					
Person	0.502	0.577	0.584	0.652	0.680
Place	0.166	0.317	0.340	0.410	0.504
Macro-F1	0.334	0.447	0.462	0.531	0.592
Micro-F1	0.454	0.532	0.541	0.610	0.650
Spacy					
Person	0.440	0.482	0.483	0.524	0.530
Place	0.086	0.188	0.209	0.254	0.338
Macro-F1	0.263	0.335	0.346	0.389	0.434
Micro-F1	0.386	0.432	0.437	0.478	0.498
BERT					
Person	0.513	0.622	0.626	0.731	0.748
Place	0.208	0.405	0.408	0.497	0.508
Macro-F1	0.360	0.513	0.517	0.614	0.628
Micro-F1	0.470	0.575	0.580	0.681	0.696

Table 3: Scores of NER systems

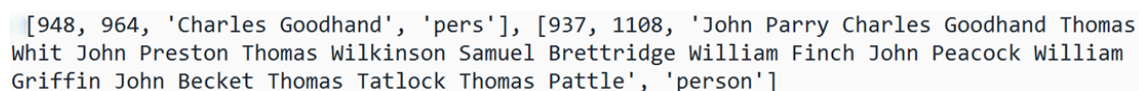
can be attributed to its context-aware architecture and large-scale pretraining, which provides robustness against the complex historical texts. The Edinburgh Geoparser performs moderately well. Its reliance on rule-based methods and gazetteer limit its general applicability to more diverse or ambiguous contexts. In contrast, spaCy shows the lowest scores. This can be attributed to its heavy reliance on models trained on contemporary, clean data, which makes it difficult to adapt to historical documents. Consistently, across all modes, all NER tools perform better on person entities compared to place entities. Person names typically follow more recognisable patterns, and the presence of titles and honorifics aids identification. As shown in Table 1 and Table 2, the scope and span of place entities are more complex and diverse, which makes NER tools harder to identify consistently. Moreover, compared to person names, the identification of place names is easier affected by OCR errors.

Next, we analysed the types and causes of errors made by spaCy, BERT, and the Edinburgh Geoparser on the historical evaluation corpus. Four broad categories of errors were identified. **Error 1** represents missed entities where systems fail to detect entities present in the gold annotations. These are commonly caused by formatting issues such as extra spaces, transcription errors, abbreviations, the use of hyphens, and the omission of honorifics or small-scale locations. **Error 2** involves spurious entities where non-entities are incorrectly identified as entities. This is mainly due to formatting issues and case sensitivity in spaCy and the Geoparser, while BERT specifically struggles with sub-word splitting errors. **Error 3** is about type misclassification where entities are detected but assigned incorrect categories, commonly due to nested entities, ambiguity and formatting issues. **Error 4** covers boundary mismatches where entity types are correct but spans are incomplete or excessive, typically resulting from entity merging errors, definite articles, nested entities, and various formatting issues.

As frequently noted above, spelling and formatting issues are the most common sources of errors. spaCy and Geoparser struggle to identify entities correctly when their forms are disrupted by elements such as extra spaces introduced during transcription, hyphens, or apostrophes, which break the expected entity patterns. In comparison, BERT is less affected by entities with formatting and spelling issues, although hyphens and abbreviations still impact its performance. Errors introduced during transcription also present a significant challenge for NER systems. In HIPE dataset, for instance, many occurrences of ‘s’ were transcribed as ‘f’ incorrectly, which disrupts

the word spelling and reduces the accuracy of recognition. Additionally, the presence of historical orthographic features, such as the long s, also disrupt recognition.

Furthermore, the recognition capabilities of NER tools are highly dependent on capitalized initials, especially in the case of Geoparser. While this may work well in well-edited modern texts, it introduces more errors when applied to complex historical texts. For example, Geoparser frequently misclassifies capitalized non-entities as entities, or fails to recognize entities that appear in lowercase. Figure 1 illustrates an extreme example where the Geoparser incorrectly merged all adjacent capitalized words into a single entity. The figure consists of two elements: the left one represents the actual annotation in the evaluation corpus, while the right one displays the result identified by the NER system.

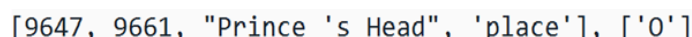


```
[948, 964, 'Charles Goodhand', 'pers'], [937, 1108, 'John Parry Charles Goodhand Thomas Whit John Preston Thomas Wilkinson Samuel Brettridge William Finch John Peacock William Griffin John Becket Thomas Tatlock Thomas Pattle', 'person']
```

Figure 1: Over-merging Capitalized Words

For person names, a significant source of score loss stems from the inconsistent recognition of honorifics. Both BERT and spaCy rarely identify honorifics as part of person entities. The Edinburgh Geoparser holds a slight advantage in this regard, as it can recognize some common honorifics, such as ‘Mr.’ This limitation also affects nested entity recognition. For example, in cases where a place name is nested within a person entity, such as ‘Prince of Wales,’ spaCy only recognizes ‘Wales’ as a place entity. Another instance is the place entity ‘Davis’s Straits,’ where spaCy incorrectly tags ‘Davis’ as a person entity.

For place names, the evaluation corpus includes many small-scale locations extracted from the Old Bailey dataset, such as streets, pubs, and churches. Both spaCy and the Edinburgh Geoparser perform poorly on these types of entities, while BERT demonstrates stronger recognition capabilities. However, BERT’s performance declines when processing entities that lack explicit place designating terms. For example, as shown in Figure 2, ‘Prince’s Head’ is the name of a pub and does not contain common locative indicators, BERT fails to correctly classify it as a place entity. In this figure, the ‘O’ on the right side means the entity was not recognized by the NER system.



```
[9647, 9661, "Prince 's Head", 'place'], ['O']
```

Figure 2: Entities without Explicit Indicators

Additionally, for BERT, most errors are concentrated in the splitting of sub-words. Due to its tokenization mechanism, BERT may split an entity into multiple sub-word tokens, and some sub-words fail to merge after prediction, which results in a large number of meaningless sub-word tokens.

Besides the aforementioned errors, NER tools also make errors in identifying the boundaries of entities with definite articles, locative specifiers and locative designators. Although this issue appears less severe—since they often capture the core element of the entity—it still impacts overall accuracy. For example, historians may prefer recognizing the full phrase ‘North of France’ rather than just ‘France.’

Overall, the errors made by the three NER tools are concentrated in several specific areas, particularly issues related to entity formatting and spelling, including transcription errors and inherent spelling variations in the texts. The performance of spaCy and Geoparser in recognizing small-scale places is poor, whereas BERT performs significantly better. However, some place entities are still missed. Additionally, there are also issues such as the failure to recognize honorifics, nested entities and locative designators. These findings reveal the limitations of NER when applied to historical texts and provide valuable insights to guide the future development of historical NER.

5 Conclusion and Future Work

In this study, we address the issue of incomparability of different NER systems on historical documents by developing a shared evaluation corpus. This corpus was constructed from four historically significant datasets with annotations normalized to ensure consistency across sources. Using this corpus, we evaluated the performance of three representative NER tools in recognizing person and place entities. The testing results provide an initial understanding of the limitations of existing NER tools and reveal the challenges inherent in applying NER to historical texts.

However, this study also has some limitations. Firstly, we only focus on the recognition of person and place names, whereas historical research may involve more entity types. In addition, the annotations in the corpus were normalized through majority voting across historical and generic datasets; whether these annotation choices truly align with historians' expectations remains uncertain. The normalization rules should be adjusted to better reflect historians' actual research needs.

In future work, we plan to communicate with historians through interviews and questionnaires to better understand their requirements. Based on the results, we aim to develop a gold standard corpus that aligns with their needs and expectations, which can be used to comprehensively evaluate the performance of NER applied to historical documents.

Acknowledgements

References

- [1] Marrero, Mónica et al. "Named Entity Recognition: Fallacies, challenges and opportunities." In: *Computer Standards Interfaces*. Vol. 35. 2012, pp. 482–489.
- [2] Chiron, Guillaume et al. "Impact of OCR Errors on the Use of Digital Libraries: Towards a Better Access to Information." In: *Computer Standards Interfaces*. 2017, pp. 1–4.
- [3] Bates, Marcia J. "The Getty End-User Online Searching Project in the Humanities: Report No. 6: Overview and Conclusions." In: *College Research Libraries*. Vol. 57. 1996, pp. 514–523.
- [4] Cunningham, Hamish. "Information extraction, automatic." In: *Encyclopedia of Language Linguistics*, ed. by Keith Brown. 2006, pp. 665–677.
- [5] Ehrmann, Maud et al. "Named Entity Recognition and Classification in Historical Documents: A Survey." In: *ACM Computing Surveys*. Vol. 56. 2023, pp. 1–47.
- [6] Humbel, Marco et al. "Named-entity recognition for early modern textual documents: a review of capabilities and challenges with strategies for the future." In: *Journal of Documentation*. Vol. 77. 2021, pp. 1223–1247.
- [7] Grishman, Ralph and Sundheim, Beth. "Design of the MUC-6 evaluation." In: *MUC6 '95: Proceedings of the 6th conference on Message understanding*. 1995, pp. 1–11.
- [8] Chinchor, Nancy and Sundheim, Beth. "MUC-5 Evaluation Metrics." In: *Proceedings of the 5th conference on Message understanding*. 1993, pp. 69–78.
- [9] Sang, Erik F. Tjong Kim and Meulder, Fien De. "Introduction to the CoNLL-2003 Shared Task: Language-Independent Named Entity Recognition." In: *Proceedings of the Seventh Conference on Natural Language Learning at HLT-NAACL 2003*. 2003, pp. 142–147.
- [10] "IREX. IREX Named Entity Task Homepage." [Online; accessed 13-Jul-2025]. URL: <https://nlp.cs.nyu.edu/irex/index-e.html>.
- [11] Nadeau, David and Sekine, Satoshi. "A survey of named entity recognition and classification." In: *Linguistic Investigations. International Journal of Linguistics and Language Resources*. Vol. 30. 2007, pp. 3–26.

- [12] Segura-Bedmar, Isabel, Martínez, Paloma, and Herrero-Zazo, María. “SemEval-2013 Task 9 : Extraction of Drug-Drug Interactions from Biomedical Texts (DDIExtraction 2013).” In: *Second Joint Conference on Lexical and Computational Semantics (*SEM)*. Vol. 2. 2013, pp. 341–350.
- [13] Doddington, George et al. “The Automatic Content Extraction (ACE) Program Tasks, Data, and Evaluation.” In: *Proceedings of the Fourth International Conference on Language Resources and Evaluation (LREC’04)*. 2004, pp. 837–840.
- [14] “Automatic Content Extraction 2008 Evaluation Plan (ACE08).” In: *Proceedings of the International Conference on Language Resources and Evaluation*. 2008.
- [15] Ehrmann, Maud et al. “Introducing the HIPE 2022 Shared Task: Named Entity Recognition and Linking in Multilingual Historical Documents.” In: *Advances in Information Retrieval: 44th European Conference on IR Research*. 2022, pp. 347–354.
- [16] Rodriguez, Kepa J. et al. “Comparison of Named Entity Recognition Tools for Raw OCR Text.” In: *Proceedings of KONVENS 2012 (LThist 2012 workshop)*. 2012, pp. 410–414.
- [17] Ortolja-Baird, Alexandra et al. “Digital Humanities in the Memory Institution: The Challenges of Encoding Sir Hans Sloane’s Early Modern Catalogues of His Collections.” In: *Open Library of Humanities*. Vol. 5. 2019.
- [18] Barker, Hannah et al. “Unlocking the Mary Hamilton Papers.” [Online; accessed 13-Jul-2025]. 2023. URL: <https://www.maryhamiltonpapers.alc.manchester.ac.uk/>.
- [19] Nyhan, Julianne et al. “Welcome to the Sloane Lab - Sloane Lab.” [Online; accessed 13-Jul-2025]. 2023. URL: <https://sloanelab.org/>.
- [20] Hitchcock, Tim et al. “The Old Bailey Proceedings Online.” [Online; accessed 13-Jul-2025]. 2023. URL: www.oldbaileyonline.org.

A Majority Voting Results

Table 4 summarizes the annotation choices for person entities in both generic and historical NER datasets respectively. A ‘Yes’ indicates that the dataset includes this type of person entity, whereas a ‘No’ indicates it does not. ‘Partial yes’ indicates that the entity type exists in the corpus, but only some of its occurrences are annotated, revealing inconsistencies in the labelling process.

	Generic Corpora				Historical Corpora			
	MUC7	ACE	CoNLL-2003	OntoNotes	Old Bailey	Sloane’s Catalogue	Mary Hamilton	HIPE2020
Names	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes
Honorifics	No	Yes	No	Partial yes	Yes	Yes	Yes	No
Honorifics with names	No	No	No	Partial yes	Yes	Yes	Yes	Yes
Conjoined Entities	No	No	No	No	Yes	No	No	No
Nominal Mentions	No	Yes	No	No	Yes	No	No	No

Table 4: Person Annotations in Generic vs. Historical Corpora

Table 5 presents the annotation choices regarding place entities in both generic and historical corpora. Empty cells indicate cases where information was not available in the official documentation.

	Generic Corpora				Historical Corpora			
	MUC7	ACE	CoNLL-2003	OntoNotes	Old Bailey	Sloane's Catalogue	Mary Hamilton	HIPE2020
Administrative Locations	Yes	No	Yes	Yes	Yes	Yes	Yes	Yes
Natural Features	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes
Facilities	No	No	Yes	No	Yes	Yes	Yes	Yes
Thoroughfares	Yes	Yes	Yes	No	Yes	Yes	No	Yes
Possessive and Descriptive	No		No	No	Yes	No	No	No
Nested and Embedded Entities	No	Yes	No	No	No	No	NO	No
Locative Designators	Yes	Yes	Yes	Yes	Partial yes	Yes	Partial yes	Partial yes
Locative Specifiers	No	Yes	No	No	No	Yes	Partial yes	No
Locative Modifier	No	Yes	No	No	No		No	No
Nominal mentions	No	Yes	No	No	No	No	No	Yes
Compound toponyms	No	Yes	No	No	Yes	Yes	No	No

Table 5: Place Annotations in Generic vs. Historical Corpora

B Code and Resources

All code used for preprocessing, normalization, and evaluation is available at: [repository link](#)