

Relation Extraction or Pattern Matching? Unravelling the Generalisation Limits of Language Models for Biographical RE

Anonymous ACL submission

Abstract

Analysing the generalisation capabilities of relation extraction (RE) models is crucial for assessing whether they learn robust relational patterns or rely on spurious correlations. Our cross-dataset experiments find that RE models struggle with unseen data, even within similar domains. Notably, higher intra-dataset performance does not indicate better transferability, instead often signaling overfitting to dataset-specific artefacts. Our results also show that data quality, rather than lexical similarity, is key to robust transfer, and the choice of optimal adaptation strategy depends on the quality of data available: while fine-tuning yields the best cross-dataset performance with high-quality data, few-shot in-context learning (ICL) is more effective with noisier data. However, even in these cases, zero-shot baselines occasionally outperform all cross-dataset results. Structural issues in RE benchmarks, such as single-relation per sample constraints and non-standardised negative class definitions, further hinder model transferability.¹

1 Introduction

Relation extraction (RE) is the core information extraction task of identifying the semantic relationship between entities in text. Traditional RE evaluations rely predominantly on in-distribution testing, but this approach often overestimates true model performance by implicitly assuming that individual datasets wholly represent the underlying task (Linzen, 2020; Kovatchev and Lease, 2024). While model generalisation has gained increasing attention in NLP, RE remains relatively unexplored in this context (§ 2).

However, understanding RE generalisation to out-of-distribution (OOD) data is crucial both for the task itself as well as for the robust application of RE systems in downstream tasks like question

answering and knowledge-base population (Bassigana and Plank, 2022a). Given the popularity of representing internal language model (LM) knowledge as relational triples (Geva et al., 2023; Hernandez et al., 2024), building robust RE systems beyond the mere memorisation of dataset-specific patterns may also be key to more interpretable and trustworthy models.

This paper systematically analyses how well RE systems generalise across datasets. Due to the limited relation overlap in popular RE datasets, we focus our experiments on biographical relations, which are pervasive in RE settings; this also allows us to include a domain-specific dataset for grounded analysis (§ 3). While prior work explored various robustness tests, including domain shift and adversarial attacks (Chen et al., 2023; Meng et al., 2024), to the best of our knowledge, this work provides the first thorough analysis of cross-dataset generalisation capabilities across sentence-level RE benchmarks.

Through our cross-dataset experiments, this paper makes the following contributions:

- We document key challenges in analysing RE generalisation, including inconsistent relation schemas and highly imbalanced class distributions (§ 3), as well as propose methods for overcoming these issues.
- We find that strong in-distribution RE performance often masks fundamental generalisation failures, with models that excel on intra-dataset evaluations frequently failing to transfer effectively (§ 5).
- Our cross-dataset analysis shows that data quality dictates the best RE adaptation method: while fine-tuning achieves superior cross-dataset performance on clean data, few-shot in-context learning (ICL) better handles

¹We will release the code upon publication.

078	noisy data. However, zero-shot prompting out-	and adversarially modified sets (Wu et al., 2019;	126
079	performs all cross-dataset methods in some	Gardner et al., 2020; Goel et al., 2021; Rusert et al.,	127
080	settings (§ 5.2).	2022).	128
081	• We identify structural issues in current RE	Recent studies have explored various ways to im-	129
082	benchmarks that lead to observed generali-	prove RE model robustness. Bassignana and Plank	130
083	sation errors, including single-relation con-	(2022a) introduce a cross-domain RE dataset with	131
084	straints, reliance on external knowledge, and	broad relation types, while Meng et al. (2024) and	132
085	coverage biases (§ 6).	Chen et al. (2023) evaluate state-of-the-art (SOTA)	133
086	These findings reveal that while current RE	document-level RE models on perturbed test sets.	134
087	systems achieve high in-distribution results, their	Chen et al. (2023) further reveal that even when	135
088	cross-dataset performance shows critical gaps in	models predict correctly, they often rely on spuri-	136
089	genuine relation understanding, limiting their ap-	ous correlations, highlighting their vulnerability to	137
090	applicability in real-world situations.	minor evaluation shifts. To reduce dependence on	138
091		mere pattern matching, Allen-Zhu and Li (2024)	139
092	2 Related Work	propose augmenting training data with synthetic	140
093	2.1 Approaches for RE	samples reformulated by an auxiliary model.	141
094	RE is traditionally framed as a classification task,	3 Methodology	142
095	tackled via either a pipeline approach—where sub-	We assess the robustness of RE systems by eval-	143
096	tasks like named entity recognition (NER), coref-	uating their performance on OOD data. Standard	144
097	erence resolution, and relation classification (RC)	evaluations on in-distribution test sets likely overes-	145
098	are performed sequentially—or a joint model that	timate RE performance (Linzen, 2020), as models	146
099	processes them simultaneously (Taillé et al., 2020;	can exploit spurious cues rather than learning gen-	147
100	Bassignana and Plank, 2022b; Saini et al., 2023).	uine relational patterns (Chen et al., 2023; Meng	148
101	It is further categorised into sentence- (Alt et al.,	et al., 2024; Arzt and Hanbury, 2024).	149
102	2020; Plum et al., 2022) and document-level RE	To systematically evaluate generalisation capa-	150
103	(Yao et al., 2019; Meng et al., 2024).	bilities of RE models, we conduct both intra- and	151
104	Since the introduction of BERT (Devlin et al.,	cross-dataset experiments. The intra-dataset ex-	152
105	2019), encoder-based models have dominated RE	periments act as a control, evaluating RE mod-	153
106	due to their bidirectional attention mechanism,	els on data drawn from the distribution used for	154
107	which effectively captures context for classification	model adaptation, while the cross-dataset experi-	155
108	tasks (Alt et al., 2020; Plum et al., 2022). How-	ments measure model robustness with OOD test	156
109	ever, the rise of autoregressive models has led to	sets derived from a different RE dataset.	157
110	increasing adoption of decoder-based architectures	For our experiments, we consider three sentence-	158
111	to RE (Wang et al., 2022; Sun et al., 2023; Xu et al.,	level RE datasets: TACRED-RE (Alt et al., 2020),	159
112	2023; Liu et al., 2024; Efeoglu and Paschke, 2024a).	NYT (Riedel et al., 2010), and Biographical (Plum	160
113	While encoder-decoder models have been explored	et al., 2022). While TACRED-RE and NYT are	161
114	(Huguet Cabot and Navigli, 2021; Li et al., 2023b),	general-purpose RE datasets, we focus exclusively	162
115	our experiments focus on the dominant encoder-	on biographical relations, or relations that describe	163
116	only and decoder-only architectures for RE.	aspects of an individual’s life (e.g., <i>place_of_birth</i> ,	164
117	2.2 Generalisation Capabilities of RE Models	<i>studied_at</i> , or <i>children</i>). Our focus on biographical	165
118	Recent work advocates for transparent evaluation	RE is motivated by two key factors. First, the rela-	166
119	(Neubig et al., 2019; Liu et al., 2021) and OOD	tion type overlap between TACRED-RE and NYT	167
120	testing (Linzen, 2020; Allen-Zhu and Li, 2024; Qi	is limited to two non-biographical relations (com-	168
121	et al., 2023) to assess model robustness. Com-	pared to the six overlapping biographical relations).	169
122	mon strategies include cross-dataset (Antypas and	In addition, focusing on biographical relations al-	170
123	Camacho-Collados, 2023; Jang and Frassinelli,	lows for additional cross-dataset evaluations with	171
124	2024) and cross-domain (Fu et al., 2017; Liu et al.,	the Biographical dataset, which only contains bi-	172
125	2020; Bassignana and Plank, 2022a; Calderon et al.,	ographical relations. This setup thus allows us to	173
	2024) experiments, as well as testing on perturbed	directly compare the generalisation of two popular	174
		RE datasets in a third, held-out evaluation setting.	175

3.1 Data

We now briefly describe three RE datasets used.

TACRED-RE (Alt et al., 2020) is a general-purpose RE dataset with 41 relation types and a ‘no_relation’ class.² It contains over 106k instances but is highly imbalanced, with $\sim 80\%$ labeled as ‘no_relation’. Built from English newswire and web text, it is a revisited version of TACRED (Zhang et al., 2017), where challenging samples were re-annotated by professional annotators to reduce noise from crowdsourced labels. Experimental results demonstrate improved performance on TACRED-RE compared to TACRED (see Tables 11 and 12 in the Appendix), leading to its use in our cross-dataset experiments. Figure 1 shows a TACRED-RE example. We focus on 26 biographical relations from TACRED-RE, including the ‘no_relation’ class (Table 4, Appendix).

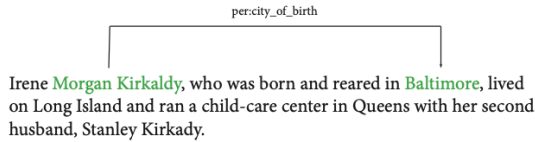


Figure 1: Example from TACRED-RE dataset. This annotation example is from Zhang et al. (2017).

NYT (Riedel et al., 2010) is a general-purpose RE dataset comprising 24 relations and a ‘None’ class. The dataset contains over 266k sentences, with 64% of the instances labeled as ‘None’ and half of positive instances containing a single dominant relation, ‘/location/location/contains’.³ The NYT dataset was constructed via distant supervision, by applying Freebase (Bollacker et al., 2008) as external supervision on the text of New York Times articles (Sandhaus, 2008). Figure 2 shows an NYT example. We focus on a subset of the NYT dataset with 7 biographical relations, including a ‘None’ class (Table 5, Appendix).

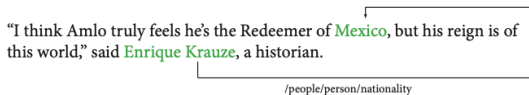


Figure 2: Example from the NYT dataset

Biographical is an RE dataset for the biographical text domain, with 10 relation types (Plum et al.,

2022). Built from Wikipedia articles on prominent individuals and containing 346,257 instances, Biographical was created using a semi-supervised approach.⁴ Named entities were automatically extracted from text using spaCy (Honnibal et al., 2020) and Stanford CoreNLP (Manning et al., 2014). Subsequently, Wikipedia sentences containing the extracted named entities were matched with Pantheon and Wikidata to automatically infer relations between the entities. Figure 3 shows an example⁵ from Biographical. The statistics for Biographical, downsampled to match the size of the TACRED-RE and NYT subsets, are presented in Table 5 in the Appendix.

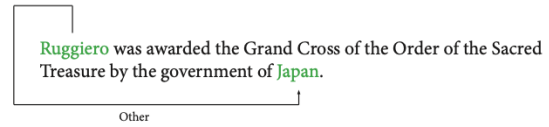


Figure 3: Example from the Biographical dataset

3.2 Cross-Dataset Comparison: Challenges

Single Relation per Sample: Both TACRED-RE and Biographical restrict each sample to at most two named entities and one relation, even when multiple relations exist within a sentence. For instance, the TACRED-RE example in Figure 1 is labeled with ‘per:city_of_birth’ but also contains relations like ‘per:stateorprovinces_of_residence’, ‘per:employee_of’, and ‘per:spouse’, all within TACRED-RE’s relation set. This constraint may potentially confuse a model trained on such data, as it enforces a single-label assignment per sentence.

While NYT better reflects real-world scenarios by allowing multiple relations per sentence, we filtered it to two-entity, single-relation samples for fair cross-dataset comparison, retaining only those like the example in Figure 2.

Unclear ‘Negative’ Class: Clear negative samples—instances with entities but no meaningful relation—are crucial for RE systems. While TACRED-RE has an explicit ‘no_relation’ class, NYT’s ‘None’ class lacks clear definition (Riedel et al., 2010), potentially confusing models about whether it indicates absence of predefined relations or any relation. Similarly, the Biographi-

²TACRED, and therefore TACRED-RE, are licensed by the Linguistic Data Consortium (LDC).

³The dataset is publicly available and accessible at <https://github.com/INK-USC/ReQuest>.

⁴It has multiple versions. We use the version without coreference resolution or first Wikipedia sentence skipping (referred to as *m2_normal_final1* in Plum et al. (2022)).

⁵Sample ID ‘mS7/1269356’ in Biographical

cal dataset lacks an explicit negative class, using instead an ‘Other’ class for unspecified relations (Plum et al., 2022), as shown in Figure 3. This inconsistency between choosing a ‘no_relation’ versus a ‘none_of_the_above’ class in RE benchmarks highlights the general challenge of consistently defining the boundaries between presence and absence of semantic relations in text (Bassignana and Plank, 2022b).

Expected Factual Knowledge: The design of RE datasets influences whether models genuinely learn RE or rely on dataset-specific cues. NYT’s distant supervision approach incorporates Freebase-derived relations not stated in text requiring external world knowledge rather than textual evidence, as shown in Figure 2 where the text lacks explicit information about Enrique Krause’s nationality—such annotations extend beyond RE’s scope and corrupt models trained on such data. Similarly, although manually curated, TACRED-RE encompasses relations like ‘per:city_of_birth’, which require factual knowledge from a model, limiting generalisation to instances seen during adaptation.

3.3 Cross-Dataset Label Overlap

To enable cross-dataset evaluation, a manual label mapping was conducted across TACRED-RE, NYT, and Biographical. Table 7 shows six overlapping biographical relations between NYT and TACRED-RE, with twelve fine-grained TACRED-RE relations mapping to six broader NYT labels (e.g., NYT’s ‘place_of_birth’ encompasses three TACRED-RE location-specific birth relations).

Treating Biographical’s ‘Other’ class as negative—supported by manual analysis of 30 random instances revealing typically negative and not unspecified relations—we find four overlapping relations across all three datasets (Table 9, Appendix). The NYT-Biographical overlap includes these same four relations (Table 6, Appendix), while TACRED-RE and Biographical share nine relations (Table 8, Appendix).

4 Experiments

4.1 Data Format and Standardisation

To facilitate cross-dataset evaluations and focus our experiments on *relation classification*, we standardise our data into a unified format providing the entity spans in each input. Thus, each example’s entities are marked with the tags <e1>head entity</e1> and <e2>tail entity</e2>, respectively.

To address class imbalance, negative instances were randomly downsampled across all three datasets to balance the number of positive and negative instances, and Biographical (~350,000 instances) was downsampled to match the size of other biographical subsets for computational efficiency (Appendix Tables 3, 4, 5). Given TACRED-RE’s fine-grained relations, we mapped them to broader NYT labels during cross-dataset evaluation (see Appendix C for details).

4.2 Model Selection, Training, and Evaluation

We consider two types of models: an encoder-only model (DeBERTa-v3-large 304M; He et al. (2021)) and a decoder-only model (an instruction-tuned LLaMA 3.1 8B model; Grattafiori et al. (2024)). We do not include SOTA systems for each dataset; SOTA approaches for the three datasets in question differ—with most SOTA systems for one dataset not evaluated on the others (Orlando et al., 2024; Efeoglu and Paschke, 2024b; Sainz et al., 2024)—and we focus our experiments on the biographical subset of relations from TACRED-RE and NYT. Thus, our results are not directly comparable to prior work on these datasets.

We employ two commonly used model adaptation strategies for RE: fine-tuning and in-context learning (ICL). Specifically, we consider direct fine-tuning with DeBERTa, fine-tuning LLaMA using low-rank adaptation (LoRA; Hu et al., 2022), and zero-shot and five-shot ICL (Brown et al., 2020) with LLaMA. For few-shot ICL, we perform five runs with different demonstration sets to account for demonstration sensitivity (Zhang et al., 2022; Webson and Pavlick, 2022; Lu et al., 2022). However, due to computational constraints, fine-tuning experiments are limited to a single run. For NYT and TACRED-RE, we conduct experiments in two adaptation settings: adaptation on all biographical relations in each dataset (Appendix Tables 4 and 5) and adaptation on only overlapping relations (Appendix Table 7). This applies to both fine-tuning and ICL, where zero-shot prompts and few-shot demonstrations are selected accordingly.

We then perform two types of evaluations: **intra-dataset**, where models are evaluated on the same dataset they were adapted to; and **cross-dataset**, where the adapted models were tested on OOD data to assess their generalisation capabilities. Further implementation details, including hyperparameter settings, full label sets, and prompting details, are provided in Appendix D.

5 Results

Overview of Reported Results: Table 1 reports intra- and cross-dataset results for NYT and TACRED-RE on six overlapping relations, using models adapted on all biographical relations. For TACRED-RE, which maps its 12 fine-grained labels to NYT’s shared label space, we report both dataset-specific and shared label results.

Table 2 shows model generalisation to Biographical across three overlapping relation sets: (1) four relations shared across all datasets, (2) same four relations shared between NYT and Biographical, and (3) nine relations shared between TACRED-RE and Biographical. Models were adapted on each dataset’s full biographical relations, with Biographical’s intra-dataset results for comparison. We focus on results with models adapted on the full overlap, as they show similar performance to those adapted only on overlap (Table 14, Appendix) while better reflecting real-world scenarios.

While we primarily focus on macro F1 to address class imbalance, additional experimental results, including per-class performance breakdowns, are provided in Appendix F.

5.1 Intra-Dataset Results

We evaluate our RE models on their training data distribution for comparison of cross-dataset generalisation. Unsurprisingly, we find that fine-tuning performs best for intra-dataset evaluations: fine-tuned LLaMA outperforms DeBERTa on TACRED-RE and NYT (Table 1), while the two models perform nearly identically on the full label overlap when fine-tuned on Biographical (Table 2).⁶

Our ICL experiments similarly show expected results, with the five-shot prompting moderately outperforming zero-shot prompting but underperforming full model fine-tuning. In the case of Biographical, this few-shot ICL gain over zero-shot is significant, increasing from 0.24 to 0.53 ± 0.05 with five demonstrations (Table 2).

We also note different performance trends within intra-dataset evaluations of TACRED-RE and NYT. Specifically, we find while fine-tuning yields higher intra-dataset performance on NYT than on TACRED-RE, this performance trend flips for the zero- and few-shot ICL settings, with prompting on NYT performing significantly *worse* than

TACRED-RE despite TACRED-RE’s finer-grained relation schema. This difference likely stems from differing data quality between datasets: the noisy labeling during NYT creation (Yaghoobzadeh et al., 2017) likely leads to over-fitting during fine-tuning (Tänzer et al., 2022) (rather than learning robust relational patterns), but harms model generalisation to NYT when not fine-tuned for that data distribution.

Comparison with SOTA: As our generalisation analysis focuses on biographical relations across datasets (and prior works rarely report per-class results), we cannot directly compare against SOTA models on the full datasets. However, our best intra-dataset results on biographical subsets (with fine-tuned LLaMA) remain close to the overall SOTA performance, trailing the reported scores by 1-2 F1 points on all three datasets (Han et al., 2022; Plum et al., 2022; Orlando et al., 2024).

5.2 Cross-Dataset Results

We now turn to examining the cross-dataset generalisation of our RE systems. Unsurprisingly, we find that performance almost always declines with cross-dataset evaluations. However, models adapted on TACRED-RE exhibit relatively strong generalisation capabilities—the few exceptions of better cross-dataset performance stemming from TACRED-RE models applied to the Biographical dataset—while those adapted on NYT struggle to transfer effectively, likely due to dataset noise.

RE Models Struggle to Generalise across Datasets Cross-dataset evaluations (almost) always perform worse than the comparable intra-dataset experiment: NYT and TACRED-RE show substantial drops of 25-27 points, while Biographical exhibits smaller decreases of 4-10 points depending on the relation setting (Tables 1 and 2). We also observe somewhat different performance trends across model and adaptation approaches from the intra-dataset experiments; while fine-tuning LLaMA with TACRED-RE achieves the best cross-dataset performance on NYT, the best TACRED-RE cross-dataset results are obtained using few-shot ICL with NYT demonstrations (rather than fine-tuning). However, these results remain below the zero-shot TACRED-RE baseline.

The cross-dataset experiments on Biographical similarly perform worse than the corresponding intra-dataset experiments in most settings (Table 2); one notable exception is LLaMA prompted

⁶This may indicate a performance ceiling due to data noise limiting further improvement (Alt et al., 2020; Arzt and Hanbury, 2024).

Model	Setting	Dataset	Intra-Dataset		Cross-Dataset	
			Shared Labels	Dataset Labels	NYT	TACRED-RE
DeBERTa-v3 large 304M	Fine-tuned on	NYT	0.67	0.67	–	0.27
		TACRED-RE	0.66	0.64	0.50	–
LLaMA 3.1 8B	Fine-tuned on	NYT	0.87	0.87	–	0.45
		TACRED-RE	0.79	0.76	0.62	–
LLaMA 3.1 8B	Zero-Shot	NYT	0.31	0.31	–	–
		TACRED-RE	0.58	0.37	–	–
LLaMA 3.1 8B	5-Shot	NYT	0.45 ± 0.07	0.45 ± 0.07	–	0.52 ± 0.06
		TACRED-RE	0.63 ± 0.06	0.43 ± 0.07	0.39 ± 0.02	–

Table 1: Macro F1-scores for intra- and cross-dataset predictions on six overlapping relations. Results show both shared and dataset-specific labels, with models adapted on all biographical relations through fine-tuning or ICL.

Model	Setting	Dataset	Full Overlap	Overlap w. NYT	Overlap w. TACRED-RE
DeBERTa-v3-large 304M	Fine-tuned on	NYT	0.47	0.47	–
		TACRED-RE	0.60	–	0.68
		Biographical	0.79	0.79	0.71
LLaMA 3.1 8B	Fine-tuned on	NYT	0.30	0.30	–
		TACRED-RE	0.69	–	0.70
		Biographical	0.79	0.79	0.74
LLaMA 3.1 8B	Zero-Shot	Biographical	0.24	0.24	0.35
LLaMA 3.1 8B	5-Shot	NYT	0.48 ± 0.04	0.48 ± 0.04	–
		TACRED-RE	0.51 ± 0.04	–	0.58 ± 0.02
		Biographical	0.53 ± 0.05	0.53 ± 0.05	0.54 ± 0.03

Table 2: Evaluation on Biographical Dataset (macro F1-scores). Models adapted on all biographical relations through fine-tuning or ICL.

with five TACRED-RE examples, which outperforms the intra-dataset few-shot experiments on the TACRED-RE/Biographical label overlap. The best Biographical cross-dataset results are achieved with fine-tuning LLaMA on TACRED-RE, though this still underperforms intra-dataset fine-tuning.

NYT Models Generalise Worse than TACRED-RE Models NYT-adapted models exhibit significantly poorer generalisation capabilities than those adapted on TACRED-RE (Tables 1 and 2). For example, fine-tuning LLaMA with TACRED-RE surpasses zero- and cross-dataset ICL on NYT (by ~ 30 and ~ 20 points, respectively), whereas all cross-dataset experiments transferring from NYT to TACRED-RE underperform zero-shot evaluations with no cross-dataset signal. This is also clear from the evaluations on the held-out Biographical dataset, where transferring from TACRED-RE always performs better than NYT (and occasionally outperforms the intra-dataset performance).

This performance gap is unlikely due to domain differences, as both datasets contain newspaper articles (with TACRED-RE including some NYT newspaper content without instance overlap), while Biographical covers Wikipedia articles. Rather, we attribute it to NYT’s distant supervision annotations, which introduce noise and limit model ro-

bustness. This is likely why LLaMA fine-tuned on NYT and evaluated on Biographical data (0.30) underperforms DeBERTa (0.47; Table 2)—the overparametrised LLaMA exhibits stronger overfitting to NYT noise and generalises poorly to unseen data, a phenomenon also highlighted by Liu et al. (2022) with corrupted training data.

Effect of Adaptation Strategy on Generalisation While prior work suggests that ICL often generalises more effectively to OOD data than fine-tuning (Awadalla et al., 2022; Song et al., 2023; Si et al., 2023), our results indicate this advantage depends heavily on data quality. With high-quality data like TACRED-RE, fine-tuning consistently achieves the best cross-dataset performance, surpassing few-shot ICL on both NYT and Biographical evaluations. In fact, TACRED-RE adaptations can even perform comparably to intra-dataset ones: DeBERTa (0.68) and LLaMA (0.70) fine-tuned on TACRED-RE achieve similar results to their Biographical intra-dataset performance (0.71 and 0.74) on the TACRED-RE/Biographical label overlap (Table 2).

However, when adaptation data are noisy, as with NYT, few-shot ICL becomes a more effective strategy: few-shot ICL via NYT consistently performs better than fine-tuning on NYT for both TACRED-

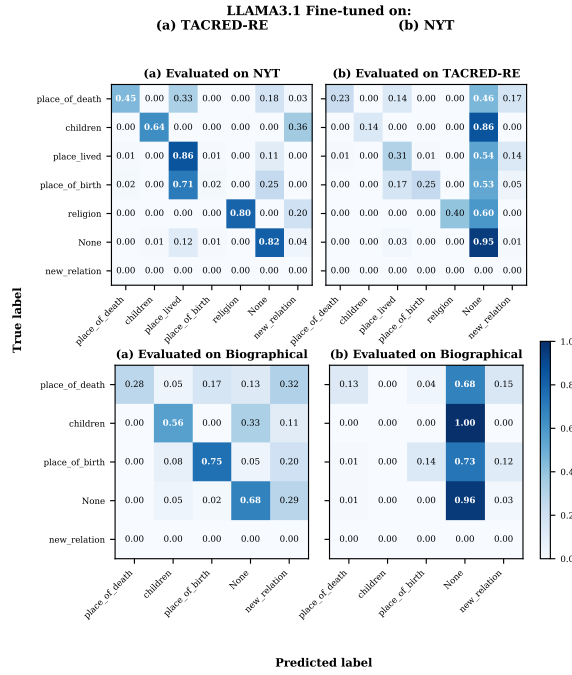


Figure 4: Confusion matrices comparing cross-dataset results for LLAMA fine-tuned on TACRED-RE/ NYT

RE and Biographical evaluations (Tables 1, 2). This is likely because ICL limits the signal from noisy training data, which in turn reduces the overfitting to dataset-specific artefacts and catastrophic forgetting compared to fine-tuning (Tran et al., 2024; Anonymous, 2024; Kotha et al., 2024).

6 Analysing RE Generalisation Failure Cases

Given RE benchmarks’ numerous relations and class imbalance, we analyse the strongest performing model, fine-tuned LLaMA, beyond aggregated metrics. We examine per-relation performance (Figure 4) and qualitatively analyse 30 random misclassifications from four evaluation settings to identify their likely underlying causes.⁷ Through these analyses, we find the following causes of RE generalisation mistakes:

Effect of Noisy Supervision on Generalisation

Confusion matrices in Figure 4 reveal that NYT-adapted models systematically overpredict the ‘None’ class on both TACRED-RE and Biographical, with manual analysis showing these misclassifications stem primarily from NYT’s *distant supervision* (Table 21, Appendix) rather than vocabulary differences (Figure 5) or domain shift. This issue is particularly evident in NYT’s location-based

relations, where reliance on external knowledge leads to conflicting annotations that hinder pattern learning—for instance, “Henryk Tomaszewski [...] died on Sunday at his home in Warsaw”⁸ is labeled as birthplace despite clear evidence of death location. Similar issues arise with Biographical’s *semi-supervised* data, where models adapted on cleaner datasets like TACRED-RE fail to replicate ground truth labels that lack textual evidence. Notably, despite higher lexical overlap between Biographical and NYT (Figure 6, Appendix), TACRED-RE-adapted models perform better on Biographical, indicating that adaptation data quality matters more than lexical similarity.

Single Relation Constraint & Negative Class

The issue extends beyond noisy supervision to fundamental constraints in RE task design. Models adapted on biographical relations often detect valid but unlabeled relations (marked as ‘new_relation’ in Figure 4), highlighting limitations of enforcing single-relation per sample. This manifests in cases like “Gross, who is himself Jewish [...] was sent to Cuba”⁹, where only ‘place_lived’ is labeled while ‘religion’ is omitted due to TACRED-RE’s constraint. Also, unclear or missing negative class affects cross-dataset evaluation: while Biographical’s ‘Other’ class is intended to cover undefined relations, our analysis reveals it contains both instances without meaningful relations and those with valid but undefined relations. This explains the high frequency of ‘new_relation’ predictions for Biographical when using LLaMA fine-tuned on TACRED-RE with their finer-grained relation schema (Figure 4) and also highlights the fundamental difficulty in defining boundaries between presence and absence of semantic relations. Despite this ambiguity, the class achieves strong cross-dataset performance when mapped to ‘None’ (0.78 and 0.72 for LLaMA fine-tuned on TACRED-RE and NYT), further underscoring the importance of distinguishing between ‘no_relation’ and ‘none_of_the_above’ cases in RE (Bassigiana and Plank, 2022b).

Reliance on External Knowledge even in Manually Curated Datasets Even with high-quality manual annotations, RE often requires external knowledge and complex reasoning capabilities. Our analysis reveals this challenge manifests in two key ways: through implicit relations requiring

⁸NYT instance ID: ‘/m/vinci8/data1/riedel/projects/relation/kb/nyt1/docstore/nyt-2005-2006.backup/1701917.xml.pb’
⁹TACRED-RE instance ID: ‘098f6f318be29eddb182’

⁷More misclassification patterns in Appendix Table 21.

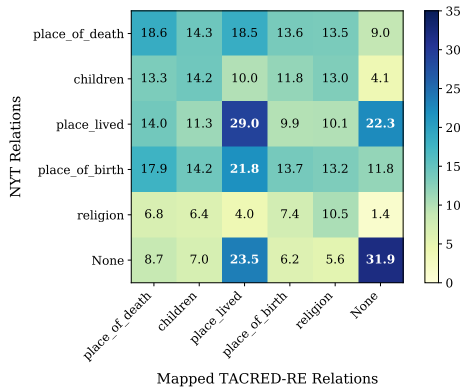


Figure 5: Vocabulary Overlap (%) per overlapping relation between NYT and TACRED-RE

inference, and through necessary world knowledge for entity interpretation. For example, in “Gross [...] was sent to Cuba as a spy”⁹, the NYT-adapted model predicts ‘None’ instead of ‘place_lived’, failing to infer that being sent somewhere as a spy implies residence. While detecting implicit relations is crucial (Geva et al., 2021), ensuring consistent and objective interpretation remains challenging.

Beyond implicit relations, models must also rely on world knowledge for basic entity understanding—as in cases like ‘Idaho businesswoman’¹⁰, where identifying entity types requires knowing Idaho as a location. TACRED-RE fine-grained relation schema further demonstrates this issue, where even with world knowledge, distinguishing between relations like city of birth and state/province of birth can be ambiguous (e.g., whether New York refers to the city or state). As Chen et al. (2023) note, even human annotators tend to rely on such prior knowledge despite the lack of rationales, motivating the need for finer-grained word evidence annotation.

Dataset Composition and Coverage Biases

Analysis of the most shared words across NYT and TACRED-RE demonstrates a strong US-centric coverage bias, likely limiting generalisation to non-US contexts (Appendix Table 22). NYT also exhibits topical skews in specific relations, such as religion being predominantly associated with Islam, potentially leading models to learn narrow, biased representations of relations.

Analysing part-of-speech distributions also reveals distinct patterns across all three datasets (Appendix Table 20). While proper nouns dominate head and tail entities in all datasets (reaching nearly 100% in NYT), TACRED-RE shows more linguis-

tic diversity with 17% of head entities as pronouns and 17% of tail entities as common nouns. Biographical, sourced from Wikipedia, contains a high proportion (26%) of numerical tail entities, primarily dates. These compositional differences, along with TACRED-RE’s longer, compound sentences and higher average entity distance (~12 tokens vs NYT’s ~8 tokens), most likely impact cross-dataset performance; NYT-adapted models struggle with these more complex patterns, which are absent from their training data (Appendix Table 21).

7 Conclusion

This work examines cross-dataset generalisation in language model-based RE systems in biographical settings. We find RE models struggle to generalise even within similar domains, with high intra-dataset performance often masking spurious overfitting rather than genuine learning of relational patterns. Furthermore, data quality is crucial for robust transfer (with the optimal adaptation strategy depending on data quality); fine-tuning yields the best cross-dataset performance when high-quality data is available, but few-shot ICL performs better in settings with noisy data. However, in some cases, a zero-shot baseline surpasses all cross-dataset results, further underscoring the limitations of current RE systems.

Our analysis also reveals several structural issues in all current RE benchmarks: (1) single-relation constraints that ignore other valid relations between entities in text, (2) the lack of a well-defined negative class with challenging samples (e.g., sentences containing commonly used tokens for relations like ‘born’ or ‘died’) to enforce deeper semantic understanding beyond pattern matching, and (3) limited diversity in data sources. These issues, compounded by inconsistent relation definitions and limited overlap across datasets, hinder meaningful evaluations of RE generalisation.

These findings thus highlight the need for more transparent evaluation beyond in-distribution testing and aggregated metrics, as limiting evaluation to these may not reflect genuine improvements in capturing relational patterns or account for class imbalance and the large number of relations in RE benchmarks. We see many promising directions for future work, including testing RE robustness on perturbed evaluation sets and applying interpretability methods to better understand how models infer relational knowledge.

¹⁰TACRED-RE instance ID: ‘098f6bd9fa786293e49d’

Limitations

Our cross-dataset analysis is limited to a particular set of biographical relations but reflects a broader challenge in RE evaluation where datasets, even covering the same domain, typically share a small relation overlap. We also constrain our analysis to single-relation examples: while, real-world scenarios often involve multiple relations per instance (and NYT allows multiple relations), we focused on single-relation setting for fair-cross dataset comparison, as TACRED-RE and Biographical are annotated with single relations. Similarly, we exclusively evaluate relation classification (RC) due to dataset constraints: TACRED-RE and Biographical assume a single relation triple per sentence, unlike real-world text where multiple relations can coexist. By focusing on RC with entity tags as guidance, we aim to minimise the prediction of other potential relations present in a sentence, but not between the specified entities.

The adaptation sets we used contain a large class imbalance due to the underlying distributions of the datasets, even after we perform data rebalancing. While this could be viewed as a limitation, it reflects real-world scenarios where models must adapt with limited training data (Bassignana and Plank, 2022a). Finally, we note that our intra-dataset results cannot be directly compared with reported SOTA performance, as most papers lack detailed relation-based metrics, reporting only aggregated results.

References

- Zeyuan Allen-Zhu and Yuanzhi Li. 2024. [Physics of language models: Part 3.1, knowledge storage and extraction](#). *Preprint*, arXiv:2309.14316.
- Christoph Alt, Aleksandra Gabryszak, and Leonhard Hennig. 2020. [TACRED revisited: A thorough evaluation of the TACRED relation extraction task](#). In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 1558–1569, Online. Association for Computational Linguistics.
- Anonymous. 2024. [Analyzing and reducing catastrophic forgetting in parameter efficient tuning](#). In *Submitted to ACL Rolling Review - April 2024*. Under review.
- Dimosthenis Antypas and Jose Camacho-Collados. 2023. [Robust hate speech detection in social media: A cross-dataset empirical evaluation](#). In *The 7th Workshop on Online Abuse and Harms (WOAH)*, pages 231–242, Toronto, Canada. Association for Computational Linguistics.
- Varvara Arzt and Allan Hanbury. 2024. [Beyond the numbers: Transparency in relation extraction benchmark creation and leaderboards](#). In *Proceedings of the 2nd GenBench Workshop on Generalisation (Benchmarking) in NLP*, pages 120–130, Miami, Florida, USA. Association for Computational Linguistics.
- Anas Awadalla, Mitchell Wortsman, Gabriel Ilharco, Sewon Min, Ian Magnusson, Hannaneh Hajishirzi, and Ludwig Schmidt. 2022. [Exploring the landscape of distributional robustness for question answering models](#). In *Findings of the Association for Computational Linguistics: EMNLP 2022*, pages 5971–5987, Abu Dhabi, United Arab Emirates. Association for Computational Linguistics.
- Elisa Bassignana and Barbara Plank. 2022a. [CrossRE: A cross-domain dataset for relation extraction](#). In *Findings of the Association for Computational Linguistics: EMNLP 2022*, pages 3592–3604, Abu Dhabi, United Arab Emirates. Association for Computational Linguistics.
- Elisa Bassignana and Barbara Plank. 2022b. [What do you mean by relation extraction? a survey on datasets and study on scientific relation classification](#). In *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics: Student Research Workshop*, pages 67–83, Dublin, Ireland. Association for Computational Linguistics.
- Kurt Bollacker, Colin Evans, Praveen Paritosh, Tim Sturge, and Jamie Taylor. 2008. [Freebase: a collaboratively created graph database for structuring human knowledge](#). In *Proceedings of the 2008 ACM SIGMOD International Conference on Management of Data, SIGMOD ’08*, page 1247–1250, New York, NY, USA. Association for Computing Machinery.
- Tom B. Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, Sandhini Agarwal, Ariel Herbert-Voss, Gretchen Krueger, Tom Henighan, Rewon Child, Aditya Ramesh, Daniel M. Ziegler, Jeffrey Wu, Clemens Winter, Christopher Hesse, Mark Chen, Eric Sigler, Mateusz Litwin, Scott Gray, Benjamin Chess, Jack Clark, Christopher Berner, Sam McCandlish, Alec Radford, Ilya Sutskever, and Dario Amodei. 2020. Language models are few-shot learners. In *Proceedings of the 34th International Conference on Neural Information Processing Systems, NIPS ’20*, Red Hook, NY, USA. Curran Associates Inc.
- Nitay Calderon, Naveh Porat, Eyal Ben-David, Alexander Chapanin, Zorik Gekhman, Nadav Oved, Vitaly Shalumov, and Roi Reichart. 2024. [Measuring the robustness of NLP models to domain shifts](#). In *Findings of the Association for Computational Linguistics: EMNLP 2024*, pages 126–154, Miami, Florida, USA. Association for Computational Linguistics.

766	Haotian Chen, Bingsheng Chen, and Xiangdong Zhou.	Karan Goel, Nazneen Fatema Rajani, Jesse Vig, Zachary	823
767	2023. Did the models understand documents?	Taschdjian, Mohit Bansal, and Christopher Ré. 2021.	824
768	benchmarking models for language understanding in	Robustness gym: Unifying the NLP evaluation land-	825
769	document-level relation extraction. In <i>Proceedings</i>	scape. In <i>Proceedings of the 2021 Conference of</i>	826
770	<i>of the 61st Annual Meeting of the Association for</i>	<i>the North American Chapter of the Association for</i>	827
771	<i>Computational Linguistics (Volume 1: Long Papers),</i>	<i>Computational Linguistics: Human Language Tech-</i>	828
772	pages 6418–6435, Toronto, Canada. Association for	<i>nologies: Demonstrations,</i> pages 42–55, Online. As-	829
773	Computational Linguistics.	sociation for Computational Linguistics.	830
774	Jacob Devlin, Ming-Wei Chang, Kenton Lee, and	Aaron Grattafiori, Abhimanyu Dubey, Abhinav Jauhri,	831
775	Kristina Toutanova. 2019. BERT: Pre-training of	and et al. 2024. The llama 3 herd of models. <i>Preprint,</i>	832
776	deep bidirectional transformers for language under-	arXiv:2407.21783.	833
777	standing. In <i>Proceedings of the 2019 Conference of</i>	Jiale Han, Shuai Zhao, Bo Cheng, Shengkun Ma, and	834
778	<i>the North American Chapter of the Association for</i>	Wei Lu. 2022. Generative prompt tuning for rela-	835
779	<i>Computational Linguistics: Human Language Tech-</i>	tion classification. In <i>Findings of the Association</i>	836
780	<i>nologies, Volume 1 (Long and Short Papers),</i> pages	<i>for Computational Linguistics: EMNLP 2022,</i> pages	837
781	4171–4186, Minneapolis, Minnesota. Association for	3170–3185, Abu Dhabi, United Arab Emirates. As-	838
782	Computational Linguistics.	sociation for Computational Linguistics.	839
783	Sefika Efeoglu and Adrian Paschke. 2024a. Relation ex-	Pengcheng He, Xiaodong Liu, Jianfeng Gao, and	840
784	traction with fine-tuned large language models in re-	Weizhu Chen. 2021. Deberta: Decoding-enhanced	841
785	trieval augmented generation frameworks. <i>Preprint,</i>	bert with disentangled attention. In <i>International</i>	842
786	arXiv:2406.14745.	<i>Conference on Learning Representations.</i>	843
787	Sefika Efeoglu and Adrian Paschke. 2024b. Retrieval-	Dan Hendrycks, Collin Burns, Steven Basart, Andy	844
788	augmented generation-based relation extraction.	Zou, Mantas Mazeika, Dawn Song, and Jacob Stein-	845
789	<i>Preprint,</i> arXiv:2404.13397.	hardt. 2021. Measuring massive multitask language	846
790	Lisheng Fu, Thien Huu Nguyen, Bonan Min, and Ralph	understanding. <i>Proceedings of the International Con-</i>	847
791	Grishman. 2017. Domain adaptation for relation ex-	<i>ference on Learning Representations (ICLR).</i>	848
792	traction with domain adversarial neural network. In	Evan Hernandez, Arnab Sen Sharma, Tal Haklay, Kevin	849
793	<i>Proceedings of the Eighth International Joint Con-</i>	Meng, Martin Wattenberg, Jacob Andreas, Yonatan	850
794	<i>ference on Natural Language Processing (Volume 2:</i>	Belinkov, and David Bau. 2024. Linearity of rela-	851
795	<i>Short Papers),</i> pages 425–429, Taipei, Taiwan. Asian	tion decoding in transformer language models. In	852
796	Federation of Natural Language Processing.	<i>The Twelfth International Conference on Learning</i>	853
797	Matt Gardner, Yoav Artzi, Victoria Basmov, Jonathan	<i>Representations.</i>	854
798	Berant, Ben Bogin, Sihao Chen, Pradeep Dasigi,	Matthew Honnibal, Ines Montani, Sofie Van Lan-	855
799	Dheeru Dua, Yanai Elazar, Ananth Gottumukkala,	degheem, and Adriane Boyd. 2020. spaCy: Industrial-	856
800	Nitish Gupta, Hannaneh Hajishirzi, Gabriel Ilharco,	strength Natural Language Processing in Python.	857
801	Daniel Khashabi, Kevin Lin, Jiangming Liu, Nel-	Edward J Hu, yelong shen, Phillip Wallis, Zeyuan Allen-	858
802	son F. Liu, Phoebe Mulcaire, Qiang Ning, Sameer	Zhu, Yuanzhi Li, Shean Wang, Lu Wang, and Weizhu	859
803	Singh, Noah A. Smith, Sanjay Subramanian, Reut	Chen. 2022. LoRA: Low-rank adaptation of large	860
804	Tsarfaty, Eric Wallace, Ally Zhang, and Ben Zhou.	language models. In <i>International Conference on</i>	861
805	2020. Evaluating models’ local decision boundaries	Learning Representations.	862
806	via contrast sets. In <i>Findings of the Association</i>	Pere-Lluís Huguet Cabot and Roberto Navigli. 2021.	863
807	<i>for Computational Linguistics: EMNLP 2020,</i> pages	REBEL: Relation extraction by end-to-end language	864
808	1307–1323, Online. Association for Computational	generation. In <i>Findings of the Association for Com-</i>	865
809	Linguistics.	<i>putational Linguistics: EMNLP 2021,</i> pages 2370–	866
810	Mor Geva, Jasmijn Bastings, Katja Filippova, and Amir	2381, Punta Cana, Dominican Republic. Association	867
811	Globerson. 2023. Dissecting recall of factual associa-	for Computational Linguistics.	868
812	tions in auto-regressive language models. In <i>Proceed-</i>	Hyewon Jang and Diego Frassinelli. 2024. Generaliz-	869
813	<i>ings of the 2023 Conference on Empirical Methods in</i>	able sarcasm detection is just around the corner, of	870
814	<i>Natural Language Processing,</i> pages 12216–12235,	course! In <i>Proceedings of the 2024 Conference of</i>	871
815	Singapore. Association for Computational Linguistics.	<i>the North American Chapter of the Association for</i>	872
816		<i>Computational Linguistics: Human Language Tech-</i>	873
817	Mor Geva, Daniel Khashabi, Elad Segal, Tushar Khot,	<i>nologies (Volume 1: Long Papers),</i> pages 4238–4249,	874
818	Dan Roth, and Jonathan Berant. 2021. Did aristotle	Mexico City, Mexico. Association for Computational	875
819	use a laptop? a question answering benchmark with	Linguistics.	876
820	implicit reasoning strategies. <i>Transactions of the</i>	Suhas Kotha, Jacob Mitchell Springer, and Aditi Raghu-	877
821	<i>Association for Computational Linguistics,</i> 9:346–	nathan. 2024. Understanding catastrophic forgetting	878
822	361.		

879	in language models via implicit inference. In <i>The Twelfth International Conference on Learning Representations</i> .	
880		
881		
882	Venelin Kovatchev and Matthew Lease. 2024. Benchmark transparency: Measuring the impact of data on evaluation . In <i>Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers)</i> , pages 1536–1551, Mexico City, Mexico. Association for Computational Linguistics.	
883		
884		
885		
886		
887		
888		
889		
890	Alina Leidinger, Robert van Rooij, and Ekaterina Shutova. 2023. The language of prompting: What linguistic properties make a prompt successful? In <i>Findings of the Association for Computational Linguistics: EMNLP 2023</i> , pages 9210–9232, Singapore. Association for Computational Linguistics.	
891		
892		
893		
894		
895		
896	Guozheng Li, Peng Wang, and Wenjun Ke. 2023a. Revisiting large language models as zero-shot relation extractors . In <i>Findings of the Association for Computational Linguistics: EMNLP 2023</i> , pages 6877–6892, Singapore. Association for Computational Linguistics.	
897		
898		
899		
900		
901		
902	Shaoshuai Li, Wenbin Yu, Zijian Chen, and Yangyang Luo. 2023b. A joint entity and relation extraction model based on encoder-decoder . In <i>2023 IEEE 3rd International Conference on Information Technology, Big Data and Artificial Intelligence (ICIBA)</i> , volume 3, pages 996–1000.	
903		
904		
905		
906		
907		
908	Tal Linzen. 2020. How can we accelerate progress towards human-like linguistic generalization? In <i>Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics</i> , pages 5210–5217, Online. Association for Computational Linguistics.	
909		
910		
911		
912		
913		
914	Pengfei Liu, Jinlan Fu, Yang Xiao, Weizhe Yuan, Shuaichen Chang, Junqi Dai, Yixin Liu, Zihuiwen Ye, and Graham Neubig. 2021. ExplainsBoard: An explainable leaderboard for NLP . In <i>Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing: System Demonstrations</i> , pages 280–289, Online. Association for Computational Linguistics.	
915		
916		
917		
918		
919		
920		
921		
922		
923	Sheng Liu, Zhihui Zhu, Qing Qu, and Chong You. 2022. Robust training under label noise by overparameterization. In <i>Proceedings of the 39th International Conference on Machine Learning</i> , volume 162 of <i>Proceedings of Machine Learning Research</i> , pages 14153–14172. PMLR.	
924		
925		
926		
927		
928		
929	Siyi Liu, Yang Li, Jiang Li, Shan Yang, and Yunshi Lan. 2024. Unleashing the power of large language models in zero-shot relation extraction via self-prompting . In <i>Findings of the Association for Computational Linguistics: EMNLP 2024</i> , pages 13147–13161, Miami, Florida, USA. Association for Computational Linguistics.	
930		
931		
932		
933		
934		
935		
	Zihan Liu, Yan Xu, Tiezheng Yu, Wenliang Dai, Ziwei Ji, Samuel Cahyawijaya, Andrea Madotto, and Pascale Fung. 2020. Crossner: Evaluating cross-domain named entity recognition . In <i>AAAI Conference on Artificial Intelligence</i> .	936
		937
		938
		939
		940
	Yao Lu, Max Bartolo, Alastair Moore, Sebastian Riedel, and Pontus Stenetorp. 2022. Fantastically ordered prompts and where to find them: Overcoming few-shot prompt order sensitivity . In <i>Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)</i> , pages 8086–8098, Dublin, Ireland. Association for Computational Linguistics.	941
		942
		943
		944
		945
		946
		947
		948
	Christopher Manning, Mihai Surdeanu, John Bauer, Jenny Finkel, Steven Bethard, and David McClosky. 2014. The Stanford CoreNLP natural language processing toolkit . In <i>Proceedings of 52nd Annual Meeting of the Association for Computational Linguistics: System Demonstrations</i> , pages 55–60, Baltimore, Maryland. Association for Computational Linguistics.	949
		950
		951
		952
		953
		954
		955
		956
	Shiao Meng, Xuming Hu, Aiwei Liu, Fukun Ma, Yawen Yang, Shuang Li, and Lijie Wen. 2024. On the robustness of document-level relation extraction models to entity name variations . In <i>Findings of the Association for Computational Linguistics: ACL 2024</i> , pages 16362–16374, Bangkok, Thailand. Association for Computational Linguistics.	957
		958
		959
		960
		961
		962
		963
	Graham Neubig, Zi-Yi Dou, Junjie Hu, Paul Michel, Danish Pruthi, and Xinyi Wang. 2019. compare-mt: A tool for holistic comparison of language generation systems . In <i>Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics (Demonstrations)</i> , pages 35–41, Minneapolis, Minnesota. Association for Computational Linguistics.	964
		965
		966
		967
		968
		969
		970
		971
	Riccardo Orlando, Pere-Lluís Huguet Cabot, Edoardo Barba, and Roberto Navigli. 2024. ReLiK: Retrieve and LinK, fast and accurate entity linking and relation extraction on an academic budget . In <i>Findings of the Association for Computational Linguistics: ACL 2024</i> , pages 14114–14132, Bangkok, Thailand. Association for Computational Linguistics.	972
		973
		974
		975
		976
		977
		978
	Alistair Plum, Tharindu Ranasinghe, Spencer Jones, Constantin Orasan, and Ruslan Mitkov. 2022. Biographical semi-supervised relation extraction dataset . In <i>Proceedings of the 45th International ACM SIGIR Conference on Research and Development in Information Retrieval, SIGIR ’22</i> , page 3121–3130, New York, NY, USA. Association for Computing Machinery.	979
		980
		981
		982
		983
		984
		985
		986
	Ji Qi, Chuchun Zhang, Xiaozhi Wang, Kaisheng Zeng, Jifan Yu, Jinxin Liu, Lei Hou, Juanzi Li, and Xu Bin. 2023. Preserving knowledge invariance: Rethinking robustness evaluation of open information extraction . In <i>Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing</i> , pages 5876–5890, Singapore. Association for Computational Linguistics.	987
		988
		989
		990
		991
		992
		993
		994

995	Sebastian Riedel, Limin Yao, and Andrew McCallum.	Komal Teru. 2023. Semi-supervised relation extraction via data augmentation and consistency-training .	1051
996	2010. Modeling relations and their mentions without	In <i>Proceedings of the 17th Conference of the Euro-</i>	1052
997	labeled text. In <i>Machine Learning and Knowledge</i>	<i>pean Chapter of the Association for Computational</i>	1053
998	<i>Discovery in Databases</i> , pages 148–163, Berlin, Hei-	<i>Linguistics</i> , pages 1112–1124, Dubrovnik, Croatia.	1054
999	delberg. Springer Berlin Heidelberg.	Association for Computational Linguistics.	1055
1000	Jonathan Rusert, Zubair Shafiq, and Padmini Srinivasan.		1056
1001	2022. On the robustness of offensive language classi-	Quyen Tran, Nguyen Xuan Thanh, Nguyen Hoang	1057
1002	fiers . In <i>Proceedings of the 60th Annual Meeting of</i>	Anh, Nam Le Hai, Trung Le, Linh Van Ngo, and	1058
1003	<i>the Association for Computational Linguistics (Vol-</i>	Thien Huu Nguyen. 2024. Preserving generalization	1059
1004	<i>ume 1: Long Papers</i>), pages 7424–7438, Dublin,	of language models in few-shot continual relation ex-	1060
1005	Ireland. Association for Computational Linguistics.	traction . In <i>Proceedings of the 2024 Conference on</i>	1061
1006	Pratik Saini, Samiran Pal, Tapas Nayak, and Indrajit	<i>Empirical Methods in Natural Language Processing</i> ,	1062
1007	Bhattacharya. 2023. 90% f1 score in relation triple	pages 13771–13784, Miami, Florida, USA. Associa-	1063
1008	extraction: Is it real? In <i>Proceedings of the 1st</i>	tion for Computational Linguistics.	1064
1009	<i>GenBench Workshop on (Benchmarking) Generalisa-</i>		
1010	<i>tion in NLP</i> , pages 1–11, Singapore. Association for	Shubham Vatsal and Harsh Dubey. 2024. A sur-	1065
1011	Computational Linguistics.	vey of prompt engineering methods in large lan-	1066
		guage models for different nlp tasks . <i>Preprint</i> ,	1067
		arXiv:2407.12994.	1068
1012	Oscar Sainz, Iker García-Ferrero, Rodrigo Agerri,	Chenguang Wang, Xiao Liu, Zui Chen, Haoyun Hong,	1069
1013	Oier Lopez de Lacalle, German Rigau, and Eneko	Jie Tang, and Dawn Song. 2022. DeepStruct: Pre-	1070
1014	Agirre. 2024. GoLLIE: Annotation guidelines im-	training of language models for structure prediction .	1071
1015	prove zero-shot information-extraction . In <i>The</i>	In <i>Findings of the Association for Computational Lin-</i>	1072
1016	<i>Twelfth International Conference on Learning Repre-</i>	<i>guistics: ACL 2022</i> , pages 803–823, Dublin, Ireland.	1073
1017	<i>sentations</i> .	Association for Computational Linguistics.	1074
1018	Evan Sandhaus. 2008. The new york times annotated	Qing Wang, Kang Zhou, Qiao Qiao, Yuepei Li, and	1075
1019	corpus. <i>Linguistic Data Consortium, Philadelphia</i> ,	Qi Li. 2023. Improving unsupervised relation extrac-	1076
1020	6(12):e26752.	tion by augmenting diverse sentence pairs . In <i>Pro-</i>	1077
1021	Chenglei Si, Zhe Gan, Zhengyuan Yang, Shuohang	<i>ceedings of the 2023 Conference on Empirical Meth-</i>	1078
1022	Wang, Jianfeng Wang, Jordan Lee Boyd-Graber, and	<i>ods in Natural Language Processing</i> , pages 12136–	1079
1023	Lijuan Wang. 2023. Prompting GPT-3 to be reli-	12147, Singapore. Association for Computational	1080
1024	able . In <i>The Eleventh International Conference on</i>	Linguistics.	1081
1025	<i>Learning Representations</i> .		
1026	Yisheng Song, Ting Wang, Puyu Cai, Subrota K. Mon-	Albert Webson and Ellie Pavlick. 2022. Do prompt-	1082
1027	dal, and Jyoti Prakash Sahoo. 2023. A comprehen-	sive survey of few-shot learning: Evolution, applica-	1083
1028	tions, challenges, and opportunities . <i>ACM Comput.</i>	<i>Proceedings of the 2022 Conference of</i>	1084
1029	<i>Surv.</i> , 55(13s).	<i>the North American Chapter of the Association for</i>	1085
1030		<i>Computational Linguistics: Human Language Tech-</i>	1086
		<i>nologies</i> , pages 2300–2344, Seattle, United States.	1087
		Association for Computational Linguistics.	1088
1031	Xiaofei Sun, Xiaoya Li, Jiwei Li, Fei Wu, Shangwei	Tongshuang Wu, Marco Tulio Ribeiro, Jeffrey Heer, and	1089
1032	Guo, Tianwei Zhang, and Guoyin Wang. 2023. Text	Daniel Weld. 2019. Errudite: Scalable, reproducible,	1090
1033	classification via large language models . In <i>Find-</i>	and testable error analysis . In <i>Proceedings of the</i>	1091
1034	<i>ings of the Association for Computational Linguis-</i>	<i>57th Annual Meeting of the Association for Compu-</i>	1092
1035	<i>tics: EMNLP 2023</i> , pages 8990–9005, Singapore.	<i>tational Linguistics</i> , pages 747–763, Florence, Italy.	1093
1036	Association for Computational Linguistics.	Association for Computational Linguistics.	1094
1037	Bruno Taillé, Vincent Guigue, Geoffrey Scuttheeten,	Xin Xu, Yuqi Zhu, Xiaohan Wang, and Ningyu Zhang.	1095
1038	and Patrick Gallinari. 2020. Let’s Stop Incorrect	2023. How to unleash the power of large language	1096
1039	Comparisons in End-to-end Relation Extraction! In	models for few-shot relation extraction? In <i>Proceed-</i>	1097
1040	<i>Proceedings of the 2020 Conference on Empirical</i>	<i>ings of The Fourth Workshop on Simple and Efficient</i>	1098
1041	<i>Methods in Natural Language Processing (EMNLP)</i> ,	<i>Natural Language Processing (SustaiNLP)</i> , pages	1099
1042	pages 3689–3701, Online. Association for Computa-	190–200, Toronto, Canada (Hybrid). Association for	1100
1043	tional Linguistics.	Computational Linguistics.	1101
1044	Michael Tänzler, Sebastian Ruder, and Marek Rei. 2022.	Yadollah Yaghoobzadeh, Heike Adel, and Hinrich	1102
1045	Memorisation versus generalisation in pre-trained	Schütze. 2017. Noise mitigation for neural entity	1103
1046	language models . In <i>Proceedings of the 60th Annual</i>	typing and relation extraction . In <i>Proceedings of</i>	1104
1047	<i>Meeting of the Association for Computational Lin-</i>	<i>the 15th Conference of the European Chapter of the</i>	1105
1048	<i>guistics (Volume 1: Long Papers)</i> , pages 7564–7578,	<i>Association for Computational Linguistics: Volume</i>	1106
1049	Dublin, Ireland. Association for Computational Lin-	<i>1, Long Papers</i> , pages 1183–1194, Valencia, Spain.	1107
1050	guistics.	Association for Computational Linguistics.	1108

- Yuan Yao, Deming Ye, Peng Li, Xu Han, Yankai Lin, Zhenghao Liu, Zhiyuan Liu, Lixin Huang, Jie Zhou, and Maosong Sun. 2019. [DocRED: A large-scale document-level relation extraction dataset](#). In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*, pages 764–777, Florence, Italy. Association for Computational Linguistics.
- Yiming Zhang, Shi Feng, and Chenhao Tan. 2022. [Active example selection for in-context learning](#). In *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing*, pages 9134–9148, Abu Dhabi, United Arab Emirates. Association for Computational Linguistics.
- Yuhao Zhang, Victor Zhong, Danqi Chen, Gabor Angeli, and Christopher D. Manning. 2017. [Position-aware attention and supervised data improve slot filling](#). In *Proceedings of the 2017 Conference on Empirical Methods in Natural Language Processing*, pages 35–45, Copenhagen, Denmark. Association for Computational Linguistics.

A Dataset Statistics: Class Distribution

Table 3: Balanced Biographical Dataset

Relation	# of Samples
Other	10,000
birthdate	2914
bplace_name	2845
dplace_name	1138
occupation	1105
deathdate	1011
parent	394
educatedAt	339
child	136
sibling	118
Positive Samples	10,000
Negative Samples	10,000
All	20,000

Table 4: Balanced TACRED-RE Subset with Biographical Relations (26 relations)

Relation	# of Samples
no_relation	14,192
per:title	3805
per:employee_of	2104
per:age	818
per:countries_of_residence	695
per:cities_of_residence	596
per:origin	652
per:stateorprovinces_of_residence	444
per:spouse	463
per:date_of_death	343
per:children	347
per:cause_of_death	318
per:parents	282
per:charges	270
per:other_family	241
per:siblings	238
per:schools_attended	219
per:city_of_death	204
per:religion	145
per:alternate_names	132
per:city_of_birth	107
per:stateorprovince_of_death	100
per:date_of_birth	99
per:stateorprovince_of_birth	77
per:country_of_death	57
per:country_of_birth	45
Positive Samples	12,801
Negative Samples	14,192
All	26,993

Table 5: Balanced NYT Subset with Biographical Relations after removal of instances with multiple labels (7 relations)

Relation	# of Samples
None	5068
/people/person/nationality	2160
/people/person/place_lived	2016
/people/person/place_of_birth	437
/people/deceased_person/place_of_death	284
/people/person/children	147
/people/person/religion	24
Positive Samples	5068
Negative Samples	5068
All	10,136

B Relation Type Overlap

Table 6: NYT/Biographical Relation Overlap

NYT	Biographical
/people/person/children	child
/people/person/place_of_birth	bplace_name
/people/deceased_person/place_of_death	dplace_name
None	Other

Table 7: NYT/TACRED-RE Relation Overlap

NYT	TACRED-RE
None	no_relation
/people/person/children	per:children
/people/person/religion	per:religion
/people/person/place_lived	per:stateorprovinces_of_residence per:countries_of_residence per:cities_of_residence
/people/person/place_of_birth	per:stateorprovince_of_birth per:country_of_birth per:city_of_birth
/people/deceased_person/place_of_death	per:stateorprovince_of_death per:country_of_death per:city_of_death

Table 8: Biographical/TACRED-RE Relation Overlap

Biographical	TACRED-RE
bplace_name	per:stateorprovince_of_birth per:country_of_birth per:city_of_birth
birthdate	per:date_of_birth
deathdate	per:date_of_death
parent	per:parents
educatedAt	per:schools_attended
dplace_name	per:stateorprovince_of_death per:country_of_death per:city_of_death
sibling	per:siblings
child	per:children
Other	no_relation

Table 9: Biographical/TACRED-RE/NYT Overlap

Biographical	TACRED-RE	NYT
child	per:children	/people/person/children
bplace_name	per:stateorprovince_of_birth per:country_of_birth per:city_of_birth	/people/person/place_of_birth
dplace_name	per:stateorprovince_of_death per:country_of_death per:city_of_death	/people/deceased_person/ place_of_death
Other	no_relation	None

C Data Formatting Details

TACRED-RE’s fine-grained relations (e.g., "per:city_of_birth" and "per:country_of_birth") were mapped to broader categories (e.g., "place_of_birth") used in NYT and Biographical datasets, as shown in Tables 7 and 8. Cross-dataset results are reported using NYT label names (Table 15).

D Implementation Details

Listing 1: System prompt for LLaMA 3.1 8B zero-shot, few-shot, and fine-tuning experiments, shown here with NYT relation inventory

```
system_message = {
  "role": "system",
  "content": (
    "You are an intelligent assistant specializing in identifying relations between entities in a sentence. "
    "Question: What is the relation between two tagged entities <e1>entity1</e1> and <e2>entity2</e2> in the following sentence? "
    "Choose one relation from the list: "
    "[ '/people/person/children', '/people/person/nationality', '/people/person/place_lived', "
    "'/people/person/place_of_birth', '/people/deceased_person/place_of_death', '/people/person/religion', "
    "'None']. "
    "Rules: Select exactly one relation from the list. If none of the listed relations apply, select 'None'. "
    "Output must strictly follow this format: <relation_type>. Provide no additional text or explanation ."
  )
}
```

We fine-tuned DeBERTa-v3-large¹¹ for 10 epochs employing early stopping. Following prior work in RE (Teru, 2023; Wang et al., 2023), we extended the deberta-v3-large tokeniser with entity marker tokens, namely, <e1> and </e1> for the head entity, and <e2> and </e2> for the tail entity. For LLAMA-3.1¹², we used LoRA fine-tuning ($r = 8$) over three epochs, applying it to attention and feedforward modules. Both models were fine-tuned using HuggingFace’s Trainer class. For evaluation of LLaMA 3.1, predictions were considered correct only if they matched ground-truth labels exactly (Hendrycks et al., 2021).

For prompting, we used vanilla prompting vanilla (Li et al., 2023a; Vatsal and Dubey,

¹¹<https://huggingface.co/microsoft/deberta-v3-large>

¹²<https://huggingface.co/meta-llama/Llama-3.1-8B-Instruct>

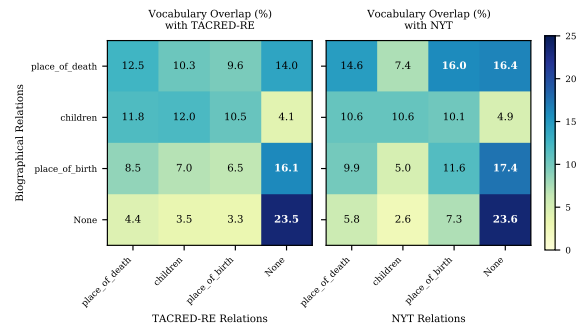


Figure 6: Vocabulary Overlap (%) Between Biographical and TACRED-RE/NYT Relations following lemmatisation and stop word removal using spaCy’s en_core_web_trf

2024) and tested several RE-specific prompt designs (Leidinger et al., 2023; Li et al., 2023a; Efeoglu and Paschke, 2024a), given LLaMA’s sensitivity to prompt formulation (Leidinger et al., 2023). The prompt 1 performed best and was used across all datasets with adapted label sets. Further prompt optimisation techniques were not considered, as they were beyond the scope of this paper. Hyperparameter settings for all experiments are detailed in Table 19.

Due to Biographical’s ambiguous ‘Other’ class (Section 3), we use it only for evaluation, excluding these relations during adaptation.

All experiments with DeBERTa-v3-large were run on a single NVIDIA® TITAN RTX 24GB GPU; all experiments with Llama-3.1-8B-Instruct were run on a single NVIDIA® A100 80GB GPU. All experiments were performed with a fixed random seed for reproducibility.

E Vocabulary Overlap with Biographical

Figure 6 depicts vocabulary overlap between Biographical and TACRED-RE as well as Biographical and NYT per overlapping relation.

F Results

F.1 Intra-Dataset Results

Model	Deberta-v3-large 304M			Llama-3.1 8B zero-shot			Llama-3.1 8B 5-shot			Llama-3.1 8B fine-tuned		
	P	R	F1	P	R	F1	P	R	F1	P	R	F1
/people/deceased_person/place_of_death	0.81	0.76	0.78	0.93	0.82	0.87	0.75	0.09	0.16	0.93	0.82	0.87
/people/person/children	0.80	0.86	0.83	0.87	0.93	0.90	0.77	0.71	0.74	0.87	0.93	0.90
/people/person/nationality	0.95	1.00	0.97	0.99	0.99	0.99	0.96	0.10	0.18	0.99	0.99	0.99
/people/person/place_lived	0.88	0.91	0.90	0.84	0.95	0.89	0.36	0.58	0.44	0.84	0.95	0.89
/people/person/place_of_birth	0.54	0.56	0.55	0.88	0.46	0.60	0.07	0.06	0.07	0.88	0.46	0.60
/people/person/religion	0.00	0.00	0.00	1.00	1.00	1.00	0.83	1.00	0.91	1.00	1.00	1.00
None	0.98	0.95	0.96	0.97	0.97	0.97	0.62	0.76	0.68	0.97	0.97	0.97
macro avg	0.71	0.72	0.71	0.92	0.87	0.89	0.62	0.47	0.45	0.92	0.87	0.89
micro avg	0.92	0.92	0.92	0.30	0.22	0.25	0.52	0.52	0.52	0.94	0.94	0.94
weighted avg	0.92	0.92	0.92	0.94	0.94	0.94	0.63	0.52	0.47	0.94	0.94	0.94

Table 10: NYT Results

Model	Deberta-v3-large 304M			Llama-3.1 8B zero-shot			Llama-3.1 8B 5-shot			Llama-3.1 8B fine-tuned		
	P	R	F1	P	R	F1	P	R	F1	P	R	F1
macro avg	0.67	0.64	0.64	0.49	0.38	0.34	0.37	0.30	0.29	0.76	0.71	0.73
micro avg	0.83	0.83	0.83	0.31	0.29	0.30	0.51	0.51	0.51	0.87	0.87	0.87
weighted avg	0.83	0.83	0.83	0.68	0.29	0.32	0.59	0.51	0.49	0.87	0.87	0.87

Table 11: TACRED Results

Model	DeBERTa-v3-large 304M			Llama-3.1 8B zero-shot			Llama-3.1 8B 5-shot			Llama-3.1 8B fine-tuned		
	P	R	F1	P	R	F1	P	R	F1	P	R	F1
no_relation	0.96	0.79	0.87	0.81	0.01	0.02	0.80	0.61	0.69	0.97	0.87	0.92
per:age	0.91	0.99	0.95	0.97	0.32	0.48	0.93	0.65	0.77	0.95	1.00	0.97
per:cause_of_death	0.74	0.74	0.74	0.46	0.26	0.33	0.53	0.24	0.33	0.65	0.95	0.77
per:charges	0.79	0.95	0.86	0.76	0.30	0.43	0.67	0.43	0.52	0.87	0.99	0.93
per:children	0.54	0.73	0.62	0.23	0.14	0.17	0.19	0.38	0.25	0.96	0.73	0.83
per:cities_of_residence	0.51	0.95	0.66	0.38	0.55	0.45	0.51	0.34	0.41	0.61	0.92	0.73
per:city_of_birth	0.75	0.50	0.60	0.50	0.67	0.57	0.43	0.50	0.46	1.00	0.50	0.67
per:city_of_death	0.78	0.44	0.56	0.33	0.31	0.32	0.28	0.56	0.38	0.43	0.56	0.49
per:countries_of_residence	0.44	0.81	0.57	0.33	0.45	0.38	0.34	0.31	0.32	0.59	0.91	0.71
per:country_of_death	0.00	0.00	0.00	0.00	0.00	0.00	0.50	0.56	0.53	0.50	0.44	0.47
per:date_of_birth	0.70	1.00	0.82	0.26	0.86	0.40	0.71	0.71	0.71	0.86	0.86	0.86
per:date_of_death	0.76	0.80	0.78	0.53	0.42	0.47	0.62	0.12	0.21	0.74	0.93	0.82
per:employee_of	0.76	0.90	0.82	0.10	0.97	0.19	0.17	0.73	0.27	0.86	0.89	0.88
per:origin	0.70	0.83	0.76	0.65	0.12	0.21	0.45	0.13	0.20	0.82	0.79	0.80
per:other_family	0.39	0.49	0.43	0.00	0.00	0.00	0.11	0.59	0.19	0.61	0.97	0.75
per:parents	0.79	0.94	0.86	0.33	0.05	0.09	0.41	0.39	0.40	0.87	0.87	0.87
per:religion	0.65	0.75	0.70	0.67	0.35	0.46	0.65	0.75	0.70	0.71	0.93	0.80
per:schools_attended	0.95	0.70	0.81	1.00	0.04	0.07	0.86	0.22	0.35	1.00	0.81	0.90
per:siblings	0.69	0.85	0.77	0.40	0.60	0.48	0.45	0.69	0.54	0.98	0.94	0.96
per:spouse	0.77	0.97	0.86	0.26	0.69	0.38	0.20	0.58	0.30	0.89	0.81	0.85
per:stateorprovince_of_birth	1.00	0.67	0.80	0.00	0.00	0.00	0.50	0.33	0.40	0.50	0.67	0.57
per:stateorprovince_of_death	0.50	0.20	0.29	0.00	0.00	0.00	1.00	0.10	0.18	0.44	0.70	0.54
per:stateorprovinces_of_residence	0.53	0.84	0.65	0.40	0.03	0.06	0.43	0.52	0.47	0.63	0.84	0.72
per:title	0.90	0.97	0.93	0.88	0.01	0.03	0.97	0.08	0.14	0.94	0.97	0.96
macro avg	0.69	0.74	0.70	0.43	0.30	0.25	0.53	0.44	0.41	0.77	0.83	0.78
micro avg	0.83	0.83	0.83	0.20	0.14	0.17	0.53	0.51	0.52	0.89	0.89	0.89
weighted avg	0.87	0.83	0.84	0.70	0.14	0.11	0.72	0.51	0.54	0.91	0.89	0.89

Table 12: TACRED-RE Results

Model	DeBERTa-v3-large 304M			Llama-3.1 8B zero-shot			Llama-3.1 8B 5-shot			Llama-3.1 8B fine-tuned		
	P	R	F1	P	R	F1	P	R	F1	P	R	F1
Other	0.96	0.89	0.92	0.78	0.29	0.42	0.94	0.54	0.69	0.92	0.93	0.93
birthdate	0.99	1.00	0.99	0.53	0.91	0.67	0.83	0.91	0.87	1.00	1.00	1.00
bplace_name	0.82	0.92	0.87	0.82	0.08	0.14	0.69	0.82	0.75	0.87	0.85	0.86
child	0.44	0.44	0.44	0.02	0.56	0.05	0.00	0.00	0.00	0.57	0.44	0.50
deathdate	0.95	0.96	0.96	0.73	0.82	0.77	0.65	0.85	0.74	0.96	0.99	0.98
dplace_name	0.64	0.71	0.67	0.39	0.25	0.31	0.50	0.55	0.52	0.56	0.81	0.66
educatedAt	0.72	0.90	0.80	0.06	1.00	0.12	0.19	1.00	0.32	0.00	0.00	0.00
occupation	1.00	0.99	1.00	0.83	0.41	0.55	0.77	0.94	0.85	1.00	0.98	0.99
parent	0.58	0.91	0.71	0.38	0.59	0.46	0.26	0.84	0.39	1.00	0.68	0.81
sibling	0.00	0.00	0.00	0.10	0.13	0.11	0.42	0.33	0.37	1.00	0.67	0.80
macro avg	0.71	0.77	0.74	0.46	0.50	0.36	0.53	0.68	0.55	0.79	0.74	0.75
micro avg	0.90	0.90	0.90	0.41	0.40	0.41	0.69	0.69	0.69	0.91	0.91	0.91
weighted avg	0.91	0.90	0.90	0.70	0.40	0.43	0.80	0.69	0.71	0.90	0.91	0.90

Table 13: Biographical Results

F.2 Cross-Dataset Results

Model	Setting	Dataset	Intra-Dataset		Cross-Dataset	
			Shared Labels	Dataset Labels	NYT	TACRED-RE
DeBERTa-v3-large 304M	Fine-tuned on	NYT	0.62	0.62	–	0.27
		TACRED-RE	0.71	0.67	0.49	–
LLaMa 3.1 8B	Fine-tuned on	NYT	0.83	0.83	–	0.45
		TACRED-RE	0.79	0.76	0.58	–
LLaMa 3.1 8B	Shot Setting					
	Zero-Shot	NYT	0.35	0.35	–	–
		TACRED-RE	0.64	0.54	–	–
LLaMa 3.1 8B	5-Shot	NYT	0.46 ± 0.04	0.46 ± 0.04	–	0.50 ± 0.06
		TACRED-RE	0.59 ± 0.03	0.45 ± 0.05	0.44 ± 0.04	–

Table 14: Macro F1-scores for intra- and cross-dataset predictions for the overlap of six relations. Results are reported for shared and dataset-specific labels in both fine-tuned and shot settings. For fine-tuning, models are fine-tuned the overlap of six relations; for shot setting, also only the overlap of six relations is used for prompt or random shot samples.

F.3 Cross-Dataset Per-Class Results

Model	DeBERTa-v3-large 304M			LLaMA-3.1 8B 5-shot			LLaMA-3.1 8B fine-tuned		
	P	R	F1	P	R	F1	P	R	F1
/people/deceased_person/place_of_death	0.75	0.36	0.49	0.54 ± 0.07	0.39 ± 0.11	0.45 ± 0.08	0.71	0.45	0.56
/people/person/children	0.50	0.43	0.46	0.47 ± 0.05	0.29 ± 0.10	0.34 ± 0.08	0.69	0.64	0.67
/people/person/place_lived	0.62	0.93	0.75	0.27 ± 0.06	0.22 ± 0.15	0.23 ± 0.11	0.63	0.86	0.73
/people/person/place_of_birth	0.33	0.06	0.11	0.12 ± 0.05	0.07 ± 0.04	0.09 ± 0.04	0.17	0.02	0.04
/people/person/religion	1.00	0.20	0.33	0.64 ± 0.03	0.72 ± 0.23	0.66 ± 0.11	1.00	0.80	0.89
None	0.94	0.78	0.85	0.84 ± 0.03	0.46 ± 0.07	0.59 ± 0.06	0.90	0.82	0.86
macro avg	0.69	0.46	0.50	0.48 ± 0.03	0.36 ± 0.04	0.39 ± 0.02	0.68	0.60	0.62
weighted avg	0.80	0.75	0.75	0.62 ± 0.04	0.37 ± 0.03	0.45 ± 0.03	0.78	0.76	0.76

Table 15: Adapted on TACRED-RE with all biographical relations present in TACRED-RE. Evaluations on NYT. The labels, although borrowed from NYT dataset, reflect the shared labels between NYT and TACRED-RE. More fine-grained TACRED-RE were mapped to broader shared labels to enable cross-dataset evaluation comparison.

Model	DeBERTa-v3-large 304M			LLaMA-3.1 8B 5-shot			LLaMA-3.1 8B fine-tuned		
	P	R	F1	P	R	F1	P	R	F1
/people/deceased_person/place_of_death	0.47	0.20	0.28	0.66 ± 0.11	0.33 ± 0.20	0.40 ± 0.17	0.62	0.23	0.33
/people/person/children	0.50	0.14	0.21	0.49 ± 0.15	0.48 ± 0.20	0.47 ± 0.17	1.00	0.14	0.24
/people/person/place_lived	0.69	0.17	0.27	0.65 ± 0.09	0.38 ± 0.06	0.48 ± 0.05	0.71	0.31	0.43
/people/person/place_of_birth	0.01	0.08	0.02	0.18 ± 0.10	0.88 ± 0.26	0.28 ± 0.11	0.38	0.25	0.30
/people/person/religion	0.00	0.00	0.00	0.86 ± 0.04	0.64 ± 0.17	0.72 ± 0.11	0.94	0.40	0.56
None	0.81	0.90	0.85	0.85 ± 0.04	0.78 ± 0.12	0.81 ± 0.05	0.79	0.95	0.86
macro avg	0.41	0.25	0.27	0.62 ± 0.05	0.58 ± 0.06	0.52 ± 0.06	0.74	0.38	0.45
weighted avg	0.73	0.66	0.66	0.79 ± 0.01	0.66 ± 0.09	0.71 ± 0.05	0.77	0.74	0.72

Table 16: Adapted on NYT with all biographical relations present in NYT. Evaluations on TACRED-RE. The labels, although borrowed from NYT dataset, reflect the shared labels between NYT and TACRED-RE. More fine-grained TACRED-RE were mapped to broader shared labels to enable cross-dataset evaluation comparison.

Model	DeBERTa-v3-large 304M			LLaMA-3.1 8B 5-shot			LLaMA-3.1 8B fine-tuned		
	P	R	F1	P	R	F1	P	R	F1
None	0.92	0.63	0.75	0.87 ± 0.10	0.49 ± 0.09	0.62 ± 0.04	0.90	0.68	0.78
/people/person/place_of_birth	0.90	0.73	0.80	0.84 ± 0.02	0.70 ± 0.08	0.76 ± 0.05	0.90	0.75	0.82
/people/person/children	0.57	0.44	0.50	0.14 ± 0.14	0.13 ± 0.14	0.14 ± 0.14	0.71	0.56	0.63
/people/deceased_person/place_of_death	0.92	0.23	0.36	0.87 ± 0.05	0.36 ± 0.05	0.51 ± 0.05	0.95	0.38	0.54
macro avg	0.83	0.51	0.60	0.68 ± 0.06	0.42 ± 0.04	0.51 ± 0.04	0.87	0.59	0.69
weighted avg	0.91	0.61	0.72	0.85 ± 0.06	0.55 ± 0.01	0.65 ± 0.01	0.90	0.67	0.76

Table 17: Adapted on TACRED-RE with all biographical relations present in TACRED-RE. Evaluations on Biographical with four biographical relations (full overlap between three datasets). The labels, although borrowed from NYT dataset, reflect the shared labels between NYT, TACRED-RE, and Biographical.

Model	DeBERTa-v3-large 304M			LLaMA-3.1 8B 5-shot			LLaMA-3.1 8B fine-tuned		
	P	R	F1	P	R	F1	P	R	F1
/people/person/place_of_birth	0.76	0.39	0.51	0.82 ± 0.02	0.71 ± 0.14	0.75 ± 0.09	0.92	0.14	0.25
/people/person/children	0.27	0.44	0.33	0.19 ± 0.06	0.51 ± 0.06	0.27 ± 0.07	0.00	0.00	0.00
/people/deceased_person/place_of_death	0.62	0.22	0.32	0.76 ± 0.09	0.11 ± 0.08	0.19 ± 0.12	0.74	0.13	0.22
None	0.65	0.83	0.73	0.85 ± 0.09	0.62 ± 0.09	0.71 ± 0.06	0.57	0.96	0.72
macro avg	0.57	0.47	0.47	0.65 ± 0.02	0.49 ± 0.04	0.48 ± 0.04	0.56	0.31	0.30
weighted avg	0.68	0.59	0.59	0.82 ± 0.04	0.59 ± 0.05	0.66 ± 0.04	0.71	0.55	0.48

Table 18: Adapted on NYT with all biographical relations present in NYT. Evaluations on Biographical with four biographical relations (full overlap between three datasets). The labels, although borrowed from NYT dataset, reflect the shared labels between NYT, TACRED-RE, and Biographical.

Setting	Parameter	DeBERTa-v3-large Fine-tuned	LLaMA 3.1 8B Zero-Shot	LLaMA 3.1 8B Five-Shot	LLaMA 3.1 8B Fine-tuned
Common	# of Epochs	10	–	–	3
	seed	42	42	42	42
	Loss	Cross-Entropy Loss	–	–	Cross-Entropy Loss
	Optimiser	AdamW	–	–	AdamW
	Batch Size	8	–	–	4
	Gradient Accumulation	4	–	–	–
	Early Stopping Patience	2	–	–	2
	Temperature	–	0.1	0.1	–
	Nucleus Sampling	–	0.9	0.9	–
	Lora Settings [†]	–	–	–	8/32/0.1
TACRED-RE	Learning Rate	5×10^{-6} / 5×10^{-5}	–	–	5×10^{-5}
	Max Length	–	–	–	800/384
	Max New Tokens	–	40	256	–
	Cross-Validation	–/5-fold	–	–	–
NYT	Learning Rate	5×10^{-6}	–	–	1×10^{-4}
	Max Length	–	–	–	384
	Max New Tokens	–	256	256	–
	Cross-Validation	–/5-fold	–	–	–
Biographical	Learning Rate	5×10^{-6}	–	–	1×10^{-4}
	Max Length	–	–	–	384
	Max New Tokens	–	40	256	–

Table 19: Hyperparameter settings across datasets. Two values (x/y) indicate *All/Overlap* relation experiment settings respectively (if these differ), where *All* indicates experiments with the whole set of biographical relations in each dataset and *Overlap* uses only the intersection. Biographical experiments are performed only with the whole set of biographical relations. [†]Lora Settings: Rank/Alpha/Dropout.

POS	TACRED-RE		NYT		Biographical	
	Head Entity	Tail Entity	Head Entity	Tail Entity	Head Entity	Tail Entity
PROPN	77.4	55.6	98.4	98.8	87.7	60.5
PRON	16.8	6.2	–	–	0.1	0.2
NOUN	2.4	16.8	0.2	0.4	1.7	5.4
ADJ	1.2	4.5	0.2	–	0.7	0.9
ADP	0.7	1.6	0.2	0.3	0.3	1.1
NUM	0.0	8.5	–	–	6.4	25.5
DET	0.3	1.0	0.2	0.1	0.8	2.3
VERB	0.4	0.7	–	–	0.1	0.1

Table 20: (Top 8) POS Distribution Across TACRED-RE, NYT, and Biographical Test Sets with all Biographical Relations (%). POS tags are obtained with spaCy’s transformer-based en_core_web_trf model.

G Misclassification Analysis

Issue	Description	Representative Example	Misclassifications on	Models Affected
Overpredicting ‘None’	Overpredicting ‘None’ and struggling with even clear relations with cues like ‘born’ or ‘died’	“My name is <e1>Ruben</e1> and I am from <e2>Holland</e2>” (GT: <i>place_lived</i> , Pred: <i>None</i> ; TACRED-RE sample ID: ‘098f6f318bc468878bbb’)	TACRED-RE and Biographical	NYT-adapted models
Failure to Capture Implicit Relations	Models struggling to detect implicit relations requiring reasoning	“<e1>Gross</e1> [...] was sent to <e2>Cuba</e2> as a spy” (GT: <i>place_lived</i> , Pred: <i>None</i> ; TACRED-RE sample ID: ‘098f6f318be29eddb182’)	TACRED-RE	NYT-adapted models
Expected world knowledge	For NYT and Biographical this issue is also frequently paired with detatable ground truth labels	“<e1>Augustus</e1> also amassed an impressive art collection and built lavish baroque palaces in Dresden and <e2>Warsaw</e2>” (GT: <i>dplace_name</i> , Pred: <i>None</i> ; Biographical sample ID: ‘mS2/247724’)	NYT, TACRED-RE, Biographical	models adapted on all 3 datasets affected
Relation Present in Sentence but Not Between Specified Entities	This issue raises concerns about the framing of the RE task itself	“Jan Malte, [...] resident of <e1>Bridgehampton</e1>, died [...] in <e2>San Francisco</e2>” (GT: <i>None</i> , Pred: <i>place_of_death</i> ; NYT article ID: ‘/m/vinci8/data1/riedel/projects/relation/kb/nyt1/docstore/nyt-2005-2006.backup/1777142.xml.pb’)	NYT, TACRED-RE, Biographical	models adapted on all 3 datasets affected
Debatable ground truth (GT) labels	Caused by distantly or semi-supervised manner in which NYT and Biographical were created	“<e1>Ida Freund</e1> was born in <e2>Austria</e2>” (GT: <i>Other</i> , Pred: <i>place_of_birth</i> ; Biographical sample ID: ‘mS10/37387826’)	TACRED-RE, Biographical	NYT-adapted models
Single-Label Annotation Limitation	Sentences labeled with a single relation may contain additional relations that remain unlabeled	“<e1>Gross</e1>, who is himself Jewish [...] was sent to <e2>Cuba</e2>” (GT: <i>place_lived</i> , Pred: <i>None</i> ; TACRED-RE sample ID: ‘098f6f318b69f98c850c’)	NYT, TACRED-RE, Biographical	models adapted on all 3 datasets affected
Relation missing in annotation schema	Lack of granularity needed to fully capture an individual’s biography	“Wen was detained in August and accused of protecting the gang operations masterminded by his sister-in-law, <e1>Xie Caiping</e1>, 46, known as the “godmother” of the <e2>Chinese</e2> underworld (GT: <i>place_lived</i> , Pred: <i>nationality</i> ; TACRED-RE sample ID: ‘098f637935e6e6d1d093’)	NYT, TACRED-RE, Biographical	—
Failure to Capture Relations in Long, Compound Sentences	Models struggling with long-term relational dependencies	“Ecoffey told jurors that he and another federal agent met with <e1>Graham</e1> in April 1994 in Yellowknife, the city in northwest <e2>Canada</e2> where Graham lived at the time” (GT: <i>place_lived</i> , Pred: <i>None</i> ; TACRED-RE sample ID: ‘098f6f318b3ea9531448’)	TACRED-RE	NYT-adapted models

Table 21: Common Misclassification Patterns Across TACRED-RE, NYT, and Biographical

Relation	NYT	TACRED-RE	Biographical
None	year, york, united, mr, like, states, president, company, work, city	year, national, president, group, include, state, percent, million, american, china	release, contract, announce, song, star, award, series, role, sign, championship
children	father, son, higgins, clark, favre, richard, mary, daughter, carol, daley	son, daughter, grandchild, survive, wife, year, child, include, gude, jr	daughter, son, child, li, father, mother, wife, give, marry, actor
religion	islam, muhammad, prophet, religion, convert, leader, school, al, church, close	jewish, al, islam, shiite, christian, group, muslim, sunni, mohammed, tantawi	–
place_lived	senator, republican, state, year, representative, gov, democrat, city, john, mr	year, state, die, home, york, city, president, live, iran, old	–
place_of_birth	city, year, orleans, chicago, bear, bill, attorney, general, mr, california	bear, grow, family, child, york, year, native, july, old, son	bear, raise, née, grow, family, youth, york, california, city, mother
place_of_death	die, year, home, city, london, los, angeles, mr, yesterday, paris	die, home, hospital, cancer, paris, wednesday, sunday, find, early, dead	die, home, paris, age, california, near, october, london, live, move

Table 22: Top 10 tokens per overlapping relation in NYT, TACRED-RE, and Biographical datasets, following lemmatisation and stop word removal using spaCy’s transformer-based en_core_web_trf model.