

Maverick: Efficient and Accurate Coreference Resolution Defying Recent Trends

Anonymous ACL submission

Abstract

Large autoregressive generative models have emerged as the cornerstone for achieving the highest performance across several Natural Language Processing tasks. However, the urge to attain superior results has, at times, led to the premature replacement of carefully designed task-specific approaches without exhaustive experimentation. The Coreference Resolution task is no exception; all recent state-of-the-art solutions adopt large generative autoregressive models that outperform encoder-based discriminative systems. In this work, we challenge this recent trend by introducing Maverick, a carefully designed – yet simple – pipeline, which enables running a state-of-the-art Coreference Resolution system within the constraints of an academic budget, outperforming models with up to 13 billion parameters with as few as 500 million parameters. Maverick achieves state-of-the-art performance on the CoNLL-2012 benchmark, training with up to 0.006x the memory resources and obtaining a 170x faster inference compared to previous state-of-the-art systems. We extensively validate the robustness of the Maverick framework with an array of diverse experiments, reporting improvements over prior systems in data-scarce, long-document, and out-of-domain settings. We release our code and models for research purposes at [omitted.link](#).

1 Introduction

As one of the core tasks in Natural Language Processing, Coreference Resolution aims to identify and group expressions (called mentions) that refer to the same entity (Karttunen, 1969). Given its crucial role in various downstream tasks, such as Knowledge Graph Construction (Li et al., 2020), Entity Linking (Kundu et al., 2018; Agarwal et al., 2022), Question Answering (Dhingra et al., 2018; Dasigi et al., 2019; Bhattacharjee et al., 2020; Chen and Durrett, 2021), Machine Translation

(Stojanovski and Fraser, 2018; Voita et al., 2018; Ohtani et al., 2019) and Text Summarization (Falke et al., 2017; Pasunuru et al., 2021; Liu et al., 2021), *inter alia*, there is a pressing need for both performance and efficiency. However, recent works in Coreference Resolution either explore methods to obtain reasonable performance optimizing time and memory efficiency (Kirstain et al., 2021; Dobrovolskii, 2021; Otmazgin et al., 2022), or strive to improve benchmark scores regardless of the increased computational demand (Bohnet et al., 2023; Zhang et al., 2023).

Efficient solutions usually rely on discriminative formulations, frequently employing the mention-antecedent classification method proposed by Lee et al. (2017). These approaches leverage relatively small encoder-only transformer architectures (Joshi et al., 2020; Beltagy et al., 2020) to encode documents and build on top of them task-specific networks that ensure high speed and efficiency. On the other hand, performance-centered solutions are nowadays dominated by general-purpose large Sequence-to-Sequence models (Liu et al., 2022; Zhang et al., 2023). A notable example of this formulation, and currently the state of the art in Coreference Resolution, is Bohnet et al. (2023), which proposes a transition-based system that incrementally builds clusters of mentions by generating coreference links sentence by sentence in an autoregressive fashion. Although these solutions achieve remarkable performance, their autoregressive nature and the size of the underlying language models (up to 13B parameters) make them dramatically slower and memory-demanding compared to traditional encoder-only approaches. This not only makes their usage for downstream applications impractical but also poses a significant barrier to their accessibility for a large number of users operating within an academic budget.

This work argues that discriminative encoder-only approaches for Coreference Resolution have

still not expressed their full potential and have been discarded too early in the urge to achieve state-of-the-art performance. By proposing Maverick, we strike an optimal balance between high performance and efficiency, a combination that was missing in previous systems. Our framework enables an encoder-only model to achieve top-tier performance while keeping the overall model size less than one-twentieth of the current state-of-the-art system, and training it with academic resources. Moreover, when further reducing the size of the underlying transformer encoder, Maverick performs in the same ballpark as encoder-only efficiency-driven solutions while improving speed and memory consumption. Finally, we propose a novel incremental Coreference Resolution method that, integrated into the Maverick framework, results in a robust architecture for out-of-domain, data-scarce, and long-document settings.

2 Related Work

We now introduce well-established approaches to neural Coreference Resolution. In particular, we first delve into the details of traditional discriminative solutions, including their incremental variations, and then present the recent paradigm shift for approaches based on large generative architectures.

2.1 Discriminative models

Discriminative approaches tackle the Coreference Resolution task as a classification problem, usually employing encoder-only architectures. The pioneering works of Lee et al. (2017, 2018) introduced the Coarse-to-Fine model, the first end-to-end discriminative system for Coreference Resolution. First, it involved a mention extraction step, in which the spans most likely to be coreference mentions are identified. This is followed by a mention-antecedent classification step where, for each extracted mention, the model searches for its most probable antecedent (i.e. the extracted span that appears before in the text). This pipeline, composed of mention extraction and mention-antecedent classification steps, has been adopted with minor modifications in many subsequent works, that we refer to as *Coarse-to-Fine* models.

Coarse-to-Fine Models Among the works that build upon the Coarse-to-Fine formulation, Lee et al. (2018), Joshi et al. (2019) and Joshi et al. (2020) experimented with changing the underlying document encoder, utilizing ELMo (Peters et al.,

2018), BERT (Devlin et al., 2019) and SpanBERT (Joshi et al., 2020) respectively, achieving remarkable score improvements on the English OntoNotes (Pradhan et al., 2012). Similarly, Kirstain et al. (2021) introduced s2e-coref that reduces the high memory footprint of SpanBERT leveraging the Longformer (Beltagy et al., 2020) sparse-attention mechanism. Based on the same architecture, Otmazgin et al. (2023) analyzed the impact of having multiple experts scoring different linguistically motivated categories (e.g., pronouns-nouns, nouns-nouns, etc.). While these works have been able to modernize the original Coarse-to-Fine formulation, training those architectures on the OntoNotes dataset still requires a considerable amount of memory.¹ This occurs because they rely on the traditional Coarse-to-Fine pipeline that, as we will cover in Section 3.1, has a large memory overhead and is based on manually-set thresholds to regulate memory usage.

Incremental Models Discriminative systems also include incremental techniques. Incremental Coreference Resolution has a strong cognitive grounding: research on the “garden-path” effect shows that humans resolve referring expressions incrementally (Altmann and Steedman, 1988).

A seminal work that proposed an incremental automatic system is Webster and Curran (2014), which introduced a clustering approach based on the shift-reduce paradigm. In this formulation, for each mention, a classifier decides whether to SHIFT it into a singleton (i.e. single mention cluster) or to REDUCE it within an existing cluster. The same approach has recently been reintroduced in ICoref (Xia et al., 2020) and longdoc (Toshniwal et al., 2021), which adopted SpanBERT and Longformer respectively. In these works the mention extraction step is identical to that of Coarse-to-Fine models. On the other hand, the mention clustering step is performed by using a linear classifier that scores each mention against a vector representation of previously built clusters, in an incremental fashion. Since cluster representations are updated with a learnable function, this method ensures constant memory usage. In Section 3.2 we present a novel performance-driven incremental method that obtains superior performance and generalization capabilities, in which we adopt a lightweight transformer architecture that retains the mention representations.

¹Training those models requires at least 32G of VRAM.

2.2 Sequence-to-Sequence models

Recent state-of-the-art Coreference Resolution systems all employ autoregressive generative approaches. However, an early example of Sequence-to-Sequence model, TANL (Paolini et al., 2021), failed to achieve competitive performance on OntoNotes. The first system to show that the autoregressive formulation was competitive is ASP (Liu et al., 2022), which outperformed encoder-only discriminative approaches. ASP is an autoregressive pointer-based model that first generates actions for mention extraction (bracket pairing) and then conditions the next step to generate coreference links. Notably, the breakthrough in ASP does not lie only in its novel formulation but in the usage of large generative models. Indeed, the success of their approach is strictly correlated with the underlying model size, since, when using models with a comparable number of parameters, the final performance is significantly lower than encoder-only approaches. The same occurs in Zhang et al. (2023), a fully-seq2seq approach where a model learns to generate a formatted sequence encoding coreference notation, in which they report a strong positive correlation between performance and model sizes.

Finally, the current state-of-the-art system on the OntoNotes benchmark is held by Link-Append (Bohnet et al., 2023), a transition-based system that incrementally builds clusters exploiting a multi-pass Sequence-to-Sequence architecture. This approach incrementally maps the mentions in previously coreference-annotated sentences to system actions for the current sentence, using the same shift-reduce incremental paradigm presented in Section 2.1. This method obtains state-of-the-art performance at the cost of using a 13B parameters model and processing one sentence at a time, drastically increasing the need for computational power. While these models ensure superior performance compared to previous discriminative approaches, using them for inference is out of reach for many users, not to mention training them from scratch.

3 Methodology

In this section, we present the Maverick framework. We propose to replace the preprocessing and training strategy of Coarse-to-Fine models with the Maverick Pipeline, improving the training and inference efficiency of Coreference Resolution systems. Furthermore, with the Maverick Pipeline, we eliminate the dependency on long-

standing manually-set hyperparameters that regulate memory usage. Finally, building on top of the Maverick Pipeline, we propose three models that adopt a mention-antecedent classification technique, namely Maverick_{s2e} and Maverick_{mes}, and a system that is based upon a novel incremental formulation, Maverick_{incr}.

3.1 Maverick Pipeline

The Maverick Pipeline is a combination of i) an efficient mention extraction method, ii) a novel mention regularization technique, and iii) a new mention pruning strategy.

Mention Extraction When it comes to extracting mentions from a document D , there are different strategies to model the probability that a span contains a mention. Several previous works follow the Coarse-to-Fine formulation presented in Section 2.1, which consists of scoring all the possible spans in D . This implies a quadratic computational cost with respect to the input length, which they mitigate by introducing several pruning techniques.

In this work, we employ a different strategy. We extract coreference mentions by first identifying all the possible starts of a mention, and then, for each start, extracting its possible end. To extract start indices, we first compute the hidden representation $(x_1, \dots, x_n)$ of the tokens $(t_1, \dots, t_n) \in D$ using a transformer encoder, and then use a fully-connected layer F to compute the probability for each t_i being the start of a mention as:

$$F_{start}(x) = W'_{start}(GeLU(W_{start}x))$$

$$p_{start}(t_i) = \sigma(F_{start}(x_i))$$

With W'_{start} , W_{start} being the learnable parameters, and σ the sigmoid function. For each start of a mention t_s , i.e. those tokens having $p_{start}(t_s) > 0.5$, we then compute the probability of its subsequent tokens t_j , with $s \leq j$, to be the end of a mention that starts with t_s . We follow the same process of the mention start classification, but we condition the prediction on the starting token by concatenating the start, x_s , and end, x_j , hidden representations before the linear classifier:

$$F_{end}(x, x') = W'_{end}(GeLU(W_{end}[x, x']))$$

$$p_{end}(t_j|t_s) = \sigma(F_{end}(x_s, x_j))$$

With W'_{end} , W_{end} being learnable parameters. This formulation considers overlapping mentions, since

for each start t_s we can find multiple t_e (i.e. those that have $p_{end}(t_j|t_s) > 0.5$) and also reduces 9 times the number of considered mentions compared to the Coarse-to-Fine pipeline (Table 1).

To further reduce the computation demand of this process, in the Maverick Pipeline we introduce the end-of-sentence (EOS) mention regularization strategy: after extracting the span start, we only consider the tokens up to the nearest EOS as possible mention end candidates.² Since annotated mentions never span across sentences, EOS mention regularization can efficiently consider all the possible spans in a document. In contrast, previous Coarse-to-Fine formulations rely on a manually-set hyperparameter that regulates maximum span length. This implies a large overhead of unnecessary computations and ignores mentions that exceed a fixed length.³

Mention Pruning After the mention extraction step, as a result of the Maverick Pipeline, we consider an 18x lower number of candidate mentions for the successive mention clustering phase (Table 1). This step consists of computing, for each mention, the probability of all its antecedents being in the same cluster, incurring a quadratic computational cost. Within the Coarse-to-Fine formulation, this high computational cost is mitigated by considering only the top k mentions according to their probability score, where k is a manually set hyperparameter. Since we obtain probabilities for a very concise number of mentions, we consider only predicted mentions (i.e. those with $p_{end} > 0.5$ and $p_{start} > 0.5$), reducing the number of considered mention-pairs by a factor of 10. In Table 1, we compare the previous Coarse-to-Fine formulation with the new Maverick Pipeline.

3.2 Mention Clustering

As a result of the Maverick Pipeline, we obtain a set of candidate mentions $M = (m_1, m_2, \dots, m_l)$, for which we propose three different clustering techniques: $Maverick_{s2e}$ and $Maverick_{mes}$, which follow the traditional Coarse-to-Fine mention-antecedent formulation, and $Maverick_{incr}$, which adopts a novel incremental technique that leverages a light transformer architecture.

²We note that all the well-established Coreference Resolution datasets are sentence-splitted.

³In previous works, max-length regularization filters out 196 correctly annotated spans when training on OntoNotes.

	Coarse-to-fine	Maverick	Δ
Ment. Extraction	Enumeration 183,577	(i) Start-End 20,565	-8,92x
(+) Regularization	(+) Span-length 14,265	(ii) (+) EOS 777	-18,3x
Ment. Clustering	Top-k 29,334	(iii) Pred-only 2,713	-10,81x

Table 1: Comparison between the Coarse-to-Fine pipeline and the Maverick Pipeline in terms of the average number of considered mentions in the mention extraction step (top) and the average number of considered mention-pairs in the mention clustering step (bottom). The statistics are computed on the OntoNotes devset, and refer to the hyperparameters proposed in (Lee et al., 2018), which were unchanged by subsequent Coarse-to-Fine works, i.e. span-len = 30, top-k = 0.4.

Mention-Antecedent models The first proposed model, $Maverick_{s2e}$, adopts a similar mention clustering strategy to Kirstain et al. (2021): given a mention $m_i = (x_s, x_e)$ and its antecedent $m_j = (x_{s'}, x_{e'})$, with their start and end token hidden states, we use two fully-connected layers to model their corresponding representations:

$$F_s(x) = W'_s(GeLU(W_s x))$$

$$F_e(x) = W'_e(GeLU(W_e x))$$

We then calculate their probability to be in the same cluster as:

$$p_c(m_i, m_j) = \sigma(F_s(x_s) \cdot W_{ss} \cdot F_s(x_{s'}) + F_e(x_e) \cdot W_{ee} \cdot F_e(x_{e'}) + F_s(x_s) \cdot W_{se} \cdot F_e(x_{e'}) + F_e(x_e) \cdot W_{es} \cdot F_s(x_{s'}))$$

With $W_{ss}, W_{ee}, W_{se}, W_{es}$ being four learnable matrices and W_s, W'_s, W_e, W'_e the learnable parameters of the two fully connected layers.

A similar formulation is adopted in $Maverick_{mes}$, where, instead of using only one generic mention-pair scorer, we use 6 different scorers that handle linguistically motivated categories, as introduced by Otmazgin et al. (2023). We detect which category k a pair of mentions m_i and m_j belongs to (e.g., if m_i is a pronoun and m_j is a proper noun, the category will be PRONOUN-ENTITY) and use a category-specific scorer to compute p_c . A complete description of the process along with the list of categories can be found in Appendix A.

Incremental model Finally, we introduce a novel approach to tackle the mention clustering step, namely $Maverick_{incr}$, which incrementally builds

clusters following the shift-reduce paradigm introduced in Section 2.1. In `Maverickincr`, in contrast to previous incremental techniques, we leverage a lightweight transformer model to attend to previous clusters, for which we retain the mentions hidden representations. Specifically, we compute the hidden representations $(h_1, \dots, h_l)$ for all the candidate mentions in M using a fully-connected layer on top of the concatenation of their start and end token representations. We first assign the first mention m_0 to the first cluster $c_0 = (m_0)$. Then, for each mention $m_i \in M$ at step i we obtain the probability of m_i to be in a certain cluster c_j by encoding h_i with all the representations of the mentions contained in the cluster c_j using a transformer architecture. In particular, we use the first special token ([CLS]) of a single-layer transformer architecture T to obtain the score $S(m_i, c_j)$ of m_i being in the cluster $c_j = (m_f, \dots, m_g)$ with $f \leq g < i$ as:

$$S(m_i, c_j) = (W_c \cdot (ReLU(T_{CLS}(h_i, h_f, \dots, h_g)))$$

Finally, we compute the probability of m_i to belong to c_j as:

$$p_c(m_i \in c_j | (m_f, \dots, m_g) \in c_j) = \sigma(S(m_i, c_j))$$

We compute this probability for each cluster c_j computed up to step i . We assign the mention m_i to the most probable cluster c_j having $p_c(m_i \in c_j) > 0.5$ if one exists, or we create a new singleton cluster containing m_i .

As we show in Section 5.3 and in Section 5.5, this formulation obtains better results than previous incremental methods, and is particularly beneficial when dealing with long-document and out-of-domain settings.

3.3 Training

To train a `Maverick` model, we optimize the sum of three binary cross-entropy losses:

$$L_{coref} = L_{start} + L_{end} + L_{clust}$$

L_{start} , L_{end} comes from the mention extraction step and L_{clust} from mention clustering. All the models we introduce are trained using teacher forcing. In particular, in the mention token end classification step, we use gold start indices to condition the end tokens prediction, and, for the mention clustering step, we consider only gold mention indices. For `Maverickincr`, at each iteration, we compare each mention only to previous gold clusters.

Dataset	# Train	# Dev	# Test	Tokens	Mentions	% Sing
OntoNotes	2802	343	348	467	56	0
LitBank	80	10	10	2105	291	19.8
PreCo	36120	500	500	337	105	52.0
GAP	-	-	2000	95	3	-
WikiCoref	-	-	30	1996	230	0

Table 2: Datasets statistics: number of documents in each dataset split, the average number of words and mentions per document, and the singletons percentage.

4 Experiments Setup

4.1 Datasets

We train and evaluate all the comparison systems on three Coreference Resolution datasets:

OntoNotes (Pradhan et al., 2012), proposed in the CoNLL-2012 shared task, is the de facto standard dataset used to benchmark Coreference Resolution systems. It consists of documents that span seven distinct genres, including full-length documents (broadcast news, newswire, magazines, weblogs, and Testaments) and multiple speaker transcripts (broadcast and telephone conversations).

LitBank (Bamman et al., 2020) contains 100 literary documents typically used to evaluate long-document Coreference Resolution.

PreCo (Chen et al., 2018) is a large-scale dataset that includes reading comprehension tests for middle school and high school students.

Notably, both LitBank and PreCo have different annotation guidelines compared to OntoNotes, and provide annotation for singletons (i.e. clusters one mention). Furthermore, we evaluate models trained on OntoNotes on three out-of-domain datasets:

- **GAP** (Webster et al., 2018) contains sentences in which, given a pronoun, the model has to choose between two candidate mentions.
- **LitBank_{ns}** and **PreCo_{ns}**, the datasets’ test-set where we filter out singleton annotations.
- **WikiCoref** (Ghaddar and Langlais, 2016), which contains Wikipedia texts, including documents with up to 9,869 tokens.

Employed dataset statistics are shown in Table 2.

4.2 Comparison Systems

Discriminative Among the discriminative systems, we consider c2f-coref (Joshi et al., 2020) and s2e-coref (Kirstain et al., 2021), which build upon the Coarse-to-Fine formulation and adopt different

document encoders. We also report the results of LingMess (Otmazgin et al., 2023), which is the previous best encoder-only solution, and f-coref (Otmazgin et al., 2022), which is a distilled version of LingMess. Furthermore, we include CorefQA (Wu et al., 2020), which casts Coreference as extractive Question Answering, and wl-coref (Dobrovolskii, 2021), which first predicts coreference links between words, then extracts mentions spans. Finally, we report the results of incremental systems, such as ICoref (Xia et al., 2020) and longdoc (Toshniwal et al., 2021).

Sequence-to-Sequence We compare our models with TANL (Paolini et al., 2021) and ASP (Liu et al., 2022), which frame Coreference Resolution as autoregressive structured prediction. We also include Link-Append (Bohnet et al., 2023), a transition-based system that builds clusters with a multi-pass Sequence-to-Sequence architecture. Finally, we report the results of seq2seq (Zhang et al., 2023), a model that learns to generate a sequence with Coreference Resolution labels.

4.3 Maverick Setup

All Maverick models use DeBERTa-v3 (He et al., 2023) as the document encoder. We use DeBERTa because it can model very long input texts⁴, and has shown to be effective in handling long sequences (He et al., 2021). On the other hand, using it to encode long documents is computationally expensive because its attention mechanism implies a quadratic computational complexity. While this further increases the computational cost of traditional Coarse-to-Fine systems, the Maverick Pipeline enables us to train models that leverage DeBERTa_{large} on the OntoNotes dataset, without any performance-lowering pruning heuristic. To train our models we use Adafactor (Shazeer and Stern, 2018) as our optimizer, with a learning rate of 3e-4 for the linear layers, and 2e-5 for the pre-trained encoder. We perform all our experiments within an academic budget, i.e. a single RTX 4090 which has 24GB of VRAM. We report more training details in Appendix B.

5 Results

5.1 English OntoNotes

We report in Table 3 the average CoNLL-F1 score of the comparison systems trained on the English

OntoNotes, along with their underlying pre-trained language models and total parameters. Compared to previous discriminative systems, we report gains of +2.2 CoNLL-F1 points over LingMess, the best encoder-only model. Interestingly, we outperform CorefQA as well, which takes advantage of training on additional Question Answering data.

Concerning Sequence-to-Sequence approaches, we report extensive improvements over systems with a similar amount of parameters compared to our large models (500M): we obtain +3.4 points with respect to ASP (770M), and the gap is even wider when taking into consideration Link-Append (3B) and seq2seq (770M), with +6.4 and +5.6, respectively. Most importantly, Maverick models surpass the performance of all sequence-to-sequence transformers even when they have several billions of parameters. Among our proposed methods, Maverick_{mes} shows the best performance, setting a new state of the art with a score of 83.6 CoNLL-F1 points on the OntoNotes benchmark. More detailed results, including a table with MUC, B³, and CEAF_{φ4} scores and an error analysis, can be found in Appendix C.

5.2 PreCo and LitBank

We further validate the robustness of the Maverick framework by training and evaluating systems on the PreCo and LitBank datasets. As reported in Table 4, our models show superior performance when dealing with long documents in a data-scarce setting such as the one LitBank poses. On this dataset, Maverick_{incr} achieves a new state-of-the-art score of 78.3, and gains +1.0 CoNLL-F1 points compared with seq2seq. On PreCo, Maverick_{incr} outperforms longdoc, but seq2seq still shows slightly better performance. Among our systems, Maverick_{incr}, leveraging its hybrid architecture, performs better on both PreCo and LitBank.

5.3 Out-of-Domain Evaluation

In Table 5, we report the performance of Maverick systems along with LingMess, the best encoder-only model, when dealing with out-of-domain texts, that is when they are trained on OntoNotes and tested on other datasets. First of all, we report considerable improvements on the GAP test set, obtaining a +1.2 F1 score with respect to the previous state of the art. We also test models on WikiCoref, PreCo_{ns} and LitBank_{ns} (Section 4.1). However, since the span annotation guidelines of these corpora differ from the ones used in OntoNotes, in

⁴This is because its attention mechanism enables its input length to grow linearly with the number of its layers.

Model	LM	Avg. F1	Params	Training		Inference	
				Time	Hardware	Time	Mem.
Discriminative							
c2f-coref (Joshi et al., 2020)	SpanBERT _{large}	79.6	-	-	1x32GB	50s	11.9
ICoref (Xia et al., 2020)	SpanBERT _{large}	79.4	377M	40h	1x1080TI-12GB	38s	2.9
CorefQA (Wu et al., 2020)	SpanBERT _{large}	83.1*	-	-	1xTPUV3-128G	-	-
s2e-coref (Kirstain et al., 2021)	LongFormer _{large}	80.3	494M	-	1x32G	17s	3.9
longdoc (Toshniwal et al., 2021)	LongFormer _{large}	79.6	-	16h	1xA6000-48G	25s	2.1
wl-coref (Dobrovolskii, 2021)	RoBERTa _{large}	81.0	360M	5h	1xRTX8000-48G	11s	2.3
f-coref (Otmazgin et al., 2022)	DistilRoBERTa	78.5*	91M	-	1xV100-32G	3s	1.0
LingMess (Otmazgin et al., 2023)	LongFormer _{large}	81.4	590M	23h	1xV100-32G	20s	4.8
Sequence-to-Sequence							
ASP (Liu et al., 2022)	FLAN-T5L	80.2	770M	-	1xA100-40G	-	-
	FLAN-T5xxl	82.5	11B	45h	6xA100-80G	20m	-
Link-Append (Bohnet et al., 2023)	mT5xl	78.0 ^d	3B	-	128xTPUV4-32G	-	-
	mT5xxl	83.3	13B	48h	128xTPUV4-32G	30m	-
seq2seq (Zhang et al., 2023)	T5-large	77.2 ^d	770M	-	8xA100-40G	-	-
	T0-11B	83.2	11B	-	8xA100-80G	40m	-
Ours (Discriminative)							
Maverick _{s2e}	DeBERTa _{base}	81.1	192M	7h	1xRTX4090-24G	6s	1.8
	DeBERTa _{large}	83.4	449M	14h	1xRTX4090-24G	13s	4.0
Maverick _{incr}	DeBERTa _{base}	81.0	197M	21h	1xRTX4090-24G	22s	1.8
	DeBERTa _{large}	83.5	452M	29h	1xRTX4090-24G	29s	3.4
Maverick _{mes}	DeBERTa _{base}	81.4	223M	7h	1xRTX4090-24G	6s	1.9
	DeBERTa _{large}	83.6	504M	14h	1xRTX4090-24G	14s	4.0

Table 3: Results on the OntoNotes benchmark. We report the Avg. CoNLL-F1 score, the number of parameters, the training time, and the hardware used to train each model. Inference time (sec) and memory (GiB) were calculated on an RTX4090. For Sequence-to-Sequence models we include statistics that are reported in the original papers, since we could not run models locally. (*) indicates models trained on additional resources. (^d) indicates scores obtained on the development set, however, Maverick systems perform always better on the development than on the test sets.

Model	PreCo	LitBank
longdoc (Toshniwal et al., 2021)	87.8	77.2
seq2seq (Zhang et al., 2023)	88.5	77.3
Maverick _{s2e}	87.2	77.6
Maverick _{incr}	88.0	78.3
Maverick _{mes}	87.4	78.0

Table 4: Results on the PreCo and LitBank test-sets.

Table 5 we also report the performance using gold mentions, i.e. skipping the mention extraction step (gold column).⁵ On the WikiCoref benchmark, we achieve a new state-of-the-art score of 67.2 CoNLL-F1, with an improvement of +4.2 points over the previous best score obtained by LingMess. On the same dataset, when using pre-identified mentions the gap increases to +5.8 CoNLL-F1 points (76.6 vs 82.4). In the same setting, our models obtain up to +7.3 and +10.1 CoNLL-F1 points on PreCo_{ns} and LitBank_{ns} compared to LingMess. These results suggest that Maverick training strategy makes it more suitable when dealing with pre-identified mentions and out-of-domain texts. This further in-

⁵We do not include autoregressive models because none of the original articles report scores on out-of-domain datasets. We could not test those models either, because they do not provide the code to perform mention clustering alone, and this methodology is not as clear as it is in encoder-only models.

creases the potential benefits that Maverick systems can bring to many downstream applications that exploit coreference as an intermediate layer, such as Entity Linking (Rosales-Méndez et al., 2020) and Relation Extraction (Xiong et al., 2023; Zeng et al., 2023), where the mentions are already identified. Among our models, on LitBank_{ns} and WikiCoref, Maverick_{incr} outperforms Maverick_{mes} and Maverick_{s2e}, confirming the superior capabilities of the incremental formulation in the long document setting. On a final note, we highlight that the performance gap between using gold mentions and performing full Coreference Resolution is wider when tested on out-of-domain datasets (on average +17) compared to testing it directly on OntoNotes (83.6 vs 93.6, +10).⁶ This result, obtained on three different out-of-domain datasets, confirms that the difference in annotation guidelines considerably contributes to lower OOD performances (7%).

5.4 Speed and Memory Usage

In Table 3, we include details regarding the training time and the hardware used by each comparison system, along with the measurement of the inference time and peak memory usage on the

⁶More on this evaluation can be found in Appendix C.

Model	GAP	WikiCoref		PreCo _{ns}		LitBank _{ns}	
		sys.	gold	sys.	gold	sys.	gold
LingMess	89.6	63.0	76.6	65.1	80.6	64.4	73.9
Maverick _{s2e}	91.1	67.2	81.5	67.2	87.9	64.8	83.1
Maverick _{incr}	91.2	66.8	82.4	66.1	86.5	65.4	84.0
Maverick _{mes}	91.1	66.8	82.1	66.1	86.9	65.1	82.8

Table 5: Comparison between LingMess and Maverick systems on GAP, WikiCoref, PreCo_{ns} LitBank_{ns}. We report scores using systems prediction (sys.) or passing gold mentions (gold).

development set. Compared to Coarse-to-Fine models, which require 32GB of VRAM, we can train Maverick systems under 18GB. At inference time both Maverick_{mes} and Maverick_{s2e}, exploiting DeBERTa_{large}, achieve competitive speed and memory consumption compared to wl-coref and s2e-coref. Furthermore, when adopting DeBERTa_{base}, Maverick_{mes} proves to be the most efficient approach⁷ among those directly trained on OntoNotes, while, at the same time, obtaining performance that are equal to the previous best encoder-only system, LingMess. The only system that shows better inference speed is f-coref, but at the cost of lower performance (-3.0).

With respect to the previous Sequence-to-Sequence state-of-the-art approach, Link-Append, we train our models with 175x less memory requirements. Comparing inference time is more complicated since we could not run models on our memory-constrained budget. For this reason, we report the inference times from the original articles, hence achieved with their high-resource settings. Interestingly, we report as much as 170x faster inference compared to seq2seq, which exploits parallel inference on multiple GPUs, and 85x faster when compared to the more efficient ASP. Among Maverick models, Maverick_{incr} is notably slower both in inference and training time, as it incrementally builds clusters using multiple steps.

5.5 Maverick Ablation

In Table 6 we compare Maverick_{s2e} and Maverick_{mes} models with s2e-coref and LingMess respectively, using different pretrained encoders. Interestingly, when using DeBERTa, Maverick systems not only achieve better speed and memory efficiency but also obtain higher performance compared to the previous systems. When using the LongFormer, instead, their scores are in the same ballpark, suggesting that the Maverick training pro-

⁷In terms of inference peak memory usage and speed.

Model	LM	Score
Maverick_{s2e}		
Maverick _{s2e}	DeBERTa _{base}	81.0
s2e-coref _t	DeBERTa _{base}	78.3
Maverick _{s2e}	LongFormer _{large}	80.6
s2e-coref	LongFormer _{large}	80.3
Maverick_{mes}		
Maverick _{mes}	DeBERTa _{base}	81.4
LingMess _t	DeBERTa _{base}	78.6
Maverick _{mes}	LongFormer _{large}	81.0
LingMess	LongFormer _{large}	81.4
Maverick_{incr}		
Maverick _{incr}	DeBERTa _{large}	83.5
Maverick _{prev-incr}	DeBERTa _{large}	79.6

Table 6: Comparison between Maverick models and previous techniques. LingMess_t and s2e-coref_t are trained using their official scripts. We use DeBERTa_{base} because the DeBERTa_{large} could not fit in hardware when training comparison systems.

cedure better exploits the capabilities of DeBERTa. To test the benefits of our novel incremental formulation, Maverick_{incr}, we also implement a Maverick model with the previously adopted incremental method used in longdoc and ICoref (Section 2.1), which we call Maverick_{prev-incr}. Compared to the previous formulation we report an increase in score of +3.9 CoNLL-F1 points. The improvement demonstrates that exploiting a transformer architecture to attend to all the previously clustered mentions is beneficial, and enables the future usage of hybrid architectures when needed.

6 Conclusion

In this work, we challenged the recent trends of adopting large autoregressive generative models to solve the Coreference Resolution task. To do so, we proposed Maverick, a new framework that enables fast and memory-efficient Coreference Resolution while obtaining state-of-the-art results. This demonstrates that the large computational overhead required by sequence-to-sequence approaches is unnecessary. Indeed, in our experiments Maverick systems can outperform large generative models and improve the speed and memory usage of previous best-performing encoder-only approaches. Furthermore, we introduced Maverick_{incr}, a robust multi-step incremental technique that obtains superior performances in the out-of-domain and long document setting. By releasing our systems, we will make high-performance models usable by a larger portion of users in different scenarios, and potentially improve downstream applications.

7 Limitations

Our experiments were limited by our resource setting i.e. a single RTX 4090. For this reason, we could not run Maverick using larger encoders, and could not properly test sequence-to-sequence models as we did with encoder-only models. Nevertheless, we believe this limitation is a common scenario in many real-world applications that would substantially benefit from our system. We also did not test our formulation on multiple languages but note that both the methodology behind Maverick and our novel incremental formulation are language agnostic, and thus could be applied to any language.

References

- Dhruv Agarwal, Rico Angell, Nicholas Monath, and Andrew McCallum. 2022. [Entity linking via explicit mention-mention coreference modeling](#). In *Proceedings of the 2022 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, pages 4644–4658, Seattle, United States. Association for Computational Linguistics.
- Gerry Altmann and Mark Steedman. 1988. [Interaction with context during human sentence processing](#). *Cognition*, 30(3):191–238.
- Amit Bagga and Breck Baldwin. 1998. [Entity-based cross-document coreferencing using the vector space model](#). In *36th Annual Meeting of the Association for Computational Linguistics and 17th International Conference on Computational Linguistics, Volume 1*, pages 79–85, Montreal, Quebec, Canada. Association for Computational Linguistics.
- David Bamman, Olivia Lewke, and Anya Mansoor. 2020. [An annotated dataset of coreference in English literature](#). In *Proceedings of the Twelfth Language Resources and Evaluation Conference*, pages 44–54, Marseille, France. European Language Resources Association.
- Iz Beltagy, Matthew E Peters, and Arman Cohan. 2020. Longformer: The long-document transformer. *arXiv preprint arXiv:2004.05150*.
- Santanu Bhattacharjee, Rejwanul Haque, Gideon Maillette de Buy Wenniger, and Andy Way. 2020. [Investigating query expansion and coreference resolution in question answering on bert](#). *Natural Language Processing and Information Systems*, 12089:47 – 59.
- Bernd Bohnet, Chris Alberti, and Michael Collins. 2023. [Coreference resolution through a seq2seq transition-based system](#). *Transactions of the Association for Computational Linguistics*, 11:212–226.

- Hong Chen, Zhenhua Fan, Hao Lu, Alan Yuille, and Shu Rong. 2018. [PreCo: A large-scale dataset in preschool vocabulary for coreference resolution](#). In *Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing*, pages 172–181, Brussels, Belgium. Association for Computational Linguistics.
- Jifan Chen and Greg Durrett. 2021. [Robust question answering through sub-part alignment](#). In *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, pages 1251–1263, Online. Association for Computational Linguistics.
- Pradeep Dasigi, Nelson F. Liu, Ana Marasović, Noah A. Smith, and Matt Gardner. 2019. [Quoref: A reading comprehension dataset with questions requiring coreferential reasoning](#). In *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)*, pages 5925–5932, Hong Kong, China. Association for Computational Linguistics.
- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2019. [BERT: Pre-training of deep bidirectional transformers for language understanding](#). In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, pages 4171–4186, Minneapolis, Minnesota. Association for Computational Linguistics.
- Bhuwan Dhingra, Qiao Jin, Zhilin Yang, William Cohen, and Ruslan Salakhutdinov. 2018. [Neural models for reasoning over multiple mentions using coreference](#). In *Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 2 (Short Papers)*, pages 42–48, New Orleans, Louisiana. Association for Computational Linguistics.
- Vladimir Dobrovolskii. 2021. [Word-level coreference resolution](#). In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, pages 7670–7675, Online and Punta Cana, Dominican Republic. Association for Computational Linguistics.
- Tobias Falke, Christian M. Meyer, and Iryna Gurevych. 2017. [Concept-map-based multi-document summarization using concept coreference resolution and global importance optimization](#). In *Proceedings of the Eighth International Joint Conference on Natural Language Processing (Volume 1: Long Papers)*, pages 801–811, Taipei, Taiwan. Asian Federation of Natural Language Processing.
- Abbas Ghaddar and Phillippe Langlais. 2016. [Wiki-Coref: An English coreference-annotated corpus of Wikipedia articles](#). In *Proceedings of the Tenth International Conference on Language Resources and*

757	<i>Evaluation (LREC'16)</i> , pages 136–142, Portorož,	Manling Li, Alireza Zareian, Ying Lin, Xiaoman Pan,	812
758	Slovenia. European Language Resources Association	Spencer Whitehead, Brian Chen, Bo Wu, Heng Ji,	813
759	(ELRA).	Shih-Fu Chang, Clare Voss, Daniel Napierski, and	814
760	Pengcheng He, Jianfeng Gao, and Weizhu Chen. 2023.	Marjorie Freedman. 2020. GAIA: A fine-grained	815
761	Debertav3: Improving deberta using electra-style pre-	multimedia knowledge extraction system . In <i>Pro-</i>	816
762	training with gradient-disentangled embedding shar-	<i>ceedings of the 58th Annual Meeting of the Associa-</i>	817
763	ing .	<i>tion for Computational Linguistics: System Demon-</i>	818
764	Pengcheng He, Xiaodong Liu, Jianfeng Gao, and	<i>strations</i> , pages 77–86, Online. Association for Com-	819
765	Weizhu Chen. 2021. Deberta: Decoding-enhanced	putational Linguistics.	820
766	bert with disentangled attention .	Tianyu Liu, Yuchen Eleanor Jiang, Nicholas Monath,	821
767	Mandar Joshi, Danqi Chen, Yinhan Liu, Daniel S. Weld,	Ryan Cotterell, and Mrinmaya Sachan. 2022. Autore-	822
768	Luke Zettlemoyer, and Omer Levy. 2020. Span-	gressive structured prediction with language models .	823
769	BERT: Improving pre-training by representing and	In <i>Findings of the Association for Computational</i>	824
770	predicting spans . <i>Transactions of the Association for</i>	<i>Linguistics: EMNLP 2022</i> , pages 993–1005, Abu	825
771	<i>Computational Linguistics</i> , 8:64–77.	Dhabi, United Arab Emirates. Association for Com-	826
772	Mandar Joshi, Omer Levy, Luke Zettlemoyer, and	putational Linguistics.	827
773	Daniel Weld. 2019. BERT for coreference reso-	Zhengyuan Liu, Ke Shi, and Nancy Chen. 2021.	828
774	lution: Baselines and analysis . In <i>Proceedings of</i>	Coreference-aware dialogue summarization . In <i>Pro-</i>	829
775	<i>the 2019 Conference on Empirical Methods in Natu-</i>	<i>ceedings of the 22nd Annual Meeting of the Special</i>	830
776	<i>ral Language Processing and the 9th International</i>	<i>Interest Group on Discourse and Dialogue</i> , pages	831
777	<i>Joint Conference on Natural Language Processing</i>	509–519, Singapore and Online. Association for	832
778	<i>(EMNLP-IJCNLP)</i> , pages 5803–5808, Hong Kong,	Computational Linguistics.	833
779	China. Association for Computational Linguistics.	Xiaoqiang Luo. 2005. On coreference resolution perfor-	834
780	Lauri Karttunen. 1969. Discourse referents . In <i>Inter-</i>	mance metrics . In <i>Proceedings of Human Language</i>	835
781	<i>national Conference on Computational Linguistics</i>	<i>Technology Conference and Conference on Empiri-</i>	836
782	<i>COLING 1969: Preprint No. 70</i> , Sönga Säby, Swe-	<i>cal Methods in Natural Language Processing</i> , pages	837
783	den.	25–32, Vancouver, British Columbia, Canada. Asso-	838
784	Yuval Kirstain, Ori Ram, and Omer Levy. 2021. Coref-	ciation for Computational Linguistics.	839
785	erence resolution without span representations . In	Takumi Ohtani, Hidetaka Kamigaito, Masaaki Nagata,	840
786	<i>Proceedings of the 59th Annual Meeting of the Asso-</i>	and Manabu Okumura. 2019. Context-aware neural	841
787	<i>ciation for Computational Linguistics and the 11th</i>	machine translation with coreference information . In	842
788	<i>International Joint Conference on Natural Language</i>	<i>Proceedings of the Fourth Workshop on Discourse in</i>	843
789	<i>Processing (Volume 2: Short Papers)</i> , pages 14–19,	<i>Machine Translation (DiscoMT 2019)</i> , pages 45–50,	844
790	Online. Association for Computational Linguistics.	Hong Kong, China. Association for Computational	845
791	Gourab Kundu, Avi Sil, Radu Florian, and Wael Hamza.	Linguistics.	846
792	2018. Neural cross-lingual coreference resolution	Shon Otmazgin, Arie Cattan, and Yoav Goldberg. 2022.	847
793	and its application to entity linking . In <i>Proceedings</i>	F-coref: Fast, accurate and easy to use coreference	848
794	<i>of the 56th Annual Meeting of the Association for</i>	resolution . In <i>Proceedings of the 2nd Conference</i>	849
795	<i>Computational Linguistics (Volume 2: Short Papers)</i> ,	<i>of the Asia-Pacific Chapter of the Association for</i>	850
796	pages 395–400, Melbourne, Australia. Association	<i>Computational Linguistics and the 12th International</i>	851
797	for Computational Linguistics.	<i>Joint Conference on Natural Language Processing:</i>	852
798	Kenton Lee, Luheng He, Mike Lewis, and Luke Zettle-	<i>System Demonstrations</i> , pages 48–56, Taipei, Taiwan.	853
799	moyer. 2017. End-to-end neural coreference reso-	Association for Computational Linguistics.	854
800	lution . In <i>Proceedings of the 2017 Conference on</i>	Shon Otmazgin, Arie Cattan, and Yoav Goldberg. 2023.	855
801	<i>Empirical Methods in Natural Language Processing</i> ,	LingMess: Linguistically informed multi expert scor-	856
802	pages 188–197, Copenhagen, Denmark. Association	ers for coreference resolution . In <i>Proceedings of the</i>	857
803	for Computational Linguistics.	<i>17th Conference of the European Chapter of the As-</i>	858
804	Kenton Lee, Luheng He, and Luke Zettlemoyer. 2018.	<i>sociation for Computational Linguistics</i> , pages 2752–	859
805	Higher-order coreference resolution with coarse-to-	2760, Dubrovnik, Croatia. Association for Computa-	860
806	fine inference . In <i>Proceedings of the 2018 Confer-</i>	tional Linguistics.	861
807	<i>ence of the North American Chapter of the Associ-</i>	Giovanni Paolini, Ben Athiwaratkun, Jason Krone,	862
808	<i>ation for Computational Linguistics: Human Lan-</i>	Jie Ma, Alessandro Achille, Rishita Anubhai, Ci-	863
809	<i>guage Technologies, Volume 2 (Short Papers)</i> , pages	cero Nogueira dos Santos, Bing Xiang, and Stefano	864
810	687–692, New Orleans, Louisiana. Association for	Soatto. 2021. Structured prediction as translation be-	865
811	Computational Linguistics.	tween augmented natural languages. <i>arXiv preprint</i>	866
		<i>arXiv:2101.05779</i> .	867

868	Ramakanth Pasunuru, Mengwen Liu, Mohit Bansal, Su-	pages 1264–1274, Melbourne, Australia. Association	925
869	jith Ravi, and Markus Dreyer. 2021. Efficiently sum-	for Computational Linguistics.	926
870	marizing text and graph encodings of multi-document		
871	clusters . In <i>Proceedings of the 2021 Conference of</i>	Kellie Webster and James R. Curran. 2014. Limited	927
872	<i>the North American Chapter of the Association for</i>	memory incremental coreference resolution . In <i>Pro-</i>	928
873	<i>Computational Linguistics: Human Language Tech-</i>	<i>ceedings of COLING 2014, the 25th International</i>	929
874	<i>nologies</i> , pages 4768–4779, Online. Association for	<i>Conference on Computational Linguistics: Technical</i>	930
875	Computational Linguistics.	<i>Papers</i> , pages 2129–2139, Dublin, Ireland. Dublin	931
		City University and Association for Computational	932
876	Matthew E. Peters, Mark Neumann, Mohit Iyyer, Matt	Linguistics.	933
877	Gardner, Christopher Clark, Kenton Lee, and Luke		
878	Zettlemoyer. 2018. Deep contextualized word repre-	Kellie Webster, Marta Recasens, Vera Axelrod, and Ja-	934
879	sentations . In <i>Proceedings of the 2018 Conference of</i>	son Baldrige. 2018. Mind the GAP: A balanced	935
880	<i>the North American Chapter of the Association for</i>	corpus of gendered ambiguous pronouns . <i>Transac-</i>	936
881	<i>Computational Linguistics: Human Language Tech-</i>	<i>tions of the Association for Computational Linguis-</i>	937
882	<i>nologies, Volume 1 (Long Papers)</i> , pages 2227–2237,	<i>tics</i> , 6:605–617.	938
883	New Orleans, Louisiana. Association for Computa-		
884	tional Linguistics.		
		Thomas Wolf, Lysandre Debut, Victor Sanh, Julien	939
885	Sameer Pradhan, Alessandro Moschitti, Nianwen Xue,	Chaumond, Clement Delangue, Anthony Moi, Pier-	940
886	Olga Uryupina, and Yuchen Zhang. 2012. CoNLL-	ric Cistac, Tim Rault, Remi Louf, Morgan Funtow-	941
887	2012 shared task: Modeling multilingual unrestricted	icz, Joe Davison, Sam Shleifer, Patrick von Platen,	942
888	coreference in OntoNotes . In <i>Joint Conference on</i>	Clara Ma, Yacine Jernite, Julien Plu, Canwen Xu,	943
889	<i>EMNLP and CoNLL - Shared Task</i> , pages 1–40, Jeju	Teven Le Scao, Sylvain Gugger, Mariama Drame,	944
890	Island, Korea. Association for Computational Lin-	Quentin Lhoest, and Alexander Rush. 2020. Trans-	945
891	guistics.	formers: State-of-the-art natural language processing .	946
		In <i>Proceedings of the 2020 Conference on Empirical</i>	947
892	Henry Rosales-Méndez, Aidan Hogan, and Barbara	<i>Methods in Natural Language Processing: System</i>	948
893	Poblete. 2020. Fine-grained entity linking. <i>Journal</i>	<i>Demonstrations</i> , pages 38–45, Online. Association	949
894	<i>of Web Semantics</i> , 65:100600.	for Computational Linguistics.	950
895	Noam Shazeer and Mitchell Stern. 2018. Adafactor:	Wei Wu, Fei Wang, Arianna Yuan, Fei Wu, and Jiwei	951
896	Adaptive learning rates with sublinear memory cost .	Li. 2020. CorefQA: Coreference resolution as query-	952
897	In <i>Proceedings of the 35th International Conference</i>	based span prediction . In <i>Proceedings of the 58th</i>	953
898	<i>on Machine Learning</i> , volume 80 of <i>Proceedings</i>	<i>Annual Meeting of the Association for Computational</i>	954
899	<i>of Machine Learning Research</i> , pages 4596–4604.	<i>Linguistics</i> , pages 6953–6963, Online. Association	955
900	PMLR.	for Computational Linguistics.	956
901	Dario Stojanovski and Alexander Fraser. 2018. Corefer-	Patrick Xia, João Sedoc, and Benjamin Van Durme.	957
902	ence and coherence in neural machine translation: A	2020. Incremental neural coreference resolution in	958
903	study using oracle experiments . In <i>Proceedings of the</i>	constant memory . In <i>Proceedings of the 2020 Con-</i>	959
904	<i>Third Conference on Machine Translation: Research</i>	<i>ference on Empirical Methods in Natural Language</i>	960
905	<i>Papers</i> , pages 49–60, Brussels, Belgium. Association	<i>Processing (EMNLP)</i> , pages 8617–8624, Online. As-	961
906	for Computational Linguistics.	sociation for Computational Linguistics.	962
907	Shubham Toshniwal, Patrick Xia, Sam Wiseman, Karen	Yiyun Xiong, Mengwei Dai, Fei Li, Hao Fei, Bobo	963
908	Livescu, and Kevin Gimpel. 2021. On generaliza-	Li, Shengqiong Wu, Donghong Ji, and Chong Teng.	964
909	tion in coreference resolution . In <i>Proceedings of the</i>	2023. Dialogue c+: An extension of dialogue to inves-	965
910	<i>Fourth Workshop on Computational Models of Refe-</i>	tigate how much coreference helps relation extraction	966
911	<i>rence, Anaphora and Coreference</i> , pages 111–120,	in dialogs. In <i>CCF International Conference on Nat-</i>	967
912	Punta Cana, Dominican Republic. Association for	<i>ural Language Processing and Chinese Computing</i> ,	968
913	Computational Linguistics.	pages 222–234. Springer.	969
914	Marc Vilain, John Burger, John Aberdeen, Dennis	Daojian Zeng, Chao Zhao, Chao Jiang, Jianling Zhu,	970
915	Connolly, and Lynette Hirschman. 1995. A model-	and Jianhua Dai. 2023. Document-level relation ex-	971
916	theoretic coreference scoring scheme . In <i>Sixth Mes-</i>	traction with context guided mention integration and	972
917	<i>sage Understanding Conference (MUC-6): Proceed-</i>	inter-pair reasoning. <i>IEEE/ACM Transactions on</i>	973
918	<i>ings of a Conference Held in Columbia, Maryland,</i>	<i>Audio, Speech, and Language Processing</i> .	974
919	<i>November 6-8, 1995</i> .		
		Wenzheng Zhang, Sam Wiseman, and Karl Stratos.	975
920	Elena Voita, Pavel Serdyukov, Rico Sennrich, and Ivan	2023. Seq2seq is all you need for coreference resolu-	976
921	Titov. 2018. Context-aware neural machine trans-	tion. <i>arXiv preprint arXiv:2310.13774</i> .	977
922	lation learns anaphora resolution . In <i>Proceedings</i>		
923	<i>of the 56th Annual Meeting of the Association for</i>		
924	<i>Computational Linguistics (Volume 1: Long Papers)</i> ,		

A Multi-Expert Scorers

In $\text{Maverick}_{\text{mes}}$, the final coreference score between two spans is calculated using 6 linguistically motivated multi-expert scorers. This approach was introduced by [Otmazgin et al. \(2023\)](#), which demonstrated that linguistic knowledge and symbolic computation can still be used to improve results on the OntoNotes benchmark. In $\text{Maverick}_{\text{mes}}$ we adopt this approach on top of the Maverick Pipeline. We use the same set of categories, namely:

1. PRON-PRON-C. Compatible pronouns based on their attributes such as gender or number (e.g. *(I, I)*, *(I, my)* (*she, her*)).
2. PRON-PRON-NC, Incompatible pronouns (e.g. *(I, he)*, *(She, my)*, *(his, her)*).
3. ENT-PRON. Pronoun and non-pronoun (e.g. *(George, he)*, *(CNN, it)*, *(Tom Cruise, his)*).
4. MATCH. Non-pronoun spans with the same content words (e.g. *Italy, Italy*).
5. CONTAINS. One contains the other (e.g. *(Barack Obama, Obama)*).
6. OTHER. The Other pairs.

To detect pronouns we use string match with a full list of English pronouns.

To perform mention clustering, we dedicate a mention-pair scorer for each of those categories. Concretely, for the mention $m_i = (x_s, x_e)$ and its antecedent $m_j = (x_{s'}, x_{e'})$, with their start and end token hidden states, we first detect their category k_g using pattern matching on their spans of texts. Then we compute their start and end representations, using the specific fully connected layers for the category k_g :

$$F_s^{k_g}(x) = W'_{k_g,s}(GeLU(W_{k_g,s}x))$$

$$F_e^{k_g}(x) = W'_{k_g,e}(GeLU(W_{k_g,e}x))$$

The probability $p_c^{k_g}$ of m_i and m_j is then calculated as:

$$p_c^{k_g}(m_i, m_j) = \sigma(F_s^{k_g}(x_s) \cdot W_{ss} \cdot F_s^{k_g}(x_{s'}) + F_e^{k_g}(x_e) \cdot W_{ee} \cdot F_e^{k_g}(x_{e'}) + F_s^{k_g}(x_s) \cdot W_{se} \cdot F_e^{k_g}(x_{e'}) + F_e^{k_g}(x_e) \cdot W_{es} \cdot F_s^{k_g}(x_{s'}))$$

With $W_{ss}, W_{ee}, W_{se}, W_{es}$ being four learnable matrices and $W'_{k_g,e}, W'_{k_g,s}, W_{k_g,e}, W_{k_g,s}$ the learnable parameters of the two fully connected layers. In this way, each mention-pair scorer learns to model the probability for his specific linguistic categories.

B Training details

B.1 Datasets

We report technical details of the adopted datasets.

- **OntoNotes** contains several metadata information for each document such as genre, speakers, and constituent graphs. Following previous works, we incorporate the speaker's name into the text whenever there is a change in speakers for datasets that include this metadata.
- **LitBank** contains 100 literary documents and is available in different 10 different cross-validation folds. Our train, dev, and test splits refer to the first cross-validation fold, LB_0 . We report comparison systems results on the same splits.
- The authors of **PreCo** have not released their official test set. To evaluate consistently our models with previous approaches, we use the official 'dev' split as our test set and retain the last 500 training examples for model validation.

B.2 Setup

All our experiments are developed using the pytorch-lightning framework.⁸ For each Maverick model, we load the pre-trained weights for the *base*⁹ and *large*¹⁰ version of DeBERTA-v3 from the Huggingface Transformers library ([Wolf et al., 2020](#)). We accumulate gradients every 4 steps and use a gradient clipping value of 1.0. We adopt a linear learning rate scheduler a warm-up of 10% of the total steps check validation scores every 50% of the total number of steps per epoch. We select our model upon validation of Avg. CoNLL-f1 score and use a patience of 20.

C Additional Results

In Table 7 we report models performance according to the standard Coreference Resolution metrics:

⁸<https://lightning.ai>

⁹<https://huggingface.co/microsoft/deberta-v3-base>

¹⁰<https://huggingface.co/microsoft/deberta-v3-large>

Model	LM	MUC			B ³			CEAF ϕ_4			Avg. F1
		P	R	F1	P	R	F1	P	R	F1	
Discriminative											
e2e-coref (Lee et al., 2017)	-	78.4	73.4	75.8	68.6	61.8	65.0	62.7	59.0	60.8	67.2
c2f-coref (Lee et al., 2018)	ELMo	81.4	79.5	80.4	72.2	69.5	70.8	68.2	67.1	67.6	73.0
c2f-coref (Joshi et al., 2019)	BERT _{large}	84.7	82.4	83.5	76.5	74.0	75.3	74.1	69.8	71.9	76.9
c2f-coref (Joshi et al., 2020)	SpanBERT _{large}	85.8	84.8	85.3	78.3	77.9	78.1	76.4	74.2	75.3	79.6
ICoref (Xia et al., 2020)	SpanBERT _{large}	85.7	84.8	85.3	78.1	77.5	77.8	76.3	74.1	75.2	79.4
CorefQA (Wu et al., 2020)	SpanBERT _{large}	88.6	87.4	88.0	82.4	82.0	82.2	79.9	78.3	79.1	83.1*
longdoc (Toshniwal et al., 2021)	LongFormer _{large}	85.5	85.1	85.3	78.7	77.3	78.0	74.2	76.5	75.3	79.6
s2e-coref Kirstain et al. (2021)	LongFormer _{large}	86.5	85.1	85.8	80.3	77.9	79.1	76.8	75.4	76.1	80.3
wl-coref (Dobrovolskii, 2021)	RoBERTa _{large}	84.9	87.9	86.3	77.4	82.6	79.9	76.1	77.1	76.6	81.0
f-coref (Otmazgin et al., 2022)	DistilRoberta	85.0	83.9	84.4	77.6	75.5	76.6	74.7	74.3	74.5	78.5*
LingMess (Otmazgin et al., 2023)	LongFormer _{large}	88.1	85.1	86.6	82.7	78.3	80.5	78.5	76.0	77.3	81.4
Sequence-to-Sequence											
TANL (Paolini et al., 2021)	T5 _{base}	-	-	81.0	-	-	69.0	-	-	68.4	72.8
ASP (Liu et al., 2022)	FLAN-T5 _{XXL}	86.1	88.4	87.2	80.2	83.2	81.7	78.9	78.3	78.6	82.5
Link-Append (Bohnet et al., 2023)	mT5 _{XXL}	87.4	88.3	87.8	81.8	83.4	82.6	79.1	79.9	79.5	83.3
seq2seq (Zhang et al., 2023)	T0 _{XXL}	86.1	89.2	87.6	80.6	84.3	82.4	78.9	80.1	79.5	83.2
Ours (Discriminative)											
Maverick _{s2e}	DeBERTa _{large}	87.1	88.6	87.9	81.7	83.8	82.7	80.8	78.7	79.7	83.4
Maverick _{incr}	DeBERTa _{large}	87.6	88.1	87.9	82.7	82.6	82.7	80.3	79.3	79.8	83.5
Maverick _{mes}	DeBERTa _{large}	87.5	88.5	88.0	82.2	83.5	82.8	80.4	79.3	79.9	83.6

Table 7: Results on the OntoNotes test set. The average CoNLL-F1 score of MUC, B³, and CEAF ϕ_4 is the main evaluation criterion. * marks models using additional/different training data.

MUC (Vilain et al., 1995), B³ (Bagga and Baldwin, 1998), CEAF ϕ_4 (Luo, 2005) and AVG CoNLL-F1. Scores for Maverick models are computed using the official CoNLL coreference scorer.¹¹

C.1 Error Analysis

To better understand the quality of Maverick predictions, we conduct an error analysis on our best system trained on OntoNotes, Maverick_{mes}. In table 8, we report the score of performing only mention extraction (F1) or mention clustering with gold mention (CoNLL-F1) with our systems. Our results highlight that our models have strong capabilities of clustering pre-identified mentions, but limited performance in the identification of correct spans. We investigated this phenomenon by conducting a qualitative evaluation of the outputs of our best system, Maverick_{mes}, and found out that OntoNotes contains several annotation errors. We report examples of errors in Table 9. The main inconsistency we found in the gold test set is that many documents have incomplete annotations, which directly correlates with the mention extraction error.

System	Ment. Clustering	Ment. Extraction
Maverick _{s2e}	89.4	93.5
Maverick _{incr}	89.2	94.2
Maverick _{mes}	89.6	93.7

Table 8: Mention extraction (F1) and mention clustering (CoNLL-F1) scores on the OntoNotes development set.

¹¹<https://conll.github.io/reference-coreference-scorers>

Type	Text
Ex. 1	
Gold	Nine people were injured in Gaza when gunmen [opened]₁ [fire]₂ on an Israeli bus. The passengers were off - duty Israeli security workers. Witnesses say [the shots]₂ came from [the Palestinian international airport]₃. Israeli Prime Minister Ehud Barak [closed]₄ down [the two - year - old airport]₃ in response to [the incident]₁. [Palestinians]₅ criticized [the move]₄. [they]₅ regard [the airport]₃ as a symbol of emerging statehood.
Output	[Nine people]₁ were injured in Gaza when gunmen opened fire on an Israeli bus. [The passengers]₁ were off - duty Israeli security workers. Witnesses say the shots came from [the Palestinian international airport]₂. Israeli Prime Minister Ehud Barak [closed]₃ down [the two - year - old airport]₂ in response to the incident. [Palestinians]₄ criticized [the move]₃. [They]₄ regard [the airport]₂ as a symbol of emerging statehood.
Ex. 2	
Gold	[Mr. Seelenfreund]₁ is [executive vice president and chief financial officer of [McKesson]]₃₂- and will continue in [those roles]₂. [PCS]₄ also named Rex R. Malson, 57, executive vice president at McKesson,- as a director, filling the seat vacated by Mr. Field. Messrs. Malson and Seelenfreund are directors of [McKesson, which has an 86% stake in [PCS]]₄₃.
Output	[Mr. Seelenfreund]₁ is [executive vice president and chief financial officer of [McKesson]]₃₂ and will continue in [those roles]₂. [PCS]₄ also named [Rex R. Malson, 57, executive vice president at [McKesson]]₃₅- as a director, filling the seat vacated by Mr. Field. Messrs. [Malson]₅ and [Seelenfreund]₁ are directors of [McKesson, which has an 86 % stake in [PCS]]₄₃.
Ex. 3	
Gold	The Second U.S. Circuit Court of Appeals opinion in the Arcadian Phosphate case - did not repudiate the position [Pennzoil Co.]₁ took in [its]₁ dispute with [Texaco]₂, - contrary to your Sept. 8 article “ Court Backs [Texaco]₂’s View in [Pennzoil]₁ Case – Too Late. ” The fundamental rule of contract law applied to [both cases]₃ was that courts will not enforce - [agreements to [which]]₄ the parties did not intend to be bound₄. In the Pennzoil / Texaco litigation, [the courts]₅ found [Pennzoil]₁ and Getty Oil intended to be bound; in Arcadian Phosphates [they]₅ found there was no intention to be bound.
Output	The Second U.S. Circuit Court of Appeals opinion in [the Arcadian Phosphate case]₁ - did not repudiate the position [Pennzoil Co.]₂ took in [[[its]]₂ dispute with [Texaco]₄₃, - contrary to your Sept. 8 article “ Court Backs [Texaco]’s View in [[Pennzoil] Case]₃₃ - Too Late . ” [[The fundamental rule of contract law]]₅ applied to both cases]₅ was that courts will not enforce - agreements to which the parties did not intend to be bound. In [the [Pennzoil] / [Texaco] litigation]₃, [the courts]₆ found [Pennzoil]₂ and Getty Oil intended to be bound; in [Arcadian Phosphates]₁ [they]₆ found there was no intention to be bound.
Ex. 4	
Gold	... [Harry]₁ has avoided all that by living in a Long Island suburb with [his]₁ wife, who ’s so addicted to soap operas and mystery novels she barely seems to notice when [her husband] disappears for drug - seeking forays into Manhattan.
Output	... [Harry]₁ has avoided all that by living in a Long Island suburb with [[his]]₁ wife, who ’s so addicted to soap operas and mystery novels [she]₂ barely seems to notice when [[her]]₂ husband]₁ disappears for drug - seeking forays into Manhattan]₂.

Table 9: Error examples.