

Incremental Processing in the Age of Non-Incremental Encoders: An Empirical Assessment of Bidirectional Models for Incremental NLU

Brielen Madureira David Schlangen

Computational Linguistics

Department of Linguistics

University of Potsdam, Germany

{madureiralasota, david.schlangen}@uni-potsdam.de

Abstract

While humans process language incrementally, the best language encoders currently used in NLP do not. Both bidirectional LSTMs and Transformers assume that the sequence that is to be encoded is available in full, to be processed either forwards and backwards (BiLSTMs) or as a whole (Transformers). We investigate how they behave under incremental interfaces, when partial output must be provided based on partial input seen up to a certain time step, which may happen in interactive systems. We test five models on various NLU datasets and compare their performance using three incremental evaluation metrics. The results support the possibility of using bidirectional encoders in incremental mode while retaining most of their non-incremental quality. The “omni-directional” BERT model, which achieves better non-incremental performance, is impacted more by the incremental access. This can be alleviated by adapting the training regime (truncated training), or the testing procedure, by delaying the output until some right context is available or by incorporating hypothetical right contexts generated by a language model like GPT-2.

1 Introduction

In “The Story of Your Life”, a science fiction short story by Ted Chiang (2002), Earth is visited by alien creatures whose writing system does not unfold in time but rather presents full thoughts instantaneously. In our world, however, language *does* unfold over time, both in speaking and in writing. There is ample evidence (Marslen-Wilson, 1975; Tanenhaus and Brown-Schmidt, 2008, *inter alia*) that it is also *processed* over time by humans, in an incremental fashion where the interpretation of a full utterance is continuously built up while the utterance is being perceived.

In Computational Linguistics and Natural Language Processing, this property is typically abstracted away by assuming that the unit to be processed (*e.g.*, a sentence) is available as a whole.¹ The return and subsequent mainstreaming of Recurrent Neural Networks (RNNs), originally introduced by Elman (1990) and repopularized *i.a.* by Mikolov et al. (2010), may have made it seem that time had found a place as a first-class citizen in NLP. However, it was quickly discovered that certain technical issues of this type of model could be overcome, for example in the application of machine translation, by encoding input sequences in reverse temporal order (Sutskever et al., 2014).

This turns out to be a special case of the more general strategy of bidirectional processing, proposed earlier in the form of BiRNNs (Schuster and Paliwal, 1997; Baldi et al., 1999) and BiLSTMs (Hochreiter and Schmidhuber, 1997), which combine a forward and a backward pass over a sequence. More recently, Transformers (Vaswani et al., 2017) also function with representations that inherently have no notion of linear order. Atemporal processing has thus become the standard again.

In this paper, we explore whether we can adapt such bidirectional models to work in incremental processing mode and what the performance cost is of doing so. We first go back and reproduce the work of Huang et al. (2015), who compare the performance of LSTMs and BiLSTMs in sequence tagging, extending it with a BERT-based encoder and with a collection of different datasets for tagging and classification tasks. Then we address the following questions:

¹An exception is the field of research on interactive systems, where it has been shown that incremental processing can lead to preferable timing behavior (Aist et al., 2007; Skantze and Schlangen, 2009) and work on incremental processing is ongoing (Žilka and Jurčiček, 2015; Trinh et al., 2018; Coman et al., 2019, *inter alia*).

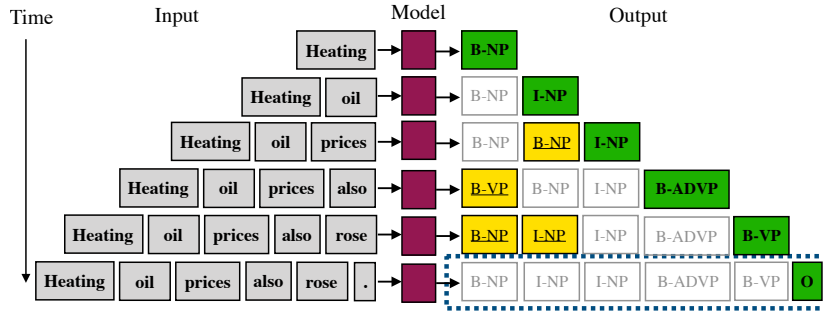


Figure 1: Incremental interface on a bidirectional tagging model (here for chunking). Each line represents the input and output at a time step. Necessary additions are green/bold, substitutions are yellow/underlined, and the dashed frame shows the output of the final time step, which is the same as the non-incremental model’s.

Q1. If we employ inherently non-incremental models in an incremental system, do we get functional representations that are adequate to build correct and stable output along the way? We examine how bidirectional encoders behave under an incremental interface, revisiting the approach proposed by Beuck et al. (2011) for POS taggers. After standard training, we modify the testing procedure by allowing the system to see only successively extended prefixes of the input available so far with which they must produce successively extended prefixes of the output, as shown in Figure 1. The evaluation metrics are described in Section 3, and the discussion is anchored on the concepts of time-liness, monotonicity, and decisiveness and their trade-off with respect to the non-incremental quality (Beuck et al., 2011; Köhn, 2018). We show that it is possible to use them as components of an incremental system (e.g. for NLU) with some trade-offs.

Q2. How can we adapt the training regime or the real-time procedure to mitigate the negative effect that the non-availability of right context (i.e., future parts of the signal) has on non-incremental models? To tackle this question, we implement three strategies that help improve the models’ incremental quality: *truncated training*, *delayed output* and *prophecies* (see Section 4).

Our results are relevant for incremental Natural Language Understanding, needed for the design of dialogue systems and more generally interactive systems, e.g. those following the incremental processing model proposed by Schlangen and Skantze (2011). These systems rely on the availability of partial results, on which fast decisions can be based. Similarly, simultaneous translation is an area where decisions need to be based on partial input with incomplete syntactical and semantic information.

2 Related Work

2.1 Bidirectionality

Language is one of the cognitive abilities that have a temporal nature. The inaugural adoption of RNNs (Elman, 1990) in NLP showed a pursuit to provide connectionist models with a dynamic memory in order to incorporate time implicitly, not as a dimension but through its effects on processing. Since then, the field has witnessed the emergence of a miscellany of neural architectures that take the temporal structure of language into account. In particular, LSTMs (Hochreiter and Schmidhuber, 1997) have been vastly used for sequence-to-sequence or sequence classification tasks, which are ubiquitous in NLP.

Bidirectional LSTMs (Schuster and Paliwal, 1997; Baldi et al., 1999) are an extension to LSTMs that exploit bidirectionality and whose basic processing units are full sentences. They achieved remarkable results in many tasks, e.g. part-of-speech tagging (Ling et al., 2015; Plank et al., 2016), chunking (Zhai et al., 2017), named entity recognition (Chiu and Nichols, 2016), semantic role labeling (He et al., 2017), slot filling and intent detection (E et al., 2019a) and opinion mining (İrsoy and Cardie, 2014). Subsequent works have confirmed that bidirectionality can afford an increase in performance (Graves and Schmidhuber, 2005; Huang et al., 2015; Zhai et al., 2017).

More recently, Vaswani et al. (2017) has consolidated the application of attention mechanisms on NLP tasks with Transformers, which are not constrained by only two directions, as BiLSTMs. Instead, complete sentences are accessed at once. The need for NLP neural networks to be grounded on robust language models and reliable word representations has become clear. The full right and

left context of words started to play a major role as in [Peters et al. \(2018\)](#), which resorts to bidirectionality to train a language model. In a combination of bidirectional word representations with the Transformer architecture, we observe the establishment of BERT ([Devlin et al., 2019](#)) as a current state-of-the-art model, on top of which an output layer can be added to solve classification and tagging tasks.

2.2 Incremental processing

The motivation to build incremental processors, as defined by [Kempen and Hoenkamp \(1982\)](#) and [Levelt \(1989\)](#), is twofold: they are more cognitively plausible and, from the viewpoint of engineering, real-time applications such as parsing ([Nivre, 2004](#)), SRL ([Konstas et al., 2014](#)), NLU ([Peldszus et al., 2012](#)), dialog state tracking ([Trinh et al., 2018](#)), NLG and speech synthesis ([Buschmeier et al., 2012](#)) and ASR ([Selfridge et al., 2011](#)) require that the input be continually evaluated based on incoming prefixes while the output is being produced and updated.

Another advantage is a better use of computational resources, as a module does not have to wait for the completion of another one to start processing ([Skantze and Schlangen, 2009](#)). In robots, linguistic processing must also be intertwined with its perceptions and actions, happening simultaneously ([Brick and Scheutz, 2007](#)).

Research on processing and generating language incrementally has been done long before the current wave of neural network models, using several different methods. For example, in ASR, a common strategy has been to process the input incrementally to produce some initial output, which was then re-scored or re-processed with a more complex model ([Vergyri et al., 2003](#); [Hwang et al., 2009](#)). While the recent accomplishments of neural encoders are cherished, bidirectional encoders drift apart from a desirable temporal incremental approach because they are trained to learn from complete sequences.

There is some cognitive resemblance underlying RNNs in the sense that they can process sequences word-by-word and build intermediary representations at every time step. This feature provides a legitimate way to employ them in incremental systems. [Trinh et al. \(2018\)](#) and [Žilka and Jurčiček \(2015\)](#) explore this, for instance, using the LSTM’s representations to predict dialogue states after each word. Recent works on simul-

taneous translation also use RNNs as incremental decoders ([Dalvi et al., 2018](#)).

Some works arouse interest in the incremental abilities of RNNs. [Hupkes et al. \(2018\)](#) use a diagnostic classifier to analyze the representations that are incrementally built by sequence-to-sequence models in disfluency detection and conclude that the semantic information is only kept encoded for a few steps after it appears in the dialogue, being soon forgotten afterwards. [Ulmer et al. \(2019\)](#) propose three metrics to assess the incremental encoding abilities of LSTMs and compare it with the addition of attention mechanisms.

According to [Beuck et al. \(2011\)](#) and [Schlangen and Skantze \(2011\)](#), incrementality is not a binary feature. Besides using inherently incremental algorithms, it is also possible to provide incremental *interfaces* to non-incremental algorithms. Such interfaces simply feed ever-increasing prefixes to what remains a non-incremental algorithm, providing some “housekeeping” to manage the potentially non-monotonic results.

To alleviate the effect of the partiality of the input, we test the use of anticipated continuations, inspired by the mechanism of predictive processing discussed in cognitive science ([Christiansen and Chater, 2016](#)) and the idea of interactive utterance completion introduced by [DeVault et al. \(2011\)](#). Related strategies to predict upcoming content and to wait for more right context are also applied in recent work on simultaneous translation ([Grissom II et al., 2014](#); [Oda et al., 2015](#); [Ma et al., 2019](#)). The use of truncated inputs during training, discussed below, aims at making intermediate structures available during learning, an issue discussed in [Köhn \(2018\)](#). This is a variation of chunked training used in [Dalvi et al. \(2018\)](#).

3 Evaluation of incremental processors

The hierarchical nature of language makes it likely that incremental processing leads to non-monotonic output due to re-analysis, as in the well-known “garden path” sentences. Incremental systems may edit the output by adding, revoking, and substituting its parts ([Baumann et al., 2011](#)). We expect an incremental system to produce accurate output as soon as possible ([Trinh et al., 2018](#)), with a minimum amount of revocations and substitutions, ideally only having correct additions, to avoid jittering that may be detrimental to subsequent processors working on partial outputs.

To assess the incremental behavior of sequence tagging and classification models, we use the evaluation metrics for incremental processors established by Schlangen et al. (2009) and Baumann et al. (2011). The latter defines three diachronic metrics: *edit overhead* ($EO \in [0, 1]$), the proportion of unnecessary edits (the closer to 0, the fewer edits were made); *correction time* ($CT \in [0, 1]$), the fraction of the utterance seen before the system commits on a final decision for a piece of the output (the closer to 0, the sooner final decisions were made); and *relative correctness* ($RC \in [0, 1]$), the proportion of outputs that are correct with respect to the non-incremental output (being close to 1 means the system outputs were most of the time correct prefixes of the non-incremental output).

The sequence tagging tasks we evaluate are massively incremental (Hildebrandt et al., 1999), meaning that a new label is always added to the output after a new word is processed. The models can also substitute any previous labels in the output sequence in the light of new input. Sequence classifiers must add one label (the sequence’s class) after seeing the first word and can only substitute that single label after each new word. In both cases, additions are obligatory and substitutions should ideally be kept as low as possible, but there can be no revocations. Moreover, our data is sequential, discrete, and order-preserving (Köhn, 2018).

Given a sequence of length n , the number of necessary edits is always the number of tokens in the sequence (all additions) for sequence taggers and we set it to 1 for sequence classifiers. All other edits (substitutions) count as unnecessary and their number is bounded by $\sum_{i=1}^{n-1} i$ for tagging, and by $(n - 1)$, for classification.

We need to slightly adapt the CT measure for sequences. It is originally defined as FD-F0, the time step of a final decision minus the time step when the output first appeared. F0 is fixed for every word in a sequence (the systems always output a new label corresponding to each new word it sees), but each label will have a different FD. In order not to penalize initial labels, which have more opportunities of being substituted than final ones, we instead sum the FD of each token and divide by the sum of the number of times each one could be modified, to get a score for the sequence as a whole. Let the sequence length be n , then here $CTscore = (\sum_{i=1}^n FD_i) / (\sum_{i=1}^n n - i)$. We define it to be 0 for sequences of one token. Again, 0 means

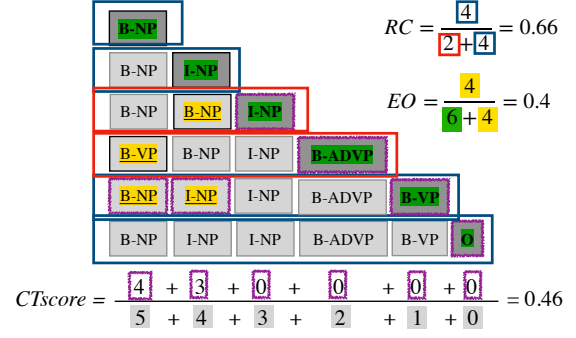


Figure 2: How we estimate the evaluation metrics for the complete sequence of outputs from Figure 1.

every label is immediately committed, 1 means all final decisions are delayed until the last time step. Figure 2 presents a concrete example of how to estimate the metrics.

Based on the trade-off between responsiveness and output quality (Skantze and Schlangen, 2009), we also estimate whether there is any improvement in the quality of the outputs if the encoder waits for some right context to appear before committing on output previously generated. For that, we use delayed EO and delayed RC (also named *discounted* in Baumann et al., 2011), which allows one or two words of the right context to be observed before outputting previous labels, named EO/RC Δ 1 and EO/RC Δ 2, respectively.

In order to concentrate on the incremental quality despite the eventual non-incremental deficiencies, we follow the approach by Baumann et al. (2011) and evaluate intermediate outputs in comparison to the processor’s final output, which may differ from the gold output but is the same as the non-incremental output. The general non-incremental correctness should be guaranteed by having high accuracy or F1 score in the non-incremental performance.

4 Models

We test the behavior of five neural networks, illustrated in Figure 3, under an incremental processing interface operating on word level and having full sentences as processing units: a) a vanilla LSTM; b) a vanilla BiLSTM; c) an LSTM with a CRF (Conditional Random Field) layer; d) a BiLSTM with a CRF layer; and e) BERT. The vanilla LSTM is the only model that works solely in temporal direction.

We choose to use the basic forms of each model to isolate the effect of bidirectionality. They per-

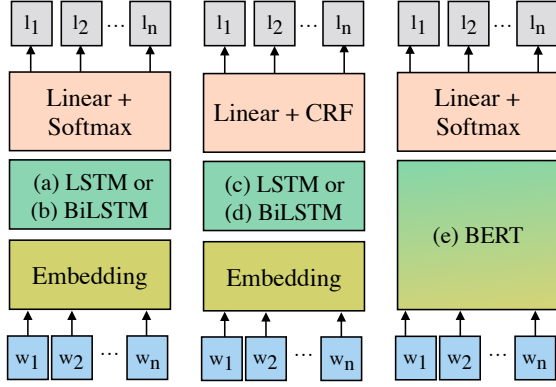


Figure 3: Models for sequence tagging, w =word and l =label. (a) is the only inherently incremental. (a), (b) and (e) can also be used for sequence classification if we consider only their final representation.

form well enough on the tasks to enable a realistic evaluation (see Table 1). Note that state-of-the-art results are typically achieved by combining them with more sophisticated mechanisms.

We use the models for both sequence tagging and classification. They use the representation at each time step to predict a corresponding label for sentence tagging, whereas for sequence classification they use the representation of the last time step (LSTM) or a combination of the last forward and backward representations (BiLSTM) or, in case of BERT, the representation at the CLS (initial) token, as suggested in Devlin et al. (2019). The two models with CRF cannot be used for classification, as there are no transition probabilities to estimate.

Sequence tagging implies a one-to-one mapping from words to labels, so that for every new word the system receives, it outputs a sequence with one extra label. In sequence classification, we map every input to a single label. In that case, the LSTM can also edit the output since it can change the chosen label as it processes more information. Because the datasets we use are tokenized and each token has a corresponding label, we follow the instructions given by Devlin et al. (2019) for dealing with BERT’s subtokenization: the scores of the first subtoken are used to predict its label, and further subtoken scores are ignored.

Except for the LSTM on sequence tagging, all models’ outputs are non-monotonic, i.e., they may reassign labels from previous words. The concept of timeliness is trivial here because we know exactly that the label for the t -th word will appear for the first time at the t -th version of the output, for all t . Even so, we can delay the output to allow

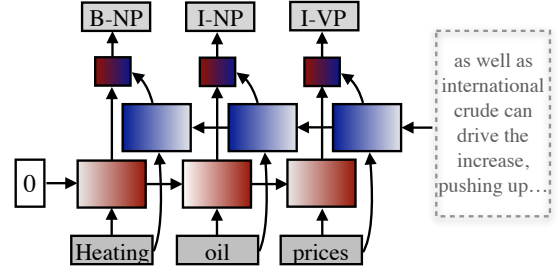


Figure 4: Incremental interface of a non-incremental bidirectional model, showing the input and output at time step 3. The context vector fed into the backward LSTM can be zero or initialized with a hypothetical right context generated by a language model.

some lookahead. In terms of decisiveness, all models commit to a single output at every time step. Figure 4 shows an example of the computation graph. BiLSTMs can recompute only the backward pass, while BERT needs a complete recomputation.

4.1 Strategies

We check the effect of three strategies: *truncated training*, *delayed output* and *prophecies*. In the first case, we modify the training regime by stripping off the endings of each sentence in the training set. We randomly sample a maximum length $l \leq n$, where n is the original sentence length, and cut the subsequent words and labels. We expect this to encourage the model to know how to deal with truncated sequences that it will have to process during testing.

The second strategy involves allowing some upcoming words to be observed before outputting a label corresponding to previous words. This is a case of lookahead described in Baumann et al. (2011), where the processor is allowed to wait for some right context before making a first decision with respect to previous time steps. We experiment with right contexts of one or two words, $\Delta 1$ and $\Delta 2$, respectively. $\Delta 1$ means the model outputs the first label for word t once it consumes word $t + 1$. Analogously, $\Delta 2$ means the model can observe words $t + 1$ and $t + 2$ before outputting the first label for word t . Figure 5 illustrates how to calculate *EO* with $\Delta 1$ delay for the same example as in Figure 2.

In the third strategy, we first feed each prefix as left context in the GPT-2 language model and let it generate a continuation up to the end of a sentence

Task	Metric	Model				
		LSTM	LSTM+CRF	BiLSTM	BiLSTM+CRF	BERT
Named Entity Recognition	F1 Score (%)	86.93 (84.23)	90.40 (88.13)	90.22 (88.07)	91.24 (89.44)	96.32 (96.11)
		70.78 (67.64)	86.30 (83.98)	88.79 (84.72)	89.29 (87.50)	93.52 (92.54)
		52.27 (49.63)	70.83 (68.71)	77.39 (73.34)	84.28 (80.88)	89.01 (87.23)
		93.82 (90.78)	95.36 (92.09)	94.84 (91.41)	95.26 (92.63)	95.57 (93.88)
		82.20 (78.09)	89.63 (85.28)	90.44 (85.41)	92.32 (87.82)	95.46 (92.93)
Intent (ATIS)	Accuracy (%)	96.86 (93.06)	-	95.74 (93.62)	-	97.31 (95.86)
Intent (SNIPS)		96.86 (97.43)	-	97.43 (97.43)	-	97.57 (97.71)
Part-of-Speech Tagging		94.98 (94.32)	96.02 (95.56)	96.44 (96.23)	96.64 (96.35)	97.87 (97.65)
Positive/Negative		82.17 (72.83)	-	83.33 (75.67)	-	93.83 (92.50)
Pros/Cons		94.51 (93.85)	-	94.40 (93.65)	-	95.74 (95.17)

Table 1: Non-incremental performance of all models on test sets (truncated training in parentheses). The results are not necessarily state-of-the-art because we use basic forms of each model in order to isolate the effect of bidirectionality and have comparable results among different tasks.

is an overall considerable improvement in the use of BERT model for all tasks. Truncated training reduces overall performance but even so BERT with truncated training outperforms all models, even with usual training, in most tasks (except for slot filling and IntentATIS).

Figure 7 presents an overview of the incremental evaluation metrics for all models and tasks. Sequence tagging has, in general, low EO and low CT score; i.e., labels are not edited much and a final decision is reached early. That does not hold for BERT, whose CT score and EO is, in general, higher. CT score and EO in sequence classification are also higher because the label in this case should capture a more global representation, which cannot reasonably be expected to be very good when only a small part of the sequence has yet been seen.

When it comes to RC (correctness relative to the final output), again BERT has worse results than other models, especially for tagging. For sequence classification, BERT’s performance is more in line with the other models. Achieving high RC is desirable because it means that, most of the time, the partial outputs are correct prefixes of the non-incremental output and can be trusted, at least to the same degree that the final result can be trusted.

This overview shows that although BERT’s non-incremental performance is normally the highest, the quality of its incremental outputs is more unstable. The next step is examining the effect of the three strategies that seek to improve the quality and stability of incremental outputs. Figure 8 shows that truncated training is always beneficial, as is delayed evaluation, with both strategies reducing EO and increasing RC. The fact that delay helps in all cases indicates that most substitutions happen

in the last or last but one label (the right frontier, given the current prefix), or, in other words, that even having a limited right context improves quality substantially.⁵

Prophecies are detrimental in classification tasks, but they help in some tagging tasks, especially for BERT. Most importantly, any of the strategies cause a great improvement to BERT’s incremental performance in sequence tagging, making its metrics be on the same level as other models while retaining its superior non-incremental quality.

Note that while CT and RC can only be measured once the final output is available, an estimate of EO may be evaluated on the fly if we consider the edits and additions up to the last output. Figure 9 shows how the mean EO evolves, breaking out the results for cases where the non-incremental final output will be correct and those where it will not with respect to the gold labels. We can observe an intriguing pattern: the mean EO grows faster for cases where the final response will be wrong; this is most pronounced for the sequence classification task. It might be possible to use this observation as an indication of how much to trust the final result: If the incremental computation was more unstable than the average, we should not expect the final result to be good. However, initial experiments on building a classifier based on the instability of partial outputs have so far not been successful in cashing in on that observation.

⁵Results for SRL are not included in Figure 8, because they go in the opposite direction of all other tasks. Since this task depends on the predicate embedding, both truncating the training sequence or adding a right context with no predicate information reduces performance in most cases, except for BERT. See Appendix for results separated by model and task.

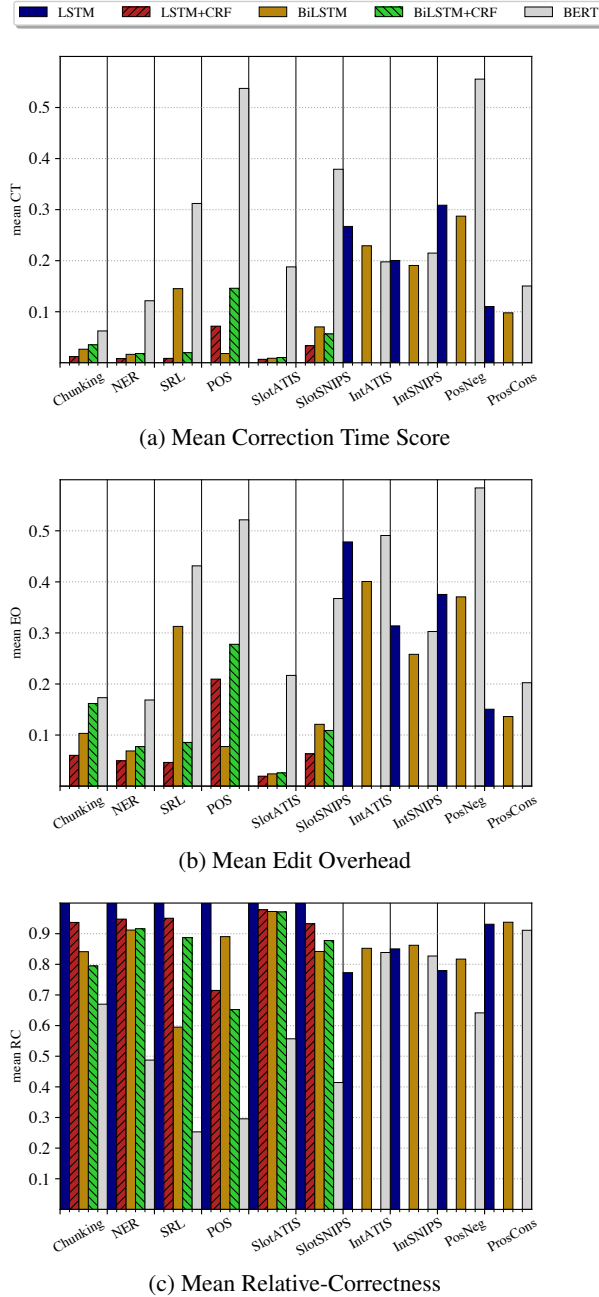


Figure 7: Comparison of evaluation metrics for all models and tasks. The incremental behavior is more stable for sequence tagging than for sequence classification. BERT takes longer to reach final decisions, and its outputs are edited more often than other models, especially in sequence tagging tasks.

7 Discussion and conclusion

We show that bidirectional encoders can be adapted to work under an incremental interface without a too drastic impact on their performance. Even though the training (being done on complete sequences) differs from the testing situation (which exposes the model to partial input), the incremental

metrics of most models are, in general, good: in sequence tagging, edit overhead is low, final decisions are taken early, and often partial outputs are a correct prefix of the complete non-incremental output. Sequence classification is more unstable because, at initial steps, there is a higher level of uncertainty on what is coming next. Our experiments show that the deficiencies of BERT in the incremental metrics can be mitigated with some adaptations (truncated training or prophecies together with delay), which make its incremental quality become as good as those of other models.

Since the semantic information is only kept encoded for a few steps in RNNs (Hupkes et al., 2018), this may be a reason why delay causes incremental metrics to be much better. If long-range dependencies are not captured, only neighboring words exert more influence in the choice of a label, so after seeing two words in the right context, the system rarely revises labels further back. BERT, having access to the whole sentence at any time step, is less stable because new input can cause it to reassess past labels more easily.

Besides, we also found evidence of different behavior of the instability of partial outputs between correct and incorrect output sequences, which could potentially be a signal of later lower quality. This could be used, for example, in dialog systems: if edit overhead gets too high, a clarification request should be made. A follow-up idea is training a classifier that predicts more precisely how likely it is that the final labels will be accurate based on the development of EO. However, our initial experiments on building such classifier were not successful. We suppose this is due to the fact that, in our datasets, incorrect final output sequences still usually have more than 90% correct labels, so the learnable signal may be too weak.

The use of GPT-2 prophecies led to promising improvements for BERT in sequence tagging. We see room for improvement, *e.g.* resorting to domain adaptation to make prophecies be more related to each genre. A natural extension is training a language model that generates the prophecies together with the encoder.

Finally, we believe that using attention mechanisms to study the grounding of the edits, similarly to the ideas in Köhn (2018), can be an important step towards understanding how the preliminary representations are built and decoded; we want to test this as well in future work.

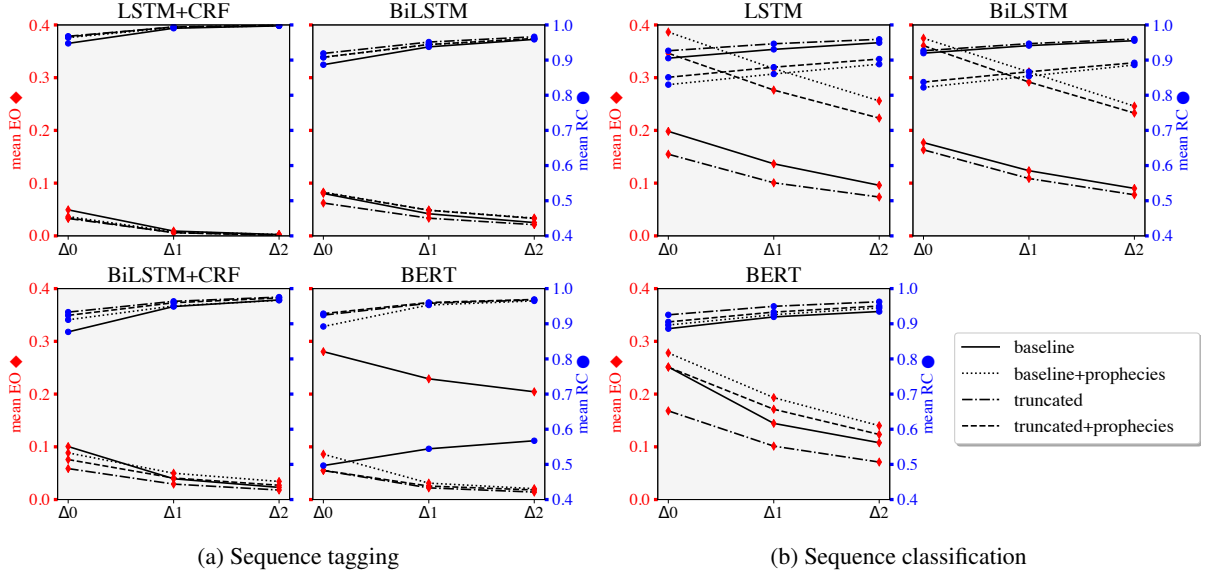


Figure 8: Comparison of mean Edit Overhead and mean Relative-Correctness on the baseline incremental interface and the three strategies using observations from all tasks except SRL.

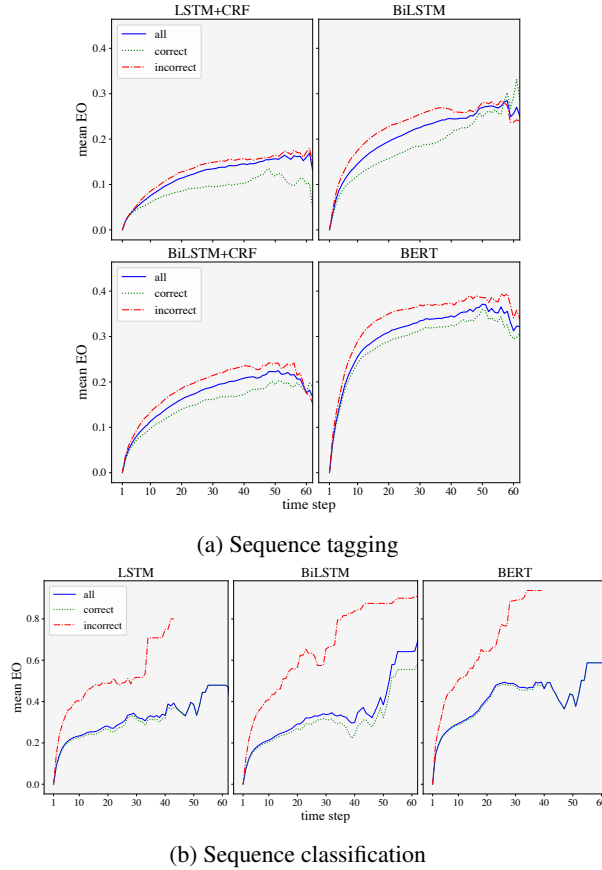


Figure 9: Development of mean Edit Overhead over time using observations from all tasks. *correct* means that all final output labels of a sentence are right and *incorrect* means that at least one label of the final output sequence is wrong. All models are more unstable when their non-incremental final output is incorrect with respect to the gold output.

Acknowledgments

We are thankful to the three anonymous reviewers and to our colleagues in the Foundations of Computational Linguistics Lab at the University of Potsdam for their valuable comments and insights. D.S. acknowledges the support by the Deutsche Forschungsgemeinschaft (DFG), grant SCHL 845/9-1.

Updates and Corrections

We have fixed a wrong step in the implementation of the delayed EO and updated Figure 5, Figure 8 and Table 10 accordingly. The text remains the same since no conclusions changed. A few citations were missing that we will add here: The GPT-2 model was proposed by [Radford et al. \(2019\)](#). The incremental interface paradigm was also defined by [Schlangen and Skantze \(2011\)](#) as *restart-incrementality*. We use the preprocessed data and splits for SNIPS and ATIS made available by [E et al. \(2019b\)](#) and relied on <https://github.com/yuchenlin/OntoNotes-5.0-NER-BIO/> and <https://github.com/ontonotes/conll-formatted-ontonotes-5.0> to preprocess OntoNotes.

References

Gregory Aist, James Allen, Ellen Campana, Carlos Gomez Gallo, Scott Stoness, Mary Swift, and Michael K. Tanenhaus. 2007. Incremental under-

- standing in human-computer dialogue and experimental evidence for advantages over nonincremental methods. In *Proceedings of Decalog 2007, the 11th International Workshop on the Semantics and Pragmatics of Dialogue*, Trento, Italy.
- Pierre Baldi, Søren Brunak, Paolo Frasconi, Giovanni Soda, and Gianluca Pollastri. 1999. Exploiting the past and the future in protein secondary structure prediction. *Bioinformatics*, 15(11):937–946.
- Timo Baumann, Okko Buß, and David Schlangen. 2011. Evaluation and optimisation of incremental processors. *Dialogue & Discourse*, 2(1):113–141.
- Niels Beuck, Arne Köhn, and Wolfgang Menzel. 2011. Decision strategies for incremental POS tagging. In *Proceedings of the 18th Nordic Conference of Computational Linguistics (NODALIDA 2011)*, pages 26–33.
- Timothy Brick and Matthias Scheutz. 2007. Incremental natural language processing for HRI. In *Proceedings of the ACM/IEEE international conference on Human-robot interaction*, pages 263–270.
- Hendrik Buschmeier, Timo Baumann, Benjamin Dosch, Stefan Kopp, and David Schlangen. 2012. [Combining incremental language generation and incremental speech synthesis for adaptive information presentation](#). In *Proceedings of the 13th Annual Meeting of the Special Interest Group on Discourse and Dialogue*, pages 295–303, Seoul, South Korea. Association for Computational Linguistics.
- T. Chiang. 2002. *Stories of Your Life and Others*. Tom Doherty Associates.
- Jason PC Chiu and Eric Nichols. 2016. Named entity recognition with bidirectional LSTM-CNNs. *Transactions of the Association for Computational Linguistics*, 4:357–370.
- Morten H Christiansen and Nick Chater. 2016. The now-or-never bottleneck: A fundamental constraint on language. *Behavioral and brain sciences*, 39.
- Andrei C. Coman, Koichiro Yoshino, Yukitoshi Murase, Satoshi Nakamura, and Giuseppe Riccardi. 2019. [An Incremental Turn-Taking Model for Task-Oriented Dialog Systems](#). In *Proc. Interspeech 2019*, pages 4155–4159.
- Alice Coucke, Alaa Saade, Adrien Ball, Théodore Bluche, Alexandre Caulier, David Leroy, Clément Doumouro, Thibault Gisselbrecht, Francesco Caltagirone, Thibaut Lavril, et al. 2018. SNIPS voice platform: an embedded spoken language understanding system for private-by-design voice interfaces. *arXiv preprint arXiv:1805.10190*.
- Fahim Dalvi, Nadir Durrani, Hassan Sajjad, and Stephan Vogel. 2018. [Incremental decoding and training methods for simultaneous translation in neural machine translation](#). 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 493–499, New Orleans, Louisiana. Association for Computational Linguistics.
- David DeVault, Kenji Sagae, and David Traum. 2011. Incremental interpretation and prediction of utterance meaning for interactive dialogue. *Dialogue & Discourse*, 2(1):143–170.
- 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.
- Haihong E, Peiqing Niu, Zhongfu Chen, and Meina Song. 2019a. [A novel bi-directional interrelated model for joint intent detection and slot filling](#). In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*, pages 5467–5471, Florence, Italy. Association for Computational Linguistics.
- Haihong E, Peiqing Niu, Zhongfu Chen, and Meina Song. 2019b. [A novel bi-directional interrelated model for joint intent detection and slot filling](#). In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*, pages 5467–5471, Florence, Italy. Association for Computational Linguistics.
- Jeffrey L Elman. 1990. Finding structure in time. *Cognitive science*, 14(2):179–211.
- Murthy Ganapathibhotla and Bing Liu. 2008. [Mining opinions in comparative sentences](#). In *Proceedings of the 22nd International Conference on Computational Linguistics (Coling 2008)*, pages 241–248, Manchester, UK. Coling 2008 Organizing Committee.
- Alex Graves and Jürgen Schmidhuber. 2005. Framewise phoneme classification with bidirectional LSTM and other neural network architectures. *Neural networks*, 18(5-6):602–610.
- Alvin Grissom II, He He, Jordan Boyd-Graber, John Morgan, and Hal Daumé III. 2014. [Don’t until the final verb wait: Reinforcement learning for simultaneous machine translation](#). In *Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 1342–1352, Doha, Qatar. Association for Computational Linguistics.
- Luheng He, Kenton Lee, Mike Lewis, and Luke Zettlemoyer. 2017. Deep semantic role labeling: What works and what’s next. In *Proceedings of the 55th*

- Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 473–483.
- Charles T Hemphill, John J Godfrey, and George R Doddington. 1990. The ATIS spoken language systems pilot corpus. In *Speech and Natural Language: Proceedings of a Workshop Held at Hidden Valley, Pennsylvania, June 24-27, 1990*.
- Bernd Hildebrandt, Hans-Jürgen Eikmeyer, Gert Rickheit, and Petra Weiß. 1999. Inkrementelle sprachrezeption. *KogWis99: Proceedings der 4. Fachtagung der Gesellschaft für Kognitionswissenschaft*.
- Sepp Hochreiter and Jürgen Schmidhuber. 1997. Long short-term memory. *Neural computation*, 9(8):1735–1780.
- Zhiheng Huang, Wei Xu, and Kai Yu. 2015. Bidirectional LSTM-CRF models for sequence tagging. *arXiv preprint arXiv:1508.01991*.
- Dieuwke Hupkes, Sanne Bouwmeester, and Raquel Fernández. 2018. [Analysing the potential of seq-to-seq models for incremental interpretation in task-oriented dialogue](#). In *Proceedings of the 2018 EMNLP Workshop BlackboxNLP: Analyzing and Interpreting Neural Networks for NLP*, pages 165–174, Brussels, Belgium. Association for Computational Linguistics.
- Mei-Yuh Hwang, Gang Peng, Mari Ostendorf, Wen Wang, Arlo Faria, and Aaron Heide. 2009. Building a highly accurate mandarin speech recognizer with language-independent technologies and language-dependent modules. *IEEE transactions on audio, speech, and language processing*, 17(7):1253–1262.
- Ozan İrsoy and Claire Cardie. 2014. [Opinion mining with deep recurrent neural networks](#). In *Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 720–728, Doha, Qatar. Association for Computational Linguistics.
- Gerard Kempen and Edward Hoenkamp. 1982. [Incremental sentence generation: Implications for the structure of a syntactic processor](#). In *Coling 1982: Proceedings of the Ninth International Conference on Computational Linguistics*, pages 151–156.
- Diederik P Kingma and Jimmy Ba. 2014. Adam: A method for stochastic optimization. *arXiv preprint arXiv:1412.6980*.
- Arne Köhn. 2018. [Incremental natural language processing: Challenges, strategies, and evaluation](#). In *Proceedings of the 27th International Conference on Computational Linguistics*, pages 2990–3003, Santa Fe, New Mexico, USA. Association for Computational Linguistics.
- Ioannis Konstas, Frank Keller, Vera Demberg, and Mirella Lapata. 2014. [Incremental semantic role labeling with tree adjoining grammar](#). In *Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 301–312, Doha, Qatar. Association for Computational Linguistics.
- Dimitrios Kotzias, Misha Denil, Nando De Freitas, and Padhraic Smyth. 2015. From group to individual labels using deep features. In *Proceedings of the 21th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pages 597–606.
- Willem J. M. Levelt. 1989. *Speaking: From intention to articulation*. London: MIT Press.
- Wang Ling, Chris Dyer, Alan W Black, Isabel Trancoso, Ramón Fernandez, Silvio Amir, Luís Marujo, and Tiago Luís. 2015. [Finding function in form: Compositional character models for open vocabulary word representation](#). In *Proceedings of the 2015 Conference on Empirical Methods in Natural Language Processing*, pages 1520–1530, Lisbon, Portugal. Association for Computational Linguistics.
- Mingbo Ma, Liang Huang, Hao Xiong, Renjie Zheng, Kaibo Liu, Baigong Zheng, Chuanqiang Zhang, Zhongjun He, Hairong Liu, Xing Li, Hua Wu, and Haifeng Wang. 2019. [STACL: Simultaneous translation with implicit anticipation and controllable latency using prefix-to-prefix framework](#). In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*, pages 3025–3036, Florence, Italy. Association for Computational Linguistics.
- William D. Marslen-Wilson. 1975. [Sentence perception as an interactive parallel process](#). *Science*, 189(4198):226–228.
- Tomáš Mikolov, Martin Karafiát, Lukáš Burget, Jan Černocký, and Sanjeev Khudanpur. 2010. Recurrent neural network based language model. In *Eleventh annual conference of the international speech communication association*.
- Joakim Nivre. 2004. [Incrementality in deterministic dependency parsing](#). In *Proceedings of the Workshop on Incremental Parsing: Bringing Engineering and Cognition Together*, pages 50–57, Barcelona, Spain. Association for Computational Linguistics.
- Yusuke Oda, Graham Neubig, Sakriani Sakti, Tomoki Toda, and Satoshi Nakamura. 2015. [Syntax-based simultaneous translation through prediction of unseen syntactic constituents](#). In *Proceedings of the 53rd Annual Meeting of the Association for Computational Linguistics and the 7th International Joint Conference on Natural Language Processing (Volume 1: Long Papers)*, pages 198–207, Beijing, China. Association for Computational Linguistics.
- Andreas Peldszus, Timo Baumann, Okko Buß, and David Schlangen. 2012. Joint satisfaction of syntactic and pragmatic constraints improves incremental spoken language understanding. In *Proceedings of the 13th Conference of the European Chapter of*

- the Association for Computational Linguistics, pages 514–523. Association for Computational Linguistics.
- Jeffrey Pennington, Richard Socher, and Christopher D. Manning. 2014. [GloVe: Global vectors for word representation](#). In *Empirical Methods in Natural Language Processing (EMNLP)*, pages 1532–1543.
- Matthew Peters, Mark Neumann, Mohit Iyyer, Matt Gardner, Christopher Clark, Kenton Lee, and Luke Zettlemoyer. 2018. [Deep contextualized word representations](#). In *Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long Papers)*, pages 2227–2237, New Orleans, Louisiana. Association for Computational Linguistics.
- Barbara Plank, Anders Søgaard, and Yoav Goldberg. 2016. [Multilingual part-of-speech tagging with bidirectional long short-term memory models and auxiliary loss](#). In *Proceedings of the 54th Annual Meeting of the Association for Computational Linguistics (Volume 2: Short Papers)*, pages 412–418, Berlin, Germany. Association for Computational Linguistics.
- Sameer Pradhan, Alessandro Moschitti, Nianwen Xue, Hwee Tou Ng, Anders Björkelund, Olga Uryupina, Yuchen Zhang, and Zhi Zhong. 2013. Towards robust linguistic analysis using ontonotes. In *Proceedings of the Seventeenth Conference on Computational Natural Language Learning*, pages 143–152.
- Alec Radford, Jeffrey Wu, Rewon Child, David Luan, Dario Amodei, Ilya Sutskever, et al. 2019. [Language models are unsupervised multitask learners](#). *OpenAI blog*, 1(8):9.
- Tjong Kim Sang, F Erik, and Sabine Buchholz. 2000. Introduction to the CoNLL-2000 shared task: chunking. In *Proceedings of CoNLL-2000, Lisbon, Portugal*, pages 127–132.
- David Schlangen, Timo Baumann, and Michaela Atterer. 2009. [Incremental reference resolution: The task, metrics for evaluation, and a Bayesian filtering model that is sensitive to disfluencies](#). In *Proceedings of the SIGDIAL 2009 Conference*, pages 30–37, London, UK. Association for Computational Linguistics.
- David Schlangen and Gabriel Skantze. 2011. A general, abstract model of incremental dialogue processing. *Dialogue & Discourse*, 2(1):83–111.
- Mike Schuster and Kuldip K Paliwal. 1997. Bidirectional recurrent neural networks. *IEEE transactions on Signal Processing*, 45(11):2673–2681.
- Ethan Selfridge, Iker Arizmendi, Peter Heeman, and Jason Williams. 2011. [Stability and accuracy in incremental speech recognition](#). In *Proceedings of the SIGDIAL 2011 Conference*, pages 110–119, Portland, Oregon. Association for Computational Linguistics.
- Gabriel Skantze and David Schlangen. 2009. [Incremental dialogue processing in a micro-domain](#). In *Proceedings of the 12th Conference of the European Chapter of the ACL (EACL 2009)*, pages 745–753, Athens, Greece. Association for Computational Linguistics.
- Ilya Sutskever, Oriol Vinyals, and Quoc V Le. 2014. Sequence to sequence learning with neural networks. In *Advances in neural information processing systems*, pages 3104–3112.
- Michael K Tanenhaus and Sarah Brown-Schmidt. 2008. Language processing in the natural world. *Philosophical Transactions of the Royal Society B: Biological Sciences*, 363(1493):1105–1122.
- Anh Duong Trinh, Robert J Ross, and John D Kelleher. 2018. A multi-task approach to incremental dialogue state tracking. In *Proceedings of The 22nd workshop on the Semantics and Pragmatics of Dialogue, SEMDIAL*, pages 132–145.
- Dennis Ulmer, Dieuwke Hupkes, and Elia Bruni. 2019. [Assessing incrementality in sequence-to-sequence models](#). In *Proceedings of the 4th Workshop on Representation Learning for NLP (RepL4NLP-2019)*, pages 209–217, Florence, Italy. Association for Computational Linguistics.
- Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. 2017. Attention is all you need. In *Advances in neural information processing systems*, pages 5998–6008.
- Dimitra Vergyi, Andreas Stolcke, Venkata Ramana Rao Gadde, Luciana Ferrer, and Elizabeth Shriberg. 2003. Prosodic knowledge sources for automatic speech recognition. In *2003 IEEE International Conference on Acoustics, Speech, and Signal Processing, 2003. Proceedings.(ICASSP'03)*, volume 1, pages I–I. IEEE.
- Ralph Weischedel, Martha Palmer, Mitchell Marcus, Eduard Hovy, Sameer Pradhan, Lance Ramshaw, Nianwen Xue, Ann Taylor, Jeff Kaufman, Michelle Franchini, et al. 2013. Ontonotes release 5.0 LDC2013t19. web download. *Linguistic Data Consortium, Philadelphia, PA*.
- Feifei Zhai, Saloni Potdar, Bing Xiang, and Bowen Zhou. 2017. Neural models for sequence chunking. In *Thirty-First AAAI Conference on Artificial Intelligence*.
- Lukáš Žilka and Filip Jurčiček. 2015. Lectrack: Incremental dialog state tracking with long short-term memory networks. In *International Conference on Text, Speech, and Dialogue*, pages 174–182. Springer.

A Reproducibility

Here we describe more details about hyperparameters and the models. Table 2 to Table 11 present more information about the hyperparameter search and the datasets as well as explicit results for each model and task, to be used for reproducibility purposes.

Data

- We use only the WSJ section of OntoNotes, with the train-valid-test splits provided by Pradhan et al. (2013), as well as their conversion to CoNLL format.
- The PosNeg and ProsCons datasets, which are not published with a train-valid-test split, were divided into 70%-10%-20% sets randomly.
- We removed the two longest sentences (>200 words) in Pros/Cons the dataset because they were infeasible to compute with BERT.
- Sentences longer than 60 words were removed in the hyperparameter search phase only.

Implementation

- Hyperparameter search is done for each dataset in both the LSTM model and BERT.
- We use Comet’s Bayes algorithm⁶, which balances exploration and exploitation and, in our experiments, tries to maximize the accuracy (for sequence classification) of F1 score (for sequence tagging) in the validation set.
- We set the maximum number of iterations to 50, but early stopping happens if no improvement is seen during 10 iterations.
- The best configuration of the LSTM model is also used for the LSTM+CRF, BiLSTM, and BiLSTM+CRF models.
- We use a five-dimensional embedding for the binary predicates in the SRL task.
- Hidden states are initialized as 0.
- All the weights and biases are initialized with PyTorch’s default (uniformly sampled from $(-\sqrt{hidden_size}, \sqrt{hidden_size})$).

- Dropout is implemented after the embedding layer and after the encoder layer with the same value.
- PyTorch’s and Numpy’s manual seeds are set to 2204 for all experiments.
- All experiments were run on a GPU GeForce GTX 1080 Ti.

Hyperparameter	LSTM
Batch size	32, 64, 128, 512
Clipping	0.25, 0.5, 1
Dropout	0.1, 0.2, 0.3, 0.5
Embedding dimension	50, 100, 200, 300
Hidden layer dimension	50, 100, 150, 200, 300
Learning rate	0.1, 0.01, 0.001
Number of layers	1, 2, 3, 4
	BERT
Batch size	16, 32
Clipping	0.25, 0.5, 1
Learning rate	5e-05, 3e-05, 2e-05, 1e-05

Table 2: Hyperparameter search space.

⁶<https://www.comet.ml/docs/python-sdk/introduction-optimizer/>

Task/Model	search trials	avrg. runtime	best configuration						
			batch size	clipping	dropout	embedding layer	hidden layer	learning rate	number of layers
SEQUENCE TAGGING									
Chunk	24	6	32	0.5	0.5	300	200	0.001	1
Named Entity Recognition	24	16	32	1	0.3	50	300	0.001	2
Part-of-Speech Tagging	26	19	32	0.25	0.5	300	150	0.001	3
Semantic Role Labeling	34	33	32	1	0.3	100	300	0.001	2
Slot Filling (ATIS)	23	2	32	0.25	0.5	50	200	0.01	1
Slot Filling (SNIPS)	39	5	32	1	0.5	50	150	0.001	2
SEQUENCE CLASSIFICATION									
Intent (ATIS)	60	1	64	0.5	0.1	200	200	0.001	2
Intent (SNIPS)	66	3	64	0.5	0.5	100	300	0.001	1
Positive/Negative	27	1	64	0.5	0.5	300	100	0.01	3
Pros/Cons	38	6	32	0.5	0.5	300	150	0.001	4

Table 3: Hyperparameter search for LSTM model. The best configuration was also used for LSTM+CRF, BiLSTM and BiLSTM+CRF. Runtime in minutes.

Task/Model	search trials	avrg. runtime	best configuration		
			batch size	clipping	learning rate
SEQUENCE TAGGING					
Chunk	8	24	16	0.25	2e-05
Named Entity Recognition	9	99	16	1	2e-05
Part-of-Speech Tagging	7	62	16	0.5	3e-05
Semantic Role Labeling	9	471	16	0.5	2e-05
Slot Filling (ATIS)	10	15	32	0.25	5e-05
Slot Filling (SNIPS)	11	33	16	0.5	2e-05
SEQUENCE CLASSIFICATION					
Intent (ATIS)	10	12	32	0.25	5e-05
Intent (SNIPS)	10	23	16	0.5	3e-05
Positive/Negative	10	4	16	0.25	2e-05
Pros/Cons	10	70	32	0.5	3e-05

Table 4: Hyperparameter search for BERT model. Runtime in minutes.

Task	Model				
	LSTM	LSTM+CRF	BiLSTM	BiLSTM+CRF	BERT
Chunking	5,775,023	5,775,598	6,181,223	6,181,798	108,327,959
NER	2,937,587	2,939,030	4,813,487	4,814,930	108,338,725
POS	11,330,748	11,333,148	12,331,548	12,333,948	108,347,184
SRL	4,829,016	4,840,464	6,791,616	6,803,064	108,391,786
SlotATIS	270,327	286,710	497,327	513,710	108,407,935
SlotSNIPS	876,522	881,850	1,369,722	1,375,050	108,365,640
IntentATIS	821,226	-	1,789,626	-	108,330,266
IntentSNIPS	1,611,007	-	2,095,507	-	108,315,655
PosNeg	1,764,402	-	2,247,002	-	108,311,810
ProsCons	4,905,302	-	6,260,402	-	108,311,810

Table 5: Number of parameters in each model.

Task	Dataset	Reference	Labels	Train	Valid	Test
SENTENCE TAGGING						
Chunking	CoNLL 2000	Sang et al. (2000)	23	7,922	1,014	2,012
Named Entity Recognition	OntoNotes 5.0	Weischedel et al. (2013)	37	30,060	5,315	1,640
Part-of-Speech Tagging	OntoNotes 5.0	Weischedel et al. (2013)	48	30,060	5,315	1,640
Semantic Role Labeling	OntoNotes 5.0	Weischedel et al. (2013)	106	83,920	15,208	4,781
Slot Filling	ATIS	Hemphill et al. (1990)	127	4,478	500	893
Slot Filling	SNIPS	Coucke et al. (2018)	72	13,084	700	700
SENTENCE CLASSIFICATION						
Intent	ATIS	Hemphill et al. (1990)	26	4,478	500	893
Intent	SNIPS	Coucke et al. (2018)	7	13,084	700	700
Sentiment	Positive/Negative	Kotzias et al. (2015)	2	2,100	300	600
Sentiment	Pros/Cons	Ganapathibhotla and Liu (2008)	2	32,088	4,602	9,175

Table 6: Tasks and datasets.

Task	Metric	Model				
		LSTM	LSTM+CRF	BiLSTM	BiLSTM+CRF	BERT
Chunk	F1 Score (%)	88.39 (88.42)	91.52 (90.79)	91.67 (90.93)	92.53 (91.76)	97.53 (97.00)
Named Entity Recognition		68.60 (67.36)	85.22 (82.08)	86.77 (83.02)	87.86 (84.78)	92.05 (89.38)
Semantic Role Labeling		52.55 (49.78)	71.48 (67.22)	77.53 (70.57)	84.16 (77.91)	89.29 (82.81)
Slot Filling (ATIS)		95.76 (94.54)	96.93 (95.91)	97.16 (96.07)	97.33 (96.62)	98.39 (96.67)
Slot Filling (SNIPS)		82.86 (80.12)	90.36 (86.12)	90.47 (84.90)	91.60 (87.26)	95.56 (88.52)
Intent (ATIS)	Accuracy (%)	98.40 (90.60)	-	98.20 (90.40)	-	98.80 (93.60)
Intent (SNIPS)		99.71 (95.14)	-	99.57 (94.00)	-	99.29 (93.57)
Part-of-Speech Tagging		94.72 (94.00)	95.75 (94.88)	96.24 (95.54)	96.27 (95.54)	97.90 (97.45)
Positive/Negative		85.33 (70.67)	-	85.67 (70.67)	-	95.67 (76.33)
Pros/Cons		94.59 (90.24)	-	94.74 (90.48)	-	96.02 (91.81)

Table 7: Non-incremental performance all models on validation sets for the purpose of reproducibility. Values in parentheses refer to using truncated samples during training.

Task	Model				
	LSTM	LSTM+CRF	BiLSTM	BiLSTM+CRF	BERT
Sentence level correctness (%)					
SEQUENCE TAGGING					
Chunk	33.95 (28.03)	44.04 (37.48)	43.74 (36.38)	48.56 (43.39)	71.42 (69.68)
Named Entity Recognition	51.04 (48.60)	70.61 (66.83)	74.02 (68.05)	75.00 (72.44)	85.61 (83.17)
Part-of-Speech Tagging	36.77 (34.02)	45.12 (42.68)	49.33 (47.38)	50.55 (49.21)	63.41 (60.49)
Semantic Role Labeling	6.82 (6.65)	24.14 (21.36)	48.90 (41.92)	56.14 (49.42)	67.85 (63.75)
Slot Filling (ATIS)	84.10 (75.92)	87.46 (78.72)	86.34 (76.48)	87.68 (79.06)	89.36 (84.88)
Slot Filling (SNIPS)	61.29 (54.71)	75.00 (66.86)	79.00 (69.14)	81.71 (71.86)	89.29 (84.43)
SEQUENCE CLASSIFICATION					
Intent (ATIS)	96.86 (93.06)	-	95.74 (93.62)	-	97.31 (95.86)
Intent (SNIPS)	96.86 (97.43)	-	97.43 (97.43)	-	97.57 (97.71)
Positive/Negative	82.17 (72.83)	-	83.33 (75.67)	-	93.83 (92.50)
Pros/Cons	94.51 (93.85)	-	94.40 (93.65)	-	95.74 (95.17)

Table 8: Sentence-level non-incremental performance of all models on test sets (the same as accuracy in sequence classification). Values in parentheses refer to using truncated samples during training.

Task/Model	Metrics			Metrics with prophecies		
	EO	CT	RC	EO	CT	RC
SEQUENCE TAGGING						
Chunk						
LSTM	0.00 (0.00)	0.00 (0.00)	1.00 (1.00)	0.00 (0.00)	0.00 (0.00)	1.00 (1.00)
LSTM+CRF	0.06 (0.03)	0.01 (0.00)	0.94 (0.97)	0.03 (0.03)	0.00 (0.00)	0.97 (0.97)
BiLSTM	0.10 (0.08)	0.03 (0.02)	0.84 (0.89)	0.09 (0.11)	0.02 (0.03)	0.90 (0.87)
BiLSTM+CRF	0.16 (0.07)	0.04 (0.02)	0.79 (0.92)	0.10 (0.09)	0.02 (0.02)	0.91 (0.91)
BERT	0.17 (0.06)	0.06 (0.02)	0.67 (0.91)	0.10 (0.06)	0.02 (0.01)	0.87 (0.92)
Named Entity Recognition						
LSTM	0.00 (0.00)	0.00 (0.00)	1.00 (1.00)	0.00 (0.00)	0.00 (0.00)	1.00 (1.00)
LSTM+CRF	0.05 (0.04)	0.01 (0.01)	0.95 (0.96)	0.04 (0.04)	0.01 (0.01)	0.96 (0.96)
BiLSTM	0.07 (0.06)	0.02 (0.01)	0.91 (0.93)	0.08 (0.08)	0.02 (0.02)	0.91 (0.92)
BiLSTM+CRF	0.08 (0.06)	0.02 (0.01)	0.92 (0.94)	0.09 (0.08)	0.02 (0.02)	0.92 (0.93)
BERT	0.17 (0.06)	0.12 (0.02)	0.49 (0.93)	0.08 (0.06)	0.02 (0.01)	0.90 (0.93)
Part-of-Speech Tagging						
LSTM	0.00 (0.00)	0.00 (0.00)	1.00 (1.00)	0.00 (0.00)	0.00 (0.00)	1.00 (1.00)
LSTM+CRF	0.05 (0.03)	0.01 (0.01)	0.95 (0.96)	0.03 (0.03)	0.01 (0.01)	0.97 (0.97)
BiLSTM	0.08 (0.06)	0.02 (0.01)	0.89 (0.93)	0.07 (0.06)	0.02 (0.02)	0.92 (0.93)
BiLSTM+CRF	0.09 (0.06)	0.02 (0.01)	0.89 (0.92)	0.07 (0.06)	0.02 (0.02)	0.92 (0.93)
BERT	0.52 (0.05)	0.54 (0.01)	0.30 (0.93)	0.07 (0.05)	0.02 (0.01)	0.90 (0.93)
Semantic Role Labeling						
LSTM	0.00 (0.00)	0.00 (0.00)	1.00 (1.00)	0.00 (0.00)	0.00 (0.00)	1.00 (1.00)
LSTM+CRF	0.21 (0.29)	0.07 (0.09)	0.71 (0.83)	0.32 (0.26)	0.10 (0.09)	0.82 (0.85)
BiLSTM	0.31 (0.37)	0.15 (0.16)	0.59 (0.60)	0.34 (0.40)	0.16 (0.19)	0.62 (0.62)
BiLSTM+CRF	0.28 (0.31)	0.15 (0.15)	0.65 (0.72)	0.28 (0.35)	0.15 (0.17)	0.71 (0.72)
BERT	0.43 (0.33)	0.31 (0.14)	0.25 (0.70)	0.22 (0.26)	0.12 (0.14)	0.69 (0.67)
Slot Filling (ATIS)						
LSTM	0.00 (0.00)	0.00 (0.00)	1.00 (1.00)	0.00 (0.00)	0.00 (0.00)	1.00 (1.00)
LSTM+CRF	0.02 (0.01)	0.01 (0.00)	0.98 (0.98)	0.01 (0.03)	0.00 (0.01)	0.99 (0.97)
BiLSTM	0.02 (0.02)	0.01 (0.01)	0.97 (0.98)	0.02 (0.02)	0.01 (0.01)	0.97 (0.98)
BiLSTM+CRF	0.03 (0.01)	0.01 (0.01)	0.97 (0.98)	0.04 (0.02)	0.01 (0.01)	0.95 (0.98)
BERT	0.22 (0.03)	0.19 (0.01)	0.56 (0.97)	0.06 (0.03)	0.02 (0.01)	0.93 (0.97)
Slot Filling (SNIPS)						
LSTM	0.00 (0.00)	0.00 (0.00)	1.00 (1.00)	0.00 (0.00)	0.00 (0.00)	1.00 (1.00)
LSTM+CRF	0.06 (0.04)	0.03 (0.02)	0.93 (0.96)	0.10 (0.07)	0.05 (0.03)	0.90 (0.93)
BiLSTM	0.12 (0.08)	0.07 (0.05)	0.84 (0.90)	0.17 (0.13)	0.10 (0.07)	0.81 (0.85)
BiLSTM+CRF	0.11 (0.08)	0.06 (0.04)	0.88 (0.91)	0.17 (0.14)	0.10 (0.07)	0.82 (0.86)
BERT	0.37 (0.08)	0.38 (0.04)	0.41 (0.91)	0.12 (0.09)	0.06 (0.05)	0.86 (0.90)
SEQUENCE CLASSIFICATION						
Intent (ATIS)						
LSTM	0.48 (0.21)	0.27 (0.12)	0.77 (0.91)	0.77 (0.78)	0.72 (0.68)	0.52 (0.54)
LSTM+CRF	-	-	-	-	-	-
BiLSTM	0.40 (0.24)	0.23 (0.14)	0.85 (0.90)	0.66 (0.77)	0.38 (0.73)	0.70 (0.50)
BiLSTM+CRF	-	-	-	-	-	-
BERT	0.49 (0.20)	0.20 (0.13)	0.84 (0.91)	0.38 (0.29)	0.19 (0.16)	0.88 (0.90)
Intent (SNIPS)						
LSTM	0.31 (0.24)	0.20 (0.14)	0.85 (0.90)	0.57 (0.49)	0.48 (0.39)	0.71 (0.77)
LSTM+CRF	-	-	-	-	-	-
BiLSTM	0.26 (0.23)	0.19 (0.13)	0.86 (0.91)	0.55 (0.49)	0.47 (0.41)	0.71 (0.76)
BiLSTM+CRF	-	-	-	-	-	-
BERT	0.30 (0.22)	0.21 (0.13)	0.83 (0.91)	0.40 (0.35)	0.26 (0.20)	0.83 (0.86)
Positive/Negative						
LSTM	0.38 (0.40)	0.31 (0.31)	0.78 (0.79)	0.65 (0.68)	0.59 (0.62)	0.74 (0.72)
LSTM+CRF	-	-	-	-	-	-
BiLSTM	0.37 (0.39)	0.29 (0.30)	0.82 (0.83)	0.63 (0.58)	0.51 (0.49)	0.76 (0.77)
BiLSTM+CRF	-	-	-	-	-	-
BERT	0.58 (0.45)	0.56 (0.30)	0.64 (0.80)	0.56 (0.55)	0.43 (0.43)	0.79 (0.79)
Pros/Cons						
LSTM	0.15 (0.13)	0.11 (0.09)	0.93 (0.94)	0.32 (0.27)	0.26 (0.21)	0.88 (0.90)
LSTM+CRF	-	-	-	-	-	-
BiLSTM	0.14 (0.14)	0.10 (0.10)	0.94 (0.94)	0.32 (0.30)	0.28 (0.24)	0.85 (0.88)
BiLSTM+CRF	-	-	-	-	-	-
BERT	0.20 (0.14)	0.15 (0.10)	0.91 (0.94)	0.24 (0.22)	0.17 (0.16)	0.91 (0.92)

Table 9: Mean values of Edit Overhead, Correction Time Score and Relative-Correctness. Values in parentheses refer to using truncated samples during training.

Task/Model	EO			EO with prophecies		
	$\Delta 0$	$\Delta 1$	$\Delta 2$	$\Delta 0$	$\Delta 1$	$\Delta 2$
SEQUENCE TAGGING						
Chunk						
LSTM	0.00 (0.00)	0.00 (0.00)	0.00 (0.00)	0.00 (0.00)	0.00 (0.00)	0.00 (0.00)
LSTM+CRF	0.06 (0.03)	0.01 (0.00)	0.00 (0.00)	0.03 (0.03)	0.00 (0.00)	0.00 (0.00)
BiLSTM	0.10 (0.08)	0.06 (0.05)	0.03 (0.03)	0.09 (0.11)	0.05 (0.07)	0.04 (0.05)
BiLSTM+CRF	0.16 (0.07)	0.06 (0.04)	0.03 (0.02)	0.10 (0.09)	0.06 (0.05)	0.04 (0.03)
BERT	0.17 (0.06)	0.10 (0.03)	0.09 (0.02)	0.10 (0.06)	0.03 (0.03)	0.03 (0.02)
Named Entity Recognition						
LSTM	0.00 (0.00)	0.00 (0.00)	0.00 (0.00)	0.00 (0.00)	0.00 (0.00)	0.00 (0.00)
LSTM+CRF	0.05 (0.04)	0.02 (0.01)	0.00 (0.00)	0.04 (0.04)	0.01 (0.01)	0.00 (0.00)
BiLSTM	0.07 (0.06)	0.04 (0.03)	0.02 (0.02)	0.08 (0.08)	0.05 (0.05)	0.04 (0.03)
BiLSTM+CRF	0.08 (0.06)	0.04 (0.03)	0.02 (0.02)	0.09 (0.08)	0.05 (0.04)	0.04 (0.03)
BERT	0.17 (0.06)	0.14 (0.03)	0.13 (0.02)	0.08 (0.06)	0.04 (0.03)	0.03 (0.02)
Part-of-Speech Tagging						
LSTM	0.00 (0.00)	0.00 (0.00)	0.00 (0.00)	0.00 (0.00)	0.00 (0.00)	0.00 (0.00)
LSTM+CRF	0.05 (0.03)	0.00 (0.00)	0.00 (0.00)	0.03 (0.03)	0.00 (0.00)	0.00 (0.00)
BiLSTM	0.08 (0.06)	0.03 (0.02)	0.02 (0.01)	0.07 (0.06)	0.04 (0.03)	0.03 (0.02)
BiLSTM+CRF	0.09 (0.06)	0.03 (0.03)	0.02 (0.02)	0.07 (0.06)	0.04 (0.03)	0.03 (0.02)
BERT	0.52 (0.05)	0.48 (0.01)	0.46 (0.01)	0.07 (0.05)	0.02 (0.02)	0.02 (0.02)
Semantic Role Labeling						
LSTM	0.00 (0.00)	0.00 (0.00)	0.00 (0.00)	0.00 (0.00)	0.00 (0.00)	0.00 (0.00)
LSTM+CRF	0.21 (0.29)	0.13 (0.23)	0.09 (0.18)	0.32 (0.26)	0.26 (0.21)	0.21 (0.17)
BiLSTM	0.31 (0.37)	0.26 (0.33)	0.23 (0.30)	0.34 (0.40)	0.31 (0.37)	0.28 (0.34)
BiLSTM+CRF	0.28 (0.31)	0.24 (0.27)	0.21 (0.23)	0.28 (0.35)	0.25 (0.32)	0.22 (0.28)
BERT	0.43 (0.33)	0.40 (0.29)	0.37 (0.25)	0.22 (0.26)	0.19 (0.23)	0.16 (0.20)
Slot Filling (ATIS)						
LSTM	0.00 (0.00)	0.00 (0.00)	0.00 (0.00)	0.00 (0.00)	0.00 (0.00)	0.00 (0.00)
LSTM+CRF	0.02 (0.01)	0.00 (0.00)	0.00 (0.00)	0.01 (0.03)	0.00 (0.00)	0.00 (0.00)
BiLSTM	0.02 (0.02)	0.01 (0.01)	0.00 (0.00)	0.02 (0.02)	0.01 (0.01)	0.00 (0.00)
BiLSTM+CRF	0.03 (0.01)	0.01 (0.00)	0.00 (0.00)	0.04 (0.02)	0.01 (0.01)	0.00 (0.00)
BERT	0.22 (0.03)	0.15 (0.01)	0.13 (0.00)	0.06 (0.03)	0.00 (0.01)	0.00 (0.00)
Slot Filling (SNIPS)						
LSTM	0.00 (0.00)	0.00 (0.00)	0.00 (0.00)	0.00 (0.00)	0.00 (0.00)	0.00 (0.00)
LSTM+CRF	0.06 (0.04)	0.02 (0.01)	0.01 (0.00)	0.10 (0.07)	0.03 (0.02)	0.02 (0.00)
BiLSTM	0.12 (0.08)	0.07 (0.05)	0.05 (0.03)	0.17 (0.13)	0.10 (0.07)	0.07 (0.04)
BiLSTM+CRF	0.11 (0.08)	0.05 (0.05)	0.03 (0.03)	0.17 (0.14)	0.09 (0.07)	0.06 (0.05)
BERT	0.37 (0.08)	0.29 (0.04)	0.23 (0.02)	0.12 (0.09)	0.05 (0.04)	0.03 (0.02)
SEQUENCE CLASSIFICATION						
Intent (ATIS)						
LSTM	0.48 (0.21)	0.34 (0.13)	0.26 (0.09)	0.77 (0.78)	0.72 (0.72)	0.65 (0.64)
LSTM+CRF	-	-	-	-	-	-
BiLSTM	0.40 (0.24)	0.32 (0.16)	0.25 (0.12)	0.66 (0.77)	0.51 (0.70)	0.36 (0.63)
BiLSTM+CRF	-	-	-	-	-	-
BERT	0.49 (0.20)	0.16 (0.12)	0.11 (0.09)	0.38 (0.29)	0.24 (0.18)	0.15 (0.13)
Intent (SNIPS)						
LSTM	0.31 (0.24)	0.23 (0.16)	0.16 (0.11)	0.57 (0.49)	0.48 (0.40)	0.40 (0.33)
LSTM+CRF	-	-	-	-	-	-
BiLSTM	0.26 (0.23)	0.22 (0.16)	0.16 (0.10)	0.55 (0.49)	0.48 (0.42)	0.41 (0.34)
BiLSTM+CRF	-	-	-	-	-	-
BERT	0.30 (0.22)	0.25 (0.14)	0.18 (0.08)	0.40 (0.35)	0.31 (0.26)	0.23 (0.16)
Positive/Negative						
LSTM	0.38 (0.40)	0.31 (0.31)	0.24 (0.25)	0.65 (0.68)	0.59 (0.62)	0.52 (0.56)
LSTM+CRF	-	-	-	-	-	-
BiLSTM	0.37 (0.39)	0.31 (0.31)	0.27 (0.27)	0.63 (0.58)	0.56 (0.52)	0.48 (0.46)
BiLSTM+CRF	-	-	-	-	-	-
BERT	0.58 (0.45)	0.49 (0.35)	0.41 (0.27)	0.56 (0.55)	0.48 (0.48)	0.41 (0.41)
Pros/Cons						
LSTM	0.15 (0.13)	0.10 (0.08)	0.07 (0.06)	0.32 (0.27)	0.25 (0.20)	0.19 (0.15)
LSTM+CRF	-	-	-	-	-	-
BiLSTM	0.14 (0.14)	0.08 (0.09)	0.06 (0.06)	0.32 (0.30)	0.26 (0.23)	0.21 (0.17)
BiLSTM+CRF	-	-	-	-	-	-
BERT	0.20 (0.14)	0.11 (0.08)	0.08 (0.06)	0.24 (0.22)	0.16 (0.14)	0.11 (0.10)

Table 10: Mean Edit Overhead and Delay of one or two time steps. Values in parentheses refer to using truncated samples during training.

Task/Model	RC			RC with prophecies		
	$\Delta 0$	$\Delta 1$	$\Delta 2$	$\Delta 0$	$\Delta 1$	$\Delta 2$
SEQUENCE TAGGING						
Chunk						
LSTM	1.00 (1.00)	1.00 (1.00)	1.00 (1.00)	1.00 (1.00)	1.00 (1.00)	1.00 (1.00)
LSTM+CRF	0.94 (0.97)	0.99 (1.00)	1.00 (1.00)	0.97 (0.97)	1.00 (1.00)	1.00 (1.00)
BiLSTM	0.84 (0.89)	0.91 (0.93)	0.94 (0.95)	0.90 (0.87)	0.94 (0.92)	0.96 (0.94)
BiLSTM+CRF	0.79 (0.92)	0.93 (0.96)	0.95 (0.97)	0.91 (0.91)	0.94 (0.95)	0.96 (0.97)
BERT	0.67 (0.91)	0.74 (0.94)	0.75 (0.95)	0.87 (0.92)	0.94 (0.95)	0.95 (0.96)
Named Entity Recognition						
LSTM	1.00 (1.00)	1.00 (1.00)	1.00 (1.00)	1.00 (1.00)	1.00 (1.00)	1.00 (1.00)
LSTM+CRF	0.95 (0.96)	0.98 (0.99)	1.00 (1.00)	0.96 (0.96)	0.99 (0.99)	1.00 (1.00)
BiLSTM	0.91 (0.93)	0.95 (0.95)	0.96 (0.97)	0.91 (0.92)	0.95 (0.95)	0.96 (0.96)
BiLSTM+CRF	0.92 (0.94)	0.95 (0.97)	0.97 (0.98)	0.92 (0.93)	0.95 (0.96)	0.97 (0.97)
BERT	0.49 (0.93)	0.51 (0.96)	0.53 (0.97)	0.90 (0.93)	0.94 (0.96)	0.96 (0.97)
Part-of-Speech Tagging						
LSTM	1.00 (1.00)	1.00 (1.00)	1.00 (1.00)	1.00 (1.00)	1.00 (1.00)	1.00 (1.00)
LSTM+CRF	0.95 (0.96)	1.00 (1.00)	1.00 (1.00)	0.97 (0.97)	1.00 (1.00)	1.00 (1.00)
BiLSTM	0.89 (0.93)	0.95 (0.96)	0.96 (0.98)	0.92 (0.93)	0.95 (0.96)	0.97 (0.98)
BiLSTM+CRF	0.89 (0.92)	0.95 (0.96)	0.97 (0.97)	0.92 (0.93)	0.96 (0.96)	0.97 (0.97)
BERT	0.30 (0.93)	0.30 (0.97)	0.29 (0.97)	0.90 (0.93)	0.96 (0.96)	0.97 (0.97)
Semantic Role Labeling						
LSTM	1.00 (1.00)	1.00 (1.00)	1.00 (1.00)	1.00 (1.00)	1.00 (1.00)	1.00 (1.00)
LSTM+CRF	0.71 (0.83)	0.83 (0.87)	0.89 (0.91)	0.82 (0.85)	0.87 (0.89)	0.90 (0.92)
BiLSTM	0.59 (0.60)	0.66 (0.65)	0.70 (0.69)	0.62 (0.62)	0.66 (0.66)	0.70 (0.70)
BiLSTM+CRF	0.65 (0.72)	0.72 (0.76)	0.76 (0.80)	0.71 (0.72)	0.75 (0.76)	0.79 (0.79)
BERT	0.25 (0.70)	0.27 (0.74)	0.28 (0.77)	0.69 (0.67)	0.74 (0.72)	0.77 (0.75)
Slot Filling (ATIS)						
LSTM	1.00 (1.00)	1.00 (1.00)	1.00 (1.00)	1.00 (1.00)	1.00 (1.00)	1.00 (1.00)
LSTM+CRF	0.98 (0.98)	1.00 (1.00)	1.00 (1.00)	0.99 (0.97)	1.00 (1.00)	1.00 (1.00)
BiLSTM	0.97 (0.98)	0.99 (0.99)	1.00 (1.00)	0.97 (0.98)	0.99 (0.99)	1.00 (1.00)
BiLSTM+CRF	0.97 (0.98)	0.99 (0.99)	1.00 (1.00)	0.95 (0.98)	0.99 (0.99)	1.00 (1.00)
BERT	0.56 (0.97)	0.65 (0.99)	0.71 (0.99)	0.93 (0.97)	0.99 (0.99)	1.00 (0.99)
Slot Filling (SNIPS)						
LSTM	1.00 (1.00)	1.00 (1.00)	1.00 (1.00)	1.00 (1.00)	1.00 (1.00)	1.00 (1.00)
LSTM+CRF	0.93 (0.96)	0.98 (0.99)	0.99 (1.00)	0.90 (0.93)	0.97 (0.99)	0.99 (1.00)
BiLSTM	0.84 (0.90)	0.91 (0.93)	0.94 (0.96)	0.81 (0.85)	0.88 (0.91)	0.92 (0.94)
BiLSTM+CRF	0.88 (0.91)	0.94 (0.95)	0.96 (0.96)	0.82 (0.86)	0.90 (0.93)	0.93 (0.95)
BERT	0.41 (0.91)	0.50 (0.95)	0.60 (0.97)	0.86 (0.90)	0.94 (0.95)	0.97 (0.97)
SEQUENCE CLASSIFICATION						
Intent (ATIS)						
LSTM	0.77 (0.91)	0.84 (0.95)	0.89 (0.96)	0.52 (0.54)	0.58 (0.59)	0.64 (0.66)
LSTM+CRF	-	-	-	-	-	-
BiLSTM	0.85 (0.90)	0.88 (0.94)	0.91 (0.95)	0.70 (0.50)	0.77 (0.55)	0.85 (0.61)
BiLSTM+CRF	-	-	-	-	-	-
BERT	0.84 (0.91)	0.93 (0.94)	0.95 (0.96)	0.88 (0.90)	0.92 (0.93)	0.95 (0.96)
Intent (SNIPS)						
LSTM	0.85 (0.90)	0.89 (0.93)	0.92 (0.95)	0.71 (0.77)	0.75 (0.80)	0.80 (0.84)
LSTM+CRF	-	-	-	-	-	-
BiLSTM	0.86 (0.91)	0.89 (0.93)	0.91 (0.95)	0.71 (0.76)	0.75 (0.80)	0.80 (0.84)
BiLSTM+CRF	-	-	-	-	-	-
BERT	0.83 (0.91)	0.86 (0.94)	0.91 (0.96)	0.83 (0.86)	0.87 (0.90)	0.90 (0.94)
Positive/Negative						
LSTM	0.78 (0.79)	0.81 (0.82)	0.84 (0.85)	0.74 (0.72)	0.76 (0.75)	0.79 (0.77)
LSTM+CRF	-	-	-	-	-	-
BiLSTM	0.82 (0.83)	0.84 (0.85)	0.86 (0.87)	0.76 (0.77)	0.79 (0.79)	0.82 (0.81)
BiLSTM+CRF	-	-	-	-	-	-
BERT	0.64 (0.80)	0.66 (0.84)	0.68 (0.86)	0.79 (0.79)	0.82 (0.82)	0.84 (0.84)
Pros/Cons						
LSTM	0.93 (0.94)	0.95 (0.96)	0.96 (0.97)	0.88 (0.90)	0.90 (0.92)	0.93 (0.94)
LSTM+CRF	-	-	-	-	-	-
BiLSTM	0.94 (0.94)	0.96 (0.96)	0.97 (0.97)	0.85 (0.88)	0.87 (0.91)	0.90 (0.93)
BiLSTM+CRF	-	-	-	-	-	-
BERT	0.91 (0.94)	0.94 (0.96)	0.95 (0.97)	0.91 (0.92)	0.94 (0.94)	0.95 (0.96)

Table 11: Mean Relative Correctness and Delay of one or two time steps. Values in parentheses refer to using truncated samples during training.