
Dependency Parsing is More Parameter-Efficient with Normalization

Paolo Gajo
University of Bologna
paolo.gajo2@unibo.it

Domenic Rosati
Dalhousie University
Domenic.Rosati@Dal.Ca

Hassan Sajjad
Dalhousie University
HSajjad@dal.ca

Alberto Barrón-Cedeño
University of Bologna
a.barron@unibo.it

Abstract

Dependency parsing is the task of inferring natural language structure, often approached by modeling word interactions via attention through biaffine scoring. This mechanism works like self-attention in Transformers, where scores are calculated for every pair of words in a sentence. However, unlike Transformer attention, biaffine scoring does not use normalization prior to taking the softmax of the scores. In this paper, we provide theoretical evidence and empirical results revealing that a lack of normalization necessarily results in overparameterized parser models, where the extra parameters compensate for the sharp softmax outputs produced by high variance inputs to the biaffine scoring function. We argue that biaffine scoring can be made substantially more efficient by performing score normalization. We conduct experiments on semantic and syntactic dependency parsing in multiple languages, along with latent graph inference on non-linguistic data, using various settings of a k -hop parser. We train N -layer stacked BiLSTMs and evaluate the parser’s performance with and without normalizing biaffine scores. Normalizing allows us to achieve state-of-the-art performance with fewer samples and trainable parameters. Code: <https://github.com/paolo-gajo/EfficientSDP>

1 Introduction

Dependency parsing (DP) consists in classifying node labels t_i (words), edges e_{ij} (relations), and edge labels r_{ij} (relation types) of a dependency graph [31]. A popular model for this task, introduced by Dozat and Manning [7], entails modeling word interactions as a fully connected graph via biaffine attention. Despite its simplicity, a number of models using this biaffine transformation [3, 7, 8, 13, 14] require more parameters than necessary, due to what we identify as a lack of normalization of its outputs. We hypothesize that this overparameterization is caused by the high variance of the outputs, which extra parameters help mitigate. After showing that variance can be reduced through normalization, we demonstrate that similar or better performance can be obtained for DP with fewer parameters and training samples.

In this task, we consider a graph $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ comprising a set of nodes $\mathcal{V}$, connected by a set of edges $\mathcal{E}$, with each edge $e_{ij} \in \mathcal{E}$ connecting pairs of nodes $(v_i, v_j) \in \mathcal{V}$. In latent graph inference (LGI) for dependency graphs, a sentence of $|\mathcal{V}|$ words is modeled as a $|\mathcal{V}| \times |\mathcal{V}|$ graph via biaffine scoring [7]:

$$XW_X^\top = X(W_QW_K^\top)X^\top = XW_Q(XW_K)^\top = QK^\top, \quad X \in \mathbb{R}^{|\mathcal{V}| \times d}.$$

Observe that this is equivalent to the unnormalized scores used in [38]’s self-attention: $QK^\top/\sqrt{d_k}$ where $\text{Attention}(Q, K, V) = \text{Softmax}(QK^\top/\sqrt{d_k})V$ and d_k is the output size of the keys.

The main difference resides in the fact that Transformer attention scores are scaled by $a = 1/\sqrt{d_k}$. As explained in [38], this scaling is done because high variance inputs to the softmax function will result in large values dominating the output ($e^{x \gg 1}$) and small values decaying ($e^{x \ll 1}$). The downstream effect is exploding and vanishing gradients. To see how this works, observe that the score s between any query and key vector is the result of their dot product $q_i k_i^\top$. If we assume that these vectors are zero mean and unit variance, then the variance is $\text{Var}(s) = d_k$. Finally, we get our normalization factor, which ensures that each entry in the score matrix has a standard deviation of 1, since $\text{Std} = \sqrt{d_k}$. Consequently, lower input variance will result in more stable outputs.

We observe that there is no consistency in the literature on the use of this scaling term for DP. Most works do not use this scaling term [3, 7, 8, 13, 14], with the exception of [37] who use [38]’s self-attention out of the box. In this paper, we carry out a series of experiments on semantic and syntactic DP tasks, along with non-linguistic LGI tasks, using a wide range of architecture ablations. We extend the work of Bhatt et al. [3], which represents the state of the art on dependency graph parsing in NLP.

Following the state of the art, we train stacked BiLSTMs and a biaffine classifier to infer the latent dependency graph connecting the words of a sentence. We find that, when not normalizing the scores produced by the biaffine transformation, model performance drops in terms of micro-averaged F_1 -measure and attachment score [26]. In particular, increasing the amount of layers produces a normalization effect by reducing the variance of the output scores. Using score normalization, we find that in some cases similar or better performance can be obtained by reducing the amount of trained BiLSTM parameters by as much as 85%.

Our contributions are thus three-fold. (i) We show the impact of normalizing the output of the biaffine scorer in relation to the architectural changes in a DP model. (ii) We show DP models can obtain better performance with substantially fewer trained parameters. (iii) We provide a new method that substantially improves scores and obtains state-of-the-art performance for semantic [21, 46] and syntactic [24, 48] dependency parsing.

The rest of the paper is structured as follows. Section 2 presents an overview of the literature concerning DP, with specific focus on how to infer the adjacency matrix of a dependency graph. Section 3 explains why layer depth can compensate for normalization. Section 4 provides an overview of the datasets, models, evaluation, and other experimental details. Section 5 illustrates how normalization mitigates overparameterization. Finally, Section 6 draws conclusions and provides avenues for future research.

2 Background

LGI involves predicting all, or a proper subset of, the edges between the nodes of a graph [15, 19, 39]. For example, one could carry out LGI over molecule graphs [33] to predict the bonds between atoms or the proximity of pixel clusters in an image [9].

In the context of NLP, DP is a task concerned with inferring the dependency graph of a sentence [31]. Depending on the nature of the nodes and relations of the graph, it can be described as semantic (SemDP) [3, 8, 14, 37] or syntactic (SynDP) [7, 13, 48]. In this work, we tackle both tasks, but focus more specifically on SemDP, which can be thought of as comprising two sub-tasks: named entity recognition (NER) [23] and relation extraction (RE) [3, 14]. They can be either approached in a pipeline as separate objectives [5, 47, 50, 53] or by jointly training a single model on both sub-tasks [3, 4, 14, 49]. In the context of deep learning models trained end-to-end, three main paradigms are used: encoder-based models [3, 7, 8, 13, 14, 41, 43], decoder-only Transformers, more commonly known as large language models (LLMs) [4, 44, 51], and seq2seq encoder-decoder Transformers [17, 20, 27, 49].

In this work, we focus on encoder-style models, since they currently achieve the best performance on DP tasks [3, 14, 37] and are much more parameter efficient than LLM-based solutions. These models approach entity (node) prediction analogously to NER, while edges and relations are handled via MLP projection of node-pair features [16, 28, 29, 54] or via attention-based edge and relation inference [3, 8, 13, 14, 37]. Current state-of-the-art models, exemplified by [3], are based on the biaffine dependency parser introduced by [7], which uses stacked BiLSTM layers on top of embeddings from a pre-trained language model.

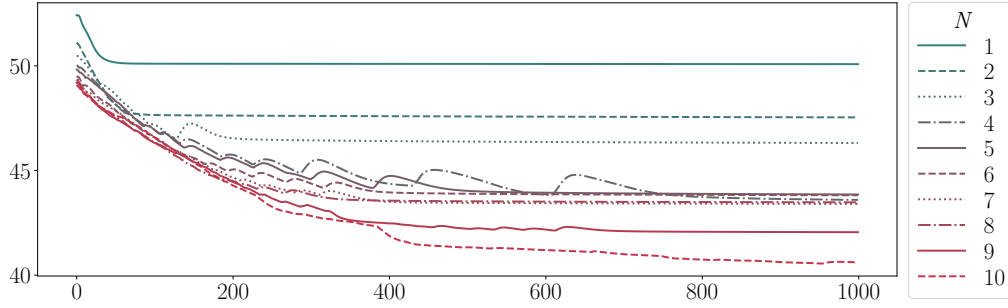


Figure 1: Effective rank $\rho(W)$ reduces over training epochs as we increase N of BiLSTM layers.

3 Layer depth can compensate for normalization

In order to reduce the number of parameters used for biaffine scoring, we observe that the softmax function is sensitive to inputs with high variance: large values dominate the output ($e^{x \gg 1}$) and small values decay ($e^{x \ll 1}$). This causes a drop in the downstream task performance, due to some values dominating the probability outputs in the score matrix, as well as exploding and vanishing gradients. Typically, contemporary architectures based on [7] do not employ normalization, but are still able to perform well on DP tasks.

We explain this discrepancy using insights from the theory of implicit regularization [1], which contains the result that, as layer depth increases, the effective rank of the weight matrices is reduced during gradient descent. Our claim is that a reduction in the rank of the weight matrices causes a corresponding decrease in the variance of the input to the softmax function.

Result 1 (Singular Value Dynamics Under Gradient Descent [1]). *Minimization of a loss function $\mathcal{L}(W)$ with gradient descent using weight matrices $W \in \mathbb{R}^{n \times n}$ (assuming a small learning rate η and initialization close to the origin). An N -layer linear neural network leads the singular values σ_r of W to evolve in the number of iterations t by:*

$$\sigma_r(t+1) \leftarrow \sigma_r(t) - \eta \cdot \langle \nabla \mathcal{L}(W(t)), \mathbf{u}_r(t) \mathbf{v}_r^\top(t) \rangle \cdot N \cdot \sigma_r(t)^{2-2/N},$$

where $\nabla \mathcal{L}(W(t)) \in \mathbb{R}^{n \times n}$, $\mathbf{u}_r, \mathbf{v}_r$ are the corresponding right and left singular vectors, and $\langle \cdot, \cdot \rangle$ is the Frobenius inner product. This implies that, as the number of layers increases, the rank of the weight matrices decreases, since the smallest singular values decay.

While Result 1 is stated for deep linear neural networks for the matrix factorization task, it has been validated empirically as applying to deep non-linear neural networks [34, 52]. We validate this finding for the biaffine scoring setting in Figure 1, which shows the inverse proportionality between the effective rank $\rho(W)$ [36] of the weights of a BiLSTM and the number N of its layers.

Claim 1 (Monotonic Increase of Output Variance with Rank). *Let $X \in \mathbb{R}^n$ be a random vector with covariance matrix $\mathbf{K}_{xx} := \text{Cov}(X) \in \mathbb{R}^{n \times n}$, and let $\mathbf{A} \in \mathbb{R}^{m \times n}$ be a fixed matrix with singular value decomposition (SVD):*

$$\mathbf{A} = \sum_{i=1}^{\min(m,n)} \sigma_i \mathbf{u}_i \mathbf{v}_i^\top,$$

where $\sigma_1 \geq \sigma_2 \geq \dots \geq 0$. Define the best rank- r approximation of $\mathbf{A}$ by truncated SVD as:

$$\mathbf{A}_r := \sum_{i=1}^r \sigma_i \mathbf{u}_i \mathbf{v}_i^\top.$$

Let $Y_r = \mathbf{A}_r X \in \mathbb{R}^m$ denote the image of X under the rank- r linear map. Then, the total variance of Y_r , as measured by the trace of its covariance matrix, increases monotonically with r :

$$\text{tr}(\text{Cov}(Y_r)) \leq \text{tr}(\text{Cov}(Y_{r+1})).$$

Proof. Let $\mathbf{K}_{xx} := \text{Cov}(X) \in \mathbb{R}^{n \times n}$ denote the covariance matrix of the random vector X . Since covariance matrices are symmetric and positive semi-definite (PSD), we have $\mathbf{K}_{xx} \geq 0$. Let $\mathbf{A}_r := \sum_{i=1}^r \sigma_i \mathbf{u}_i \mathbf{v}_i^\top$ be the rank- r truncated SVD of $\mathbf{A}$, the covariance matrix of $Y_r := \mathbf{A}_r X$ is:

$$\text{Cov}(Y_r) = \mathbf{A}_r \mathbf{K}_{xx} \mathbf{A}_r^\top.$$

To evaluate the total variance we compute the trace:

$$\text{tr}(\text{Cov}(Y_r)) = \text{tr}(\mathbf{A}_r \mathbf{K}_{xx} \mathbf{A}_r^\top).$$

Using the cyclic property of the trace ($\text{tr}(ABC) = \text{tr}(BCA)$) and symmetry of $\mathbf{K}_{xx}$ we get:

$$\text{tr}(\mathbf{A}_r \mathbf{K}_{xx} \mathbf{A}_r^\top) = \text{tr}(\mathbf{K}_{xx} \mathbf{A}_r^\top \mathbf{A}_r).$$

Now define the matrix $\mathbf{M}_r := \mathbf{A}_r^\top \mathbf{A}_r \in \mathbb{R}^{n \times n}$, which is PSD. As r increases, $\mathbf{A}_r$ includes more terms in its truncated SVD, so we have:

$$\mathbf{M}_r = \sum_{i=1}^r \sigma_i^2 \mathbf{v}_i \mathbf{v}_i^\top, \quad \mathbf{M}_{r+1} = \mathbf{M}_r + \sigma_{r+1}^2 \mathbf{v}_{r+1} \mathbf{v}_{r+1}^\top.$$

Thus:

$$\mathbf{M}_{r+1} \geq \mathbf{M}_r.$$

Since $\mathbf{K}_{xx} \geq 0$, and since the trace of a product of PSD matrices respects Loewner order (i.e., if $A \leq B$, then $\text{tr}(CA) \leq \text{tr}(CB)$ for all $C \geq 0$), we conclude:

$$\text{tr}(\mathbf{K}_{xx} \mathbf{M}_r) \leq \text{tr}(\mathbf{K}_{xx} \mathbf{M}_{r+1}),$$

which implies:

$$\text{tr}(\text{Cov}(Y_r)) \leq \text{tr}(\text{Cov}(Y_{r+1})).$$

Therefore, as the rank r increases, the total variance Y_r increases. $\square$

The implication of Claim 1 is that the pre-softmax score matrix from a shallow network without normalization will have a higher variance than a deeper network because of Result 1. Based on this finding, we propose to remove BiLSTM layers and use normalization—rather than having greater layer depth—in order to develop a more parameter-efficient method.

4 Experimental setting

4.1 Data

To train and evaluate our models, we use four SemDP datasets,¹ along with the 2.2 version of the Universal Dependencies English EWT treebank (enEWT) [24] and SciDTB [48] for SynDP. Table 1 summarizes datasets statistics and reports their entity and relation class annotations.

As regards SemDP, **ADE** [12] is a medical-domain dataset comprising reports of drug adverse-effect reactions. Each sample contains a single relation, “adverseEffect,” which links a disease to a drug. **CoNLL04** [35] contains news texts and is annotated with the classic named entities $e_i \in \{\text{per, org, loc}\}$ and relations $r_j \in \{\text{workFor, kill, orgBasedIn, liveIn, locIn}\}$. ADE and CoNLL04 are characterized by relations which are not complex enough to form connected graphs of considerable size. **SciERC** [21] is a dataset compiled from sentences extracted from artificial intelligence literature. It is arguably the most challenging dataset used in this study, since most of the entities in the validation and testing partitions are not assigned an entity class. This makes it hard to train tag embeddings that can help to infer edges. **ERFGC** [46] is a dataset comprising “flow graphs” parsed from culinary recipes. The semantic dependency graphs of these recipes are directed and acyclic, with a single root. Differently from [3], and as advised directly by [46] (personal correspondence), we ignore the “-” edge labels present in the corpus, since they link discontinuous parts of phrasal verbs and are thus irrelevant. For ADE, CoNLL04, and SciERC we use the splits provided in [4].² ERFGC is not available online; we obtained it by contacting the authors of [46].

As regards SynDP, we use **enEWT** [24] to compare directly against [13]’s results. While many works evaluate SynDP on the Penn Treebank [30], it is closed-access and prohibitively expensive. Instead, we use **SciDTB** [48], a discourse analysis dataset comprising 798 abstracts extracted from the ACL Anthology.³ It was processed for the syntax dependency parsing task using Stanza [32].

¹We neglect ACE04 [22] and ACE05 [42], two popular datasets for RE, due to their prohibitive cost.

²<https://drive.google.com/drive/folders/1vVKJIUzK4hIipfdEGmSOCCoFmUmZw0QV>

³<https://aclanthology.org>

Table 1: Size of the train/dev/test splits and entity/relation classes for each dataset.

Data		Entities	Relations
ADE	[12]	disease, drug	adverseEffect
2,563 / 854 / 300			
CoNLL04	[35]	organization, person, location	kill, locatedIn, workFor, orgBasedIn, liveIn
922 / 231 / 288			
SciERC	[21]	generic, material, method, metric, otherSciTerm, task	usedFor, featureOf, hyponymOf, evaluateFor, partOf, compare, conjunction
1,366 / 187 / 397			
ERFGC	[46]	food, tool, duration, quantity, action-ByChef, discountAction, actionBy-Food, actionByTool, foodState, tool-State	agent, target, indirectObject, tool-Complement, foodComplement, foodEq, foodPartOf, foodSet, toolEq, toolPartOf, actionEq, timingHeadVerb, other
242 / 29 / 29			
enEWT	[24]	xPOS tags	UD relations
10,098 / 1,431 / 1,427			
SciDTB	[48]	xPOS tags	UD relations
2,567 / 814 / 817			

From these two datasets, we consider the xPOS tags (e.g., noun, preposition, determiner) and the syntactic dependencies (e.g., sentence root, object, adverbial modifier) to respectively be the entities and relations of the dependency graphs to infer.

In order to show that the normalization effect is language-independent, in Appendix A we also carry out SynDP experiments using Universal Dependencies [25] datasets in six other languages. Finally, we also conduct experiments on three non-linguistic datasets: PCQM-Contact [10], CIFAR10 Superpixel [9], and QM9 [33]. PCQM-Contact and QM9 are datasets which contain graphs of molecules, while CIFAR10 Superpixel contains superpixel clusters constructed from CIFAR10 images. We use these datasets to further show that the observed effects of normalization are not limited to linguistic dependency parsing, but are a general phenomenon of carrying out inference of fully connected graphs. We report the results of this additional experiment in Appendix B.

4.2 Model

We adopt the architecture of [3] in our work, schematized in Figure 2. It can be subdivided into four main components: encoder, tagger, parser, and decoder. The input is tokenized and passed through a BERT-like encoder, where token representations are averaged into $|\mathcal{V}|$ word-level features $\mathbf{x}_i \in \mathbb{R}^{d_f}$.⁴ Optionally, additional features can be obtained by predicting the entity classes of each word with a tagger, composed by a single-layer BiLSTM ϕ , followed by a classifier:

$$\mathbf{h}_i^{tag} = \phi(\mathbf{x}_i), \quad \mathbf{h}_i^{tag} \in \mathbb{R}^{d_h},$$

$$\mathbf{y}_i^{tag} = \text{Softmax}(\text{MLP}^{tag}(\mathbf{h}_i^{tag})), \quad \mathbf{y}_i^{tag} \in \mathbb{R}^{|T|},$$

where T is the set of word tag classes. The tagger’s predictions are then converted into one-hot vectors and projected into dense representations by another MLP, such that $\mathbf{e}_i^{tag} = \text{MLP}^{emb}(\mathbf{1}_T(\mathbf{y}_i^{tag}))$. These new tag embeddings are concatenated with the original BERT output and sent to the parser.

In the parser, an optional N -layered BiLSTM ψ produces new representations $\mathbf{h}_i = \psi(\mathbf{e}_i^{tag} \oplus \mathbf{x}_i)$, which are then projected into four different representations:

$$\mathbf{e}_i^h = \text{MLP}^{(edge-head)}(\mathbf{h}_i), \quad \mathbf{e}_i^d = \text{MLP}^{(edge-dept)}(\mathbf{h}_i),$$

$$\mathbf{r}_i^h = \text{MLP}^{(rel-dept)}(\mathbf{h}_i), \quad \mathbf{r}_i^d = \text{MLP}^{(rel-head)}(\mathbf{h}_i).$$

The edge scores s_i^{edge} and relation scores s_i^{rel} are then calculated with the biaffine function f :

$$f(\mathbf{x}_1, \mathbf{x}_2; W) = \mathbf{x}_1^\top W \mathbf{x}_2 + \mathbf{x}_1^\top \mathbf{b},$$

⁴Using token-level representations resulted in much lower performance in preliminary experiments.

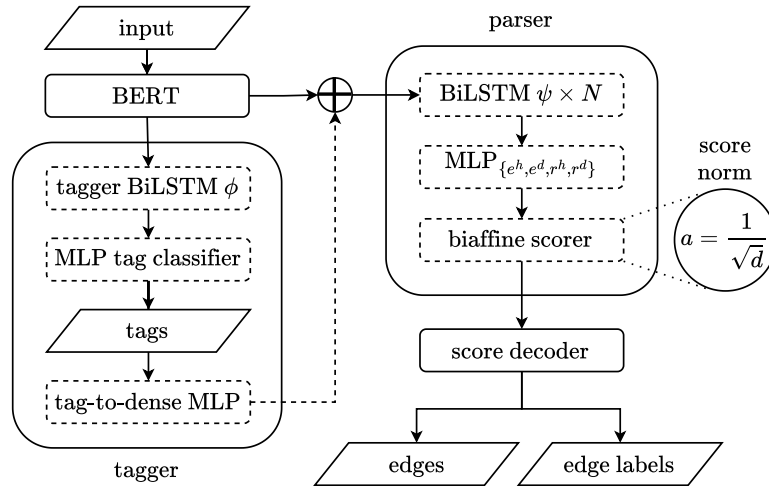


Figure 2: Dependency parsing diagram. Dashed components are the targets of ablation experiments.

$$s_i^{edge} = f^{(edge)}(\mathbf{e}_i^h, \mathbf{e}_i^d; W_e), \quad W_e \in \mathbb{R}^{d \times 1 \times d},$$

$$s_i^{rel} = f^{(rel)}(\mathbf{r}_i^h, \mathbf{r}_i^d; W_r), \quad W_r \in \mathbb{R}^{d \times |R| \times d},$$

where R is the set of relation classes, i.e., the possible labels applied to an edge. In Appendix C, following [13], we extend the architecture and instead feed the representations produced by the BERT encoder into $L_\gamma \in \{0, 1, 2, 3\}$ pairs of biaffine and graph attention network (GAT) [40] layers. This allows us to see the influence of normalization when using a k -hop parser.

Finally, in the decoder, the edge scores are used in conjunction with the relation representations $\mathbf{r}_i^h$ and $\mathbf{r}_i^d$ to obtain the final predictions. During training, we do greedy decoding, while during inference, we use Chu-Liu/Edmonds' maximum spanning tree (MST) algorithm [11] to ensure the predictions are well-formed trees. This is especially useful with big dependency graphs, since greedy decoding is more likely to produce invalid trees as size increases. When doing greedy decoding, an edge index (i.e., an adjacency matrix) $a_i = \arg \max_j s_{ij}^{edge}$ is produced by taking the argmax of the attention scores s_i^{edge} across the last dimension. The edge index is then used to select which head relation representations $\mathbf{r}_i^h$ to use to calculate the relation scores $s_i^{rel} = f(\mathbf{r}_i^h, \mathbf{r}_i^d; W)$, $W \in \mathbb{R}^{d \times |R| \times d}$. The relations are then predicted as $r_i = \arg \max_j s_{ij}^{rel}$. When using MST decoding, edge and relation scores are combined into a single energy matrix where each entry represents the score of a specific head-dependent pair with its most likely relation type. This energy matrix is then used in the MST algorithm, producing trees with a single root and no cycles. Prior to energy calculation, edge scores and relation scores are scaled so that low values are squished and high values are increased, making the log softmax produce a hard adjacency matrix.

The model is trained end-to-end jointly on the entity, edge, and relation classification objectives:

$$\mathcal{L}_{tag} = -\frac{1}{|\mathcal{V}|} \sum_{i=1}^{|\mathcal{V}|} \sum_{t=1}^{|\mathcal{T}|} y_{i,t}^{tag} \log p(y_{i,t}^{tag}),$$

$$\mathcal{L}_{edge} = -\sum_{i,j=1}^{|\mathcal{V}|} \log p(y_{i,j}^{edge} = 1),$$

$$\mathcal{L}_{rel} = -\sum_{i=1}^{|\mathcal{V}|} \sum_{j=1}^{|\mathcal{V}|} \mathbb{1}(y_{i,j}^{edge} = 1) \sum_{\ell=1}^{|R|} y_{i,j,\ell}^{rel} \log p(y_{i,j,\ell}^{rel}),$$

$$\mathcal{L} = \lambda_1 \mathcal{L}_{tag} + \lambda_2 (\mathcal{L}_{edge} + \mathcal{L}_{rel}).$$

Losses are calculated based on the gold tags, edges, and relations. We set $\lambda_1 = 0.1$ and $\lambda_2 = 1$ as hyperparameters because the tagging task is much simpler than predicting the edges, since the same top performance is always achieved regardless of any other selected architecture hyperparameters.

Table 2: Main setting hyperparameter ranges. ψ = Parser BiLSTM. f = biaffine layer.

Component	Hyperparameter	Values
Encoder	Freeze BERT	$\nabla_{\text{BERT}} \in \{\checkmark, \times\}$
Tagger	Tagger BiLSTM ϕ	$\phi \in \{\checkmark, \times\}$
	Concat. tag embeds.	$\mathbf{e}_i^{\text{tag}} \in \{\checkmark, \times\}$
Parser	ψ num. layers	$N \in \{0, 1, 2, 3\}$
	ψ hidden dim.	$h_\psi \in \{100, 200, 300, 400\}$
	MLP ^(edge) output dim.	$d_{\text{MLP}} \in \{100, 300, 500\}$
	ψ LayerNorm	$\text{LN}_\psi \in \{\checkmark, \times\}$
	MLP ^(edge) and $f^{(\text{edge})}$ init.	$I_{\text{par}} \sim \{\mathcal{U}, \mathcal{N}\}$
	Score scaling	$a \in \{1, \frac{1}{\sqrt{d}}\}$

Following the usual approach for SynDP [7, 13, 14], when training on enEWT and SciDTB we use an oracle—the gold tags—and do not predict the POS tags ourselves. Since in this case we only focus on training the edge and relation classification tasks, we set $\lambda_1 = 0$.

4.3 Hyperparameters

We experiment with a range of hyperparameters for the encoder, tagger, and parser, as listed in Table 2. We use BERT_{base} [6] as our pre-trained encoder, which we keep frozen throughout the whole training run in our main setting.

As regards the tagger, we set $L_\phi = 1$ and $h_\phi = 100$, as in [3], with weights initialized with a Xavier uniform distribution. In Appendix D.1, we ablate the ϕ BiLSTM and the use of the tag embeddings $\mathbf{e}_i^{\text{tag}}$ to assess their impact on overall performance. With relation to the parser, for all hyperparameter combinations of $\{N, h_\psi, d_{\text{MLP}}\}$, we run our experiments by initializing its weights with a Xavier uniform distribution ($I_{\text{par}} \sim \mathcal{U}$) and no LayerNorm LN_ψ . In Appendices D.2 to D.5, we conduct ablations over the parser hyperparameters indicated in Table 2.

We train our models for $2k$ steps and evaluate on the development partition of each dataset every 100 steps. For enEWT and SciDTB, we train for $5k$ steps with $1k$ -step validation intervals to make our results more comparable with the state of the art [13]. We apply early stopping at 30% of the total steps without improvement and choose the best model based on the top performance on the development split. In Appendix D.6, we extend the training to $10k$ steps and fully fine-tune a variety of small and large pre-trained language models to show time-wise test performance trends more clearly. We set the learning rate at $\eta = 1 \times 10^{-3}$ when the encoder is kept frozen. In all settings, including ablations, we use AdamW [18] as the optimizer and a batch size of 8.

We use [3]’s original architecture as our baseline. It uses a frozen BERT_{base} model as encoder and trains all of the components showed in Figure 2. Following the best results obtained by [7], they use three BiLSTM layers in the parser with a hidden size of 400, while the four MLPs following the stacked BiLSTM have an output size of 500 for the edge representations and 100 for the relations.

4.4 Evaluation

Following [3, 8, 14], we measure tagging and parsing performance on ADE, CoNLL04, SciERC, and ERFGC in terms of micro-averaged F_1 -measure. In addition, for enEWT and SciDTB we use unlabeled (UAS) and labeled (LAS) attachment score [26]. To corroborate the validity of our results, we train and evaluate each setting, including ablations, with five random seeds. We report mean and standard deviation for the F_1 , UAS, and LAS metrics, averaged over the five runs. To test the significance of our results, we use the one-tailed Wilcoxon signed-rank test [45].

5 Results and discussion

Table 3 shows the results of our experiments for the labeled edge prediction task, with the first row using the same hyperparameters as [3] ($h_\psi = 400, d_{\text{MLP}} = 500$). We also use these hyperparameters

Table 3: Micro-F₁ (SemDP) and LAS (SynDP) for the labeled edge prediction task. Best in bold.

Model	a	N	ADE	CoNLL04	SciERC	ERFGC	enEWT	SciDTB
[3]	1	3	0.653 $\pm$ 0.018	0.566 $\pm$ 0.019	0.257 $\pm$ 0.024	0.701 $\pm$ 0.009	0.804 $\pm$ 0.006	0.915 $\pm$ 0.002
Ours	1	0	0.541 $\pm$ 0.021	0.399 $\pm$ 0.024	0.147 $\pm$ 0.049	0.548 $\pm$ 0.010	0.559 $\pm$ 0.005	0.729 $\pm$ 0.004
		1	0.657 $\pm$ 0.011	0.556 $\pm$ 0.021	0.282 $\pm$ 0.009	0.676 $\pm$ 0.010	0.771 $\pm$ 0.006	0.892 $\pm$ 0.002
		2	0.667 $\pm$ 0.011	0.573 $\pm$ 0.025	0.273 $\pm$ 0.010	0.694 $\pm$ 0.010	0.796 $\pm$ 0.006	0.910 $\pm$ 0.002
		3	0.662 $\pm$ 0.027	0.562 $\pm$ 0.021	0.299 $\pm$ 0.023	0.705 $\pm$ 0.011	0.804 $\pm$ 0.006	0.915 $\pm$ 0.002
	$\frac{1}{\sqrt{d}}$	0	0.567 $\pm$ 0.014	0.438 $\pm$ 0.033	0.181 $\pm$ 0.027	0.612 $\pm$ 0.008	0.646 $\pm$ 0.002	0.796 $\pm$ 0.002
		1	0.668 $\pm$ 0.017	0.597 $\pm$ 0.015	0.299 $\pm$ 0.019	0.692 $\pm$ 0.009	0.789 $\pm$ 0.003	0.904 $\pm$ 0.002
		2	0.676 $\pm$ 0.019	0.596 $\pm$ 0.014	0.312 $\pm$ 0.011	0.699 $\pm$ 0.009	0.805 $\pm$ 0.003	0.916 $\pm$ 0.002
		3	0.686 $\pm$ 0.025	0.602 $\pm$ 0.017	0.320 $\pm$ 0.013	0.708 $\pm$ 0.008	0.807 $\pm$ 0.005	0.919 $\pm$ 0.001

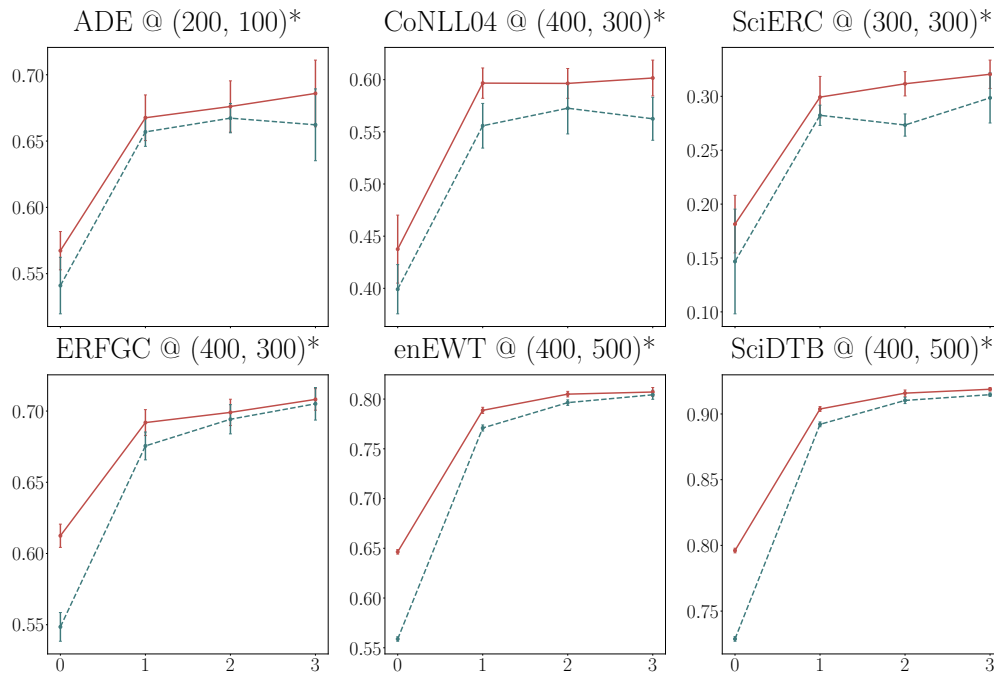


Figure 3: Micro-F₁ (SemDP) and LAS (SynDP) vs $N \in \{0, 1, 2, 3\}$ at $2k$ training steps. Red = norm; blue = raw. *Performance increase with normalization is statistically significant ($p < 0.01$).

for enEWT and SciDTB, since our ablation studies only concerns SemDP due to SynDP being less challenging. The lower half uses the best combinations for ADE (200, 100), CoNLL04 (400, 300), SciERC (300, 300), and ERFGC (400, 300), chosen based on mean performance across these four datasets. For brevity, the performance on the tagging and unlabeled edge prediction tasks are reported in Appendix E. Details on the used computational resources can be found in Appendix F.

Overall, *normalizing biaffine scores provides an evident performance boost at all BiLSTM depths*. In particular, for ERFGC we beat the state-of-the-art performance achieved by [3]. Most of the performance gain is obtained by adding the first layer, with additional ones yielding diminishing returns. In other words, the performance boost provided by score normalization is highest in the absence of the implicit normalization provided by extra parameters ($N > 0$). In general, the top performance obtained by using three BiLSTM layers and no score normalization can be matched or surpassed *with a single BiLSTM layer, when using score normalization*. Taking into account the lower values for h_ψ and d_{MLP} , this represents a reduction in trained parameters of up to 85%.

As laid out in Section 3, score variance tends to decay with deeper BiLSTM stacks. This in turn produces a converging trend as N increases. When looking at Figure 3, this is especially evident

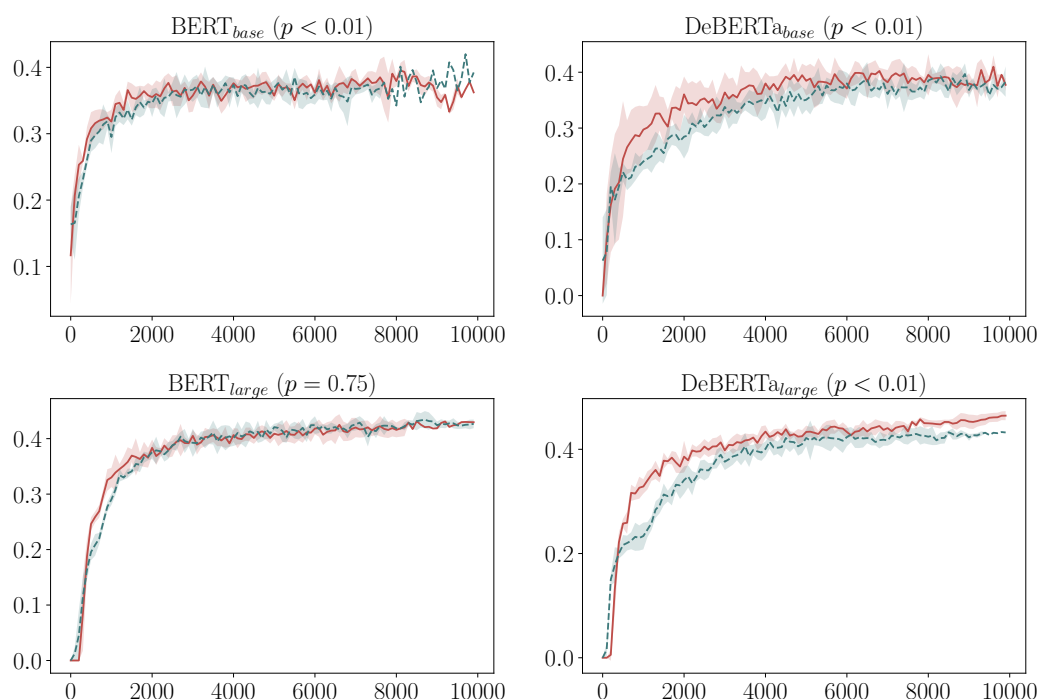


Figure 4: Test performance on SciERC (micro-averaged F_1 -measure vs number of training steps). The p -values indicate greater performance with normalization (one-tailed Wilcoxon signed-rank test).

for ERFGC, enEWT, and SciDTB, for which the beneficial effect of score normalization shrinks smoothly with higher values of N . Although the same cannot be said for ADE, CoNLL04, and SciERC, the performance boost is still statistically significant across all layer depths ($p < 0.01$).

As regards SciERC, normalizing scores without any BiLSTM layers ($N = 0$) produces a 23% increase in performance with a strong reduction in standard deviation. SciERC arguably has the hardest dependency graphs to parse, due to the little overlap between training and testing entities, which makes it difficult to leverage tag embeddings. In addition, it is characterized by complex semantic dependencies. Therefore, we reckon normalizing biaffine scores to be especially important when dealing with hard tasks. In particular, normalizing scores helps mitigate the higher variance of the model’s predictions for this challenging dataset. Conversely, the lack of score normalization exacerbates already uncertain predictions. For SciERC, this makes the trend observed in Figure 3 more unstable. Indeed, the performance first drops at $N = 2$ and then increases once again at $N = 3$, which does not happen for the other datasets.

Compared to existing approaches which do not scale the biaffine scores, fewer parameters can thus be trained to obtain similar performance. Therefore, our results indicate that parsers using raw biaffine scores are likely overparameterized, compared to the performance they could achieve by using score normalization.

In addition, score normalization increases sample efficiency by accelerating convergence. Figure 4 displays the performance on SciERC’s test set for four models during full fine-tuning. As the figure shows, the speedup in convergence is statistically significant when normalizing scores. This shows how score normalization can be beneficial even when fully fine-tuning models with $\sim 10^8$ parameters, given a hard task which forces the model to make uncertain predictions.

6 Conclusions

In this work, we explored the effect of scaling the scores produced by biaffine transformations when predicting the edges of a dependency graph. We have demonstrated, both theoretically and empirically, that the score variance produced by a lack of score scaling hurts model performance

when predicting edges and relations. In addition, our theoretical work and experiments highlight a strong relationship between the number of trained layers and their intrinsic normalization effect.

Departing from a state-of-the-art architecture for semantic dependency parsing, we were able to improve its performance on both semantic and syntactic dependency parsing on six datasets. On ERFGC, a dataset of directed acyclic semantic dependency graphs compiled from culinary recipes, our approach beats the state-of-the-art performance achieved by Bhatt et al. [3]. Moreover, our results showed that a single BiLSTM layer can be sufficient to match or surpass the results of state-of-the-art architectures with a decrease in trained parameters of up to 85%. In the case of SciERC, a challenging dataset for semantic dependency parsing, we found that the performance boost was particularly great when training the parser with three BiLSTM layers. Moreover, for this challenging dataset, we find that scaling the predictions of the biaffine scorer can accelerate convergence speed even when fully fine-tuning models in the 100 to 400M parameter range. In addition, for three of the datasets we also observed that the performance obtained with normalized and raw scores converged smoothly as the number of trained layers increased. This supported our claim that stacking BiLSTM layers mainly serves the purpose of producing an implicit regularization, and that this effect can be obtained by normalizing the scores, without any extra parameters. In additional experiments, we corroborated our findings using data in languages other than English [25], as well as non-linguistic data, for example using a dataset of molecule graphs [33]. Using GAT layers, we further confirmed our hypothesis that GNN-based iterative adjacency refinement [13] also benefits from normalization.

In future work, we plan to extend our experiments to large-scale graph inference tasks which have been limited by the lack of parameter-efficient methods, such as long-form discourse parsing tasks. With regard to score scaling, we also wish to verify whether the positive effects of this normalization can carry over to other tasks, beyond latent graph inference. Finally, further explorations in model efficiency are warranted, given our findings, e.g., via pruning of model parameters.

Acknowledgments

We acknowledge the support of the Killam Foundation, the Natural Sciences and Engineering Research Council of Canada (NSERC), RGPIN-2022-03943, Canada Foundation of Innovation (CFI) and Research Nova Scotia. Advanced computing resources are provided by ACENET, the regional partner in Atlantic Canada, the Digital Research Alliance of Canada, and the Vector Institute.

Paolo Gajo carried out this work while visiting Dalhousie University and working with the Hypermatrix team. His work is funded by the EU (NextGenerationEU) with funds made available by the National Recovery and Resilience Plan (NRRP), Mission 4, Component 1, Investment 4.1 (M.D. 118/2023).

References

- [1] Sanjeev Arora, Nadav Cohen, Wei Hu, and Yuping Luo. Implicit Regularization in Deep Matrix Factorization. In H. Wallach, H. Larochelle, A. Beygelzimer, F. d'Alché-Buc, E. Fox, and R. Garnett, editors, *Advances in Neural Information Processing Systems*, volume 32. Curran Associates, Inc., 2019.
- [2] Jimmy Lei Ba, Jamie Ryan Kiros, and Geoffrey E Hinton. Layer normalization. *arXiv preprint arXiv:1607.06450*, 2016.
- [3] Dhaivat J. Bhatt, Seyed Ahmad Abdollahpouri Hosseini, Federico Fancellu, and Afsaneh Fazly. End-to-end Parsing of Procedural Text into Flow Graphs. In Nicoletta Calzolari, Min-Yen Kan, Veronique Hoste, Alessandro Lenci, Sakriani Sakti, and Nianwen Xue, editors, *Proceedings of the 2024 Joint International Conference on Computational Linguistics, Language Resources and Evaluation (LREC-COLING 2024)*, pages 5833–5842, Torino, Italia, May 2024. ELRA and ICCL.
- [4] Zhen Bi, Jing Chen, Yinuo Jiang, Feiyu Xiong, Wei Guo, Huajun Chen, and Ningyu Zhang. CodeKGC: Code Language Model for Generative Knowledge Graph Construction, January 2024. arXiv:2304.09048 [cs].

- [5] Yee Seng Chan and Dan Roth. Exploiting syntactico-semantic structures for relation extraction. In *Proceedings of the 49th annual meeting of the association for computational linguistics: human language technologies*, pages 551–560, 2011.
- [6] Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. BERT: Pre-training of deep bidirectional transformers for language understanding. In Jill Burstein, Christy Doran, and Thamar Solorio, editors, *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, June 2019. Association for Computational Linguistics.
- [7] Timothy Dozat and Christopher D. Manning. Deep Biaffine Attention for Neural Dependency Parsing. In *International Conference on Learning Representations*. arXiv, March 2017. arXiv:1611.01734 [cs].
- [8] Timothy Dozat and Christopher D. Manning. Simpler but More Accurate Semantic Dependency Parsing. In Iryna Gurevych and Yusuke Miyao, editors, *Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (Volume 2: Short Papers)*, pages 484–490, Melbourne, Australia, July 2018. Association for Computational Linguistics.
- [9] Vijay Prakash Dwivedi, Chaitanya K. Joshi, Anh Tuan Luu, Thomas Laurent, Yoshua Bengio, and Xavier Bresson. Benchmarking graph neural networks. *Journal of Machine Learning Research*, 24(43):1–48, 2023.
- [10] Vijay Prakash Dwivedi, Ladislav Rampásek, Mikhail Galkin, Ali Parviz, Guy Wolf, Anh Tuan Luu, and Dominique Beaini. Long range graph benchmark. In *Advances in Neural Information Processing Systems 35 (NeurIPS 2022)*, 2022.
- [11] Jack Edmonds. Optimum branchings. *Journal of Research of the National Bureau of Standards Section B Mathematics and Mathematical Physics*, 71B(4):233, October 1967.
- [12] Harsha Gurulingappa, Abdul Mateen Rajput, Angus Roberts, Juliane Fluck, Martin Hofmann-Apitius, and Luca Toldo. Development of a benchmark corpus to support the automatic extraction of drug-related adverse effects from medical case reports. *Journal of Biomedical Informatics*, 45(5):885–892, 2012. Text Mining and Natural Language Processing in Pharmacogenomics.
- [13] Tao Ji, Yuanbin Wu, and Man Lan. Graph-based Dependency Parsing with Graph Neural Networks. In Anna Korhonen, David Traum, and Lluís Màrquez, editors, *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*, pages 2475–2485, Florence, Italy, July 2019. Association for Computational Linguistics.
- [14] Shu Jiang, Zuchao Li, Hai Zhao, and Weiping Ding. Entity-Relation Extraction as Full Shallow Semantic Dependency Parsing. *IEEE/ACM Transactions on Audio, Speech, and Language Processing*, 32:1088–1099, 2024.
- [15] Anees Kazi, Luca Cosmo, Seyed-Ahmad Ahmadi, Nassir Navab, and Michael M. Bronstein. Differentiable Graph Module (DGM) for Graph Convolutional Networks. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 45(2):1606–1617, February 2023. Conference Name: IEEE Transactions on Pattern Analysis and Machine Intelligence.
- [16] Eliyahu Kiperwasser and Yoav Goldberg. Simple and Accurate Dependency Parsing Using Bidirectional LSTM Feature Representations. *Transactions of the Association for Computational Linguistics*, 4:313–327, July 2016. _eprint: https://direct.mit.edu/tac1/article-pdf/doi/10.1162/tac1_a_00101/1567410/tac1_a_00101.pdf.
- [17] Tianyu Liu, Yuchen Eleanor Jiang, Nicholas Monath, Ryan Cotterell, and Mrinmaya Sachan. Autoregressive structured prediction with language models. In Yoav Goldberg, Zornitsa Kozareva, and Yue Zhang, editors, *Findings of the Association for Computational Linguistics: EMNLP 2022*, pages 993–1005, Abu Dhabi, United Arab Emirates, December 2022. Association for Computational Linguistics.

- [18] Ilya Loshchilov and Frank Hutter. Decoupled Weight Decay Regularization, January 2019. arXiv:1711.05101 [cs, math].
- [19] Jianglin Lu, Yi Xu, Huan Wang, Yue Bai, and Yun Fu. Latent graph inference with limited supervision. In *Thirty-seventh Conference on Neural Information Processing Systems*, 2023.
- [20] Yaojie Lu, Qing Liu, Dai Dai, Xinyan Xiao, Hongyu Lin, Xianpei Han, Le Sun, and Hua Wu. Unified structure generation for universal information extraction. In Smaranda Muresan, Preslav Nakov, and Aline Villavicencio, editors, *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 5755–5772, Dublin, Ireland, May 2022. Association for Computational Linguistics.
- [21] Yi Luan, Luheng He, Mari Ostendorf, and Hannaneh Hajishirzi. Multi-task identification of entities, relations, and coreference for scientific knowledge graph construction. In Ellen Riloff, David Chiang, Julia Hockenmaier, and Jun’ichi Tsujii, editors, *Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing*, pages 3219–3232, Brussels, Belgium, October–November 2018. Association for Computational Linguistics.
- [22] Alexis Mitchell et al. Ace 2004 multilingual training corpus ldc2005t09. Web Download, 2005.
- [23] David Nadeau and Satoshi Sekine. A survey of named entity recognition and classification. In Satoshi Sekine and Elisabete Ranchhod, editors, *Recognition, classification and use*, pages 3–28. John Benjamins Publishing Company, 2009.
- [24] Joakim Nivre, Mitchell Abrams, Željko Agić, Lars Ahrenberg, and Lene Antonsen. Universal dependencies 2.2, 2018. LINDAT/CLARIAH-CZ digital library at the Institute of Formal and Applied Linguistics (ÚFAL), Faculty of Mathematics and Physics, Charles University.
- [25] Joakim Nivre, Marie-Catherine de Marneffe, Filip Ginter, Jan Hajič, Christopher D. Manning, Sampo Pyysalo, Sebastian Schuster, Francis Tyers, and Daniel Zeman. Universal Dependencies v2: An evergrowing multilingual treebank collection. In Nicoletta Calzolari, Frédéric Béchet, Philippe Blache, Khalid Choukri, Christopher Cieri, Thierry Declerck, Sara Goggi, Hitoshi Isahara, Bente Maegaard, Joseph Mariani, Héléne Mazo, Asuncion Moreno, Jan Odijk, and Stelios Piperidis, editors, *Proceedings of the Twelfth Language Resources and Evaluation Conference*, pages 4034–4043, Marseille, France, May 2020. European Language Resources Association.
- [26] Joakim Nivre and Chiao-Ting Fang. Universal dependency evaluation. In *Proceedings of the NoDaLiDa 2017 Workshop on Universal Dependencies (UDW 2017)*, pages 86–95, 2017.
- [27] Giovanni Paolini, Ben Athiwaratkun, Jason Krone, Jie Ma, Alessandro Achille, RISHITA ANUBHAI, Cicero Nogueira dos Santos, Bing Xiang, and Stefano Soatto. Structured prediction as translation between augmented natural languages. In *International Conference on Learning Representations*, 2021.
- [28] Wenzhe Pei, Tao Ge, and Baobao Chang. An Effective Neural Network Model for Graph-based Dependency Parsing. In Chengqing Zong and Michael Strube, editors, *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 313–322, Beijing, China, July 2015. Association for Computational Linguistics.
- [29] Hao Peng, Sam Thomson, and Noah A. Smith. Deep multitask learning for semantic dependency parsing. In Regina Barzilay and Min-Yen Kan, editors, *Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 2037–2048, Vancouver, Canada, July 2017. Association for Computational Linguistics.
- [30] Rashmi Prasad, Nikhil Dinesh, Alan Lee, Eleni Miltsakaki, Livio Robaldo, Aravind Joshi, and Bonnie Webber. The Penn Discourse TreeBank 2.0. In Nicoletta Calzolari, Khalid Choukri, Bente Maegaard, Joseph Mariani, Jan Odijk, Stelios Piperidis, and Daniel Tapias, editors, *Proceedings of the Sixth International Conference on Language Resources and Evaluation (LREC’08)*, Marrakech, Morocco, May 2008. European Language Resources Association (ELRA).

- [31] Peng Qi, Timothy Dozat, Yuhao Zhang, and Christopher D Manning. Universal dependency parsing from scratch. *arXiv preprint arXiv:1901.10457*, 2019.
- [32] Peng Qi, Yuhao Zhang, Yuhui Zhang, Jason Bolton, and Christopher D. Manning. Stanza: A Python natural language processing toolkit for many human languages. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics: System Demonstrations*, 2020.
- [33] Raghunathan Ramakrishnan, Pavlo O. Dral, Matthias Rupp, and O. Anatole von Lilienfeld. Quantum chemistry structures and properties of 134 kilo molecules. *Scientific Data*, 1(1):140022, August 2014. Publisher: Nature Publishing Group.
- [34] Noam Razin and Nadav Cohen. Implicit regularization in deep learning may not be explainable by norms. *Advances in neural information processing systems*, 33:21174–21187, 2020.
- [35] Dan Roth and Wen-Tau Yih. A linear programming formulation for global inference in natural language tasks. In *Proceedings of the Eighth Conference on Computational Natural Language Learning (CoNLL-2004) at HLT-NAACL 2004*, pages 1–8, Boston, Massachusetts, USA, May 6 - May 7 2004. Association for Computational Linguistics.
- [36] Olivier Roy and Martin Vetterli. The effective rank: A measure of effective dimensionality. In *2007 15th European signal processing conference*, pages 606–610. IEEE, 2007.
- [37] Wei Tang, Benfeng Xu, Yuyue Zhao, Zhendong Mao, Yifeng Liu, Yong Liao, and Haiyong Xie. UniRel: Unified representation and interaction for joint relational triple extraction. In Yoav Goldberg, Zornitsa Kozareva, and Yue Zhang, editors, *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing*, pages 7087–7099, Abu Dhabi, United Arab Emirates, December 2022. Association for Computational Linguistics.
- [38] Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. Attention is All you Need. In *Advances in Neural Information Processing Systems*, volume 30. Curran Associates, Inc., 2017.
- [39] Petar Veličković, Lars Buesing, Matthew Overlan, Razvan Pascanu, Oriol Vinyals, and Charles Blundell. Pointer Graph Networks. In *Advances in Neural Information Processing Systems*, volume 33, pages 2232–2244. Curran Associates, Inc., 2020.
- [40] Petar Veličković, Guillem Cucurull, Arantxa Casanova, Adriana Romero, Pietro Liò, and Yoshua Bengio. Graph attention networks. In *International Conference on Learning Representations*, 2018.
- [41] David Wadden, Ulme Wennberg, Yi Luan, and Hannaneh Hajishirzi. Entity, relation, and event extraction with contextualized span representations. In Kentaro Inui, Jing Jiang, Vincent Ng, and Xiaojun Wan, editors, *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 5784–5789, Hong Kong, China, November 2019. Association for Computational Linguistics.
- [42] Christopher Walker et al. Ace 2005 multilingual training corpus ldc2006t06. Web Download, 2006.
- [43] Jue Wang and Wei Lu. Two are better than one: Joint entity and relation extraction with table-sequence encoders. In Bonnie Webber, Trevor Cohn, Yulan He, and Yang Liu, editors, *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 1706–1721, Online, November 2020. Association for Computational Linguistics.
- [44] Xiang Wei, Xingyu Cui, Ning Cheng, Xiaobin Wang, Xin Zhang, Shen Huang, Pengjun Xie, Jinan Xu, Yufeng Chen, Meishan Zhang, Yong Jiang, and Wenjuan Han. ChatIE: Zero-Shot Information Extraction via Chatting with ChatGPT, May 2024. arXiv:2302.10205 [cs].
- [45] Robert F Woolson. Wilcoxon signed-rank test. *Encyclopedia of biostatistics*, 8, 2005.

- [46] Yoko Yamakata, Shinsuke Mori, and John Carroll. English Recipe Flow Graph Corpus. In Nicoletta Calzolari, Frédéric Béchet, Philippe Blache, Khalid Choukri, Christopher Cieri, Thierry Declerck, Sara Goggi, Hitoshi Isahara, Bente Maegaard, Joseph Mariani, H  l  ne Mazo, Asuncion Moreno, Jan Odijk, and Stelios Piperidis, editors, *Proceedings of the Twelfth Language Resources and Evaluation Conference*, pages 5187–5194, Marseille, France, May 2020. European Language Resources Association.
- [47] Zhiheng Yan, Chong Zhang, Jinlan Fu, Qi Zhang, and Zhongyu Wei. A partition filter network for joint entity and relation extraction. In Marie-Francine Moens, Xuanjing Huang, Lucia Specia, and Scott Wen-tau Yih, editors, *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, pages 185–197, Online and Punta Cana, Dominican Republic, November 2021. Association for Computational Linguistics.
- [48] An Yang and Sujian Li. SciDTB: Discourse dependency TreeBank for scientific abstracts. In *Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (Volume 2: Short Papers)*, pages 444–449, Melbourne, Australia, July 2018. Association for Computational Linguistics.
- [49] Urchade Zaratiana, Nadi Tomeh, Pierre Holat, and Thierry Charnois. An autoregressive text-to-graph framework for joint entity and relation extraction. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 38, pages 19477–19487, 2024. Issue: 17.
- [50] Dmitry Zelenko, Chinatsu Aone, and Anthony Richardella. Kernel methods for relation extraction. *Journal of machine learning research*, 3(Feb):1083–1106, 2003.
- [51] Bowen Zhang and Harold Soh. Extract, Define, Canonicalize: An LLM-based Framework for Knowledge Graph Construction, April 2024. arXiv:2404.03868 [cs].
- [52] Dan Zhao. Combining explicit and implicit regularization for efficient learning in deep networks. *Advances in Neural Information Processing Systems*, 35:3024–3038, 2022.
- [53] Zexuan Zhong and Danqi Chen. A frustratingly easy approach for entity and relation extraction. In Kristina Toutanova, Anna Rumshisky, Luke Zettlemoyer, Dilek Hakkani-Tur, Iz Beltagy, Steven Bethard, Ryan Cotterell, Tanmoy Chakraborty, and Yichao Zhou, editors, *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, pages 50–61, Online, June 2021. Association for Computational Linguistics.
- [54] Hao Zhu, Yankai Lin, Zhiyuan Liu, Jie Fu, Tat-Seng Chua, and Maosong Sun. Graph neural networks with generated parameters for relation extraction. In Anna Korhonen, David Traum, and Llu  s M  rquez, editors, *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*, pages 1331–1339, Florence, Italy, July 2019. Association for Computational Linguistics.

Table 4: Samples per train/dev/test partition for each non-English dataset.

Dataset	Train	Dev	Test	URL
Arabic	6,075	909	680	https://github.com/UniversalDependencies/UD_Arabic-PADT
Chinese	3,996	500	500	https://github.com/UniversalDependencies/UD_Chinese-GSD
Italian	12,161	538	467	https://github.com/UniversalDependencies/UD_Italian-ISDT
Japanese	6,824	493	518	https://github.com/UniversalDependencies/UD_Japanese-GSD
Spanish	13,821	1,607	1,666	https://github.com/UniversalDependencies/UD_Spanish-AnCora
Wolof	1,149	436	453	https://github.com/UniversalDependencies/UD_Wolof-WTB

Table 5: SynDP results for different languages, with and without score normalization.

a	N	AR	CH	IT	JP	ES	WO
1	0	0.538 $\pm$ 0.005	0.395 $\pm$ 0.007	0.563 $\pm$ 0.006	0.493 $\pm$ 0.010	0.554 $\pm$ 0.004	0.252 $\pm$ 0.007
	1	0.723 $\pm$ 0.008	0.653 $\pm$ 0.007	0.792 $\pm$ 0.003	0.812 $\pm$ 0.003	0.775 $\pm$ 0.003	0.525 $\pm$ 0.006
	2	0.745 $\pm$ 0.004	0.710 $\pm$ 0.005	0.826 $\pm$ 0.003	0.844 $\pm$ 0.005	0.807 $\pm$ 0.002	0.587 $\pm$ 0.007
	3	0.748 $\pm$ 0.004	0.717 $\pm$ 0.007	0.832 $\pm$ 0.002	0.849 $\pm$ 0.003	0.808 $\pm$ 0.002	0.614 $\pm$ 0.005
$\frac{1}{\sqrt{d}}$	0	0.609 $\pm$ 0.003	0.479 $\pm$ 0.007	0.633 $\pm$ 0.002	0.585 $\pm$ 0.005	0.620 $\pm$ 0.002	0.305 $\pm$ 0.005
	1	0.737 $\pm$ 0.007	0.693 $\pm$ 0.003	0.820 $\pm$ 0.004	0.838 $\pm$ 0.004	0.801 $\pm$ 0.003	0.556 $\pm$ 0.006
	2	0.758 $\pm$ 0.003	0.736 $\pm$ 0.005	0.842 $\pm$ 0.004	0.859 $\pm$ 0.004	0.822 $\pm$ 0.001	0.613 $\pm$ 0.008
	3	0.759 $\pm$ 0.005	0.742 $\pm$ 0.006	0.845 $\pm$ 0.002	0.859 $\pm$ 0.003	0.823 $\pm$ 0.002	0.633 $\pm$ 0.005

A Non-English data results

In this appendix, we carry out experiments on the SynDP task on six non-English datasets, all downloaded from Universal Dependencies.⁵ The datasets are listed in Table 4 and comprise the following languages: Arabic, Chinese, Italian, Japanese, Spanish, and Wolof. The results of the experiments are reported in Table 5. Since the data is in different languages, as encoder we use mBERT, the multilingual version of BERT.⁶ All model hyperparameters are kept the same as in the main setting.

Similarly to the English results observed for enEWT and SciDTB, the addition of BiLSTM layers increases the performance. Furthermore, additional layers have a diminished boosting effect on the performance. This shows that the benefits of normalization are independent of the language of which the dependencies are being parsed.

B Non-linguistic results

In this appendix we lay out the unlabeled latent graph inference experiments we carried out on three non-linguistic datasets: PCQM-Contact [10], CIFAR10 Superpixel [9], and QM9 [33]. PCQM-Contact and QM9 are molecule datasets, while CIFAR10 Superpixel is a dataset of superpixel clusters constructed from CIFAR10. Nodes in QM9 and PCQM-Contact respectively contain 9 and 11 features describing each atom in a molecule. In CIFAR10 Superpixel, nodes contain the 3D RGB features of each pixel cluster. In QM9 and CIFAR10 Superpixel, we concatenate the 3D and 2D spatial coordinates of the nodes to their features. Therefore, the resulting node feature dimensionalities of the three datasets are: $d_{\text{PCQM}} = 9$, $d_{\text{CIFAR10}} = 5$, and $d_{\text{QM9}} = 14$.

We feed the graph node features into $L_\gamma \in \{1, 2, 3\}$ GAT layers and provide them with a fully-connected adjacency matrix without any edge attributes. Although the adjacency is fully connected, meaning each node can attend to all others in a single step, stacking L_γ layers refines the globally aggregated representations over 1–3 iterations. The processed node features are passed into the biaffine layer to produce edge predictions. In this case, we pass the scores into a sigmoid function

⁵<https://universaldependencies.org>

⁶<https://huggingface.co/google-bert/bert-base-multilingual-cased>

Table 6: Results for the non-linguistic datasets at varying depths of GAT layers.

a	L_γ	PCQM-Contact	CIFAR10 Superpixel	QM9
1	1	0.241 $\pm$ 0.044	0.407 $\pm$ 0.110	0.861 $\pm$ 0.008
	2	0.217 $\pm$ 0.019	0.528 $\pm$ 0.134	0.876 $\pm$ 0.006
	3	0.245 $\pm$ 0.034	0.751 $\pm$ 0.014	0.877 $\pm$ 0.008
$\frac{1}{\sqrt{d}}$	1	0.288 $\pm$ 0.090	0.531 $\pm$ 0.038	0.886 $\pm$ 0.001
	2	0.209 $\pm$ 0.028	0.585 $\pm$ 0.121	0.880 $\pm$ 0.004
	3	0.237 $\pm$ 0.019	0.703 $\pm$ 0.046	0.881 $\pm$ 0.013

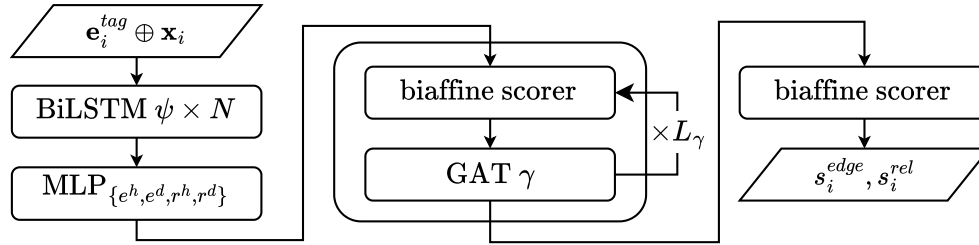


Figure 5: Extended parser architecture, featuring $L_\gamma \in \{0, 1, 2, 3\}$ pairs of biaffine and GAT layers.

because we are not trying to predict an arborescence with the maximum spanning tree algorithm like in dependency parsing. We use $N = 0$ BiLSTM layers to isolate the effect of the GAT layers with respect to any other network components. Scores are averaged over three different seeds.

The unlabeled F_1 scores are reported in Table 6. In all cases, the results show that the gap in performance is large at $L_\phi = 1$ and decreases (or flips) with additional layers. This corroborates the findings from the other settings and is an especially important result considering the phenomenon holds even for such a small network, consisting of just 100–160k parameters.

C GAT results

In this appendix, we extend the main-setting parser, following [13], and carry out k -hop parsing experiments by passing the node representations through $L_\gamma \in \{0, 1, 2, 3\}$ pairs of biaffine and GAT layers. These allow the model to encode higher-order dependencies before the final biaffine layer. Figure 5 depicts the extended architecture. We keep $N = 0$ to isolate the effect of the GAT layers.

As Table 7 shows, the results of this architecture are consistent with the ones obtained using the main-setting 1-hop parser, with the performance increasing across the board when using normalization. As regards the number of GAT layers, performance increases when using $L_\gamma \in \{1, 2\}$ layers, compared to $L_\gamma = 0$, but drops again for all datasets at $L_\gamma = 3$. In general, the best performance is obtained with normalization and $L_\gamma \in \{0, 1\}$ layers. It is particularly interesting to notice that for enEWT the top performance without normalization is obtained at $L_\gamma = 0$, while $L_\gamma = 1$ provides better performance when normalizing biaffine scores. Overall, the usefulness of multi-hop dependency parsing for SemDP is scarce if compared with the SynDP results, where higher-order dependencies provide more useful information for the predictions.

D Ablations

This appendix provides a collection of ablation experiments for various components and settings of the model used in the main experiments.

D.1 Tagger

Table 8 reports the results for the best values of h_ψ and d_{MLP} (chosen as described in Section 4), ablating over $\phi \in \{\checkmark, \times\}$ and $\mathbf{e}_i^{\text{tag}} \in \{\checkmark, \times\}$. The first quarter of the table is equivalent to Table 3.

Table 7: Micro-F₁ (SemDP) and LAS (SynDP) for the labeled edge prediction task at varying depths of $L_\gamma \in \{0, 1, 2, 3\}$ GAT layers. Best in bold.

a	L_γ	ADE	CoNLL04	SciERC	ERFGC	enEWT	SciDTB
1	0	0.517 ± 0.038	0.526 ± 0.046	0.123 ± 0.050	0.536 ± 0.013	0.610 ± 0.009	0.727 ± 0.006
	1	0.509 ± 0.022	0.493 ± 0.037	0.039 ± 0.020	0.599 ± 0.008	0.583 ± 0.011	0.731 ± 0.009
	2	0.509 ± 0.025	0.485 ± 0.033	0.029 ± 0.041	0.611 ± 0.015	0.540 ± 0.012	0.710 ± 0.013
	3	0.487 ± 0.040	0.481 ± 0.087	0.000 ± 0.000	0.585 ± 0.009	0.511 ± 0.015	0.700 ± 0.010
$\frac{1}{\sqrt{d}}$	0	0.556 ± 0.037	0.574 ± 0.028	0.156 ± 0.031	0.607 ± 0.009	0.671 ± 0.007	0.778 ± 0.006
	1	0.587 ± 0.015	0.550 ± 0.036	0.148 ± 0.028	0.639 ± 0.007	0.696 ± 0.005	0.809 ± 0.008
	2	0.534 ± 0.023	0.563 ± 0.029	0.128 ± 0.029	0.637 ± 0.007	0.657 ± 0.007	0.795 ± 0.008
	3	0.543 ± 0.024	0.514 ± 0.026	0.071 ± 0.021	0.616 ± 0.006	0.596 ± 0.010	0.770 ± 0.007

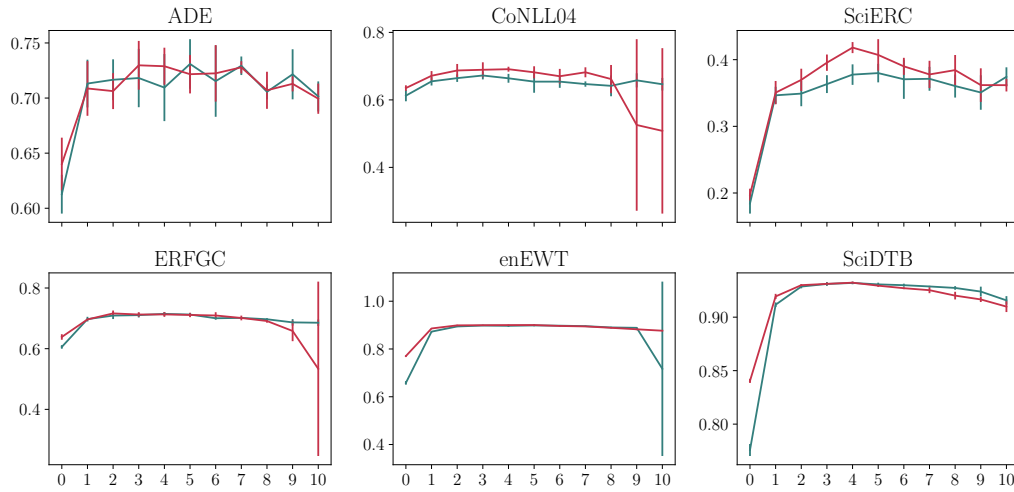


Figure 6: Micro-F₁ (SemDP) and LAS (SynDP) vs $N \in \{0, \dots, 10\}$ at $20k$ training steps. Best per-dataset hyperparameters (h_ψ, d_{MLP}) with $\phi = \checkmark$ and $\mathbf{e}_i^{\text{tag}} = \checkmark$. Red = norm; blue = raw.

Using both the tagger BiLSTM and tag embeddings on average produces the best results. This is sensible, since a better tagger produces better tag embeddings, which in turn help inform the edge and relation classification tasks. The model in this case obtains the best performance on ADE, CoNLL04, and ERFGC. Its top performance for SciERC (0.320) is also close to the overall peak performance (0.324), obtained without using the tagger BiLSTM. On average, the mean performance is considerably higher when combining the positive contributions of the BiLSTM and the tag embeddings.

D.2 Additional BiLSTM layers

Figure 6 depicts performance against the number of BiLSTM layers N at $20k$ training steps. Compared with the main setting’s $2k$ steps, we increase the amount of training because lower amounts resulted in very high F₁ standard deviations with as few as $N = 6$ layers in preliminary experiments. With a shallow stack of BiLSTM layers ($N < 6$), performance is greater when normalizing scores. However, at deeper depths performance becomes very unstable for CoNLL04, ERFGC, and enEWT. Nonetheless, it is evident that score normalization is beneficial, with deeper stacks of BiLSTM layers reducing the effect of normalization, supporting our main-setting claims and experiments.

As observed in Section 5, adding trainable layers seemingly allows the parameters to scale the variance of the scores to compensate for the lack of explicit normalization. As reported in Appendix D.6, a similar behavior can be observed with fully fine-tuned models, where normalization yields diminishing benefits the more parameters we train. This behavior points to a trade-off between the use of our

Table 8: Tagger ablations with best hyperparameters (h_ψ, d_{MLP}). Best in bold, second-best underlined.

a	N	ϕ	$\mathbf{e}_i^{\text{tag}}$	ADE	CoNLL04	SciERC	ERFGC	Mean
$(h_\psi, d_{\text{MLP}}) =$				(200, 100)	(400, 300)	(300, 300)	(400, 300)	
1	0	■	■	0.541 ± 0.021	0.399 ± 0.024	0.147 ± 0.049	0.548 ± 0.010	0.515
	1	■	■	0.657 ± 0.011	0.556 ± 0.021	0.282 ± 0.009	0.676 ± 0.010	
	2	■	■	0.667 ± 0.011	0.573 ± 0.025	0.273 ± 0.010	0.694 ± 0.010	
	3	■	■	0.662 ± 0.027	0.562 ± 0.021	0.299 ± 0.023	0.705 ± 0.011	
$\frac{1}{\sqrt{d}}$	0	■	■	0.567 ± 0.014	0.438 ± 0.033	0.181 ± 0.027	0.612 ± 0.008	0.541
	1	■	■	0.668 ± 0.017	<u>0.597</u> ± 0.015	0.299 ± 0.019	0.692 ± 0.009	
	2	■	■	0.676 ± 0.019	0.596 ± 0.014	0.312 ± 0.011	0.699 ± 0.009	
	3	■	■	0.686 ± 0.025	0.602 ± 0.017	0.320 ± 0.013	0.708 ± 0.008	
1	0	■	□	0.543 ± 0.013	0.402 ± 0.023	0.162 ± 0.019	0.554 ± 0.008	0.517
	1	■	□	0.674 ± 0.011	0.527 ± 0.026	0.275 ± 0.013	0.677 ± 0.005	
	2	■	□	0.672 ± 0.019	0.575 ± 0.009	0.293 ± 0.012	0.689 ± 0.011	
	3	■	□	0.657 ± 0.027	0.583 ± 0.011	0.301 ± 0.012	0.697 ± 0.007	
$\frac{1}{\sqrt{d}}$	0	■	□	0.563 ± 0.013	0.443 ± 0.019	0.188 ± 0.006	0.612 ± 0.003	0.534
	1	■	□	0.653 ± 0.023	0.565 ± 0.027	0.299 ± 0.020	0.687 ± 0.006	
	2	■	□	0.672 ± 0.017	0.592 ± 0.020	0.307 ± 0.008	0.702 ± 0.005	
	3	■	□	0.661 ± 0.022	0.593 ± 0.017	0.305 ± 0.014	0.708 ± 0.007	
1	0	□	■	0.550 ± 0.026	0.387 ± 0.030	0.157 ± 0.012	0.553 ± 0.006	0.514
	1	□	■	0.667 ± 0.022	0.532 ± 0.036	0.273 ± 0.013	0.673 ± 0.006	
	2	□	■	0.665 ± 0.021	0.565 ± 0.024	0.278 ± 0.027	0.694 ± 0.013	
	3	□	■	0.676 ± 0.021	0.558 ± 0.051	0.288 ± 0.012	<u>0.706</u> ± 0.006	
$\frac{1}{\sqrt{d}}$	0	□	■	0.578 ± 0.022	0.437 ± 0.030	0.187 ± 0.014	0.611 ± 0.008	0.534
	1	□	■	0.651 ± 0.032	0.559 ± 0.029	0.301 ± 0.007	0.684 ± 0.003	
	2	□	■	0.673 ± 0.029	0.582 ± 0.017	0.310 ± 0.014	0.701 ± 0.008	
	3	□	■	0.659 ± 0.019	0.586 ± 0.026	0.324 ± 0.019	0.708 ± 0.012	
1	0	□	□	0.547 ± 0.019	0.402 ± 0.033	0.156 ± 0.029	0.549 ± 0.015	0.516
	1	□	□	0.651 ± 0.017	0.552 ± 0.012	0.288 ± 0.008	0.680 ± 0.007	
	2	□	□	0.664 ± 0.031	0.566 ± 0.012	0.288 ± 0.017	0.693 ± 0.011	
	3	□	□	0.663 ± 0.008	0.561 ± 0.014	0.291 ± 0.012	0.703 ± 0.007	
$\frac{1}{\sqrt{d}}$	0	□	□	0.566 ± 0.014	0.431 ± 0.015	0.182 ± 0.023	0.606 ± 0.012	0.534
	1	□	□	0.655 ± 0.015	0.567 ± 0.010	0.286 ± 0.018	0.678 ± 0.013	
	2	□	□	0.678 ± 0.014	0.591 ± 0.026	0.307 ± 0.007	0.702 ± 0.007	
	3	□	□	<u>0.684</u> ± 0.022	0.595 ± 0.017	<u>0.321</u> ± 0.016	0.698 ± 0.015	

technique and the amount of stacked BiLSTM layers, on a given dataset and at a given amount of training steps.

D.3 MLP output dimension

As shown in Table 9, without normalization, increasing the output dimension of the two MLPs responsible for projecting edge representations leads to a decrease in performance. This is in contrast with the behavior observed when using normalization, where performance is essentially independent from the output dimension. This is in line with our claim that higher variance in the edge scores causes lower performance. Normalizing the scores reduces the variance, which makes performance stable even at high output dimensions. Once again, these results show the relationship between score variance and performance, with equivalent performance being achievable with fewer parameters.

D.4 BiLSTM hidden size

As Table 10 shows, the hidden size of the BiLSTMs does not have a visible effect on performance. This is especially the case when applying score normalization, which produces smaller standard

Table 9: Performance in terms of F₁-measure when ablating over different output dimensions for the parser’s MLPs ($\phi = \checkmark$, $\mathbf{e}_i^{tag} = \checkmark$). Best in bold.

a	d_{MLP}	ADE	CoNLL04	SciERC	ERFGC
$(N, h_\psi) =$		(3, 200)	(3, 400)	(3, 300)	(3, 400)
1	100	0.662 ± 0.027	0.579 ± 0.030	0.302 ± 0.018	0.709 ± 0.008
	300	0.658 ± 0.018	0.562 ± 0.021	0.299 ± 0.023	0.705 ± 0.011
	500	0.657 ± 0.018	0.566 ± 0.019	0.281 ± 0.016	0.701 ± 0.009
$\frac{1}{\sqrt{d}}$	100	0.686 ± 0.025	0.586 ± 0.013	0.318 ± 0.017	0.712 ± 0.008
	300	0.680 ± 0.018	0.602 ± 0.017	0.320 ± 0.013	0.708 ± 0.008
	500	0.685 ± 0.019	0.600 ± 0.015	0.311 ± 0.009	0.707 ± 0.004

Table 10: Performance in terms of F₁-measure when ablating over different hidden sizes for the parser’s stacked BiLSTMs ($\phi = \checkmark$, $\mathbf{e}_i^{tag} = \checkmark$). Best in bold.

a	h_ψ	ADE	CoNLL04	SciERC	ERFGC
$(N, d_{\text{MLP}}) =$		(3, 100)	(3, 300)	(3, 300)	(3, 300)
1	100	0.682 ± 0.021	0.580 ± 0.031	0.291 ± 0.021	0.693 ± 0.011
	200	0.662 ± 0.027	0.582 ± 0.006	0.286 ± 0.032	0.703 ± 0.007
	300	0.674 ± 0.029	0.585 ± 0.026	0.299 ± 0.023	0.703 ± 0.006
	400	0.663 ± 0.032	0.562 ± 0.021	0.289 ± 0.038	0.705 ± 0.011
$\frac{1}{\sqrt{d}}$	100	0.685 ± 0.020	0.600 ± 0.018	0.302 ± 0.013	0.698 ± 0.008
	200	0.686 ± 0.025	0.610 ± 0.018	0.314 ± 0.019	0.702 ± 0.008
	300	0.678 ± 0.012	0.599 ± 0.012	0.320 ± 0.013	0.711 ± 0.005
	400	0.674 ± 0.011	0.602 ± 0.017	0.306 ± 0.016	0.708 ± 0.008

deviations. As a result, not only do models perform better with biaffine score normalization, but performance is also less dependent on the hidden size of the BiLSTM encoders. This supports our claim that score normalization is a useful technique to obtain the same performance with lower parameter counts, since the models tend to perform comparably, despite lower BiLSTM hidden sizes.

D.5 LayerNorm and parameter initialization

In this section, we analyze the effects of using a Xavier normal initialization ($I_{par} \sim \mathcal{N}$)⁷ for the weights of the two projections $\text{MLP}^{(edge-head)}$ and $\text{MLP}^{(edge-dept)}$, along with the edge biaffine layer $f^{(edge)}(\cdot; W_e)$. We also apply a LayerNorm function LN_ψ [2] for each of the BiLSTM layers.

As reported in Table 11, on average the best results are obtained with the base setting, i.e., with $I_{par} \sim \mathcal{U}$ and $\text{LN}_\psi = \times$. Using LayerNorm can provide a performance boost for some datasets, especially with uniform initialization. However, it makes performance unstable for some. As a matter of fact, when using LayerNorm with three BiLSTM layers on CoNLL04 and SciERC, the model often breaks and is not able to converge. Based on average performance, then, the best models are obtained by scaling biaffine scores ($a = 1/\sqrt{d}$) and initializing the parser with a uniform weight distribution, without any LayerNorm in between the stacked BiLSTM layers.

D.6 Full fine-tuning and sample efficiency

In this section, we present the results obtained when fine-tuning BERT_{base} , DeBERTa_{base} , BERT_{large} , and DeBERTa_{large} over 10k steps. In our previous settings, we only used a frozen BERT_{base} , whose last hidden states we fed as input to the tagger and parser. This evaluation allows us to test the effectiveness of our approach in a setting in which the number of learnable parameters is unconstrained. In this experiment, we use different learning rates for the base models ($\eta = 1 \times 10^{-4}$) and the large models ($\eta = 3 \times 10^{-5}$) and we apply gradient norm clipping with $\|\nabla\|_{max} = 1.0$. In

⁷https://docs.pytorch.org/docs/stable/nn.init.html#torch.nn.init.xavier_normal_

Table 11: Ablation over the best hyperparameters (h_ψ, d_{MLP}) and $\phi = \checkmark$ and $\mathbf{e}_i^{\text{tag}} = \checkmark$, using LayerNorm layers $\text{LN}_\psi \in \{\checkmark, \times\}$ and Xavier uniform/normal initialization $I_{\text{par}} \sim \{\mathcal{U}, \mathcal{N}\}$ for the parser. The first quarter of the table is equivalent to Table 3. Best in bold, second-best underlined.

I_{par}	LN_ψ	a	N	ADE	CoNLL04	SciERC	ERFGC	Mean
$(h_\psi, d_{\text{MLP}}) =$				(200, 100)	(400, 300)	(300, 300)	(400, 300)	
$\mathcal{U}$	☐	1	0	0.541 ± 0.021	0.399 ± 0.024	0.147 ± 0.049	0.548 ± 0.010	0.515
	☐		3	0.662 ± 0.027	0.562 ± 0.021	0.299 ± 0.023	0.705 ± 0.011	
	☐		1	0.657 ± 0.011	0.556 ± 0.021	0.282 ± 0.009	0.676 ± 0.010	
	☐		2	0.667 ± 0.011	0.573 ± 0.025	0.273 ± 0.010	0.694 ± 0.010	
	☐	$\frac{1}{\sqrt{d}}$	0	0.567 ± 0.014	0.438 ± 0.033	0.181 ± 0.027	0.612 ± 0.008	0.541
	☐		1	0.668 ± 0.017	0.597 ± 0.015	0.299 ± 0.019	0.692 ± 0.009	
	☐		2	0.676 ± 0.019	0.596 ± 0.014	0.312 ± 0.011	0.699 ± 0.009	
	☐		3	0.686 ± 0.025	0.602 ± 0.017	0.320 ± 0.013	0.708 ± 0.008	
	☒	1	0	0.545 ± 0.018	0.399 ± 0.024	0.151 ± 0.014	0.548 ± 0.010	0.468
	☒		1	0.680 ± 0.014	0.554 ± 0.042	0.265 ± 0.028	0.686 ± 0.007	
	☒		2	0.679 ± 0.011	0.578 ± 0.033	0.260 ± 0.020	0.715 ± 0.012	
	☒		3	0.680 ± 0.013	0.117 ± 0.261	0.060 ± 0.133	0.569 ± 0.318	
	☒	$\frac{1}{\sqrt{d}}$	0	0.567 ± 0.014	0.433 ± 0.029	0.181 ± 0.027	0.614 ± 0.004	0.503
	☒		1	0.664 ± 0.017	0.564 ± 0.037	0.282 ± 0.029	0.686 ± 0.004	
	☒		2	0.697 ± 0.022	0.623 ± 0.019	0.300 ± 0.042	0.702 ± 0.011	
	☒		3	0.675 ± 0.019	0.239 ± 0.327	0.111 ± 0.152	0.703 ± 0.006	
$\mathcal{N}$	☐	1	0	0.545 ± 0.017	0.415 ± 0.014	0.155 ± 0.019	0.558 ± 0.009	0.520
	☐		1	0.667 ± 0.014	0.543 ± 0.017	0.275 ± 0.020	0.680 ± 0.015	
	☐		2	0.674 ± 0.025	0.576 ± 0.023	0.272 ± 0.014	0.699 ± 0.006	
	☐		3	0.672 ± 0.028	0.580 ± 0.022	0.297 ± 0.019	0.705 ± 0.006	
	☐	$\frac{1}{\sqrt{d}}$	0	0.570 ± 0.013	0.454 ± 0.006	0.181 ± 0.023	0.613 ± 0.007	<u>0.538</u>
	☐		1	0.671 ± 0.015	0.578 ± 0.031	0.299 ± 0.030	0.685 ± 0.010	
	☐		2	0.671 ± 0.031	0.593 ± 0.016	0.301 ± 0.019	0.704 ± 0.007	
	☐		3	0.681 ± 0.024	0.590 ± 0.014	<u>0.315</u> ± 0.009	0.702 ± 0.011	
	☒	1	0	0.545 ± 0.017	0.415 ± 0.014	0.155 ± 0.019	0.558 ± 0.009	0.469
	☒		1	0.679 ± 0.012	0.582 ± 0.012	0.259 ± 0.010	0.680 ± 0.016	
	☒		2	0.668 ± 0.015	0.580 ± 0.041	0.284 ± 0.024	<u>0.711</u> ± 0.005	
	☒		3	0.679 ± 0.018	0.000 ± 0.000	0.000 ± 0.000	<u>0.711</u> ± 0.009	
	☒	$\frac{1}{\sqrt{d}}$	0	0.570 ± 0.013	0.454 ± 0.006	0.181 ± 0.023	0.613 ± 0.007	0.512
	☒		1	0.675 ± 0.019	0.578 ± 0.022	0.280 ± 0.022	0.682 ± 0.006	
	☒		2	<u>0.687</u> ± 0.021	<u>0.609</u> ± 0.012	0.308 ± 0.012	0.702 ± 0.011	
	☒		3	0.678 ± 0.009	0.476 ± 0.266	0.000 ± 0.000	0.701 ± 0.006	

Table 12: Results for the fully fine-tuned models ($h_{\psi} = 400$, $h_{out} = 500$, $\phi = \checkmark$, and $\mathbf{e}_i^{tag} = \checkmark$).

Model	a	ADE	CoNLL04	SciERC	ERFGC	Mean
BERT _{base}	1	0.748 ± 0.028	0.613 ± 0.019	0.414 ± 0.011	0.726 ± 0.006	0.321
	$\frac{1}{\sqrt{d}}$	0.731 ± 0.025	0.629 ± 0.022	0.412 ± 0.016	0.726 ± 0.006	0.321
BERT _{large}	1	0.748 ± 0.016	0.700 ± 0.019	0.473 ± 0.011	0.750 ± 0.006	0.340
	$\frac{1}{\sqrt{d}}$	0.777 ± 0.011	0.697 ± 0.014	0.446 ± 0.025	0.748 ± 0.010	0.341
DeBERTa _{base}	1	0.754 ± 0.013	0.674 ± 0.014	0.429 ± 0.015	0.751 ± 0.002	0.332
	$\frac{1}{\sqrt{d}}$	0.761 ± 0.021	0.700 ± 0.013	0.425 ± 0.019	0.751 ± 0.010	0.338
DeBERTa _{large}	1	0.786 ± 0.010	0.739 ± 0.016	0.478 ± 0.019	0.763 ± 0.006	0.352
	$\frac{1}{\sqrt{d}}$	0.794 ± 0.015	0.747 ± 0.013	0.476 ± 0.010	0.763 ± 0.007	0.353

the case of the large models, we also use a cosine schedule with warm-up over 6% of the steps. We adopt these measures because during early trials we experienced sudden mid-run gradient explosions. In this case, we do not use any downstream BiLSTMs and only vary whether we use biaffine score normalization in the parser.

Table 12 reports the performance in terms of micro-averaged F_1 -measure for the best models, picked based on top validation performance, evaluated every 100 steps. As the table shows, normalizing the scores generally does not increase top performance when fully fine-tuning these models. This is in line with our theoretical results, since tuning all of their $\sim 100 - 400M$ parameters can compensate for the lack of biaffine score normalization.

In order to have a stronger indication of whether score normalization also works in this setting, we study the performance on the test set versus the amount of training steps. Figure 7 and Figure 8 respectively visualize the performance of the base and large models on the test set in terms of micro-averaged F_1 -measure over $10k$ training steps. We still use early stopping, which is why some of the series are cut short before reaching $10k$ steps.

As Figure 7 shows, a one-tailed Wilcoxon signed-rank test finds the difference in performance throughout the training to be significantly higher when normalizing scores. Since the average test performance of the models is similar with and without normalization, the gap between the two curves being statistically significant indicates faster convergence with normalization.

For both BERT_{base} and DeBERTa_{base}, the difference in convergence speed and performance is statistically significant on ADE, CoNLL04, and SciERC. The same is true for BERT_{large} only on ADE. However, the effect is statistically significant with DeBERTa_{large} for all datasets. This indicates our score normalization approach can produce positive effects in terms of sample efficiency also when fine-tuning models with hundreds of millions of trainable parameters.

E Full main-setting results

Table 13 reports the main-setting results for the best hyperparameter combinations for all three tasks of predicting tags, edges, and relations of the semantic dependency graphs. The first part of Table 13 is identical to the bottom of Table 3. Since we already discussed the performance for labeled edges in Section 5, we forgo discussing them again in this appendix.

As regards unlabeled edges, using score normalization improves mean performance for all datasets, similarly to relations. This is expected, as unlabeled edges directly influence the prediction of the edge labels, together with entity class prediction. As we have already observed for relations, the performance boost provided by normalization is particularly great for SciERC at $N = 0$.

It is also interesting to notice that for CoNLL04, the performance for the unlabeled edge prediction task is only slightly higher than that of the relations. In fact, as regards ADE, the opposite is true, with the labeled performance seemingly being higher. Due to the high variances, it would therefore seem that for these two datasets there is essentially no difference in difficulty between the labeled and unlabeled edge prediction tasks. In the case of ADE, this makes perfect sense, since there is

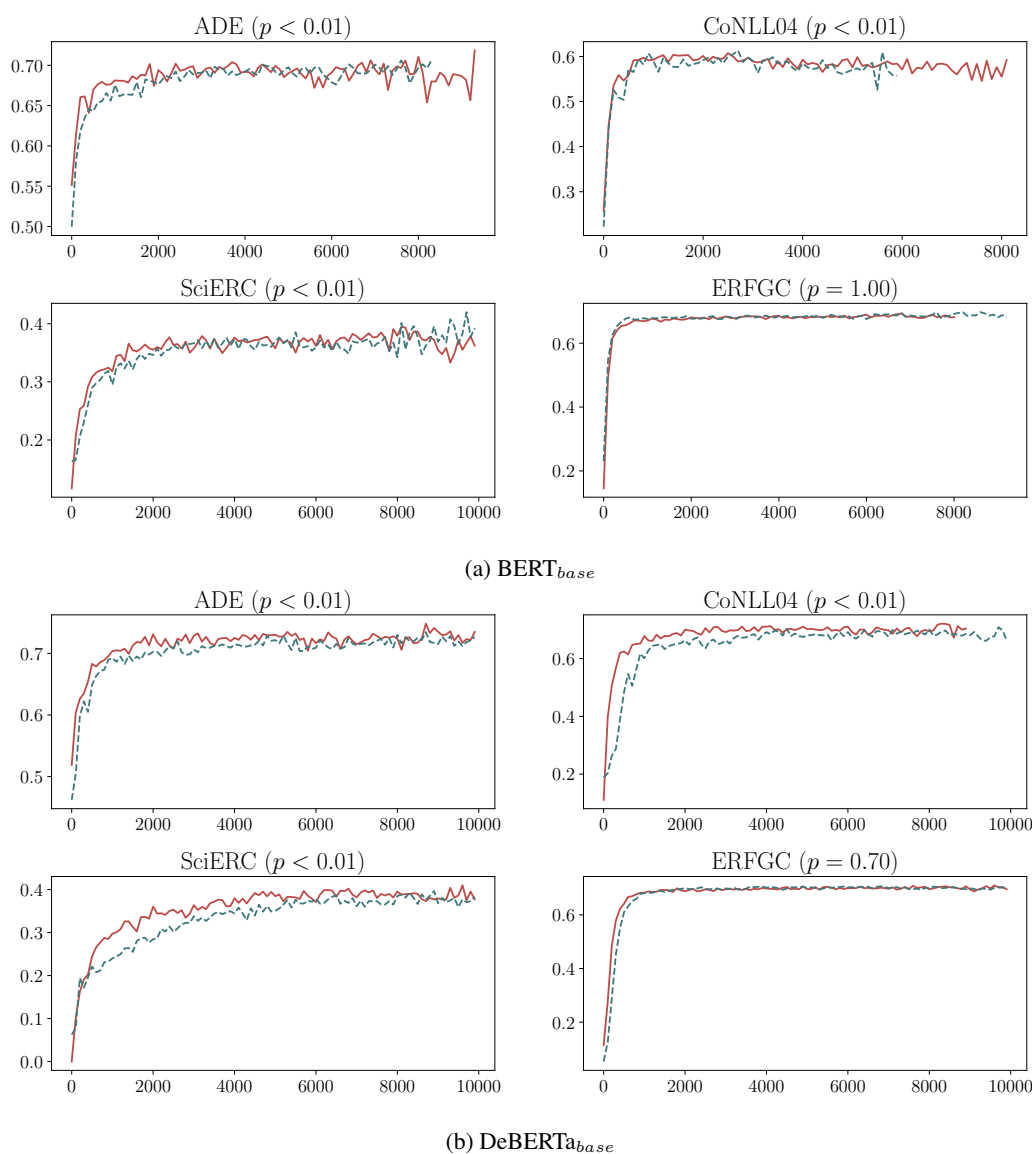


Figure 7: Performance in terms of F_1 -measure vs the number of training steps for the base models on the SemDP datasets. Red = norm; blue = raw. The p -values refer to the performance being greater with score normalization (one-tailed Wilcoxon signed-rank test).

only one possible relation. For CoNLL04, the amount of relation types is rather limited as well; therefore, it is sensible for the performance to also be similar in this case. Indeed, when looking at SciERC and ERFGC, the performance gap between labeled and unlabeled tasks is considerable. This is a strong indication that CoNLL04 is thus found somewhere in the middle, where the number of possible relations only slightly affects labeled performance.

As regards enEWT and SciDTB, the performance on the prediction of edges (UAS) and relations (LAS) is also rather similar. Since in this case the predictions involve syntactic rather than semantic relations, the similar performance hints at relations being easier to predict in SynDP than in SemDP.

In the tagging task, very high standard deviations can be observed. In the case of ADE and CoNLL04, this could be caused by our choice of $\lambda_1 = 0.1$ being too low a value. Regarding SciERC, as already mentioned, most of the entities in the validation and test sets are not found in the training set. This means it is not possible to train the model to recognize them, which results in the abysmal tagging performance reported in Table 13. This makes it challenging to train good tag embeddings which can

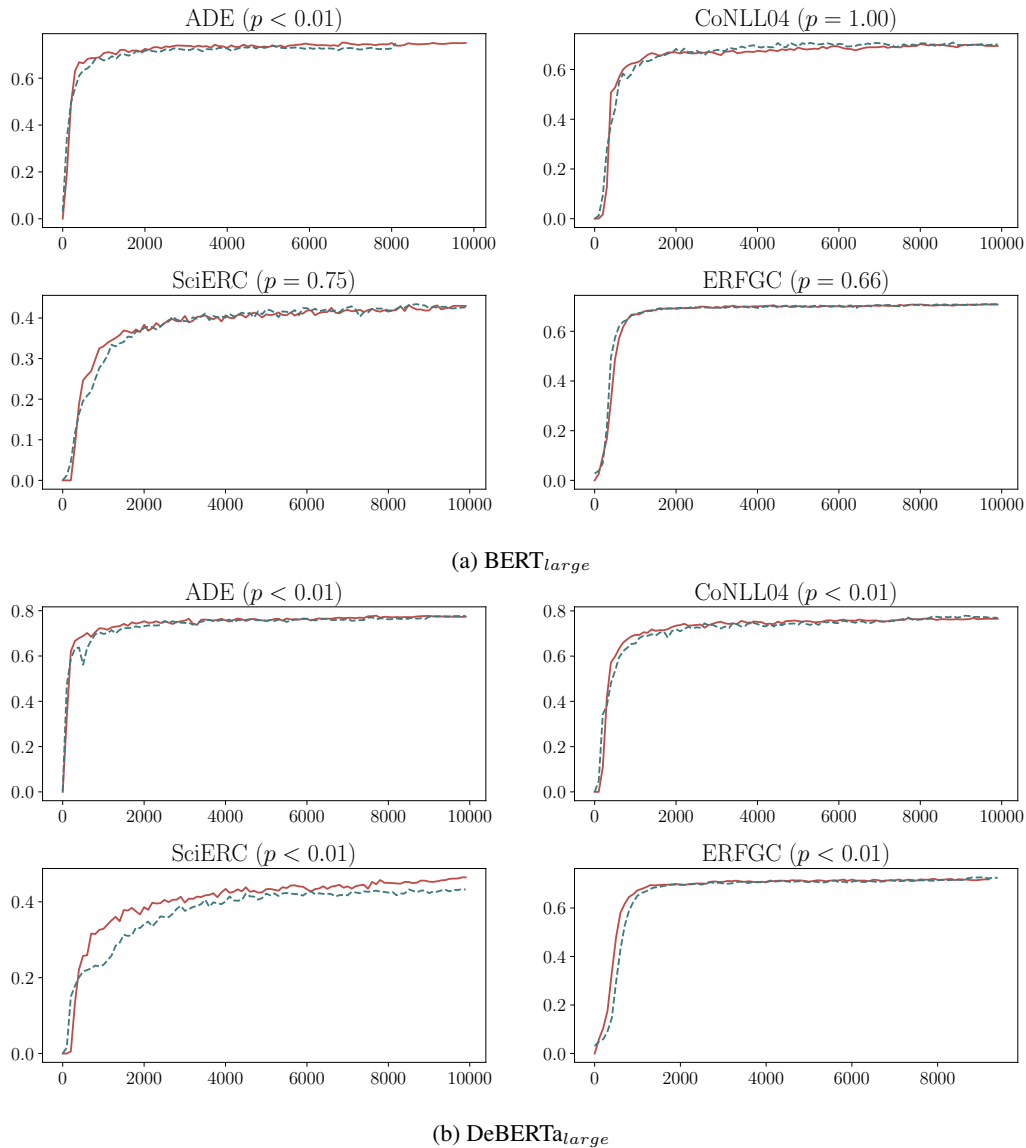


Figure 8: Performance in terms of F₁-measure vs the number of training steps for the large models on the SemDP datasets. Red = norm; blue = raw. The p -values refer to the performance being greater with score normalization (one-tailed Wilcoxon signed-rank test).

help inform the edge and relation prediction tasks. Surprisingly, we notice that for CoNLL04 and ERFGC the tagging performance is also seemingly higher with score normalization. However, the overlaps of the standard deviations are too high to be able to make any claims on the matter.

F Computational resources

We ran all of our experiments on a cluster of NVIDIA H100 (96GB of VRAM) and NVIDIA L40 (48GB of VRAM) GPUs, one run per single GPU. When freezing the BERT_{base} encoder and only training the BiLSTMs and the classifiers, each training and evaluation run took ~5-7 minutes, depending on the number of BiLSTM layers. For the main setting, finding the best hyperparameters involved training and testing 6,120 models, for a total of ~600 GPU hours. In the full fine-tuning setting (Appendix D.6), training and evaluation took ~1 hour for each of the 40 base models and ~2-3 hours for each of the 40 large models, for an additional ~140 GPU hours.

Table 13: Micro-F₁ (SemDP) and LAS (SynDP) on all tasks ($\phi = \checkmark$, $\mathbf{e}_i^{tag} = \checkmark$). Best in bold.

Metric	a	N	ADE	CoNLL04	SciERC	ERFGC	enEWT	SciDTB
$(h_\psi, d_{\text{MLP}}) =$			(200, 100)	(400, 300)	(300, 300)	(400, 300)	(400, 500)	(400, 500)
rels.	1	0	0.541 ± 0.021	0.399 ± 0.024	0.147 ± 0.049	0.548 ± 0.010	0.559 ± 0.005	0.729 ± 0.004
		1	0.657 ± 0.011	0.556 ± 0.021	0.282 ± 0.009	0.676 ± 0.010	0.771 ± 0.006	0.892 ± 0.002
		2	0.667 ± 0.011	0.573 ± 0.025	0.273 ± 0.010	0.694 ± 0.010	0.796 ± 0.006	0.910 ± 0.002
		3	0.662 ± 0.027	0.562 ± 0.021	0.299 ± 0.023	0.705 ± 0.011	0.804 ± 0.006	0.915 ± 0.002
	$\frac{1}{\sqrt{d}}$	0	0.567 ± 0.014	0.438 ± 0.033	0.181 ± 0.027	0.612 ± 0.008	0.646 ± 0.002	0.796 ± 0.002
		1	0.668 ± 0.017	0.597 ± 0.015	0.299 ± 0.019	0.692 ± 0.009	0.789 ± 0.003	0.904 ± 0.002
		2	0.676 ± 0.019	0.596 ± 0.014	0.312 ± 0.011	0.699 ± 0.009	0.805 ± 0.003	0.916 ± 0.002
		3	0.686 ± 0.025	0.602 ± 0.017	0.320 ± 0.013	0.708 ± 0.008	0.807 ± 0.005	0.919 ± 0.001
edges	1	0	0.536 ± 0.009	0.415 ± 0.020	0.173 ± 0.050	0.601 ± 0.013	0.589 ± 0.006	0.745 ± 0.005
		1	0.652 ± 0.009	0.566 ± 0.022	0.321 ± 0.009	0.744 ± 0.013	0.793 ± 0.006	0.901 ± 0.002
		2	0.657 ± 0.021	0.586 ± 0.017	0.322 ± 0.017	0.769 ± 0.014	0.819 ± 0.006	0.919 ± 0.002
		3	0.649 ± 0.027	0.578 ± 0.011	0.355 ± 0.028	0.782 ± 0.007	0.827 ± 0.006	0.924 ± 0.002
	$\frac{1}{\sqrt{d}}$	0	0.549 ± 0.014	0.453 ± 0.032	0.203 ± 0.030	0.674 ± 0.007	0.682 ± 0.003	0.815 ± 0.001
		1	0.654 ± 0.020	0.598 ± 0.021	0.351 ± 0.019	0.762 ± 0.009	0.810 ± 0.003	0.913 ± 0.002
		2	0.660 ± 0.015	0.600 ± 0.008	0.367 ± 0.015	0.775 ± 0.008	0.827 ± 0.003	0.925 ± 0.002
		3	0.666 ± 0.028	0.610 ± 0.022	0.377 ± 0.021	0.786 ± 0.004	0.829 ± 0.004	0.928 ± 0.002
tags	1	0	0.615 ± 0.115	0.285 ± 0.150	0.014 ± 0.015	0.662 ± 0.073		
		1	0.665 ± 0.147	0.671 ± 0.018	0.049 ± 0.014	0.794 ± 0.057		
		2	0.746 ± 0.011	0.648 ± 0.074	0.060 ± 0.012	0.838 ± 0.016		
		3	0.768 ± 0.022	0.669 ± 0.033	0.062 ± 0.005	0.853 ± 0.013		
	$\frac{1}{\sqrt{d}}$	0	0.604 ± 0.134	0.464 ± 0.115	0.046 ± 0.005	0.735 ± 0.066		
		1	0.734 ± 0.048	0.702 ± 0.031	0.058 ± 0.014	0.857 ± 0.010		
		2	0.724 ± 0.064	0.688 ± 0.080	0.060 ± 0.007	0.864 ± 0.018		
		3	0.762 ± 0.017	0.690 ± 0.062	0.057 ± 0.012	0.850 ± 0.022		

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [Yes]

Justification: Our abstract and introduction faithfully report the claims made in the rest of the paper, i.e. we show theoretically and empirically that normalizing biaffine scores for dependency parsing in NLP provides higher parsing performance.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [Yes]

Justification: In our paper, we discuss that we limit ourselves to the task of dependency parsing in NLP. In addition, in the conclusions we mention that our approach does not use k -hop parsers, for example with GNNs or triaffine transformations.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [Yes]

Justification: In Section 3 we provide a lengthy introduction to Result 1, with additional information right after it, which provides more context.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [Yes]

Justification: All of the hyperparameters and methodology are provided to the reader and the experiments can be easily reproduced through the provided GitHub repository.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility.

In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [\[Yes\]](#)

Justification: The code is provided. The repository includes a README.md file which guides the reader through its usage. The data for ADE, CoNLL04, and SciERC can be downloaded from the link provided in Section 4.1. As regards ERFGC, [46] provide access to the corpus upon request. UD 2.2 enEWT and SciDTB are openly available.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyperparameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [\[Yes\]](#)

Justification: Yes, we report information about the data we use in Section 4.1. Training and evaluation information is provided in Sections 4.2 to 4.4.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [\[Yes\]](#)

Justification: Yes, all of our tables and diagrams provide standard deviations and error bars, calculated from model training runs carried out with five different seeds.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g., negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [\[Yes\]](#)

Justification: Appendix F provides all relevant information with relation to the hardware used, its characteristics, and the time necessary to reproduce the experiments.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [\[Yes\]](#)

Justification: Our work does not raise any ethics concerns with regard to any of the guidelines listed within the NeurIPS Code of Ethics. We have no human participants, we use publicly available data, and we believe our work does not have any broader potentially unintended societal impact.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [NA]

Justification: As mentioned at point #9, we do not believe that the Code of Ethics is relevant for our work, reason for which we do not discuss any societal impacts in relation with our study.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: We do not release any new data and/or pre-trained language models.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [Yes]

Justification: We explicitly cite all of the used datasets and models and discuss their availability.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. New assets

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [\[Yes\]](#)

Justification: We provide a well-documented GitHub repo, which is linked to in the abstract.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. Crowdsourcing and research with human subjects

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [\[NA\]](#)

Justification: Our work does not involve any human subjects in any manner.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. Institutional review board (IRB) approvals or equivalent for research with human subjects

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

Justification: Our work does not involve the participation of any human subjects in any manner.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

16. **Declaration of LLM usage**

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigor, or originality of the research, declaration is not required.

Answer: [NA]

Justification: We only use LLMs for the construction of pandas tables and matplotlib visualizations.

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.