

La reconnaissance automatique des relations de cohérence RST en français

Martial Pastor¹ Erik Bran Marino² Nelleke Oostdijk¹

(1) Centre for Language Studies, Radboud University, 6525 HT Nimègue, Pays-Bas

(2) CIDEHUS, Universidade de Évora, 7004-516 Évora, Portugal

{martial.pastor|nelleke.oostdijk}@ru.nl, erik.marino@uevora.pt

RÉSUMÉ

Les parseurs de discours ont suscité un intérêt considérable dans les récentes applications de traitement automatique du langage naturel. Cette approche dépasse les limites traditionnelles de la phrase et peut s'étendre pour englober l'identification de relation de discours. Il existe plusieurs parseurs spécialisés dans le traitement automatique du discours, mais ces derniers ont été principalement évalués sur des corpus anglais. Par conséquent, il n'est pas évident de bien cerner les éléments linguistiques importants sur lesquels les parseurs se basent pour classifier les relations de discours en dehors de l'anglais. Cet article évalue les performances du parseur DMRST sur le corpus RST-DT traduit en français. Nous constatons que les performances de classification des relations de discours en français sont comparables à celles obtenues pour d'autres langues. En analysant les succès et échecs de la classification des relations, nous soulignons l'impact des marqueurs de discours et des structures syntaxiques sur la précision du parseur.

ABSTRACT

RST relation label classification for French.

Discourse parsers have attracted considerable interest in recent natural language processing applications. This approach goes beyond the conventional scope of sentences and can extend to encompass discourse relation identification. There are several parsers specialized in Discourse Parsing, but these have primarily been evaluated on English corpora. Consequently, it is not straightforward to capture the important linguistic elements upon which parsers rely to classify discourse relations outside of English. This article evaluates the performance of the DMRST parser on the RST-DT corpus translated into French. We find that parser performance for the classification of coherence relations in French is comparable to those obtained for other languages. By analyzing the successes and failures of relation classification, we highlight the impact of discourse markers and syntactic structures on the parser's accuracy.

MOTS-CLÉS : parseurs de discours, rhetorical structure theory (RST), relations de discours, traitement automatique du discours, marqueurs de discours, structures syntaxiques.

KEYWORDS: discourse parsers, rhetorical structure theory (RST), discourse relations, discourse parsing, discourse markers, syntactic structures.

1 Introduction

Les parseurs de discours ont suscité un intérêt considérable dans les récentes applications de traitement automatique du langage naturel (Chernyavskiy *et al.*, 2024; Braud *et al.*, 2023). Cette tâche consiste à identifier les relations qu’entretiennent des unités de textes à l’intérieur d’un texte plus large. Cette approche dépasse les limites traditionnelles de la phrase et peut s’étendre pour englober l’identification des Relations de Cohérence au niveau du discours. L’un des formalismes les plus populaires pour représenter les Relations de Cohérence est la Rhetorical Structure Theory (RST; Mann & Thompson, 1988), qui a encouragé la création de jeux de données maintenant utilisés pour l’analyse automatique du discours. Cette dernière tâche est complexe et les parseurs de discours n’ont pas atteint le même niveau de succès que pour d’autres tâches au niveau de la phrase.

Il existe plusieurs parseurs spécialisés dans les structures de type RST (Nguyen *et al.*, 2021, Guz & Carenini; 2020), mais ces derniers ont été principalement entraînés sur le corpus anglais RST-DT (Carlson *et al.*, 2003). Par conséquent, en ce qui concerne les questions d’explicabilité des parseurs orientés deep learning, il n’est pas évident de bien cerner les éléments linguistiques importants sur lesquels les parseurs se basent pour classifier les relations de discours en dehors de l’anglais. Toutefois, des corpus RST sont disponibles pour d’autres langues que l’anglais, ce qui a donné lieu au développement de parseurs multilingues. Cela dit, à notre connaissance, l’analyse des performances des parseurs et les questions d’explicabilité restent limitées à des analyses effectuées sur de l’anglais. Par exemple, certaines études montrent que l’efficacité des parseurs dépend de la présence de marqueurs de discours¹ (voir exemple (1)²) qui rendrait la classification plus facile pour les relations de discours signalées par ces derniers (Pitler *et al.*, 2008).

(1) « [If I sell now,] $\xrightarrow[\text{gold : condition}]{\text{pred : condition}}$ [I’ll take a big loss.] » **wsj_2386**

Par ailleurs, selon les analyses effectuées sur des parseurs plus récents (Liu *et al.*, 2023), il semble que certaines structures syntaxiques, comme des *modifications nominales* (2), jouent un rôle déterminant dans la capacité des parseurs à identifier de manière plus efficace les relations de discours.

(2) « [Negotiable, bank-backed business credit instruments] $\xleftarrow[\text{gold : elaboration}]{\text{pred : elaboration}}$ [typically financing an import order.] » **wsj_0602**

Cet article propose une évaluation puis une analyse des réussites et des échecs du parseur DMRST sur le corpus RST-DT traduit en français. Nous commençons par reproduire les expériences avec le parseur multilingue DMRST, que nous évaluons ensuite sur le corpus test du RST-DT. Nous constatons initialement que les performances en termes de classification des relations de cohérence en français sont comparables à celles obtenues pour d’autres langues. Ensuite, nous procédons à une analyse qualitative des cas de succès et d’échecs dans la classification des relations, en mettant en lumière des cas de figure où les relations sont signalées soit par des marqueurs de discours, soit

1. La définition à laquelle nous faisons référence dans cet article concerne les *Discourse Markers* en anglais, qui sont des connecteurs, généralement des conjonctions, entre propositions indépendantes.

2. "pred" ici correspond à l’étiquette prédite par le parseur DMRST présenté dans la section 2, et "gold" correspond à l’étiquette annotée dans le jeu de données RST-DT. La direction de la flèche pointe vers le noyau. La différenciation entre le noyau et le satellite repose sur le fait qu’une unité textuelle agissant comme satellite peut être retirée sans altérer la cohérence du discours, tandis que le retrait d’une unité textuelle agissant comme noyau rendrait le texte incohérent (Mann & Thompson, 1988).

par des structures syntaxiques. Les résultats révèlent que certains signaux syntaxiques facilitent la classification, tandis que certains marqueurs de discours contribuent à une confusion accrue entre différents types de relations.

2 Méthodologie

La Rhetorical Structure Theory. La RST est un modèle d'analyse textuelle des relations de cohérence. Avec celui-ci, les relations entre les segments de texte sont annotées avec différentes classes de relations de cohérence telles que l'élaboration, le contraste, le causal, le temporel, etc. (voir Table 2 pour une liste complète des relations utilisées dans cet article). Les segments de texte d'un arbre RST sont des "unités discursives élémentaires" (EDUs), qui sont des ensembles contigus de tokens approximativement similaires à des propositions indépendantes. Les relations se font non seulement entre les segments de texte mais aussi entre des groupes de segments de texte, ce qui signifie que la représentation finale de RST d'un texte complet (livre, chapitre, article, commentaire, etc.) est un arbre hiérarchique de segments de texte connectés par des relations de cohérence comme sur la Figure 1.

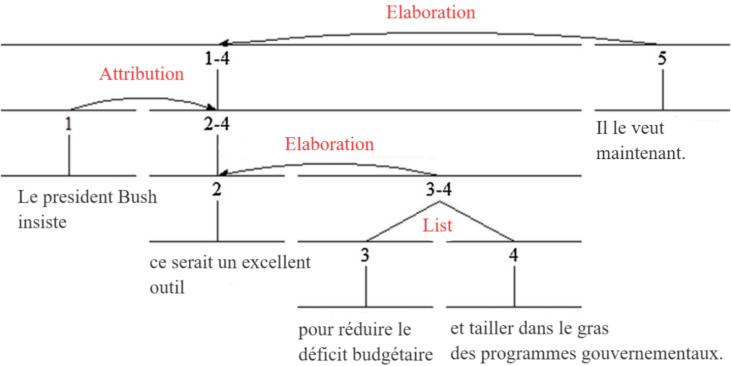


FIGURE 1 – Analyse RST traduite en français issue du document wsj_609 dans le corpus RST-DT qui décrit le modèle d'analyse textuelle des relations de cohérence.

Le corpus RST Discourse Treebank. Le corpus RST-DT (Carlson *et al.*, 2003) a été largement utilisé pour l'analyse RST en anglais et est devenu un choix standard pour évaluer les parseurs RST. Il est reconnu pour ses structures arborescentes hiérarchiques et a été initialement annoté avec 76 relations de cohérence. Les relations examinées ici proviennent de l'ensemble de test RST-DT, qui contient un total de 38 documents, comprenant environ 21 600 tokens. En ce qui concerne les étiquettes de relations, nous utilisons actuellement l'ensemble harmonisé de 18 étiquettes (Braud *et al.*, 2017).

Le parseur multilingue DMRST. Le parseur DMRST développé par Liu *et al.*, 2021 est basé sur XLM-ROBERTA-BASE (Conneau *et al.*, 2020) et est un système multilingue TOP-DOWN qui gère simultanément la segmentation des EDUs et l'analyse des arbres RST. Sa pertinence pour notre étude réside dans ses performances à l'état de l'art en matière de classification des relations de cohérence.

Les auteurs ont fourni un accès à un modèle optimisé pour l'inférence disponible en ligne³. Ce modèle particulier a été entraîné sur une collection multilingue de corpus RST, offrant une prise en charge native pour six langues : l'anglais, le portugais, l'espagnol, l'allemand, le néerlandais et le basque. Il est important de noter que, bien que le parseur puisse prédire la structure arborescente et les relations directement à partir de texte brut, notre étude choisit d'utiliser la segmentation des unités discursives (EDU) de référence. Dans notre configuration expérimentale, nous utilisons à la fois le texte brut de l'ensemble de test RST-DT d'origine (traduit en français) et la segmentation des EDUs de référence.

Traduction du corpus RST-DT en français avec ChatGPT. Afin d'évaluer les performances du système DMRST en français, nous avons traduit l'ensemble test du corpus RST-DT en français en utilisant ChatGPT. La traduction avait pour objectif de rester aussi fidèle que possible au texte original : chaque unité discursive des 38 documents (soit 2308 EDUs) a été traduite individuellement pour préserver la structure RST de chaque texte et éviter la génération de paraphrases. De plus, ChatGPT prend en compte le contexte de chaque EDU dans sa traduction ; bien que les unités discursives soient traduites individuellement, les accords grammaticaux, de sujets-verbes ou encore de références sont préservés. Par exemple, la deuxième unité ici accorde correctement au pluriel : [sans les frais de rachat] [qui sont courants pour les rentes.].

La version de ChatGPT utilisée est GPT-4-Turbo. Le prompt qui a été utilisé pour obtenir les traductions est le suivant :

```
message : {  
    role.system : "Vous êtes traducteur expert de l'anglais vers le français et un linguiste  
        spécialisé dans le domaine de la Rhetorical Structure Theory (RST)."  
    role.user : "Traduire le texte ci-dessous en Français segment par segment, tout en veillant  
        à conserver le contexte global du texte dans son ensemble.  
        Texte à traduire :  
            segment 1  
            segment 2  
            ...  
        Afficher le résultat avec chaque segment traduit sur une ligne distincte."  
}
```

FIGURE 2 – Prompt utilisé pour traduire le corpus de test RST-DT en français. Ce prompt a été utilisé 38 fois pour les 38 documents de ce corpus.

Encore une fois, il était nécessaire de contraindre le système à traduire EDU par EDU afin de préserver les structures textuelles annotées en relations de cohérence du corpus RST-DT. La question qui se pose dès lors est la suivante : le texte produit est-il bien du français ? Afin d'analyser les biais introduits par une telle contrainte, nous avons comparé, pour un ensemble de 10 documents du corpus de test RST-DT choisis au hasard, une traduction automatique des documents traduits en entier avec la traduction EDU par EDU. Ces 10 mêmes documents ont ensuite été relus et évalués par 3 locuteurs natifs français.

3. https://github.com/seq-to-mind/DMRST_Parser

En ce qui concerne la comparaison entre le texte traduit par segments et le texte traduit en entier, étant donné que les traducteurs automatiques traduisent souvent phrase par phrase, les traductions étaient similaires, à l’exception de légères variations dans le vocabulaire utilisé et la génération de paraphrases où l’ordre de quelques groupes prépositionnels a été inversé. Nous avons ensuite comparé ces deux traductions automatiques pour chacun des 10 documents avec les métriques BLEU, ROUGE, TER et WER et avons obtenu les scores moyens suivants : 0.65, 0.82, 20.0 et 0.35.

Quant à l’évaluation manuelle des 3 locuteurs, les 10 documents traduits en français ont été jugés intelligibles et corrects dans leur ensemble. Toutefois les participants ont noté un manque de fluidité, une absence totale d’idiomaticité et la présence, rare mais non négligeable, de calques syntaxiques venus de l’anglais produisant des imprécisions grammaticales, comme « an equity position in Leaseway » traduit par « une position en capitaux dans Leaseway ».

Les résultats de ces analyses nous permettent de conclure que, bien que la contrainte de traduction EDU par EDU n’introduise pas de biais trop importants par rapport à la traduction du texte entier (comme le montrent les scores obtenus), la traduction du français obtenue n’est pas la plus naturelle qui soit. Toutefois, comme nous le verrons dans la partie 4, les éléments linguistiques analysés restent pertinents et correspondent à des usages naturels en français.

3 Expérimentations et résultats

Reconnaissance des relations de cohérence RST en français. Les 38 textes du corpus RST-DT traduits en français avec ChatGPT ont été parsés avec le système multilingue DMRST. Nous présentons ci-dessous les résultats évalués avec la métrique RST-Parseval (Marcu, 2000) pour la classification de relation.

	Précision	Rappel	F1-score
Français	0.54	0.57	0.54
Anglais	0.65	0.67	0.65

TABLE 1 – Résultats globaux évalués avec la métrique RST-Parseval (Marcu, 2000) pour la classification de relation en français et en anglais.

Sans surprise, les résultats sont moins bons que ceux obtenus pour l’anglais, mais sont comparables aux résultats obtenus sur d’autres corpus RST-DT dans différentes langues (Avec un F1-score de 0.62 en portugais, 0.63 en espagnol, 0.47 en allemand, 0.52 en néerlandais et 0.48 en basque, conformément aux résultats signalés par Liu *et al.*, 2021). Comme nous pouvons le voir avec la matrice de confusion sur la Figure 3 ci-dessous, nous remarquons également que la confusion entre les différentes classes de relations est similaire à ce que nous obtenons pour l’anglais concernant les classes ELABORATION et JOINT, cela étant dû à un déséquilibre des classes dans le jeu de données RST-DT.

Relation	rel. frq. (%)	abs. frq. (N)
elaboration	34.43	794
attribution	14.87	343
joint	9.19	212
contrast	6.37	147
same-unit	5.51	127
explanation	4.77	110
enablement	3.82	88
cause	3.56	82
evaluation	3.47	80
temporal	3.17	73
background	2.95	68
topic-comment	1.69	39
summary	1.39	32
comparison	1.26	29
condition	1.21	28
manner-means	1.17	27
topic-change	0.78	18
textualorganization	0.39	9
TOTAL	100.0	2306

TABLE 2 – Fréquence absolue et relative des étiquettes de relation en or dans RST-DT pour un total de 2306 relations.

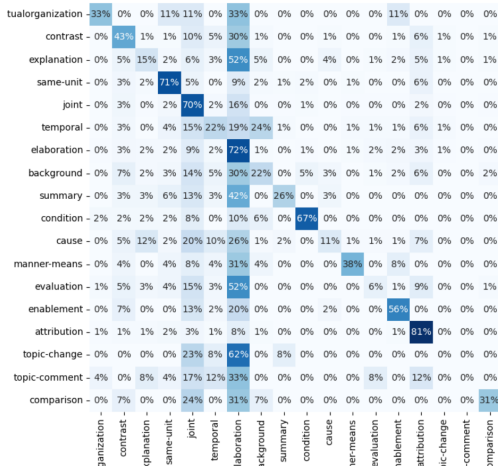


FIGURE 3 – Matrice de confusion de 18 étiquettes harmonisées pour le français.

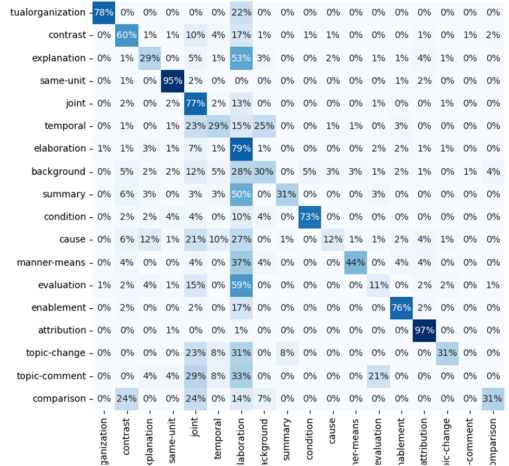


FIGURE 4 – Matrice de confusion pour l'anglais.

4 Discussion et analyse des cas d'échecs / succès

Dans cette section, nous procédons à une analyse qualitative d'un sous-ensemble de relations. Les cas sélectionnés pour l'annotation visent à illustrer en quoi la présence de marqueurs de discours ou de certaines structures syntaxiques influence les erreurs de classification du parseur DMRST. Les cas analysés ci-dessous proviennent d'un ensemble de 301 relations. Ces 301 cas ont été sélectionnés à partir du RST Signaling corpus (Das & Taboada, 2018), où chacune des relations du RST-DT

a été annotée en signaux de relations de cohérence pour de l'anglais. Nous avons extrait les cas correspondant aux catégories décrites ci-dessous et les avons réannotés pour vérifier si ces signaux étaient également applicables au français. À la suite de cette procédure, 21 relations ont été exclues. L'accord inter-annotateur, mesuré par le coefficient Kappa de Cohen, est de 0,83 pour cette phase d'annotations.

Présences de marqueurs de discours. Dans certaines situations, les marqueurs de discours jouent un rôle décisif et contribuent efficacement à prédire la relation. Par exemple, la présence du marqueur de discours *si*, signalant une relation de CONDITION, permet une reconnaissance précise dans 73% des cas pour 41 relations annotées. Cependant, nous constatons que les marqueurs de discours posent souvent des problèmes en raison de leur présence dans de nombreuses relations. Par exemple, nous observons que le marqueur de discours *et* est largement utilisé dans la relation JOINT, mais il est également présent dans de nombreuses autres relations telles que la relation TEMPORAL. Par conséquent, seulement 17% (sur 26 annotations) des relations TEMPORAL signalées par ce marqueur sont correctement prédites. Nous observons une tendance similaire en ce qui concerne la confusion entre les relations TEMPORAL et les relations BACKGROUND, souvent signalées par des marqueurs tels que *après que*, *depuis que*, *tant que*, *pendant que*, *alors que* ou d'autres formes + QUE comme dans l'exemple (3).

(3) « [dit-il] $\leftarrow \begin{smallmatrix} \text{pred : background} \\ \text{gold : temporal} \end{smallmatrix} \right.$ [alors que le boucher s'éloigne ».] **wsj_1146**

La relation BACKGROUND se produit entre une unité de texte principal (nucleus) et un satellite qui est nécessaire à la compréhension de la première unité. Bien que cette relation n'implique pas d'aspect temporel, elle est souvent signalée par des marqueurs de discours tels que *depuis que*, *tant que* ou tout simplement par un *après*, comme le montre l'exemple 4, qui sont également utilisés pour signaler la relation TEMPORAL.

(4) « [Mme Crump a déclaré que le portefeuille de son club d'investissement Ashwood avait perdu environ un tiers de sa valeur] $\leftarrow \begin{smallmatrix} \text{pred : temporal} \\ \text{gold : background} \end{smallmatrix} \right.$ [après le crash du Black Monday] **wsj_2386**

Présences de structures syntaxique. D'un autre côté, on remarque que les relations signalées par des structures syntaxiques ont plus de chances d'être systématiquement prédites. Nous avons annoté ici les structures telles que les *constructions syntaxiques parallèles* (voir exemple (5)) pour la relation JOINT et les *propositions relatives* pour la relation ELABORATION). Nous notons que 79% des relations annotées avec des *constructions syntaxiques parallèles*, soit 72 cas annotés, sont correctement prédites. De même, nous notons un taux de réussite de 85% pour les 141 cas annotés des *propositions relatives*.

(5) « [Parlez-nous de la retenue en matière de dépenses] $\leftarrow \begin{smallmatrix} \text{pred : temporal} \\ \text{gold : background} \end{smallmatrix} \right.$ [Parlez-nous des scandales du HUD] » **wsj_0623**

Ambiguïté et Spécificité. L'ambiguïté et la spécificité jouent un rôle crucial dans les performances différenciées du parseur en fonction de la présence de marqueurs de discours ou de structures syntaxiques. Les marqueurs de discours ont tendance à être plus ambigus, car ils peuvent avoir une forme plus ou moins lexicalisée, comme *alors que* ou *après*, qui peuvent être utilisés dans différents contextes. Cependant, cela ne s'applique pas aux marqueurs tels que *si*, qui sont très spécifiques à la relation CONDITION. D'autre part, on observe que les structures syntaxiques sont beaucoup plus spécifiques à certaines relations et, par conséquent, elles induisent moins de confusion pour le parseur.

5 Conclusion

Actuellement, il est important de noter que nous dépendons d'un seul point de contrôle de modèle pour l'expérimentation, ce qui introduit le potentiel d'erreurs influencées par des certaines variations lors de l'entraînement. De plus, nous tenons à souligner que le corpus est limité aux données issues de la presse, et explorer des données provenant de différents genres serait susceptible de fournir des perspectives supplémentaires. Par ailleurs, l'annotation manuelle effectuée ici s'est concentrée sur des cas particuliers inspirés par des analyses menées sur de l'anglais. La création d'un corpus annoté avec des signaux de relations de cohérence, à l'instar de ce qu'ont fait [Das & Taboada, 2018](#) pour l'anglais, nous permettrait d'avoir une vue plus complète des analyses des cas d'échecs et de succès.

En conclusion, cet article a évalué les performances du parseur DMRST sur le corpus RST-DT traduit en français. Nous avons constaté que les performances de ce parseur pour la classification des relations de cohérence en français sont comparables à celles obtenues pour d'autres langues. En analysant les succès et les échecs de la classification des relations, nous avons mis en évidence l'impact des marqueurs de discours et des structures syntaxiques sur la précision du parseur. Cela est attribuable à la spécificité des structures syntaxiques, qui sont étroitement liées à des relations individuelles et ne sont pas aussi ambiguës que les marqueurs de discours.

Remerciements

Le travail présenté dans cet article a été réalisé dans le cadre du projet HYBRIDS, un réseau doctoral Marie Skłodowska-Curie financé par l'Union européenne sous le numéro de subvention 101073351 et par le UK Research and Innovation (UKRI) Horizon Funding Guarantee.

Références

- BRAUD C., COAVOUX M. & SØGAARD A. (2017). Cross-lingual RST discourse parsing. In *Proceedings of the 15th Conference of the European Chapter of the Association for Computational Linguistics : Volume 1, Long Papers*, p. 292–304, Valencia, Spain : Association for Computational Linguistics.
- BRAUD C., LIU Y. J., METHENITI E., MULLER P., RIVIÈRE L., RUTHERFORD A. & ZELDES A. (2023). The DISRPT 2023 shared task on elementary discourse unit segmentation, connective detection, and relation classification. In C. BRAUD, Y. J. LIU, E. METHENITI, P. MULLER, L. RIVIÈRE, A. RUTHERFORD & A. ZELDES, Éd., *Proceedings of the 3rd Shared Task on Discourse Relation Parsing and Treebanking (DISRPT 2023)*, p. 1–21, Toronto, Canada : The Association for Computational Linguistics. DOI : [10.18653/v1/2023.disrpt-1.1](https://doi.org/10.18653/v1/2023.disrpt-1.1).
- CARLSON L., MARCU D. & OKUROWSKI M. E. (2003). *Building a Discourse-Tagged Corpus in the Framework of Rhetorical Structure Theory*, In J. VAN KUPPEVELT & R. W. SMITH, Éd., *Current and New Directions in Discourse and Dialogue*, p. 85–112. Springer Netherlands : Dordrecht. DOI : [10.1007/978-94-010-0019-2_5](https://doi.org/10.1007/978-94-010-0019-2_5).
- CHERNYAVSKIY A., ILVOVSKY D. & NAKOV P. (2024). Unleashing the power of discourse-enhanced transformers for propaganda detection. In *Proceedings of the 18th Conference of the*

European Chapter of the Association for Computational Linguistics (Volume 1 : Long Papers), p. 1452–1462.

CONNEAU A., KHANDLWAL K., GOYAL N., CHAUDHARY V., WENZKE G., GUZMÁN F., GRAVE E., OTT M., ZETTLEMOYER L. & STOYANOV V. (2020). Unsupervised cross-lingual representation learning at scale. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, p. 8440–8451, Online : Association for Computational Linguistics. DOI : [10.18653/v1/2020.acl-main.747](https://doi.org/10.18653/v1/2020.acl-main.747).

DAS D. & TABOADA M. (2018). Signalling of coherence relations in discourse, beyond discourse markers. *Discourse Processes*, **55**(8), 743–770. DOI : [10.1080/0163853X.2017.1379327](https://doi.org/10.1080/0163853X.2017.1379327).

GUZ G. & CARENINI G. (2020). Coreference for discourse parsing : A neural approach. In C. BRAUD, C. HARDMEIER, J. J. LI, A. LOUIS & M. STRUBE, Édts., *Proceedings of the First Workshop on Computational Approaches to Discourse*, p. 160–167, Online : Association for Computational Linguistics. DOI : [10.18653/v1/2020.codi-1.17](https://doi.org/10.18653/v1/2020.codi-1.17).

LIU Y. J., AOYAMA T. & ZELDES A. (2023). What’s hard in English RST parsing ? predictive models for error analysis. In *Proceedings of the 24th Meeting of the Special Interest Group on Discourse and Dialogue*, p. 31–42, Prague, Czechia : Association for Computational Linguistics.

LIU Z., SHI K. & CHEN N. (2021). DMRST : A joint framework for document-level multilingual RST discourse segmentation and parsing. In *Proceedings of the 2nd Workshop on Computational Approaches to Discourse*, p. 154–164, Punta Cana, Dominican Republic and Online : Association for Computational Linguistics. DOI : [10.18653/v1/2021.codi-main.15](https://doi.org/10.18653/v1/2021.codi-main.15).

MANN W. C. & THOMPSON S. A. (1988). Rhetorical structure theory : Toward a functional theory of text organization. *Text-interdisciplinary Journal for the Study of Discourse*, **8**(3), 243–281. DOI : [/doi.org/10.1515/text.1.1988.8.3.243](https://doi.org/10.1515/text.1.1988.8.3.243).

MARCU D. (2000). The Rhetorical Parsing of Unrestricted Texts : A Surface-based Approach. *Computational Linguistics*, **26**(3), 395–448. DOI : [10.1162/089120100561755](https://doi.org/10.1162/089120100561755).

NGUYEN T.-T., NGUYEN X.-P., JOTY S. & LI X. (2021). RST parsing from scratch. In K. TOUTANOVA, A. RUMSHISKY, L. ZETTLEMOYER, D. HAKKANI-TUR, I. BELTAGY, S. BETHARD, R. COTTERELL, T. CHAKRABORTY & Y. ZHOU, Édts., *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics : Human Language Technologies*, p. 1613–1625, Online : Association for Computational Linguistics. DOI : [10.18653/v1/2021.naacl-main.128](https://doi.org/10.18653/v1/2021.naacl-main.128).

PITLER E., RAGHUPATHY M., MEHTA H., NENKOVA A., LEE A. & JOSHI A. (2008). Easily identifiable discourse relations. In *Coling 2008 : Companion volume : Posters*, p. 87–90, Manchester, UK : Coling 2008 Organizing Committee.