



Analyse sémantique robuste par apprentissage antagoniste pour la généralisation de domaine

Gabriel Marzinotto, Géraldine Damnati, Frédéric Béchet, Benoît Favre

► To cite this version:

Gabriel Marzinotto, Géraldine Damnati, Frédéric Béchet, Benoît Favre. Analyse sémantique robuste par apprentissage antagoniste pour la généralisation de domaine. 6e conférence conjointe Journées d'Études sur la Parole (JEP, 33e édition), Traitement Automatique des Langues Naturelles (TALN, 27e édition), Rencontre des Étudiants Chercheurs en Informatique pour le Traitement Automatique des Langues (RÉCITAL, 22e édition). Volume 4: Démonstrations et résumés d'articles internationaux, 2020, Nancy, France. pp.71-72. hal-02768521v3

HAL Id: hal-02768521

<https://hal.science/hal-02768521v3>

Submitted on 23 Jun 2020

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

Analyse sémantique robuste par apprentissage antagoniste pour la généralisation de domaine

Gabriel Marzinotto^{1,2} Géraldine Damnati¹ Frédéric Béchet² Benoît Favre²

(1) Orange Labs / Lannion France

(2) Aix Marseille Univ, CNRS, LIS / Marseille France

gabriel.marzinotto@orange.com, geraldine.damnati@orange.com

frederic.bechet@lis-lab.fr, benoit.favre@lis-lab.fr

RÉSUMÉ

Nous présentons des résumés en français et en anglais de l'article ([Marzinotto et al., 2019](#)) présenté à la conférence *North American Chapter of the Association for Computational Linguistics : Human Language Technologies* en 2019.

ABSTRACT

Robust Semantic Parsing with Adversarial Learning for Domain Generalization

We present French and English abstracts of the article ([Marzinotto et al., 2019](#)) that was presented at the 2019 North American Chapter of the Association for Computational Linguistics: Human Language Technologies.

MOTS-CLÉS : Analyse sémantique, adaptation de domaine, apprentissage antagoniste.

KEYWORDS: Semantic parsing, domain adaptation, adversarial learning.

1 Résumé en français

Cet article étudie l'amélioration de la capacité de généralisation des modèles d'analyse sémantique à travers de techniques d'apprentissage antagoniste. La création de modèles plus robustes à la variabilité inter-documents est cruciale pour l'intégration des technologies d'analyse sémantique dans les applications réelles. La question sous-jacente de cette étude est de savoir si l'apprentissage antagoniste peut être utilisé pour entraîner des modèles à un niveau d'abstraction plus élevé afin d'augmenter leur robustesse aux variations lexicales et stylistiques.

Nous proposons d'effectuer l'analyse sémantique avec une tâche adverse de classification de domaine sans connaissance explicite du domaine. La stratégie est d'abord évaluée sur un corpus français de documents encyclopédiques, annotés avec FrameNet, dans une perspective de recherche d'informations, puis sur la tâche PropBank Semantic Role Labelling sur le benchmark CoNLL-2005. Nous montrons que l'apprentissage contradictoire augmente toutes les capacités de généralisation des modèles à la fois sur les données du domaine et hors domaine.

2 English Abstract

This paper addresses the issue of generalization for Semantic Parsing in an adversarial framework. Building models that are more robust to inter-document variability is crucial for the integration of Semantic Parsing technologies in real applications. The underlying question throughout this study is whether adversarial learning can be used to train models on a higher level of abstraction in order to increase their robustness to lexical and stylistic variations.

We propose to perform Semantic Parsing with a domain classification adversarial task without explicit knowledge of the domain. The strategy is first evaluated on a French corpus of encyclopedic documents, annotated with FrameNet, in an information retrieval perspective, then on PropBank Semantic Role Labeling task on the CoNLL-2005 benchmark. We show that adversarial learning increases all models generalization capabilities both on in and out-of-domain data.

Références

MARZINOTTO G., DAMNATI G., BÉCHET F. & FAVRE B. (2019). Robust semantic parsing with adversarial learning for domain generalization. In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics : Human Language Technologies, Volume 2 (Industry Papers)*, p. 166–173, Minneapolis, Minnesota : Association for Computational Linguistics. DOI : [10.18653/v1/N19-2021](https://doi.org/10.18653/v1/N19-2021).