



Offre n°2024-07410

PhD Position F/M Spoken language detection and clustering

Type de contrat : Fixed-term contract

Niveau de diplôme exigé : Graduate degree or equivalent

Fonction : PhD Position

Contexte et atouts du poste

Inria Défense&Sécurité (Inria D&S) was created in 2020 to federate Inria's actions for the benefit of military forces. The PhD will be carried out within the audio processing research team of Inria D&S, under the supervision of Jean-François Bonastre and co-supervised by Raphaël Duroselle.

The thesis is part of a project aimed at explainable and frugal voice profiling. Voice profiling consists in extracting information from an audio recording, such as identity, spoken language, age, geographical and ethnic origin, or socio/patho/physiological marks in the voice. The aim of this project is to make voice profiling systems explainable without degrading performance. Explainability maintains the central role of operators in the process, giving them the means to make informed decisions.

Mission confiée

Approach

The considered approach is based on the definition of a set of generic vocal attributes, which are shared by a group of individuals. Only the presence or absence of an attribute within a given speech utterance is used to make a decision. Thus the system uses binary representations. This approach was introduced for the speaker verification task [1,2].

The proposed thesis aims at developing this methodology working on the spoken language recognition task [3]. The system aims to group segments belonging to the same language and determine whether they belong to a recognized set of languages or if they are an unknown language. In the latter case, the system should rely on the attributes it knows to analyse the proximity of the new unknown language to known languages.

Since the emergence of iVector models [4] (initially for speaker recognition) in spoken language recognition, the general scheme remains the same. The system relies on an embedding extractor, trained on a large dataset and responsible of representing an acoustic sequence of any duration by a fixed size vector. Then, 1:1 classifiers, comparing two languages, or 1:N classifiers, comparing N languages, are built, and a decision-making system relies on these classifiers to perform the various tasks. Neural networks, like bottleneck features have delivered very significant gains [5]. Subsequently, embeddings derived from neural models, known as "xVectors", replaced iVectors and made it possible both to increase model size (and performance) and to simplify the training recipe of the models [6]. More recently, pre-trained models such as WavLM [7] or MMS [8] have been used [9]. These generic models enable interesting gains, especially when the training dataset is small for some languages, but at the cost of a significant increase of the number of parameters.

These approaches have common limitations: they cannot explain their decision, they suffer a drop in performance over domain change, they have difficulty managing the imbalance between datasets of different languages and they are cumbersome to adapt or retrain. Finally, they offer little or nothing in the case of unknown languages.

In this project, we propose to start from a state-of-the-art model and to adapt the vocal attributes approach to the spoken language recognition task. A language could be represented by a binary vector corresponding to the presence/absence of attributes in that language, or by a scalar vector, indicating the frequency of attributes in the language. The attributes themselves can incorporate higher-level information, such as phonotactic and linguistic levels. With this architecture, an unknown language (in the sense that no data corresponding to this language is present in the training data) can be recognized and compared to known languages, for instance exploiting geo linguistic knowledge. A model of this new language can thus be built as soon as the first recording of that language is available, and then adapted without computational cost each time an additional recording is added. If necessary, the attribute extractor can be adapted by adding one or more attributes from the new data, without need to relearn the whole model. Therefore we expect significant gains, in terms of explainability, handling of unknown languages and context adaptation.

Goals

1. Apply the vocal attribute approach to spoken language recognition.
2. Study the capacity to learn or extend (new language or new attributes) a model from unlabelled or weakly labelled data (for instance only labelled with a region), optimizing the ration between quantity and quality of weak labels.
3. Explore this approach to analyse unknown languages.
4. Exploit this approach for the language clustering of audio documents, even when some languages are unknown.

Principales activités

- Bibliography, development and evaluation of spoken language recognition systems.
- Deep learning, adaptation of self-supervised speech processing models such as WavLM [7] or MMS [8] ;
- Semi-supervised learning ;
- Explainability of deep learning models.

Compétences

- Master level in computer science, mathematics or phonetics.
- Strong interest in applied research.
- Written and spoken English
- Signal processing
- Machine learning and deep learning
- Experience with deep learning toolkits such as pytorch or keras
- Speech processing experience, knowledge of open source toolkits such as kaldi or speechbrain

References

- [1] Ben-Amor, I., & Bonastre, J. F. (2022, April). BA-LR: Binary-Attribute-based Likelihood Ratio estimation for forensic voice comparison. In 2022 International workshop on biometrics and forensics (IWBF) (pp. 1-6). IEEE.
- [2] Ben-Amor, I., Bonastre, J. F., O'Brien, B., & Bousquet, P. M. (2023, August). Describing the phonetics in the underlying speech attributes for deep and interpretable speaker recognition. In Interspeech 2023.
- [3] Li, Haizhou, Bin Ma, et Kong Aik Lee. « Spoken Language Recognition: From Fundamentals to Practice ». *Proceedings of the IEEE* 101, n° 5 (mai 2013): 1136-1159. <https://doi.org/10.1109/JPROC.2012.2237151>.
- [4] Dehak, Najim, Patrick J Kenny, Réda Dehak, Pierre Dumouchel, et Pierre Ouellet. « Front-End Factor Analysis for Speaker Verification ». *IEEE Transactions on Audio, Speech, and Language Processing* 19, n° 4 (mai 2011): 788-798. <https://doi.org/10.1109/TASL.2010.2064307>.
- [5] Fér, Radek, Pavel Matějka, František Grézl, Oldřich Plchot, Karel Veselý, et Jan Honza Černocký. « Multilingually trained bottleneck features in spoken language recognition ». *Computer Speech & Language* 46 (1 novembre 2017): 252-267. <https://doi.org/10.1016/j.csl.2017.06.008>.
- [6] Snyder, David, Daniel Garcia-Romero, Alan McCree, Gregory Sell, Daniel Povey, et Sanjeev Khudanpur. « Spoken Language Recognition Using X-Vectors ». In *The Speaker and Language Recognition Workshop (Odyssey 2018)*, 105-111. ISCA, 2018. <https://doi.org/10.21437/Odyssey.2018-15>.
- [7] Chen, Sanyuan, Chengyi Wang, Zhengyang Chen, Yu Wu, Shujie Liu, Zhuo Chen, Jinyu Li, et al. « WavLM: Large-Scale Self-Supervised Pre-Training for Full Stack Speech Processing ». *IEEE Journal of Selected Topics in Signal Processing* 16, no 6 (octobre 2022): 1505-1518. <https://doi.org/10.1109/JSTSP.2022.3188113>.
- [8] Pratap, Vineel, Andros Tjandra, Bowen Shi, Paden Tomasello, Arun Babu, Sayani Kundu, Ali Elkahky, et al. 2023. « Scaling Speech Technology to 1,000+ Languages ». arXiv. <http://arxiv.org/abs/2305.13516>.
- [9] Alumäe, Tanel, et Kunnar Kukk. 2022. « Pretraining Approaches for Spoken Language Recognition: TalTech Submission to the OLR 2021 Challenge ». arXiv. <http://arxiv.org/abs/2205.07083>.

Avantages

- Subsidized meals,
- Partial reimbursement of public transport costs,
- Leave: 7 weeks of annual leave + 10 extra days off due to RTT (statutory reduction in working hours) + possibility of exceptional leave (sick children, moving home, etc.),
- Possibility of teleworking and flexible organization of working hours,
- Professional equipment available (videoconferencing, loan of computer equipment, etc.),
- Social, cultural and sports events and activities,

Rémunération

- 1st and 2nd year : 2082 € bruts - gross /month
- 3rd year : 2190 € bruts - gross /month

Informations générales

- Ville : PARIS
- Centre Inria : [Siège](#)
- Date de prise de fonction souhaitée : 2024-05-01
- Durée de contrat : 3 years

Contacts

- Équipe Inria : MIS-DEFENSE (DIRECTION)
- Directeur de thèse :
Maillet Florence / florence.maillet@inria.fr

A propos d'Inria

Inria est l'institut national de recherche dédié aux sciences et technologies du numérique. Il emploie 2600 personnes. Ses 215 équipes-projets agiles, en général communes avec des partenaires académiques, impliquent plus de 3900 scientifiques pour relever les défis du numérique, souvent à l'interface d'autres disciplines. L'institut fait appel à de nombreux talents dans plus d'une quarantaine de métiers différents. 900 personnels d'appui à la recherche et à l'innovation contribuent à faire émerger et grandir des projets scientifiques ou entrepreneuriaux qui impactent le monde. Inria travaille avec de nombreuses entreprises et a accompagné la création de plus de 200 start-up. L'institut s'efforce ainsi de répondre aux enjeux de la transformation numérique de la science, de la société et de l'économie.

Attention: Les candidatures doivent être déposées en ligne sur le site Inria. Le traitement des candidatures adressées par d'autres canaux n'est pas garanti.

Consignes pour postuler

Sécurité défense :

Ce poste est susceptible d'être affecté dans une zone à régime restrictif (ZRR), telle que définie dans le décret n°2011-1425 relatif à la protection du potentiel scientifique et technique de la nation (PPST). L'autorisation d'accès à une zone est délivrée par le chef d'établissement, après avis ministériel favorable, tel que défini dans l'arrêté du 03 juillet 2012, relatif à la PPST. Un avis ministériel défavorable pour un poste affecté dans une ZRR aurait pour conséquence l'annulation du recrutement.

Politique de recrutement :

Dans le cadre de sa politique diversité, tous les postes Inria sont accessibles aux personnes en situation de handicap.