



Offre n°2024-07569

Post-Doctoral Research Visit F/M Deep generative models for robust and generalizable audio-visual speech enhancement

Le descriptif de l'offre ci-dessous est en Anglais

Type de contrat : CDD

Niveau de diplôme exigé : Thèse ou équivalent

Fonction : Post-Doctorant

Contexte et atouts du poste

This postdoctoral research is part of the **REAVISE** project: "Robust and Efficient Deep Learning based Audiovisual Speech Enhancement" (2023-2026) funded by the French National Research Agency (ANR). The general objective of REAVISE is to develop a unified audio-visual speech enhancement (AVSE) framework. This will leverage recent breakthroughs in statistical signal processing, machine learning, and deep neural networks to create a robust and efficient AVSE system.

The postdoctoral researcher will be supervised by [Mostafa Sadeghi](#) (researcher, Inria), [Romain Serizel](#) (associate professor, University of Lorraine), as members of the [Multispeech team](#), and [Xavier Alameda-Pineda](#) (Inria Grenoble), member of the [RobotLearn team](#). The team has access to powerful computational resources, including efficient GPUs and CPUs, required for the experiments planned in this project.

Work environment: [Multispeech team](#), Inria Nancy, France.

Starting date & duration: October 2024 (flexible), for two years.

Mission confiée

Background. Audio-visual speech enhancement (AVSE) aims to improve the intelligibility and quality of noisy speech signals by utilizing complementary visual information, such as the lip movements of the speaker [1]. This technique is especially useful in highly noisy environments. The advent of deep neural network (DNN) architectures has led to significant advancements in AVSE, prompting extensive research into the area [1]. Existing DNN-based AVSE methods are divided into supervised and unsupervised approaches. In supervised approaches, a DNN is trained on a large audiovisual corpus, like AVSpeech [2], which includes a wide range of noise conditions. This training enables the DNN to transform noisy speech signals and corresponding video frames into a clean speech estimate. These models are typically complex, containing millions of parameters.

On the other hand, unsupervised methods [3-5] employ statistical modeling combined with DNNs. These methods use deep generative models, such as variational autoencoders (VAEs) [6] and diffusion models [7], trained on clean datasets like TCD-TIMIT [8], to probabilistically estimate clean speech signals. Since these models do not train on noisy data, they are generally lighter than supervised models and may offer better generalization capabilities and robustness to visual noise, as indicated by their probabilistic nature [3-5]. Despite these advantages, unsupervised methods remain less explored compared to their supervised counterparts.

Principales activités

Objectives. In this project, we aim to develop a robust and efficient AVSE framework by thoroughly exploring the integration of recent deep-learning architectures designed for speech enhancement, encompassing both supervised and unsupervised approaches. Our goal is to leverage the strengths of both strategies alongside cutting-edge generative modeling techniques to bridge their gap. This includes the implementation of computationally efficient multimodal (latent) diffusion models, dynamical VAEs [9], temporal convolutional networks (TCNs) [10], and attention-based methods [11]. The main objectives of the project are outlined as follows:

1. Develop a neural architecture that assesses the reliability of lip images—whether they are frontal, non-frontal, occluded, in extreme poses, or missing—by providing a normalized reliability score at the output [12];
2. Design deep generative models that efficiently exploit the sequential nature of data and effectively fuse audio-visual features;
3. Integrate the visual reliability analysis network within the deep generative model to selectively use

visual data. This will enable a flexible and robust framework for audio-visual fusion and enhancement.

References:

- [1] D. Michelsanti, Z. H. Tan, S. X. Zhang, Y. Xu, M. Yu, D. Yu, and J. Jensen, "An overview of deep learning-based audio-visual speech enhancement and separation," in *IEEE/ACM Transactions on Audio, Speech, and Language Processing*, vol. 29, 2021.
- [2] A. Ephrat, I. Mosseri, O. Lang, T. Dekel, K. Wilson, A. Hassidim, W.T. Freeman, M. Rubinstein, "Looking-to-Listen at the Cocktail Party: A Speaker-Independent Audio-Visual Model for Speech Separation," in *SIGGRAPH* 2018.
- [3] M. Sadeghi, S. Leglaive, X. Alameda-Pineda, L. Girin, and R. Horaud, "Audio-visual speech enhancement using conditional variational auto-encoders," in *IEEE/ACM Transactions on Audio, Speech and Language Processing*, vol. 28, pp. 1788 –1800, 2020.
- [4] A. Golmakani, M. Sadeghi, and R. Serizel, "Audio-visual Speech Enhancement with a Deep Kalman Filter Generative Model," in *IEEE International Conference on Acoustics Speech and Signal Processing (ICASSP)*, Rhodes Island, June 2023.
- [5] B. Nortier, M. Sadeghi, and R. Serizel, "Unsupervised Speech Enhancement with Diffusion-based Generative Models," in *IEEE International Conference on Acoustics Speech and Signal Processing (ICASSP)*, Seoul, Korea, April 2024.
- [6] D. P. Kingma and M. Welling, "An introduction to variational autoencoders," in *Foundations and Trends in Machine Learning*, vol. 12, no. 4, 2019.
- [7] Y. Song, J. Sohl-Dickstein, D. P. Kingma, A. Kumar, S. Emon, and B. Poole, "Score-based generative modeling through stochastic differential equations," in *International Conference on Learning Representations (ICLR)*, 2021.
- [8] N. Harte and E. Gillen, "TCD-TIMIT: An Audio-Visual Corpus of Continuous Speech," in *IEEE Transactions on Multimedia*, vol.17, no.5, pp.603-615, May 2015.
- [9] L. Girin, S. Leglaive, X. Bie, J. Diard, T. Hueber, and X. Alameda-Pineda, "Dynamical variational autoencoders: A comprehensive review," in *Foundations and Trends in Machine Learning*, vol. 15, no. 1-2, 2021.
- [10] C. Lea, R. Vidal, A. Reiter, and G. D. Hager. "Temporal convolutional networks: A unified approach to action segmentation," in *European Conference on Computer Vision (ECCV)*, pp. 47-54. Springer, Cham, 2016.
- [11] A. Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N Gomez, L. Kaiser, and I. Polosukhin, "Attention is all you need," in *Advances in neural information processing systems (NeurIPS)* 2017, pp. 5998–6008.
- [12] Z. Kang, M. Sadeghi, R. Horaud, and X. Alameda-Pineda, "Expression-preserving Face Frontalization Improves Visually Assisted Speech Processing," in *International Journal of Computer Vision (IJCV)*, 2022.

Compétences

The preferred profile is described below.

- Master's degree, or equivalent, in the field of speech/audio processing, computer vision, machine learning, or in a related field;
- Ability to work independently as well as in a team;
- Solid programming skills (Python, PyTorch) and deep learning knowledge;
- Good level of written and spoken English.

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 (after 6 months of employment) and flexible organization of working hours
- Professional equipment available (videoconferencing, loan of computer equipment, etc.)
- Social, cultural and sports events and activities
- Access to vocational training
- Social security coverage

Rémunération

2788€ gross/month

Informations générales

- **Thème/Domaine** : Langue, parole et audio
Statistiques (Big data) (BAP E)
- **Ville** : Villers lès Nancy
- **Centre Inria** : [Centre Inria de l'Université de Lorraine](#)
- **Date de prise de fonction souhaitée** : 2024-10-01
- **Durée de contrat** : 2 ans
- **Date limite pour postuler** : 2024-08-18

Contacts

- **Équipe Inria** : [MULTISPEECH](#)
- **Recruteur** :
Sadeghi Mostafa / mostafa.sadeghi@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.

L'essentiel pour réussir

Interested candidates are encouraged to contact Mostafa Sadeghi (mostafa.sadeghi@inria.fr), Xavier Alameda-Pineda (xavier.alameda-pineda@inria.fr), and Romain Serizel (romain.serizel@loria.fr), and upload the required documents (CV, transcripts, motivation letter, and recommendation letters) to the dedicated Inria Job platform.

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.