



Automatic Detection and Sub-typing of Primary Progressive Aphasia from Speech: Integrating Task-Specific Features and Spatio-Semantic Graphs

Fritz Peters¹, W Richard Bevan-Jones¹, Grace Threlfall¹, Jenny M Harris³, Julie S Snowden²,
Matthew Jones², Jennifer C Thompson², Daniel J Blackburn¹, Heidi Christensen¹

¹University of Sheffield, UK; ²University of Manchester, UK; ³University of Exeter, UK

fpeters3@sheffield.ac.uk, heidi.christensen@sheffield.ac.uk

Abstract

Primary progressive aphasia (PPA) describes a group of neurodegenerative diseases that predominantly affect language abilities. Its diagnostic process typically requires experienced clinicians, often available only in specialised hospital departments. Patients with PPA frequently display changes in speech and language early in the disease progression. In this study, we extracted acoustic, linguistic, and task-specific features from audio recordings and evaluated their utility for PPA classification. Using a subset of task-specific features, we detected PPA with 97% accuracy. For sub-typing, models trained on the full feature set achieved 74% accuracy in a three-way classification of PPA variants. Our results highlight the added value of task-specific features, which complement traditional approaches. Additionally, their visualisation offers an intuitive representation of task execution, improving clinical interpretability and potential diagnostic utility.

Index Terms: clinical applications of speech technology, Primary Progressive Aphasia, PPA

1. Introduction

In recent years, there has been an increased interest in the automatic detection and tracking of voice-based biomarkers for several neurological conditions, including cognitive decline associated with dementia. The potential impact of this innovation on the early detection and diagnosis could be transformational, not least because it gives the opportunity to provide diagnostic aid to specialists for very rare conditions like Primary Progressive Aphasia (PPA). PPA is a syndrome caused by several underlying diseases, including frontotemporal degeneration, Lewy body disease or Alzheimer's disease. It affects only 1 in 100,000, and there is no accurate diagnostic test for PPA. Instead, people undergo an often protracted diagnostic process. An automatic tool could speed up this anxious wait and provide objective expert-level identification of significant cues in a person's speech and language. This paper presents a study exploring the feasibility of detecting and sub-typing PPA based on acoustic, linguistic and task-specific features extracted from the speech signal.

1.1. Primary Progressive Aphasia

Primary progressive aphasia (PPA) describes a clinically heterogeneous group of disorders in which the primary problem is a progressive impairment of language affecting activities of daily living. The disorders are defined by their predominant clinical features and, in particular, findings on examination of language. Whilst these disorders have been described since the 1890s [1], they have had various names, and their diagnostic criteria have been debated and revised several times. There are

now three subtypes of PPA defined in consensus diagnostic criteria [2]. They are the non-fluent variant of PPA (nfvPPA), the semantic variant of PPA (svPPA) (or semantic dementia), and the logopenic variant of PPA (lvPPA). The core clinical features of the language disorder in each syndrome vary. In lvPPA, the core features are anomia, impaired sentence repetition and relative non-fluency of speech, sometimes with phonologic errors. In svPPA, the core deficit is loss of semantic memory underpinning knowledge of objects and concepts. In language examination this manifests most obviously as anomia, impaired single word and/or object comprehension and regularisation errors when pronouncing irregular words (surface dyslexia). The core features of nfvPPA in varying combinations are non-fluency, agrammatism and impaired syntactic comprehension.

The objective of this study is to pick up these changes in participants with PPA by extracting a set of acoustic, linguistic, and task-specific features for the automatic detection and sub-typing of PPA.

1.2. Automatic detection and sub-typing from speech

Previous research has explored the automatic extraction of speech biomarkers in Primary Progressive Aphasia and its various subtypes. For instance, Nevler et al. [3] demonstrated that participants with nfvPPA ($N = 15$) and lvPPA ($N = 23$) exhibit an increased speech pause rate. Additionally, individuals with the non-fluent variant show a restricted fundamental frequency (i.e., F0) range. Themistocleous et al. [4] extracted a set of acoustic features from audio recordings of participants performing a picture description task. Machine learning models trained on this feature set showed a relatively high classification accuracy for nfvPPA (82%) compared to svPPA (66%) and lvPPA (57%). For all classifiers, pause and F0-related measures ranked among the most important features. The lower classification accuracy for svPPA and lvPPA highlights the need to incorporate additional features, such as linguistic and task-specific features, to improve the automatic detection of these PPA variants. For instance, Cho et al. [5] extracted a set of lexical features from transcribed participant speech to differentiate between different variants of frontotemporal dementia (FTD) and PPA ($N = 64$). They were able to show statistically significant group differences on a number of these features, including the number of total words, nouns, and verbs in a participant's transcript. Moreover, Zimmerer et al. [6] achieved an accuracy of 90% in a binary classification task distinguishing between healthy controls (HC) and individuals from a diagnostic group. Their classifier was based on linguistic features. However, when extended to a more complex five-way classification task differentiating between HCs ($N = 20$), the three PPA variants ($N_{\text{nfvPPA}} = 34$; $N_{\text{lvPPA}} = 25$; $N_{\text{svPPA}} = 29$), and the behavioural variant of

FTD ($N = 14$), the accuracy dropped to 59%. In later work, Themistocleous et al. [7] extracted a combination of acoustic and linguistic features for the automatic sub-typing of different PPA variants: nfPPA ($N = 19$) vs lvPPA ($N = 16$) vs svPPA ($N = 9$). They trained a Deep Neural Network, which for this three-way classification achieved an accuracy of 80%, even outperforming trained clinicians. This highlights the importance of combining features from different domains for the automatic detection of Primary Progressive Aphasia, especially for the more complex task of sub-typing different diagnostic groups.

Besides these previously employed task-invariant acoustic and linguistic features, speech elicited by the Cookie Theft picture description task allows for additional task-specific features to be extracted. Such features cannot only help to distinguish between HCs and diagnostic groups but also to understand and visualise diagnosis specific patterns in the execution of the task. The Cookie Theft picture contains various aspects that may be identified by a participant, these are also called content information units (CIUs). Extracting and quantifying the CIUs can help to understand which aspects of the picture a participant is describing. Favaro et al. [8] demonstrated the robustness of CIU related features by showing significant group differences between participants with Alzheimer’s Disease and HCs across multiple corpora. Additionally, the automatic extraction of spatio-semantic graphs and features was proposed to further analyse the narrative path during picture description [9]. This approach showed promising results for effectively distinguishing between cognitively impaired and unimpaired speakers.

In general, the automatic extraction of acoustic and linguistic features from speech has shown promising results for detecting and sub-typing PPA. While these features are useful across various speech-elicitation tasks, they do not fully leverage task-specific information [3–7]. In the context of picture description, the automatic extraction of CIUs and spatio-semantic features has proven effective in detecting cognitive impairment, highlighting the potential benefits of incorporating task-specific features [8, 9]. This paper expands on these findings and explores the automatic extraction of acoustic, linguistic, and task-specific features for the detection and sub-typing of PPA. We then compare and visualise these task-specific features to understand diagnosis specific patterns in the task execution. All features were extracted from audio recordings of the most widely used picture description task, the Cookie Theft picture.

2. Dataset

The dataset consists of audio recordings of 59 participants completing the Cookie Theft picture description task from the Boston Diagnostic Aphasia Examination (Figure 1). The task was administered by a trained clinician who prompted the participant to describe the picture in as much detail as possible. Throughout the task, further prompts were given if necessary (e.g., if the participant seemed to be making an extended pause before giving a sufficient description of the picture.). The mean length of the resulting audio recordings is 48.0 seconds ($SD = 25.43s$). Data collection was carried out at a research institution in the UK between 2013 and 2014.

The Cookie Theft picture description task was first introduced as part of the Boston Diagnostic Aphasia Examination [10]. It requires participants to describe a kitchen scene in as much detail as possible. Ever since its introduction, the pic-

ture has become a routine task in clinical assessments of various cognitive and language disorders. Its widespread use extends to research on automatic dementia diagnosis, where it serves as a staple way to elicit patient speech [11, 12].

Participants were categorised into two diagnostic groups: HC ($N = 14$) and individuals with PPA. The PPA group was further subdivided into specific clinical subtypes, including the nfPPA ($N = 10$), lvPPA ($N = 13$), svPPA ($N = 8$), and a mixed Alzheimer’s Disease (mixed AD) ($N = 14$) group. The mixed AD group consists of participants presenting with shared clinical features of FTD and AD, reflecting the variability in the neuropathological causes of language and cognitive decline. For the purposes of binary classification, all clinical subtypes, including mixed AD, were grouped together as PPA for comparison against HCs.

2.1. Data preprocessing

Open AI’s whisper model (`‘medium.en’`)¹ was used to transcribe all audio files. To differentiate between participant and clinician speech, the generated transcripts were manually diarised and aligned with their respective audio files by a single human annotator. Subsequently, clinician speech was excluded from both the audio files and transcripts. This was achieved by trimming the original recordings and removing text prefixed with the clinician identifier in the transcripts. As a result, the original audio recordings and respective diarised transcripts, as well as the trimmed participant-only audio recordings and respective transcripts, were available.

3. Classification

3.1. Feature extraction

A total number of 363 features were extracted from the participants’ transcripts and audio files. These can be grouped into three domain-specific feature sets: acoustic ($N = 88$), linguistic ($N = 220$), and task-specific ($N = 55$) which are further subdivided into CIU ($N = 44$) and spatio-semantic ($N = 11$) features. The selection of the extracted features was based on previous research in the field and consultations with clinicians; further description below.

3.1.1. Acoustic features

All acoustic features were extracted from the audio files containing only participant speech using the eGeMAPSv02 feature set [13]. This feature set includes spectral, cepstral, and prosodic descriptors and is widely used in speech analysis to assess cognitive and affective states (e.g., [14]).

3.1.2. Linguistic features

The Linguistic Feature Toolkit (LFTK) is an open source library for the extraction of handcrafted linguistic features [15]. All of its 220 features were extracted from the participant only transcripts to form the linguistic feature set used in this work. These features can be grouped into four different classes. Surface features ($N = 24$) are general high-level descriptors, such as character, word, and sentence counts, that do not align with the more specific linguistic domains. Lexico-semantic features ($N = 70$) are word-level features and describe aspects such as

¹A. Radford et al., “Robust Speech Recognition via Large-Scale Weak Supervision,” arXiv preprint arXiv:2212.04356, 2022. Available: <https://arxiv.org/abs/2212.04356>.

lexical variation and the difficulty and frequency of the words used. Discourse features ($N = 57$) capture broader relationships between words and sentences and are mainly concerned with the use of named entities. Finally, syntax features ($N = 69$) reflect structural properties of speech and are constructed from part of speech tags and metrics that capture the readability of the text.

3.1.3. Task-specific features

Content information units. The majority of task-specific features is concerned with the CIUs of the Cookie Theft picture that a participant correctly identifies. We measured whether and how often a participant names each of the 20 CIUs (see Figure 1C for a list of the CIUs), or a respective synonym, outlined by [8]. Moreover, the total number of CIUs (including repetitions), the CIU-to-word ratio, and the percentage of correctly identified CIUs were measured. In addition, the number of clinician interventions was measured by counting the instances of clinician speech during the recording.

Spatio-semantic features. For each participant, a spatio-semantic graph was constructed, and corresponding features were extracted following the approach outlined in [9, 16]. First, the 20 Conceptual CIUs from [8] were manually assigned two-dimensional coordinates based on their respective locations in the Cookie Theft image. For each participant, the CIUs mentioned in their transcript were identified, stored in the order of occurrence, and encoded with their corresponding coordinates. Using the NetworkX toolkit [17], a directed graph was generated to represent the sequence in which the participant described the image. Each node in the graph corresponds to a correctly named CIU, while edges capture the temporal transitions between them. Euclidean distance measures the distance between two connected nodes. Additionally, the image was divided into four quadrants to analyse shifts in attention throughout the narrative.

Once the graph was constructed, a set of spatio-semantic features was derived to characterise the structural and spatial properties of the description. These features summarise the semantic sequencing and visuospatial organisation of CIUs. A total of 11 features were extracted, omitting one from the original set due to overlap with the CIU feature set. These graphs and features provide insights into how participants navigate the visual scene and structure their descriptions.

3.2. Classification

Random Forest models were employed for the detection and sub-typing of PPA. The detection task involved a binary classification to distinguish between HCs and participants from any diagnostic group. The sub-typing task included a five-class classification to differentiate between HCs and the four diagnostic categories (mixed AD, lvPPA, nfPPA, and svPPA). Additionally, a three-way classification task was performed to distinguish between the traditional diagnostic PPA categories: nfPPA, lvPPA, and svPPA.

A Random Forest (RF) classifier (N trees = 100) was used for all tasks, with two different approaches. In one approach, Recursive Feature Elimination with Cross-Validation (RFECV) was applied to identify an optimal subset of features. RFECV iteratively removes less informative features based on model performance across stratified 5-fold cross-validation. The final model was trained on the selected feature set, and classification performance was assessed using cross-validation predictions. In the second approach, the classifier was trained on the feature set

without feature selection to evaluate whether feature reduction improved classification performance. Both approaches were evaluated on the domain specific feature sets individually as well as the whole concatenated feature set for the binary, three-way, and five-way classification tasks. Moreover, predictions from the RF classifier without RFECV, trained on the domain-specific feature sets, were aggregated using hard majority voting, where the final prediction was based on the majority vote from the individual models trained on each domain-specific feature set.

Performance was evaluated primarily based on accuracy, with additional assessment of classification metrics such as F1-score. The models were implemented using the scikit-learn toolkit [18].

4. Results

4.1. Detection and sub-typing of PPA

Binary classification. The RF classifier without Recursive Feature Elimination, trained on the CIU feature set, achieved the highest performance for the binary detection of PPA, with an accuracy of 97%. This model outperformed both classifiers trained on the full feature set and the majority voting approach. A summary of the results can be found in Table 1.

3-way classification. The two most successful models for classifying the three traditional PPA variants achieved an accuracy of 74%. One of these models was the RF classifier without RFECV, trained on the complete feature set, while the other utilised RFECV for feature selection, also trained on the full feature set. Notably, the model with RFECV identified a substantial set of optimal features ($N = 341$), incorporating features from all four domain-specific feature sets.

5-way classification. The most complex sub-typing task, which involved distinguishing between HCs and all diagnostic groups, was best performed by the RF classifier with RFECV, trained on the full feature set. This model achieved an accuracy of 63% and selected an optimal feature set of 46 features, which also drew from all four domain-specific feature sets.

4.2. Task-specific features

4.2.1. Content Information Units

We conducted further analyses to examine differences between diagnostic groups and HCs in terms of which CIUs they mentioned. An ANOVA revealed statistically significant group differences in the naming of multiple CIUs. For instance, while many HCs identified that the water is overflowing, this detail was mentioned far less frequently by individuals in the diagnostic groups. The heatmap in Figure 1 visualises the naming patterns for each CIU and highlights those that showed significant group differences.

Additionally, the CIUs can be categorised into four groups: animate objects, inanimate objects, exterior elements, and actions. We analysed the proportion of named CIUs within each category and found significant group differences in all but the exterior category. In each case, HCs named more CIUs than any of the diagnostic groups.

4.2.2. Spatio-semantic graphs and features

Further analysis of the spatio-semantic features revealed significant group differences across multiple features. An ANOVA comparing the five diagnostic groups showed significant differences for several features, including Total Path Distance, Self

Table 1: Overview of classification results.

Classifier	Feature Set	Binary				3-Way				5-Way			
		Acc.	Precision	Recall	F1	Acc.	Precision	Recall	F1	Acc.	Precision	Recall	F1
RF RFECV	<i>all</i>	0.93	0.93	0.88	0.90	0.74	0.77	0.70	0.71	0.63	0.58	0.59	0.57
	<i>acoustic</i>	0.90	0.86	0.86	0.86	0.61	0.63	0.61	0.62	0.59	0.61	0.56	0.56
	<i>linguistic</i>	0.88	0.88	0.77	0.81	0.68	0.65	0.64	0.63	0.49	0.48	0.49	0.48
	<i>CIU</i>	0.92	0.88	0.90	0.89	0.65	0.63	0.60	0.59	0.53	0.52	0.49	0.49
	<i>spatio-semantic</i>	0.86	0.82	0.79	0.80	0.42	0.42	0.41	0.42	0.39	0.38	0.37	0.37
RF	<i>all</i>	0.92	0.95	0.82	0.86	0.74	0.75	0.72	0.72	0.53	0.44	0.49	0.46
	<i>acoustic</i>	0.83	0.78	0.72	0.74	0.58	0.62	0.56	0.56	0.49	0.41	0.44	0.42
	<i>linguistic</i>	0.88	0.93	0.75	0.80	0.68	0.66	0.65	0.65	0.44	0.42	0.44	0.42
	<i>CIU</i>	0.97	0.95	0.95	0.95	0.58	0.55	0.53	0.52	0.54	0.54	0.52	0.52
	<i>spatio-semantic</i>	0.83	0.78	0.72	0.74	0.35	0.37	0.36	0.36	0.39	0.38	0.37	0.37
	<i>majority voting</i>	0.93	0.96	0.86	0.90	0.58	0.62	0.54	0.55	0.56	0.57	0.53	0.53

Note. Precision, recall, and F1 scores are macro averages. Acc, Accuracy.

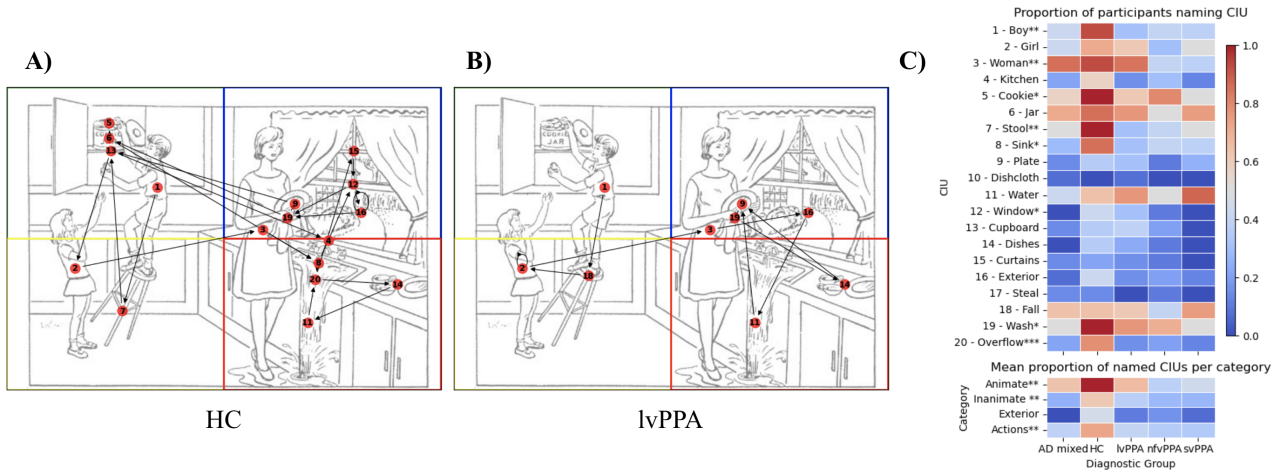


Figure 1: Cookie Theft picture from the Boston Diagnostic Aphasia Examination [10], highlighting the location of the 20 CIUs and their respective quadrants. A) and B) show spatio-semantic graphs for an HC and a participant with lvPPA. C) Heatmap of group differences in CIU identification. ANOVA was used to test significance, with * $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$ next to variable names.

Cycles, Cycles, Unique Nodes, Self Cycles (quadrants), and Cross Ratio (quadrants). Further explanation of these features is provided in [9]. Spatio-semantic graphs were constructed for all participants to visualise these differences. Figure 1 illustrates an example of a graph for a HC and a participant with lvPPA. Notably, the HC's graph displays distinct patterns compared to the lvPPA participant. The HC's graph shows a higher number of correctly identified CIUs (numbered, red circles), whereas the lvPPA participant identifies fewer CIUs. Both graphs exhibit a consecutive repetition of an individual CIU, but the total path length for the HC is much greater. The shorter path length, reduced number of identified CIUs, and repetitive CIU usage in the lvPPA participant reflect difficulties in generating a broader range of descriptive language, which is indicative of the anomia commonly seen in individuals with lvPPA. Overall, the spatio-semantic graphs provide a clear and intuitive means of interpreting how participants perform on the task, offering valuable insights into their cognitive and speech patterns.

5. Discussion and conclusions

Previous approaches to the automatic detection and sub-typing of PPA from speech have primarily relied on task-invariant

acoustic and linguistic features. In this study, we built upon these methods by additionally incorporating task-specific features that have been previously shown to be effective in detecting cognitive impairment. A random forest classifier trained on a subset of these features (CIU features) achieved the highest PPA detection accuracy (97%). Additionally, the best-performing models for the two sub-typing tasks (3-way accuracy = 74%; 5-way accuracy = 63%) were trained on the full feature set. Recursive feature elimination identified an optimal feature set that spanned all domains, highlighting the complementary role of task-invariant and task-specific features in distinguishing between PPA subtypes. Moreover, visualisations of the task-specific features provide an intuitive overview of how a person executes the Cookie Theft picture description task. In a clinical setting, these visual representations can enhance interpretability for clinicians, allowing for a quicker and more accessible assessment of a participant's speech patterns.

6. Acknowledgements

This work was supported by the UKRI AI Centre for Doctoral Training in Speech and Language Technologies (SLT) and their Applications funded by UK Research and Innovation [grant number EP/S023062/1].

7. References

References

- [1] A. Pick, "Über die Beziehungen der senilen Hirnatrophie zur Aphasie," *Prag Med Wchnschr*, vol. 17, pp. 165–167, 1892.
- [2] M. Gorno-Tempini, A. Hillis, S. Weintraub, A. Kertesz, M. Mendez, S. Cappa, J. Ogar, J. Rohrer, S. Black, B. Boeve, F. Manes, N. Dronkers, R. Vandenberghe, K. Rascovsky, K. Patterson, B. Miller, D. Knopman, J. Hodges, M. Mesulam, and M. Grossman, "Classification of primary progressive aphasia and its variants," *Neurology*, vol. 76, no. 11, pp. 1006–1014, Mar. 2011.
- [3] N. Nevler, S. Ash, D. J. Irwin, M. Liberman, and M. Grossman, "Validated automatic speech biomarkers in primary progressive aphasia," *Annals of Clinical and Translational Neurology*, vol. 6, no. 1, pp. 4–14, Jan. 2019.
- [4] C. Themistocleous, B. Ficek, B. Ficek, K. Webster, H. Wendt, A. Hillis, D. Den Ouden, and K. Tsapkini, "Acoustic markers of PPA variants using machine learning," *Frontiers in Human Neuroscience*, vol. 12, 2018.
- [5] S. Cho, N. Nevler, S. Ash, S. Shellikeri, D. J. Irwin, L. Massimo, K. Rascovsky, C. Olm, M. Grossman, and M. Liberman, "Automated analysis of lexical features in frontotemporal degeneration," *Cortex*, vol. 137, pp. 215–231, Apr. 2021.
- [6] V. C. Zimmerer, C. J. Hardy, J. Eastman, S. Dutta, L. Varnet, R. L. Bond, L. Russell, J. D. Rohrer, J. D. Warren, and R. A. Varley, "Automated profiling of spontaneous speech in primary progressive aphasia and behavioral-variant frontotemporal dementia: An approach based on usage-frequency," *Cortex*, vol. 133, pp. 103–119, Dec. 2020.
- [7] C. Themistocleous, B. Ficek, K. Webster, D.-B. Den Ouden, A. E. Hillis, and K. Tsapkini, "Automatic Subtyping of Individuals with Primary Progressive Aphasia," *Journal of Alzheimer's Disease*, vol. 79, no. 3, pp. 1185–1194, Feb. 2021.
- [8] A. Favaro, N. Dehak, T. Thebaud, J. Villalba, E. Oh, and L. Moro-Velázquez, "Discovering Invariant Patterns of Cognitive Decline Via an Automated Analysis of the Cookie Thief Picture Description Task," in *The Speaker and Language Recognition Workshop (Odyssey 2024)*. ISCA, Jun. 2024, pp. 201–208.
- [9] P. S. Ambadi, K. Basche, R. L. Kosciak, V. Berisha, J. M. Liss, and K. D. Mueller, "Spatio-Semantic Graphs From Picture Description: Applications to Detection of Cognitive Impairment," *Frontiers in Neurology*, vol. 12, p. 795374, Dec. 2021.
- [10] H. Goodglass, E. Kaplan, and S. Weintraub, *BDAE: The Boston diagnostic aphasia examination*. Lippincott Williams & Wilkins Philadelphia, PA, 2001.
- [11] S. Luz, F. Haider, S. D. L. Fuente, D. Fromm, and B. MacWhinney, "Alzheimer's Dementia Recognition Through Spontaneous Speech: The ADReSS Challenge," in *Interspeech 2020*. ISCA, Oct. 2020, pp. 2172–2176.
- [12] F. Tao, B. Mirheidari, M. Pahar, S. Young, Y. Xiao, H. Elghazaly, F. Peters, C. Illingworth, D. Braun, R. O'Malley, S. Bell, D. Blackburn, F. Haider, S. Luz, and H. Christensen, "Early Dementia Detection Using Multiple Spontaneous Speech Prompts: The PROCESS Challenge," in *ICASSP 2025 - 2025 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, 2025, pp. 1–2.
- [13] F. Eyben, M. Wöllmer, and B. Schuller, "Opensmile: the munich versatile and fast open-source audio feature extractor," in *Proceedings of the 18th ACM international conference on Multimedia*. Firenze Italy: ACM, Oct. 2010, pp. 1459–1462.
- [14] J. Chen, J. Ye, F. Tang, and J. Zhou, "Automatic Detection of Alzheimer's Disease Using Spontaneous Speech Only," in *Interspeech 2021*. ISCA, Aug. 2021, pp. 3830–3834.
- [15] B. W. Lee and J. Lee, "LFTK: Handcrafted Features in Computational Linguistics," in *Proceedings of the 18th Workshop on Innovative Use of NLP for Building Educational Applications (BEA 2023)*. Toronto, Canada: Association for Computational Linguistics, 2023, pp. 1–19.
- [16] S.-I. Ng, P. S. Ambadi, K. D. Mueller, J. Liss, and V. Berisha, "Automated Extraction of Spatio-Semantic Graphs for Identifying Cognitive Impairment," in *ICASSP 2025 - 2025 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, 2025, pp. 1–5.
- [17] A. Hagberg, P. Swart, and D. Chult, "Exploring Network Structure, Dynamics, and Function Using NetworkX," in *Proceedings of the 7th Python in Science Conference*, Jun. 2008.
- [18] F. Pedregosa, G. Varoquaux, A. Gramfort, V. Michel, B. Thirion, O. Grisel, M. Blondel, P. Prettenhofer, R. Weiss, V. Dubourg, J. Vanderplas, A. Passos, D. Cournapeau, M. Brucher, M. Perrot, and Duchesnay, "Scikit-learn: Machine Learning in Python," *Journal of Machine Learning Research*, vol. 12, no. 85, pp. 2825–2830, 2011.