

Abstract

Un modèle de la facilité d'écoute des documents audios pour les apprenants du français langue étrangère

Minami Ozawa

L'objectif de cette thèse est de proposer un modèle de prédiction automatique de la facilité d'écoute des documents audios en français. Les résultats du modèle de prédiction sont discutés afin de développer l'enseignement et l'apprentissage de la compréhension orale pour les apprenants de la langue française.

Les données utilisées dans cette thèse sont composées de documents audios et de transcriptions issus de 25 manuels de français langue étrangère. Nous analysons donc un modèle dans l'optique des variables linguistiques à l'aide de la transcription ainsi que des variables acoustiques à l'aide de l'audio. Les données sont classées en fonction des styles de parole – dialogue et monologue – et des modèles spécifiques à chacun d'entre eux sont envisagés.

L'analyse réalisée dans cette thèse au sujet des modèles prédictifs de la facilité d'écoute comporte deux étapes principales. Un premier modèle a été créé en se basant sur les données corrigées manuellement. Ensuite, un second modèle a été créé automatiquement, sans aucune correction manuelle des données. En comparant les résultats de ces deux modèles, nous avons mesuré la possibilité d'un modèle prédictif automatique de la facilité d'écoute. Pour chaque étape d'analyse, trois approches de prédiction ont été considérées : Machines à vecteur de support (SVM), wav2vec et CamemBERT.

En conséquence, il a été constaté qu'un modèle utilisant des variables linguistiques issues de la transcription était efficace pour prédire la facilité d'écoute. Concernant les variables acoustiques obtenues à partir de l'audio, leur corrélation avec la facilité d'écoute était plus élevée dans les dialogues que dans les monologues. De plus, l'analyse du modèle a révélé que ces variables acoustiques contribuaient, quoique légèrement, à améliorer la performance du modèle. Lors de l'automatisation du modèle, il a été identifié que l'alignement automatique de la parole

constituait un point faible. En tenant compte de cette faiblesse, il est nécessaire d'approfondir l'examen des variables acoustiques et d'explorer la mise en œuvre d'un modèle automatique de prédiction de la facilité d'écoute.

Cette thèse se compose de sept chapitres. Elle commence par une introduction et une explication de la problématique. Comme cette étude porte sur les compétences de la compréhension orale, le chapitre 1 résume les recherches antérieures en ce domaine. Il débute par une définition de la compréhension orale, puis identifie le processus de la compréhension orale et ses composantes. Ensuite, il examine les difficultés rencontrées par les apprenants de langues étrangères en la matière et mentionne les méthodes d'enseignement et d'apprentissage qui ont été testées jusqu'à présent. Le chapitre 1 se conclut en discutant des avantages de l'utilisation de matériaux authentiques et en confirmant l'utilité de mener des recherches sur la facilité d'écoute.

Le chapitre 2 donne une vue d'ensemble de la facilité d'écoute. Tout d'abord, il définit la facilité d'écoute, puis décrit, de manière chronologique, l'historique des recherches menées dans le domaine, en montrant les points d'inflexion. Il souligne que la recherche sur la facilité d'écoute s'est jusqu'à présent principalement concentrée sur la langue anglaise et que l'analyse n'a pas été assez effectuée dans le contexte du son.

Le chapitre 3 décrit la méthodologie de cette étude. Il clarifie les questions de recherche et explique les données et les méthodes. Il passe en revue la procédure d'analyse, en expliquant les outils et les moyens utilisés pour créer le modèle. Il fournit également une explication des variables utilisées dans cette thèse. Le niveau du Cadre européen commun de référence pour les langues (CECR) est défini comme variable dépendante tandis que les variables de lisibilité et de fluidité le sont comme variables indépendantes, en particulier pour le SVM.

Le chapitre 4 présente les résultats de l'analyse du premier modèle prédictif de la facilité d'écoute. Il constate que la performance prédictive augmente davantage avec le dialogue qu'avec le monologue, que le signal vocal ne contribue pas beaucoup à la performance du modèle et que la performance prédictive varie en fonction du niveau du CECR.

Le chapitre 5 détaille les résultats du second modèle. Nous montrons, notamment, que la performance prédictive de la facilité d'écoute ne varie pas considérablement en fonction du style de parole. En outre, à l'instar des résultats décrits au chapitre précédent, nous constatons que le signal vocal ne contribue pas nettement aux performances du modèle et que la performance de prédiction diffère en fonction du niveau du CECR.

Le chapitre 6 compare les résultats des chapitres 4 et 5 et discute de l'automatisation du modèle prédictif de la facilité d'écoute. Sur cette base, les points de difficulté liés à la création d'un modèle de prédiction entièrement automatisé sont identifiés.

Enfin, le chapitre 7 résume les résultats de l'analyse effectuée dans cette thèse. Il apporte les réponses aux questions de recherche et décrit la manière dont nos résultats peuvent être appliqués à l'enseignement, avant de mettre en perspective nos contributions.