

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

Predicting Second Language Reading Ability from Eye Fixation Patterns During Subtitled Video Viewing

Permalink

<https://escholarship.org/uc/item/9431m46t>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 48(0)

Authors

Li, Xiaokun

Qi, Yuxuan

Li, Yu

Publication Date

2026

Copyright Information

This work is made available under the terms of a Creative Commons Attribution License, available at <https://creativecommons.org/licenses/by/4.0/>

Peer reviewed

Predicting Second Language Reading Ability from Eye Fixation Patterns During Subtitled Video Viewing

Xiaokun Li

Department of Life Sciences, Beijing Normal-Hong Kong Baptist University, Zhuhai, China

Yuxuan Qi

Department of Life Sciences, Beijing Normal-Hong Kong Baptist University, Zhuhai, China

Yu Li

Department of Life Sciences, Beijing Normal-Hong Kong Baptist University, Zhuhai, China

Abstract

Reading requires eye movements. This study examines how second-language (L2) reading ability predicts eye fixation patterns on subtitles during L2 video viewing. Thirty-six native Chinese speakers watched English (L2) videos under three conditions: no subtitle, Chinese subtitle, and English subtitle. Eye movements were recorded using WebGazer. L2 reading ability was assessed in terms of reading fluency and reading accuracy. Results showed compared with the no-subtitle condition, there were significantly increased subtitle dwell time proportion and fixation counts in both Chinese and English subtitle conditions. Furthermore, participants with lower L2 reading ability, particularly lower reading accuracy, showed greater reliance on English subtitles, even when eye movements on Chinese subtitles were controlled for. Reading fluency was negatively associated with dwell time and fixation counts during English subtitle viewing. These results suggest that individual differences in L2 reading skills shape visual attention allocation and cognitive load during multimodal language comprehension.