

# A Study of the POS Keyword Caption Effect on Listening Comprehension

Jie Chi YANG<sup>a</sup>, Chia Ling CHANG<sup>b</sup>, Yi Lung LIN<sup>a</sup> & Mei Jen Audrey SHIH<sup>a</sup>

<sup>a</sup>*Graduate Institute of Network Learning Technology, National Central University, Taiwan*

<sup>b</sup>*Department of Information Management, National Central University, Taiwan*

yang@cl.ncu.edu.tw

**Abstract:** Listening comprehension is the most important part of English communication. Watching English films helps getting both audio and visual messages which can improve the learners' listening comprehension. In this study, we examined the full text captions films and the keyword captions with nouns and verbs to help understand the different effects of listening comprehension. One hundred English majored high school students participated in this experiment. This research indicated that with only nouns in the keyword captions the participants' listening comprehension was greatly assisted and reduced both the audio and visual senses from information overloading. On the other hand, with only verbs in the keyword captions, the participants' listening comprehension was excessively low. Moreover, the effects of listening comprehension for different types of films which include natural science and cultural history films were discussed. Natural science films provided more information without a necessary of a great deal of background knowledge; therefore it was easier for the participants to understand. However, it was difficult to comprehend cultural history films without having a certain level of knowledge of the background information. From the results, due to the participants' positive attitude towards the system satisfaction and learning approach, the participants had more self-confidence and higher motivation to operate the system and interface afterwards.

**Keywords:** Keyword Captions; Listening Comprehension; Part-of-Speech (POS) Tagger; Information Processing; Computer-Assisted Language Learning.

## 1. Introduction

Listening skill is the essential element in English learning [1] [2]. We can learn from many materials; especially video-based language learning is a good method to promote listening comprehension. This method integrates both the visual and audio messages which can help improve listening skills. Pavio [3] brought up the dual-coding theory which pointed out that while accepting the outside messages, human brains would make information encoding logogens and imagens. The information from both systems would produce correspondence and relation to help stimulate the strengthening of the memory. However, during the message processing, human brains tend to neglect the audio messages when there are too many visual messages to decode [4]. The learning theory revealed that too much noise would destroy the assisted learning effects [5]. Therefore, it is better to use the dual-coding method to help assist to decode messages for listening comprehension and prevent information overflowing. During the limited time of film watching, plenty of visual and audio messages would likely interfere with message decoding.

Most foreign films are being played with full text captions. Now we want to know how much content people understand while watching films with fewer texts or without any texts. The amount of captions will influence the level of listening comprehension while watching films. Schlesinger [6] pointed out that 2.5 to 3 chunking (the numbers of words that appear on the screen as captions each time) is a suitable quantity for assisting reading comprehension and also good for listening comprehension. Some researchers mention that

partial captions such as keyword captions are good for assisting listening comprehension and can help reduce information overload [7][8]. To this end, this study focuses on the effect of part-of-speech of keyword captions while watching films.

For English learning, Markham [9] indicated that people could improve listening and reading comprehensions through watching foreign films with captions. Reese [10] pointed out that text captions would only interfere in decoding the message and dispersing the learning effect with news that was rich in information as pictures and voices. Garza [11] evaluated whether the text captions could assist people to understand film contents. Participants were English native speakers learning Russian via films watching to compare the learning difference within the full text captions and non-text captions. The result showed a positive correlation between listening comprehension and full text captions.

This study focused on the different types of captions and effects on listening comprehension. There were three kinds of captions, nouns with the keyword captions, verbs with the keyword captions, and full text captions. We also examined the usages of different captions for different types of films for assisting listening comprehension.

## **2. System Implementation**

The system is divided into three areas: User Interface Area, POS Tagging Area, and Fetch Area. The User Interface Area has two interfaces that are for film processing and operator interfacing. The second area, The POS Tagging Area, is part-of-speech tagger and subtitles processor. This POS tagger rule is based on Brill's tagger [12]. After being processed in the part-of-speech tagger, the text captions are filtered to retain noun and verb to assist in text captions, and allow stop words to filter noise message. Then, POS keyword captions are processed to be the final transcripts. The third area is "Fetch Area" that POS KW Films Database combines the final transcripts, film database, and test models.

This system is called "Film-Based Computer-assisted English Listening comprehension system". There are three interfaces in this system. The System Interface includes Login in Interface, Demo Interface, and Test Interface. The operations of the system process are as follows: Part1. Login Interface: This provides participants to login into the system. Part2. DEMO Interface: this is the same as System Interface to allow familiarity of interface operation. Part3. Test Interface: When participants enter a film catalogue, they would have to watch one film and take an exam in fifteen minutes.

## **3. Experiment Design**

There are eighty-four girls and eighteen participants who are English majored high school students. The first experimental group used nouns with the keyword captions while the second experimental group used only verbs with the keyword captions. The control group used full text captions. Two types of films: natural science and cultural history, and six films in each category respectively. Every film endures for about five minutes.

The dependent variable is the preference of film content test which was called comprehension of the film contents. All test questions are from the content of films. There were three experts to co-edit the test questions. There were 25 questions from the natural science films and 27 questions from the cultural history films. The scale (7-point Likert's scale) shows the two questionnaires of English learning attitudes (the explanatory power is 76.37% & reliability is .913) and satisfaction (the explanatory power is 76.96%, reliability is .912) with the system. The five dimensions in the questionnaires are Confidence Dimension,

English Learning Attitudes scale, System Usability, Quality of provided Information, Interface Quality of System Satisfaction scale.

We processed the experiment for three weeks. TOEIC listening test was done in the first week and Film content-focused and understanding test for the second and third week. Participants had to write the TOEIC listening test and we then divided them into different level groups according to the result of the test. We arranged the different levels groups to our three caption groups and used the three caption groups as the same English listening test to process our experiment. The three groups were experimental groups: nouns with the keyword captions group (N) and verbs with the keyword captions (V) group. Both groups were attended by thirty-three participants. Thirty-four participants were in the control group with the full text captions group (F). In this experiment, experimental and control groups had to watch six subjects from the natural and science films and six subjects from the human and history films under the same conditions. Each film was played for approximately five minutes. In order to understand the learning results, we immediately performed the Film content-focused and understanding test after watching the films. Participants had to fill in the questionnaire of English learning attitudes after the completion of first stage. The questionnaire of system satisfaction was done after the second stage of the experiment.

## 4. Results

### 4.1 Results of Content-focused Understanding Test

The results of content-focused understanding test in films shows that using the three kinds of captions to enhance listening comprehension were significantly different ( $F_M=34.06$ ,  $F_{SD}=5.371$ ;  $V_M=30.94$ ,  $V_{SD}=3.061$ ;  $N_M=32.58$ ,  $N_{SD}=3.734$ ), with ANOVA test of association  $F=4.657^*$ ,  $p=.012<.05$ . With Post hoc tests of mean score different= $3.119^*$ ,  $p=.012<.05$  for full text captions group and verbs with the keyword captions group. But, it was not significantly different, with Post hoc tests of mean score different= $1.483$ ,  $p=.353>.05$  for full text captions group and nouns with the keyword captions group. This result showed that the effect for listening comprehension was almost the same while learners watched the films with full text captions or with only nouns in the keyword captions. But the effects for listening comprehension were different between full text captions group and with only verbs in the keyword captions group. The listening comprehension from watching films with only verbs in the keyword captions was less effecting than with only nouns in the keyword captions.

### 4.2 Results of Listening Comprehension on the Different Type of Films

The analysis result of keyword captions help improve listening comprehension on the different type of films: Natural science films were significantly different ( $F_M=16.12$ ,  $F_{SD}=3.436$ ;  $V_M=14.21$ ,  $V_{SD}=2.205$ ;  $N_M=15.00$ ,  $N_{SD}=2.487$ ), with ANOVA test of association  $F=4.015^*$ ,  $p=.021<.05$ , and the cultural history film was not significantly different ( $F_M=17.94$ ,  $F_{SD}=2.628$ ;  $V_M=16.73$ ,  $V_{SD}=2.066$ ;  $N_M=17.58$ ,  $N_{SD}=2.319$ ), with ANOVA test of association  $F=2.339$ ,  $p=.102>.05$ . This result indicates that the effect to enhance listening comprehension while watching natural science films was different with all three types of captions. ANOVA test was not significantly different for improving listening comprehension while watching cultural history films with the three different kinds of captions. Hence, no matter what kind of captions the learners selected, it did not affect any listening comprehension while watching cultural history films. On the other hand, there were the differences of the effect for helping listening comprehension with

three different types of captions while watching natural science films. With Post hoc tests of mean score different=1.906\*,  $p=.022<.05$  for full text captions group vs. verbs with the keyword captions group. But there was not much difference shown with Post hoc tests of mean score different=1.118,  $p=.260>.05$  for full text captions group vs. nouns with the keyword captions group. These results showed that there was a different effect for listening comprehension while watching natural science films but the effect was to be equal when learners watch films with full text captions or with only nouns in the keyword captions.

#### 4.3 Results of Investigations on the Questionnaire of English Learning Attitudes and System Satisfaction

There are five dimensions in the questionnaire, including confidence, English learning motivation, system usability, information quality and interface quality. The results of these five dimensions on full text captions group, only nouns with the keyword captions group and only verbs with the keyword captions group were analyzed.

Table 1 shows that confidence dimension was significantly differently, with ANOVA test of association  $F=7.465^{**}$ ,  $p=.001<.05$ . This result shows that these three types of keyword captions groups represent the different level of confidence in English learning. We compare the mean score of these three types of keyword captions groups on confidence dimension, the mean score with full text captions group was 5.64, the mean score with only nouns with the keyword captions group was 5.01 and the mean score with only verbs with the keyword captions group was 4.74. The mean score of the confidence dimension on full text captions group was higher than only nouns with the keyword captions group, and the mean score of only verbs with the keyword captions group was the lowest in these three groups. In addition, we compared the scores of confidence for these three types of keyword caption groups and uses Sheffe post hoc tests to test. Table 1 displays the result of the score of confidence between full text captions group and only verbs with the keyword captions group. Post hoc tests of association  $p=.001<.05$ , and the score of confidence between full text captions group and only nouns with the keyword captions group was significantly different, with post hoc tests of association  $p=.036<.05$ . When we look at the differences of the average score between full text captions group and only verbs with the keyword captions group was more different than the average score between full text captions group and only nouns with the keyword captions group. This result shows that participants have more confidence while watching films in full text captions, but they have less confidence while watching films with only verb with keyword captions.

Table 1. The result of investigation on caption groups

| Group(N)                                                  | Confidence | Eng-Learning motivation | System Usability  | Information Quality | Interface Quality |
|-----------------------------------------------------------|------------|-------------------------|-------------------|---------------------|-------------------|
| <i>F (34)</i>                                             | 5.64       | 5.47                    | 6.16              | 5.90                | 6.07              |
| <i>V (33)</i>                                             | 4.74       | 4.95                    | 5.92              | 5.76                | 5.60              |
| <i>N (33)</i>                                             | 5.01       | 5.04                    | 6.16              | 5.90                | 6.06              |
| <i>F test</i>                                             | 7.465      | 2.014                   | .954              | .229                | 2.631             |
| <i>P-value</i>                                            | .001**     | .139                    | .389              | .796                | .077              |
| <i>Confidence Post Hoc Test Mean Difference (P-value)</i> |            |                         | F-V               | F-N                 | V-N               |
|                                                           |            |                         | 3.619<br>(.001**) | 2.529<br>(.036*)    | 1.091<br>(.533)   |

In addition, Table 1 shows the different results for English learning motivation dimension ( $F=2.014$ ,  $p=.139>.05$ ), system usability dimension ( $F=0.954$ ,  $P=.389>.05$ ), information quality dimension ( $F=0.229$ ,  $p=.796>.05$ ) and interface quality dimension ( $F=2.631$ ,  $p=.077>.05$ ). According to Table 1, the mean score of English learning motivation

dimension, system usability dimension, information quality dimension and interface quality dimension were higher than the medium value of scale ( $M=3.5$ ).

These results indicate that using the three types of captions groups gives the participants a positive learning attitude and satisfaction in English learning motivation dimension, system usability dimension, information quality dimension and interface quality dimension.

## 5. Conclusions

This study showed that the participants in full text caption groups pay more attention to read the text captions. Over saturation of the visual sense channel at all times makes the participants have difficulty decoding audio information. Hence, the full text caption group usually neglects listening and focuses more on reading. Kintsch [13] explained films with captions could help viewers to understand the content, but the listening stimulation would be reduced. In this study, using only nouns in the keyword captions in films is more helpful for listening comprehension than full text captions and only verbs in the keyword captions. Moreover, these two kinds of films with different kinds of captions influence listening comprehension differently. Also, gaining knowledge from the film's content is difficult for the participants who do not know enough vocabulary. From the results, the cultural history films are difficult to watch without having enough background knowledge. Even if there is vocabulary to help the viewers understand the film content, the result on understanding the content is still limited. With nature science films in our study, it is easy for the participants to catch the main points and also to understand the content. Therefore, the captions have a different effect on understanding film content. In future study, the focus on learners' different backgrounds and increasing other types of film materials to classify film contents to understand participants' listening ability are recommended.

## References

- [1] Decker, M. A. (2004). Incorporating guided self-study listening into the language curriculum. *The Language Teacher*, 28(6), 5-9.
- [2] Verdugo, D. R., & Belmonte, I. A. (2007). Using digital stories to improve listening comprehension with Spanish young learners of English. *Language Learning and Technology*, 11(1), 87-101.
- [3] Paivio, A. (1971). *Imagery and verbal processes*. New York: Holt, Rinehart and Winston.
- [4] Seo, K. (2002). Research note: The effect of visuals on listening comprehension: A study of Japanese learners. *International Journal of Listening*, 16, 57-81.
- [5] Mayer, R. E., & Anderson, R. B. (1992). The instructive animation: Helping students build connections between words and pictures in multimedia learning. *Journal of Educational Psychology*, 84(4), 444-452.
- [6] Schlesinger, I. M. (1968). Sentence structure and the reading process. Mouton, The Hague.
- [7] Guillory, H. G. (1998). The effects of keyword captions to authentic French video on learner comprehension. *CALICO Journal*, 15, 89-108.
- [8] Jones, L. C., & Plass, J. L. (2002). Supporting listening comprehension and vocabulary acquisition in French with multimedia annotations. *Modern Language Journal*, 86(4), 546-561.
- [9] Markham, P. (1999). Captioned videotapes and second-language listening word recognition. *Foreign Language Annals*, 32(3), 321-328.
- [10] Reese, S. D. (1983). Improving audience learning from television news through between-channel redundancy. Presented at the Annual Meeting of the Association for Education in Journalism and Mass Communication (66<sup>th</sup>, Corvallis).
- [11] Garza, T. J. (1991). Evaluating the use of captioned video materials in advanced foreign language learning. *Foreign Language Annals*, 24(3), 239-258.
- [12] Brill, E. (1995). Transformation-based error-driven learning and natural language processing: A case study in part-of-speech tagging. *Computational Linguistics*, 21(4), 543-565.
- [13] Kintsch, W. (1988). The role of knowledge in discourse comprehension: A construction- integration model. *Psychological Review*, 95(2), 163-182.