

Digital engagement in pronunciation learning: effects on learning performance and language anxiety

URL	https://weko.wou.edu.my/?action=repository_uri&item_id=569
-----	---



Digital Engagement in Pronunciation Learning: Effects on Learning

Performance and Language Anxiety

Fei Ping Por^{1*}

Wawasan Open University

fppor@wou.edu.my

Soon Fook Fong²

Universiti Malaysia Sabah

sffong@ums.edu.my

Abstract The aim of this study was to investigate the effects of digital engagement towards the pronunciation performance of high, medium and low anxiety learners. The ePronounceTM was designed and developed in this study to enable the learners to engage with three presentation modes: Text+Sound+Phonetic Symbols(TSP), Text+Sound+Phonetic Symbols+Mouth Movements(TSPM), and Text+Sound+Phonetic Symbols+Face Gestures(TSPF). The Foreign Language Classroom Anxiety Scales (FLCAS) was employed to measure different levels of language anxiety, and the Pronunciation Competence Test was used as pretest and posttest to evaluate the pronunciation performance. The sample consisted of 329 Primary Five students from three different schools in Malaysia. Analyses of covariance (ANCOVA) and pairwise comparisons were conducted to examine the main effects and the interaction effects. The findings showed that there was no significant difference in the achievement scores attained by learners with different levels of language anxiety in the three presentation modes. Seemingly digital engagement is able to bring the low and high language anxiety students to medium language anxiety level for optimal learning under optimal learning condition as explained in the curvilinear relationship between anxiety and performance.

Keywords: digital engagement, presentation modes, pronunciation, learning performance, language anxiety

1. Introduction

In the design and development of multimedia-based learning, digital engagement is particularly relevant because it involves various presentation modes of representing information through text, audio, video, animation and any other media to learners. The information is represented in multiple formats via multiple sensory modalities, and then stored, transmitted and processed digitally. Similar to face-to-face teaching that depends on physical layout of the classroom, multimedia-based learning depends on the interface presentation modes. The importance of

interface presentation contributes to digital engagement to engage the learners on the focus area, active participation in learning and time on task. There is a strong correlation between engagement and achievement. When the interface presentation is able to engage learners digitally, the learners will have higher tendency to be behaviourally, emotionally and cognitively involved in learning activities. Consequently, compared to less engaged learners, engaged learners demonstrate more effort, experience more positive emotions and pay more attention in their learning (Fredricks, Blumenfeld, & Paris, 2004).

The interface presentation permits the demonstration of complicated processes in an interactive manner that instructional materials can be interconnected in a more natural and intuitive way. The audio and/or video production enhances learners' interaction with the digital instructional materials through less bridging effort between the learners and the information being processed and it provides autonomy in the learning process. The pedagogical agent in the animated human-like character or in the form of real human character provide instruction through verbal and non-verbal modes of communication which creates simulated connections between the digital instruction materials and the learners (Velu & Kaur, 2018).

Past literature has documented the importance of interface presentation modes in the presentation of digital instruction, reinforcement and assessment, particularly in pronunciation learning related to this study. For instance, the various aspects of pronunciation, such as vowels, pronunciation quality of individual words and general segments, have been considerably improved with the use of multimedia applications (Mich, Neri & Giuliani, 2006; Neri, Cucchiarini, & Strik, 2006; Pennington & Ellis, 2000; Por & Fong, 2013; Seferoğlu, 2005; Tanner & Landon, 2009; Wang & Munro, 2004). It makes possible presentation of speech sounds coordinated with written text and with other visuals on a screen, such as still graphics,

animations and full motion video. This in fact stimulates the auditory, visual, and kinesthetic channels of the learners. The various inputs increase learners' interest and motivation, and help establish connections between the abstract and the concrete (Boyd & Murphrey, 2002; Wald, 2008). In learning and teaching pronunciation particularly, the interface presentation in multimedia-based learning makes the invisible sound become visible, and concrete graphics appear in front of the pronunciation learners. The learners learn to pronounce the sound not only by listening, imitating and repeating, but also seeing the phonetic symbols and the movements of the articulatory organs that demonstrate to learners how closely their own pronunciation approximates model utterances (Boyd & Murphrey, 2002). This frequent practice through "listening discrimination and focused repetition exercises, automatic visual support" enhance learners' pronunciation performance and also train them to be active, independent and critical during information processing procedures (Levis, 2007, p. 184).

In traditional classroom with high teacher-learner ratio, learners often face anxieties when they have to pronounce the words publicly in class. For instance, learners with low pronunciation abilities may feel intimidated to practise the sounds orally and publicly. They worry about their mispronunciation. In addition, shy or introverted pronunciation learners are usually reluctant to speak out in class (Lacina, 2004). By giving learners a chance to learn individually, multimedia-based learning leads to a reduction of foreign language classroom anxiety and thus indirectly favour learning. Neri, Cucchiaroni, Strik, and Boves (2002) also said that learning must take place in a stress-free environment in which learners can be exposed to considerable and meaningful input and are stimulated to actively practise oral skills. By engaging in a non-threatening learning environment, self-confidence is also built through improvement of their pronunciation. The study of Stepp-Greany (2002) affirmed that learners gained confidence in

their abilities through using technology in their learning process without having to suffer embarrassment in front of others.

Therefore, the interface presentation modes can be used as an extremely valid tool for suitable learning tasks since it can be designed and developed with many features that are particularly appropriate for pronunciation learning. This study was designed to investigate the effects of epronounce™ with three modes of interface presentation on learners with different levels of language anxiety in the learning of pronunciation.

The three modes of interface presentation designed and developed for evaluation are as follows:

(i) Text + Sound + Phonetic Symbols (TSP) (as illustrated in Figure 1);

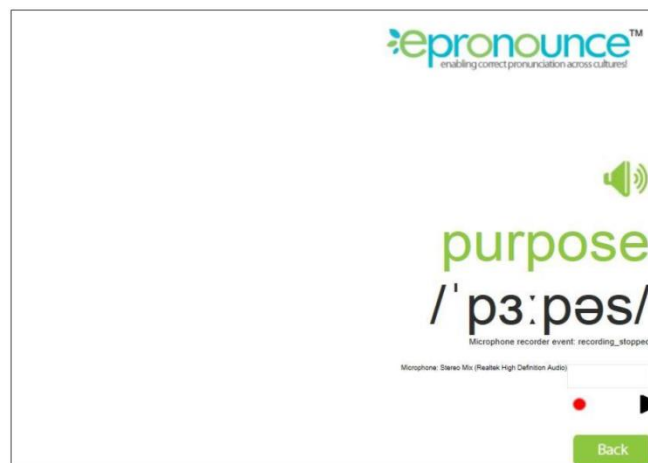


Figure 1. Text + Sound + Phonetic Symbols (TSP)

(ii) Text + Sound + Phonetic Symbols + Mouth Movements (TSPM) (as illustrated in Figure 2);

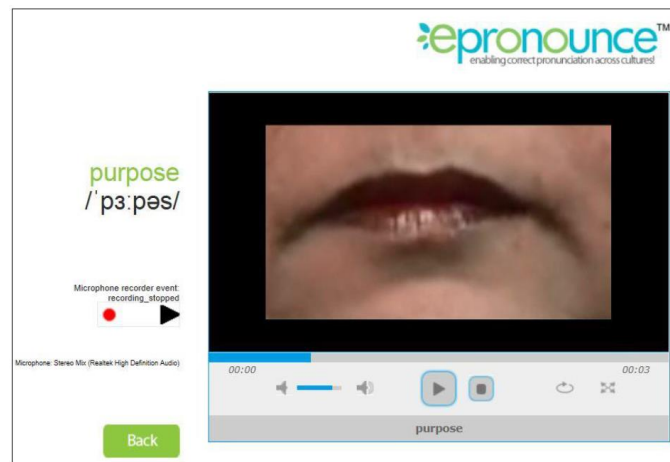


Figure 2. Text + Sound + Phonetic Symbols + Mouth Movements (TSPM)

(iii) Text + Sound + Phonetic Symbols + Face Gestures (TSPF) (as illustrated in Figure 3).

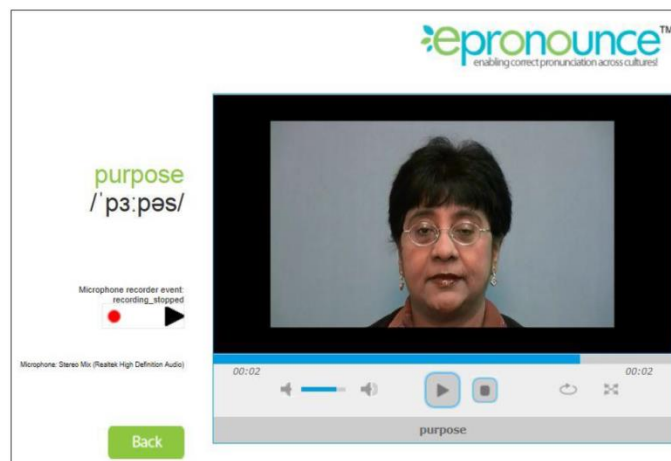


Figure 3. Text + Sound + Phonetic Symbols + Face Gestures (TSPF)

2. Theoretical Framework

The pedagogy-led epronounceTM was designed and developed based on established theoretical framework. It is important for instructional designers to develop innovative pronunciation

teaching strategies and complement the learning of English pronunciation in the English language curriculum. The results of this study provide an evident basis for instructional designers to design and develop presentation modes that best fit learners' identified needs as this study focused on the relationship between different presentation modes and different levels of language anxiety of learners to optimise the achievement of pronunciation learning.

The theories underlying this study are discussed as follows:

- (i) Second Language Acquisition Theory (Krashen, 1981)
- (ii) Baddeley's Model of Working Memory (Baddeley, 2000)
- (iii) Cognitive Theory of Multimedia Learning (Mayer, 2001)
- (iv) Cognitive Load Theory (Sweller, 1999)

Second Language Acquisition Theory (Krashen, 1981)

Stephen Krashen (1981) transformed language teaching and learning by developing Second Language Acquisition Theory, and he further developed the natural approach to language teaching together with Tracey Terrell (Krashen & Terrell, 1983). The design and development of epronounce™ imply the framework of Krashen's theory of second language acquisition by utilising phonetic symbols to monitor and correct pronunciation errors, and exposes pronunciation with phonetic symbols to the young learners at their early stage of growth in consistent with the natural order principle. In accordance with Krashen's formulation of $i + 1$ level, epronounce™ introduces pronunciation with phonetic symbols which is one step further beyond their current knowledge of phonics. To ensure the affective filter is low, epronounce™ offers a non-threatening learning environment for the learners to increase their comprehension and retention, minimise their language anxiety as well as maximise their self-confidence.

Baddeley's Model of Working Memory (Baddeley, 2000)

Baddeley's model of working memory comprises a central executive that interacts with the subsystems: the phonological loop and the visuospatial sketchpad, and also the episodic buffer. Each subsystem has its own limited capacity, which enables the subsystems to act relatively independent from each other, as shown by brain research that the subsystems are associated with different brain regions (Baddeley, 1998; Smith & Jonides, 1997). In this study, the learners were presented with the verbal stimuli which is processed in the phonological loop along with the mouth movements or face gestures for TSPM mode and TSPF mode respectively, engaging the visuospatial sketchpad.

Cognitive Theory of Multimedia Learning (Mayer, 2001)

The Cognitive Theory of Multimedia Learning (Mayer & Moreno, 2002) is formulated according to how human mind works in processing multimedia information to produce meaningful learning. The epronounceTM in this study provides various types of verbal and visual inputs which take advantage of both visual and verbal working memories without overloading one or the other (Mayer, 2001). Referring to the cognitive processes of the theory, the on-screen text, phonetic symbols, mouth movements/face gestures in epronounceTM are initially processed in the visual channel because they are brought in through the eyes, and the sounds of the word pronunciation are initially processed in the verbal channel as they are brought in through the ears. Learners will then cognitively select relevant text and graphics presented in epronounceTM and hold the corresponding verbal and visual representations in working memory. The connections will be built to organise the text and graphics in coherent mental representations. Finally, both verbal and visual mental models integrate with the learners' prior knowledge of phonics from long term memory to construct new knowledge to acquire correct pronunciation. It is noteworthy to understand that these processes are not necessarily linear because one or several

of the processes may happen simultaneously and they will occur iteratively until new knowledge is constructed (Mayer & Moreno, 2002).

Cognitive Load Theory (Sweller, 1999)

The Cognitive Load Theory proposes that learning can be enhanced by effective learning contents by directing cognitive resources toward activities that are relevant to learning (Chandler & Sweller, 1991). Initially, learners attempting to understand unfamiliar phonetic symbols may have a high level of intrinsic cognitive load. As learners understand how the phonetic symbols correlate with the sounds, low cognitive capacity is required. Although intrinsic cognitive load is a necessity in the learning process, techniques can be used to alleviate the load. For instance, the entire learning modules of epronounceTM are segmented into three main units giving the learners flexibility to digest the information in smaller chunks. Besides, sounds of phonetic symbols are learned first to prepare the learners for word pronunciation and minimal pairs. The interactivity element in epronounceTM is within the manageable effort, such as the record-play function is voluntary-based according to the learners' preference, and there is only one interactive record-play function for every word pronunciation.

A visual representation that relates all the theories in the theoretical framework is shown in Figure 4. In essence, the theoretical framework in relation to this study proposes the importance of pedagogy concerned when designing and developing epronounceTM to generate effective learning outcomes.

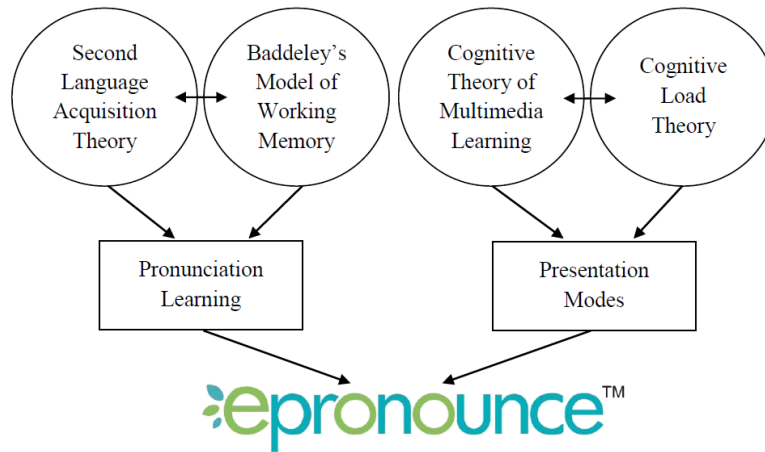


Figure 4. Visual Representation of Theoretical Framework

3. epronounce™ - Multimedia Pronunciation Learning Management System

The epronounce™ in this study is an interactive multimedia pronunciation learning management system, specially designed for young learners from non-native English speaking background to improve their pronunciation using phonetic symbols. The epronounce™ is a dynamic website with database management system and web applications. It goes beyond the conventional approaches by innovatively digitising the universally agreed system of phonetic symbols. The digitised phonetic symbols of epronounce™ with clickable sounds for each phonetic symbol, syllable and word accompanied with mouth movements and face gestures make a profound impact on the curriculum of learning and teaching pronunciation. The learning modules of epronounce™ are developed based on Mastery Learning Approach. The modules begin with laying the foundation on the basic sounds of phonemes, followed by combined sounds in words, and then proceed to minimal pairs for comparison and contrast.

Text and Sound Features of epronounce™

Learning to pronounce a word, to speak a new language, it depends primarily on hearing. By hearing the sound, the learners imitate and reproduce it. One or two vague hearings of the pronunciation of a word is insufficient to ensure good performance. Repetitive aural-oral drill is needed to build up a store of sound-memory which forms a library for the learners to acquire the sound system (Mukalel, 2007). With the text and sound features of epronounce™, the learners look at the word and simultaneously listen closely to the model pronunciation repeatedly, and then pronounce out the word. In this practice, epronounce™ supports ear recognition trainings and oral drillings which enables learners to hear and remember, recall and reproduce. The research of Iba (2008) demonstrated that the production of the participants who listened and repeated after the model pronunciation achieved higher scores than the production before listening. This approach is surprisingly simple as it does not demand a special knowledge of linguistics.

Listening to acquire pronunciation of new language involves a larger number of new skills, especially recognition skills. In order to listen to the new sound accurately, the learners must respond to a whole new sound structure. In fact, native language transfer is often observed to influence negatively the acquisition of the sounds of the second/foreign language (Celce-Murcia et al., 1996). Hence, the hearing of the learners is not adequately reliable as they are strongly influenced by the “phonological matrix of their native languages” (Schütz, 2008, p. 116) and they may unknowingly imitate these inaccuracies.

In understanding the needs and addressing the issues, epronounce™ in this study has incorporated the digitised phonetic symbols enabling learners to become active, independent and critical without mere reliance on ear.

Digitised Phonetic Symbols of epronounce™

The International Phonetic Association (IPA) was established in 1886 in Paris, in response to the inconsistencies of English orthography. The chief principle of the IPA in providing one unique symbol for one discrete sound and the symbol is used consistently for all languages (The International Phonetic Association, 2003) is meant to be easier for the pronunciation learners of non-native English speaking background to understand. As there is no overlapping of sounds, the IPA reduces the ambiguities in pronunciation learning. It is a useful aid for young learners to perceive sounds correctly. All the word pronunciation in epronounce™ is transcribed with phonetic symbols. Phonetic transcription is a system for writing the pronunciation of words using phonetic symbols in sequence to represent the speech sound of a word (The International Phonetic Association, 2003).

Stand on the promises of phonetic symbols, the interface design of epronounce™ is wholly featured with the IPA symbols which is a lack in some other pronunciation software. To further enhance epronounce™, presentation modes with mouth movements and face gestures are also designed and developed to visually and verbally guide learners through the pronunciation learning process in supplementing the digitised phonetic symbols.

Mouth Movements and Face Gestures of epronounce™

With the total dependence on sound imitation through hearing, it does not suffice to form new speech habit for non-native language. “The function of the ear is to perceive the finished sound product” (Reichmann, 1967, p. 398). Therefore, for young learners, they are incapable to just use their ears to analyse the motions of speech organs involved in producing the sounds for accurate imitation when learning the new language. Observation and imitation of lip, jaw and tongue movements are to be included to support the aural-oral approach.

According to the social agency theory, multimedia-based learning can be designed to foster virtual relationships between computers and learners by using visual social cues, or namely human agent (Atkinson, Mayer, & Merrill, 2005; Schroeder, Adesope, & Gilbert, 2013). The visual social cues of the human agent, such as facial expressions, gestures, and gaze, will engage learners in human-computer interaction as substitutes for authentic human-to-human interactions. Social cues thus result in learners being more motivated and investing more effort to understand the spoken words. Social agency theory stipulated that the life-like characteristics of a human agent prompt the social engagement of the learners, thus allowing the learners to form a simulated human bond with the agent (Atkinson, Mayer, & Merrill, 2005). Once this social partnership is established, learners will try to understand and deeply process the pronunciations produced by the friendly on-screen narrator which will improve the learners' schema activation, levels of cognitive processing, quality of learning, and ultimately increase the probability of positive knowledge transfer (Atkinson et al., 2005). Findings of Atkinson (2002) and Li (2008) indicated that the participants who were exposed to the narrator in combination with narrated instructions achieved higher scores than the control participants who were not exposed to the narrator. In view of the potential benefits of employing full face human narrator, this study incorporates full face gestures in one of the presentation modes of epronounceTM to visually and verbally guide learners through the pronunciation learning process, and to investigate its effectiveness.

In sum, auditory-visual feature implies practical applications in language learning. However, comparably less work has been done to find out which presentation mode, either mouth movements or face gestures, will yield better pronunciation competence among young learners. Hence, this study attempts to determine the effects of using the three presentation modes (TSP, TSPM, TSPF) in the learning of pronunciation.

4. Individual Differences and Pronunciation Learning

Individual differences among pronunciation learners play a noticeable role in learning as it will affect how any individual learns. These differences deserve great attention, particularly in multimedia-based learning because technology allows for the development of adaptive systems that support the learner's differences, which in turn enhance learning.

Different Levels of Language Anxiety

Horwitz et al. (1986) perceived language anxiety as “a distinct complex of self-perceptions, beliefs, feelings, and behaviours related to classroom language learning arising from the uniqueness of the language learning process” (p. 128). Steinberg and Horwitz (1986) revealed in their research that foreign language anxiety inhibits a learner’s ability to elaborate on thoughts and thus inhibit practice of the target language.

To identify the reasons behind language anxiety, Horwitz et al. (1986) noted that “anxious language learners complain difficulties in discriminating the sounds and structures of a target language message” (p. 126). They were also anxious as to whether they could pronounce correctly, speak fluently, and produce language grammatically correctly in public. learners also spoke of ‘freezing up’ when putting on the spot (Horwitz et al., 1986). In fact, the nature of pronunciation learning is a source of language anxiety. Finding a more efficient and less anxiety-producing means to learn pronunciation may, in turn, improve learners’ confidence when they practise pronunciation or speak in class. Creating a secure learning atmosphere and providing opportunities for the learners to make choices about their learning pace are feasible alternative to help reduce language anxiety. This is aligned with one of the purposes of designing and developing epronounceTM in this study. The study examines the use of epronounceTM as a tool in

the reduction of language anxiety to particularly address the needs of high language anxiety learners, while also examines its viability as a tool for pronunciation improvement by determining whether there is any significant difference in achievement scores among learners with different levels of language anxiety in using TSP, TSPM, and TSPF modes.

5. Method

To investigate the effects of TSP, TSPM and TSPF on learners with different levels of language anxiety, this study employed quasi-experimental factorial design. It is designed to investigate the effects of the independent variable on the dependent variable at each level of the moderator variables. The factors of the design in this study were the three presentation modes (TSP, TSPM, TSPF) and one moderator variable (language anxiety levels).

Research Samples and Sampling

The study was conducted on 373 Primary Five students (aged 11) but 44 students from the overall number did not manage to complete the experiment and tests required in the study. Therefore, the final total sample size calculated for analysis purposes in the study was 329. All the samples were taken from their normal intact classes, and there were a total of eleven classes involved in the study. They were randomly assigned to one of the three modes of epronounceTM (TSP, TSPM and TSPF).

The samples were sorted according to their language anxiety levels based on their scores on Foreign Language Class Anxiety Scale (FLCAS). Samples with FLCAS scores 1 standard deviation ($SD=0.72$) below the sample mean ($=2.77$) were categorised as low language anxiety, while samples with FLCAS scores in between 1 standard deviation ($SD=0.72$) above or equal to the sample mean ($=2.77$) and 1 standard deviation below or equal to the sample mean were

categorised as medium language anxiety. For samples with FLCAS scores 1 standard deviation ($SD=0.72$) above the sample mean ($=2.77$) were categorised as high language anxiety.

Instruments

In this study, there were two instruments used in collecting data. The instruments were:

- (i) Pronunciation Competence Test (Pretest and Posttest), and
- (ii) Foreign Language Classroom Anxiety Scale (FLCAS).

Pronunciation Competence Test (Pretest and Posttest)

There were 30 English words in the pretest and posttest, and the phonetic transcriptions were placed beneath the words. Items for both pretest and posttest were the same in terms of content to maintain consistency but the sequence was randomised to reduce item memory practice. To assess the Pronunciation Competence Test, the recording of individual participant's pronunciation was segmented using Praat acoustic analysis software. To quantify the pronunciation scoring objectively, each of the consonant phonemes of the 30 words was placed according to the position of Syllable Initial Word Initial (SIWI), Syllable Initial Within Word (SIWW), Syllable Final Within Word (SFWW), and Syllable Final Word Final (SFWF). For vowel and diphthong phonemes, each of them was classified into closed syllable or open syllable.

Foreign Language Classroom Anxiety Scale (FLCAS)

This study employed the Foreign Language Classroom Anxiety Scale (FLCAS) to assess the participants' language anxiety degree in affecting their performance in using epronounce™ for English pronunciation learning. This instrument was used particularly to determine whether there

was any significant difference in achievement scores among learners with different levels of language anxiety in using TSP, TSPM, and TSPF modes.

The FLCAS was scored by assigning a value of one to five points to the circled Likert response with single answer for each item. Responses indicating ‘strongly disagree’ received one point, and those indicating ‘strongly agree’ received five points. Thus the possible range of scores for the FLCAS was 33 to 165. In the case of negatively worded items (such as no. 2, 5, 8, 11, 14, 18, 22, 28, and 32), the values were reversed. The language anxiety score was gained by summing the ratings of the thirty-three items. On this instrument, a high score reflects a high level of language anxiety; whereas a low score indicates a low level of language anxiety. In this study, participants with FLCAS scores 1 standard deviation ($SD=0.72$) below the sample mean ($=2.77$) were categorised as low language anxiety, while participants with FLCAS scores in between 1 standard deviation ($SD=0.72$) above or equal to the sample mean ($=2.77$) and 1 standard deviation below or equal to the sample mean were categorised as medium language anxiety. For participants with FLCAS scores 1 standard deviation ($SD=0.72$) above the sample mean ($=2.77$) were categorised as high language anxiety.

Statistical Analysis

ANCOVA were conducted to test the hypotheses in this study to determine if there were statistically significant differences in the adjusted mean scores of the dependent variable (achievement scores of posttest) among the three presentation modes with different levels of language anxiety. The pretest scores were used as covariate. Prior to this, assumptions of ANCOVA were checked to ensure there was no violation of the assumptions of linearity, homogeneity of variance, and homogeneity of regression slopes. When the one-way ANCOVA yielded statistically significant result and there were more than two levels for the independent

variable, follow-up post-hoc pair wise comparisons were conducted to evaluate pair wise differences among the adjusted means. A two-way ANCOVA were carried out to determine if any interaction existed between the presentation modes and learners' and language anxiety levels.

6. Results and Discussion

The major research question addressed in this study concerned whether there is any significant difference in achievement scores among learners with different levels of language anxiety in using TSP, TSPM, and TSPF modes.

The two-way ANCOVA was conducted to examine the effects of language anxiety levels on the achievement scores of posttest according to presentation modes using pretest as covariate.

Referring to Table 1, there was no significant interaction effect between language anxiety level and presentation mode (FLCAS*Mode), $F(4, 319)=1.261$ at $p=0.285$. The p-value is greater than the 0.05 statistical significance cut-off level. When p-value is greater than the significance cut-off level ($p>0.05$), the interaction is considered not statistically significant (Aschengrau & Seage, 2008). This indicated that learners' language anxiety levels did not affect the posttest achievement scores among the three presentation modes. In other words, the effect of presentation modes on the achievement scores did not depend on the language anxiety levels. Due to the between-subjects effect was not significant, the follow-up analysis of pairwise comparisons was not needed to be conducted.

Table 1

Two-Way ANCOVA for Posttest Scores by Presentation Mode and Language Anxiety Level with Pretest as Covariate

Dependent Variable: Posttest

Source	Type III	Df	Mean Square	F	Sig.	Partial Eta	Observed
	Sum of Squares					Squared	Power ^b
Corrected Model	16731.667 ^a	9	1859.074	31.944	.000	.474	1.000
Intercept	13228.556	1	13228.556	227.301	.000	.416	1.000
Pretest	13010.026	1	13010.026	223.546	.000	.412	1.000
FLCAS	7715.003	2	3857.502	66.282	.000	.294	1.000
Mode	461.973	2	230.987	3.969	.020	.024	.710
FLCAS * Mode	293.571	4	73.393	1.261	.285	.016	.394
Error	18565.312	319	58.198				
Total	1392755.000	329					
Corrected Total	35296.979	328					

*a. R Squared = .474 (Adjusted R Squared = .459)**b. Computed using alpha = .05*

Table 2 presented the estimated marginal means and standard errors of the dependent variable by language anxiety levels in the three presentation modes. Estimated Marginal Means are the adjusted means with the effect of the covariate has been statistically removed. The findings demonstrated that learners with medium language anxiety attained the highest achievement scores (adjusted $M=67.500$), followed by learners with low language anxiety (adjusted $M=60.333$), and students with high language anxiety attained the lowest achievement scores (adjusted $M=52.802$).

Table 2

Estimated Marginal Means by Language Anxiety Level

Dependent Variable: Posttest

Language Anxiety Level	Mean	Std. Error	95% Confidence Interval	
			Lower Bound	Upper Bound
Low	60.333 ^a	1.060	58.247	62.418
Medium	67.500 ^a	.518	66.480	68.520
High	52.802 ^a	1.202	50.438	55.167

a. Covariates appearing in the model are evaluated at the following values: Pretest = 44.48.

The results of the two-way ANCOVA shown in Table 3 provided the adjusted means on the dependent variable for each group, split according to the level of language anxiety separately. Adjusted means refers to the fact that the effect of the covariate has been statistically removed. The findings demonstrated the adjusted means for the three presentation modes by low, medium and high language anxiety levels. For low language anxiety level, the adjusted means were reported as 58.832 for TSP mode, 60.712 for TSPM mode, and 61.454 for TSPF mode; while for medium language anxiety level, the adjusted means were reported as 64.820 for TSP mode, 66.672 for TSPM mode, 71.007 for TSPF mode. As for high visualisation level, the adjusted means were reported as 53.528 for TSP mode, 50.522 for TSPM mode), and 54.358 for TSPF mode.

Table 3

Estimated Marginal Means by Language Anxiety Level and Presentation Mode

Dependent Variable: Posttest

Language Anxiety Level	Presentation Mode	Mean	Std. Error	95% Confidence Interval	
				Lower Bound	Upper Bound
Low	TSP	58.832 ^a	1.805	55.280	62.384
	TSPM	60.712 ^a	1.923	56.928	64.496
	TSPF	61.454 ^a	1.688	58.134	64.775
Medium	TSP	64.820 ^a	.881	63.087	66.554
	TSPM	66.672 ^a	.852	64.996	68.348
	TSPF	71.007 ^a	.930	69.177	72.837
High	TSP	53.528 ^a	2.345	48.914	58.142
	TSPM	50.522 ^a	2.049	46.491	54.552
	TSPF	54.358 ^a	1.680	51.052	57.663

a. Covariates appearing in the model are evaluated at the following values: Pretest = 44.48.

H1 By using epronounceTM, the learners with different levels of language anxiety will attain significantly different achievement scores in the three presentation modes.

Referring to Table 1, there was no significant interaction effect between language anxiety level and presentation mode (FLCAS*Mode), $F(4, 319)=1.261$ and $p=0.285$. When p-value is greater than the 0.05 significance cut-off level, the interaction is considered not statistically significant. Therefore, this hypothesis was not supported.

H2 Learners with medium language anxiety (ML) will attain significantly higher achievement scores (AS) than learners with low language anxiety (LL) in the three presentation modes.

$$ASML > ASLL$$

Referring to Table 2, the achievement scores for medium language anxiety level (adjusted $M=67.500$) were higher than the achievement scores for low language anxiety level (adjusted $M=60.333$) in the three presentation modes, but $p=0.285$ ($p>0.05$) as shown in Table 1. This

indicated the differences were not significant among the achievement scores. Therefore, this hypothesis was not supported.

H3 Learners with medium language anxiety (ML) will attain significantly higher achievement scores (AS) than learners with high language anxiety (HL) in the three presentation modes.

$$ASML > ASHL$$

Referring to Table 2, the achievement scores for medium language anxiety level (adjusted $M=67.500$) were higher than the achievement scores for high language anxiety level (adjusted $M=52.802$) in the three presentation modes, but $p=0.285$ ($p>0.05$) as shown in Table 1. This indicated the differences were not significant among the achievement scores. Therefore, this hypothesis was not supported.

H4 Learners with low language anxiety (LL) will attain significantly higher achievement scores (AS) than learners with high language anxiety (HL) in the three presentation modes.

$$ASLL > ASHL$$

Referring to Table 2, the achievement scores for low language anxiety level (adjusted $M=60.333$) were higher than the achievement scores for high language anxiety level (adjusted $M=52.802$) in the three presentation modes, but $p=0.285$ ($p>0.05$) as shown in Table 1. This indicated the differences were not significant among the achievement scores. Therefore, this hypothesis was not supported.

H5 Learners with high language anxiety (HL) using the Text + Sound + Phonetic Symbols + Face Gestures (TSPF) mode will attain significantly higher achievement scores (AS) than learners with high language anxiety (HL) using the Text + Sound + Phonetic Symbols + Mouth Movements (TSPM) mode.

ASHL-TSPF > ASHL-TSPM

Referring to Table 3, the achievement scores for learners with high language anxiety using TSPF mode (adjusted $M=54.358$) were higher than the achievement scores for learners with high language anxiety using the TSPM mode (adjusted $M=50.522$), but $p=0.285$ ($p>0.05$) as shown in Table 1. This indicated the differences were not significant among the achievement scores. Therefore, this hypothesis was not supported.

H6 Learners with high language anxiety (HL) using the Text + Sound + Phonetic Symbols + Face Gestures (TSPF) mode will attain significantly higher achievement scores (AS) than learners with high language anxiety (HL) using the Text + Sound + Phonetic Symbols (TSP) mode.

ASHL-TSPF > ASHL-TSP

Referring to Table 3, the achievement scores for learners with high language anxiety using TSPF mode (adjusted $M=54.358$) were higher than the achievement scores for learners with high language anxiety using TSP mode (adjusted $M=53.528$), but $p=0.285$ ($p>0.05$) as shown in Table 1. This indicated the differences were not significant among the achievement scores. Therefore, this hypothesis was not supported.

H7 Learners with high language anxiety (HL) using the Text + Sound + Phonetic Symbols + Mouth Movements (TSPM) will attain significantly higher achievement scores (AS) than

learners with high language anxiety (HL) using the Text + Sound + Phonetic Symbols (TSP) mode.

ASHL-TSPM > ASHL-TSP

Referring to Table 3, the achievement scores for learners with high language anxiety using TSPM mode (adjusted $M=50.522$) were lower than the achievement scores for learners with high language anxiety using TSP mode (adjusted $M=53.528$), and $p=0.285$ ($p>0.05$) as shown in Table 1. This indicated the differences were not significant among the achievement scores. Therefore, this hypothesis was not supported.

Effects of Language Anxiety Levels with Presentation Modes on Pronunciation Learning

In this study it was hypothesised that the learners with different levels of language anxiety will attain significantly different achievement scores in the three presentation modes of epronounceTM. The results of this study however did not support these hypotheses. There are no significant interaction effects between language anxiety levels and presentation modes of epronounceTM. This clearly indicates that the learners of different language anxiety levels do not respond differently to epronounceTM. Their performance is at par with each other though the low and high language anxiety learners were initially expected not to perform as good as the medium language anxiety learners. Seemingly epronounceTM is able to bring the low and high language anxiety learners to medium language anxiety level for optimal learning under optimal learning condition as explained in the curvilinear relationship between anxiety and performance as well as the Yerkes-Dodson law (Keeley, Zayac, & Correia, 2008; Yerkes & Dodson, 1908).

Beauvois (1997, 1998) suggested the results are due to the fact that in multimedia-based learning learners are usually engaged more actively because of the low threatening atmosphere. This is in line with the Affective Filter principle of Krashen's Second Language Acquisition Theory

(Krashen, 2005). Krashen claimed that the best language acquisition takes place in an environment where anxiety level is low and defensiveness absent, or in another term where the affective filter is low. A low filter is associated with relaxation, confidence to take risks and a conducive learning environment which has been created by epronounce™ in this study. Krashen showed that learners whose anxiety level is low are much more likely to be successful language acquirers. Learning with epronounce™, the learners are more willing to practise their pronunciation because the mistakes made would not cause them to feel embarrassed in front of others. This situation motivates the learners to practise more and improve gradually. As a result, even reticent learners who tend not to participate in oral classroom discourse often become active contributors in the multimedia-based learning setting (Beauvois, 1998; Kelm, 1992; Kern, 1995; Meunier, 1998; Warschauer, 1996). It appears that multimedia-based learning setting provides enough practice and positive experiences for learners to become generally more engaged in language learning (Huang & Hwang, 2013; Rahimi & Yadollahi, 2011). Findings of this study suggest that epronounce™ functions as a practice platform for pronunciation learning not only in terms of pronunciation competence but also with regard to learners' affective state in which learners are seemingly more confident and engaged during learning sessions with epronounce™. The epronounce™ has also shown promise in bringing learners to medium language anxiety level for optimal learning by providing them learner-centred learning approach.

The learners also feel more at ease with epronounce™ because it is a forgiving and patient tutor (Lai, 2006) of willingly repeating the sounds for the learners ad infinitum with reliable quality in the sense of being the same every time (Pennington, 1999). Contrary, in traditional formal class setting, the learners experience fear when attempting to ask the human teachers to repeat the sounds many times because teachers may become impatient and other learners may also get irritated. In the context of this study, with epronounce™, the language anxiety of the learners is

addressed as the learners get more chance to immerse themselves in a language learning environment without fear and their pronunciation competence is enhanced. By increasing the frequency of listening to correct pronunciation with phonetic symbols, watching the videos of mouth movements or face gestures as many times as the learners desire, the learners are more engaged in sound discrimination and sound production during the information processing procedures.

The efficacy of epronounceTM with multi channels of media to transmit information has tremendously enhanced comprehension and engaged the learners actively. Therefore, it brings the learners' language anxiety to medium level which optimises their learning. The epronounceTM with the innovative use of texts, graphics, animations, videos and audios, and interactivity gives the impetus to learners to be more engaged to learning and therefore pay more attention to pronunciation learning. This in fact stimulates the verbal and visual channels of the learners. The various inputs increase learners' interest, and help establish connections between the abstract and the concrete (Boyd & Murphrey, 2002; Wald, 2008). The epronounceTM makes the invisible sound become visible, and concrete graphics appear in front of the learners. The learners learn to pronounce the sound not only by listening, imitating and repeating, but also seeing the phonetic symbols and the mouth movements as well as the face gestures.

The interactive real-time record-play function which allows the learners to record their own pronunciation and play back for listening to compare with the model pronunciation helps the learners engaged in the world of pronunciation learning and changes the role of the learners from passive contemplation to active participation which is, in turn, an essential factor for effective pronunciation learning.

Hence, pronunciation learning involves not only a cognitive process, but also a psychological process. The epronounceTM has seemingly brought the low and high language anxiety learners to medium language anxiety level for optimal learning under optimal learning condition. In regard to the private learning environment provided by epronounceTM, the high language anxiety learners manage to reduce their anxiety level by not having to practise their pronunciation publicly. Using the interactive record-play function of epronounceTM, the learners are allowed to keep practising their pronunciation privately and unlimitedly. Moreover, the learner-centred learning approach in epronounceTM helps the high language anxiety learners from being frustrated and the low language anxiety learners from getting bored. The learners can learn at a pace most effective to them.

7. Conclusion

The results of this study indicated epronounceTM is seemingly able to bring the learners to medium language anxiety level by engaging them digitally and hence optimising pronunciation learning. This is in line with the Affective Filter principle of Krashen's Second Language Acquisition Theory. Krashen claimed that the best language acquisition takes place in an environment where the affective filter is low. A low filter is associated with relaxation, confidence to take risks and a pleasant learning environment, as created by epronounceTM in this study. Krashen showed that learners who are highly motivated are much more likely to be successful language acquirers. With the innovative use of texts, graphics, animations, videos and audios, and interactivity gives the impetus to learners to be more attracted to learning and therefore pay more attention to pronunciation learning. The various inputs increase learners' interest and motivation, and help establish connections between the abstract and the concrete (Boyd & Murphrey, 2002; Wald, 2008). The epronounceTM makes the invisible sound become visible, and concrete graphics of face gestures appear in front of the learners. In accordance with

the Second Language Acquisition Theory (Krashen, 1981, 1985, 1999, 2005), Krashen proposed that learners can learn a large amount of language unconsciously where there is ample comprehensible input. In other words, language acquisition only takes place when comprehensible input is delivered sufficiently. This is another important theoretical implication of this study denotes the combination of various digital media types into an integrated multisensory interactive application ease learners' understanding and engaging in non-anxiety-provoking learning environments helps learners to enjoy the learning process and lowers the inhibition.

8. References

- Aschengrau, A., & Seage, G. R. (2008). *Essentials of epidemiology in public health* (2nd ed.). Sudbury, MA: Jones and Bartlett Publishers.
- Atkinson, R. (2002). Optimizing learning from examples using animated pedagogical agents. *Journal of Educational Psychology*, 94, 416-427.
- Atkinson, R. K., Mayer, R. E., & Merrill, M. M. (2005). Fostering social agency in multimedia learning: Examining the impact of an animated agent's voice. *Contemporary Educational Psychology*, 30(1), 117-139.
- Baddeley, A. D. (1998). *Human memory*. Boston: Allyn & Bacon.
- Baddeley, A. D. (2000). The episodic buffer: A new component of working memory? *Trends in Cognitive Sciences*, 4, 417-423.
- Beauvois, M. H. (1997). Computer-mediated communication (CMC): Technology for improving speaking and writing. In M. D. Bush & R. M. Terry (Eds.), *Technology-enhanced language learning* (pp. 165-184). Lincolnwood, IL: National Textbook Company.
- Beauvois, M. H. (1998). Conversations in slow motion: Computer-mediated communication in the foreign language classroom. *The Canadian Modern Language Review*, 54, 198-217.
- Boyd, B. L., & Murphrey, T. P. (2002). Evaluation of a computer based, asynchronous activity on student learning of leadership concepts. *Journal of Agricultural Education*, 43(1), 35-37.
- Celce-Murcia, M., Brinton, D. M., & Goodwin, J. M. (1996). *Teaching pronunciation*. Cambridge: Cambridge University Press.
- Chandler, P., & Sweller, J. (1991). Cognitive Load Theory and the format of instruction. *Cognition and Instruction*, 8, 293-332.

Fredericks, J., Blumenfeld, P., & Paris, A. (2004). School engagement: Potential of the concepts, state of the evidence. *Review of Educational Research*, 74, 59-109.

Horwitz, E. K., Horwitz, M. B., & Cope, J. A. (1986). Foreign language classroom anxiety. *The Modern Language Journal*, 70(2), 125-132.

Huang, P., & Hwang, Y. (2013). An exploration of EFL learners' anxiety and e-learning environments. *Journal of Language Teaching and Research*, 4(1), 27-35.

Iba, M. (2008). The influence of model sounds on the speech production of Japanese learners of English. *Language and Culture: The Journal of the Institute for Languages and Culture*, 12, 45-66.

Keeley, J., Zayac, R., & Correia, C. (2008). Curvilinear relationships between statistics anxiety and performance among undergraduate students: Evidence for optimal anxiety. *Statistics Education Research Journal*, 7(1), 4-15.

Krashen, S. D. (1981). *Second Language Acquisition and Second Language Learning*. Oxford: Pergamon Press Inc.

Krashen, S. D. (2005). *Explorations in language acquisition and use: The Taipei lectures*. Portsmouth, NH: Heinemann.

Krashen, S. D., & Terrell, T. (1983). *The natural approach: Language acquisition in the classroom*. Oxford: Pergamon Press Inc.

Lacina, J. (2004). Promoting language acquisitions: Technology and English language learners. *Childhood Education*, 81(2), 113-116.

Levis, J. (2007). Computer technology in teaching and researching pronunciation. *Annual Review of Applied Linguistics*, 27, 184-202.

Li, X. (2008). Effects of human agent and the application of the modality principle on the learning of Chinese idioms and the attitudes among students with different levels of intelligence. Master's Thesis (Unpublished). Universiti Sains Malaysia, Penang.

Mayer, R. E. (2001). *Multimedia learning*. New York: Cambridge University Press.

Mayer, R. E., & Moreno, R. (2002). Aids to computer-based multimedia learning. *Learning and Instruction*, 12, 107-119.

Mich, O., Neri, A., & Giuliani, D. (2006). *The effectiveness of a computer assisted pronunciation training system for young foreign language learners*. Paper presented at the Proceedings of CALL 2006 (pp. 135-143), Antwerp, Belgium.

Mukalel, J. C. (2007). *Approaches to English language teaching*. New Delhi: Discovery Publishing House.

Neri, A., Cucchiarini, C., Strik, H., & Boves, L. (2002). The pedagogy-technology interface in computer assisted pronunciation training. *Computer Assisted Language Learning*, 15(5), 441-467.

Neri, A., Cucchiarini, C., & Strik, H. (2006). *ASR-based corrective feedback on pronunciation: Does it really work?* Paper presented at the Proceedings of ICSLP 2006 (pp. 1982–1985), Pittsburgh, PA.

Pennington, M. C., & Ellis, N. C. (2000). Cantonese speakers' memory for English sentences with prosodic clues. *The Modern Language Journal*, 84, 372-389.

Por, F. P., & Fong, S. F. (2013). The Use of Mouth Movement Video and Digitised Phonetic Symbols on Pronunciation Learning among Learners with Different Psychological Profiles. *International Journal of Academic Research in Business and Social Sciences*, 3(2), 289-297.

Rahimi, M., & Yadollahi, S. (2011). Foreign language learning attitude as a predictor of attitudes towards computer-assisted language learning. *Procedia Computer Science*, 3(2011), 167–174.

Reichmann, E. (1967). Some visual aspects of pronunciation. *The German Quarterly*, 40(3), 398-403.

Schroeder, N. L., Adesope, O. O., & Gilbert, R. B. (2013). How Effective are Pedagogical Agents for Learning? A Meta-Analytic Review. *Journal of Educational Computing Research*, 49(1), 1–39. <https://doi.org/10.2190/EC.49.1.a>

Schütz, R. (2008). The importance of pronunciation. Retrieved from <http://www.sk.com.br/pronunciation.ppt>

Seferoğlu, G. (2005). Improving students' pronunciation through accent reduction software. *British Journal of Educational Technology*, 36(2), 303-316.

Smith, E. E., & Jonides, J. (1997). Working memory: A view from neuroimaging. *Cognitive Psychology*, 33, 5-42.

Steinberg, F. S., & Horwitz, E. K. (1986). The effects of induced anxiety on the denotative and interpretive content of second language speech. *Tesol Quarterly*, 20, 131-136.

Stepp-Greany, J. (2002). Student perceptions of language learning in a technological environment: Implications for the new millennium. *Language Learning & Technology*, 6(1), 165-180.

Sweller, J. (1999). *Instructional design in technical areas*. Camberwell, Australia: ACER Press.

Tanner, M. W., & Landon, M. M. (2009). The effects of computer-assisted pronunciation readings on ESL learners' use of pausing, stress, intonation, and overall comprehensibility. *Language Learning & Technology*, 13(3), 51-65.

The International Phonetic Association. (2003). *Handbook of the International Phonetic Association: A guide to the use of the International Phonetic Alphabet*. UK: Cambridge University Press.

Velu, U. R. A. R., & Kaur, S. (2018). Convergence of visual interpretation through collective practices of masculinity in a Malaysian televised show. *SEARCH: The Journal of the South East Asia Research Centre*, 10(2), 115-136.

Wang, X., & Munro, M. (2004). Computer-based training for learning English vowel contrasts. *System*, 32, 539–552.

Wald, M. (2008). Learning through multimedia: Speech recognition enhancing accessibility and interaction. *Journal of Educational Multimedia and Hypermedia*, 17, 215-233.

Yerkes, R. M., & Dodson, J. D. (1908). The relation of strength of stimulus to rapidity of habit-formation. *Journal of Comparative Neurology and Psychology*, 18, 459-482.

Authors' Biography

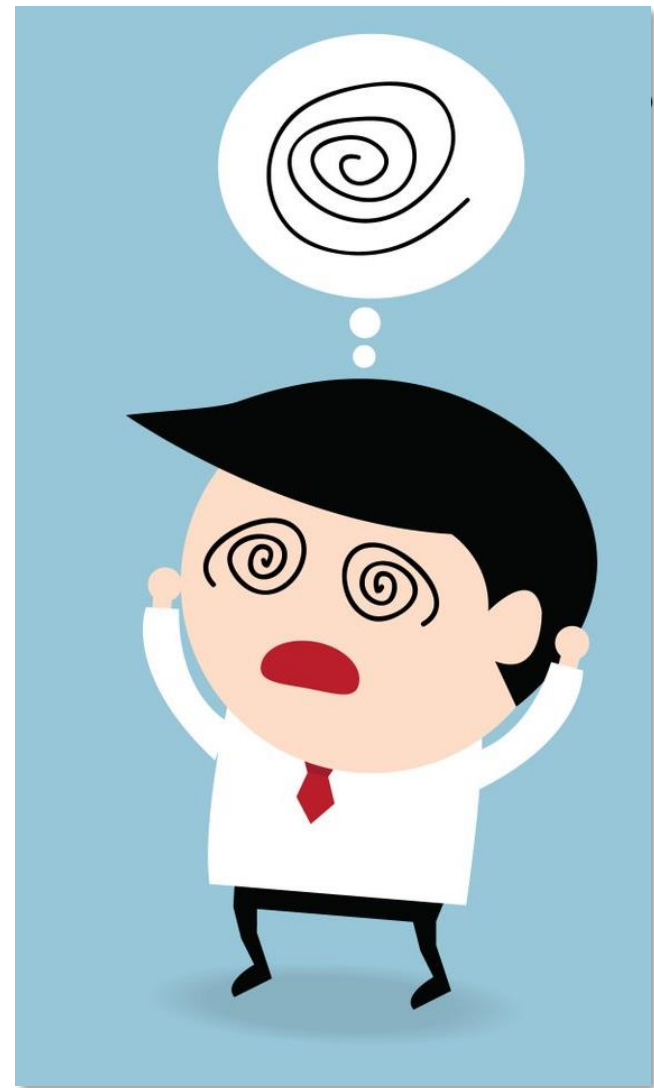
Dr Por Fei Ping is the Programme Coordinator for Master of Education and the Deputy Programme Leader for PhD (Arts and Humanities) in School of Education, Languages and Communications, Wawasan Open University. Fei Ping has shown tremendous passion in research related to educational technology and has won numerous awards.

Professor Dr Fong Soon Fook is the Director, Centre for e-Learning Universiti Malaysia Sabah. He is the Coordinator for OER@UMS and the Open-licensed Repository. His dedication as an innovative and bold researcher and educator has been recognised with various awards. He has successfully implemented consultation projects for Intel Corp., UNESCO, World Bank Institute, SEAMEO-RECSAM and various MOE agencies.

Digital Engagement in Pronunciation Learning: Effects on Learning Performance and Language Anxiety

Introduction

- Speaking in front of a crowd is anxiety-provoking.
- In traditional classroom with **high teacher-learner ratio**, learners often face anxiety when they have to pronounce the words publicly in class.
- Learners with low pronunciation abilities may feel intimidated to practise the sounds publicly. They worry about their mispronunciation.
- In addition, shy or introverted pronunciation learners are usually reluctant to speak out in class (Lacina, 2004).

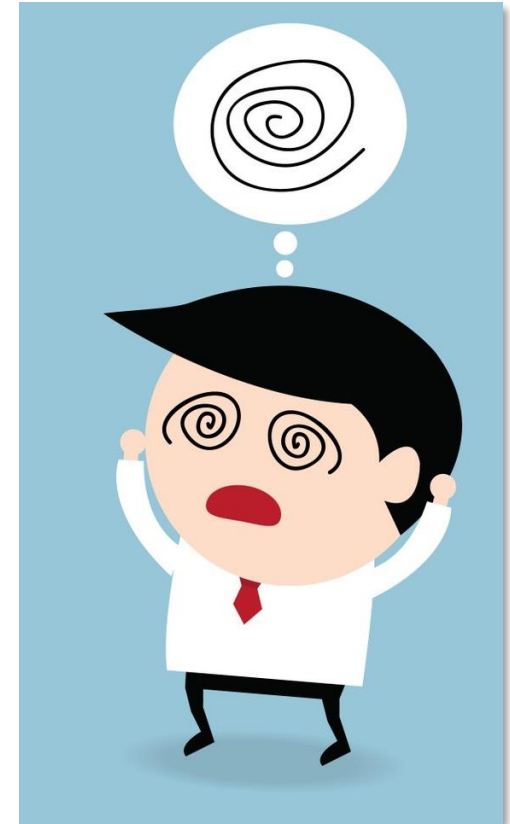


Language Anxiety

a distinct complex of **self-perceptions**, **beliefs**, **feelings**, and **behaviours** related to classroom language learning arising from the language learning process

Horwitz et al. (1986)

Steinberg and Horwitz (1986) revealed in their research that foreign language anxiety inhibits a learner's ability to elaborate on thoughts and thus inhibit practice of the target language.



- A more efficient and less anxiety-provoking means
 - pronunciation learning
 - improve learners' confidence when they practise pronunciation in class
- Creating a secure learning atmosphere and providing opportunities for the learners to make choices about their learning pace are feasible alternative to help reduce language anxiety.



- **Multimedia-based learning** - digital engagement is particularly relevant - involves **various presentation modes** of representing information



- The information is represented in **multiple formats via multiple sensory modalities**, and then stored, transmitted and processed digitally.
- Similar to face-to-face teaching that depends on physical layout of the classroom, multimedia-based learning depends on the interface presentation modes.



- Interface presentation - digital engagement
 - to engage the learners on the focus area, active participation in learning and time on task.
- Strong correlation between engagement and achievement

When the interface presentation is able to engage learners digitally, the learners will have higher tendency to be behaviourally, emotionally and cognitively involved in learning activities.

Compared to less engaged learners, engaged learners demonstrate more effort, experience more positive emotions and pay more attention in their learning.



Affordable • Accessible

(Fredricks, Blumenfeld, & Paris, 2004; Mich, Neri & Giuliani, 2006; Neri, Cucchiarini, & Strik, 2006; Pennington & Ellis, 2000; Seferoğlu, 2005; Tanner & Landon, 2009; Wang & Munro, 2004)

epronounceTM

Flexible • Affordable • Accessible





purpose
/'pɜ:pəs/

Microphone recorder event: recording_stopped

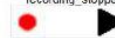
Microphone: Stereo Mix (Realtek High Definition Audio)



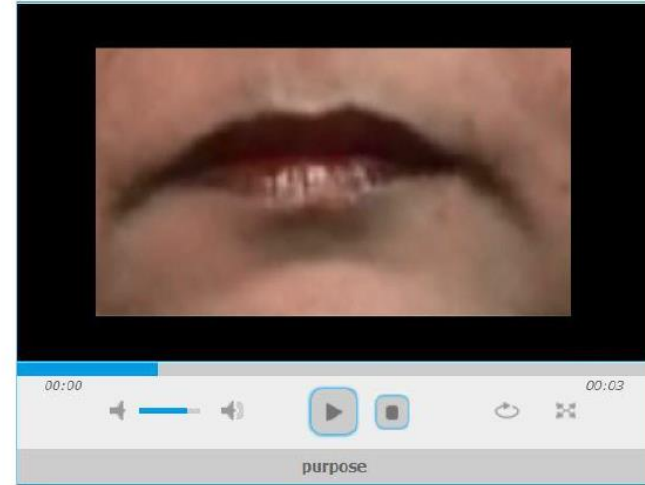
Back

purpose
/'pɜ:pəs/

Microphone recorder event: recording_stopped



Microphone: Stereo Mix (Realtek High Definition Audio)



Back

Text + Sound + Phonetic Symbols (TSP)

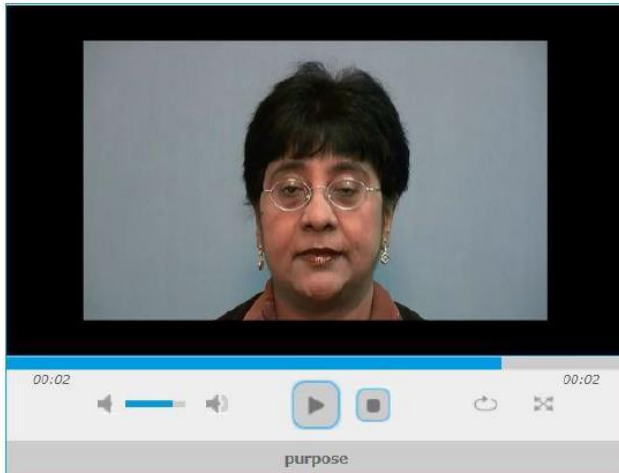
Text + Sound + Phonetic Symbols +
Mouth Movements (TSPM)

purpose
/'pɜ:pəs/

Microphone recorder event: recording_stopped



Microphone: Stereo Mix (Realtek High Definition Audio)

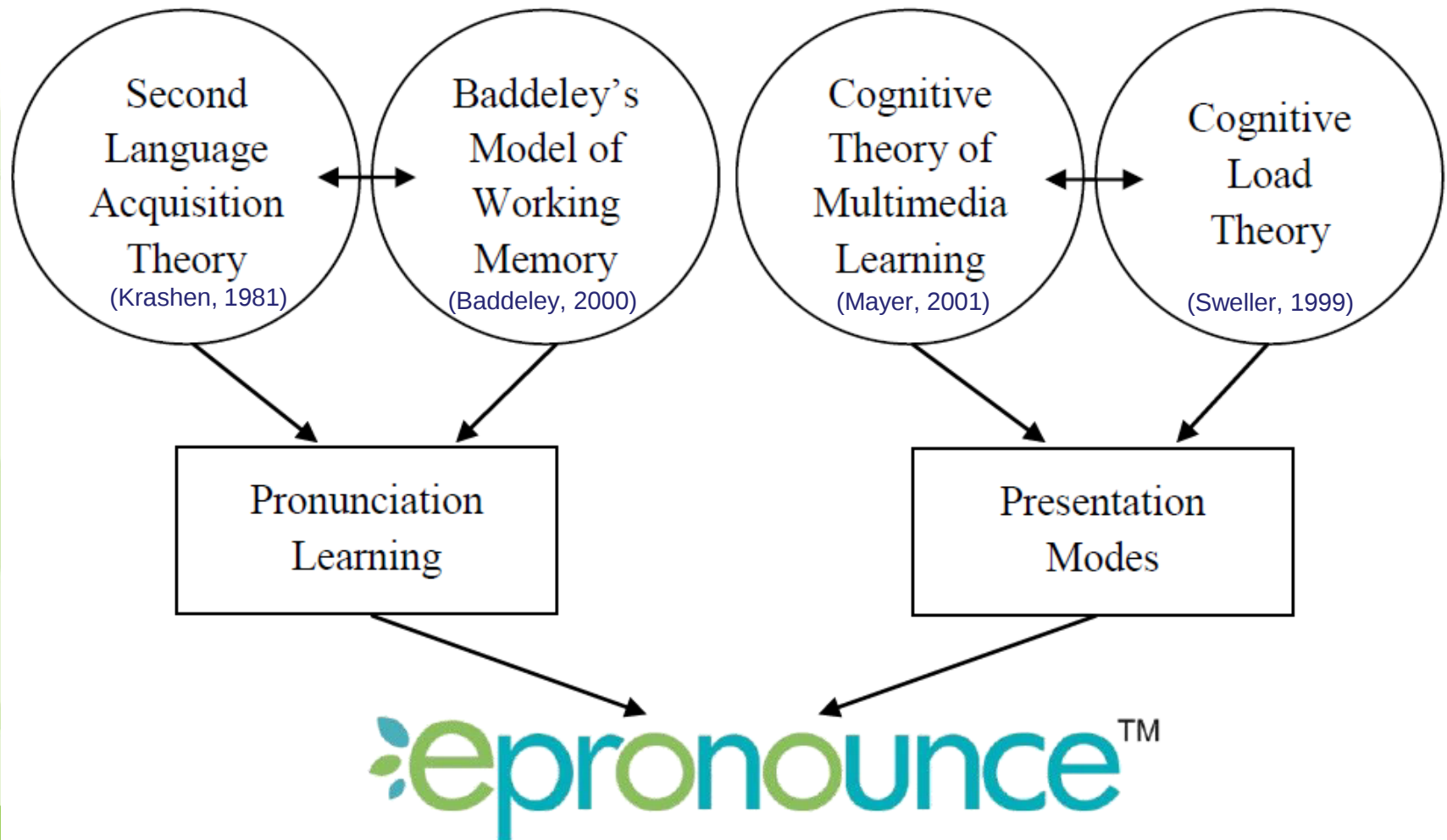


Back

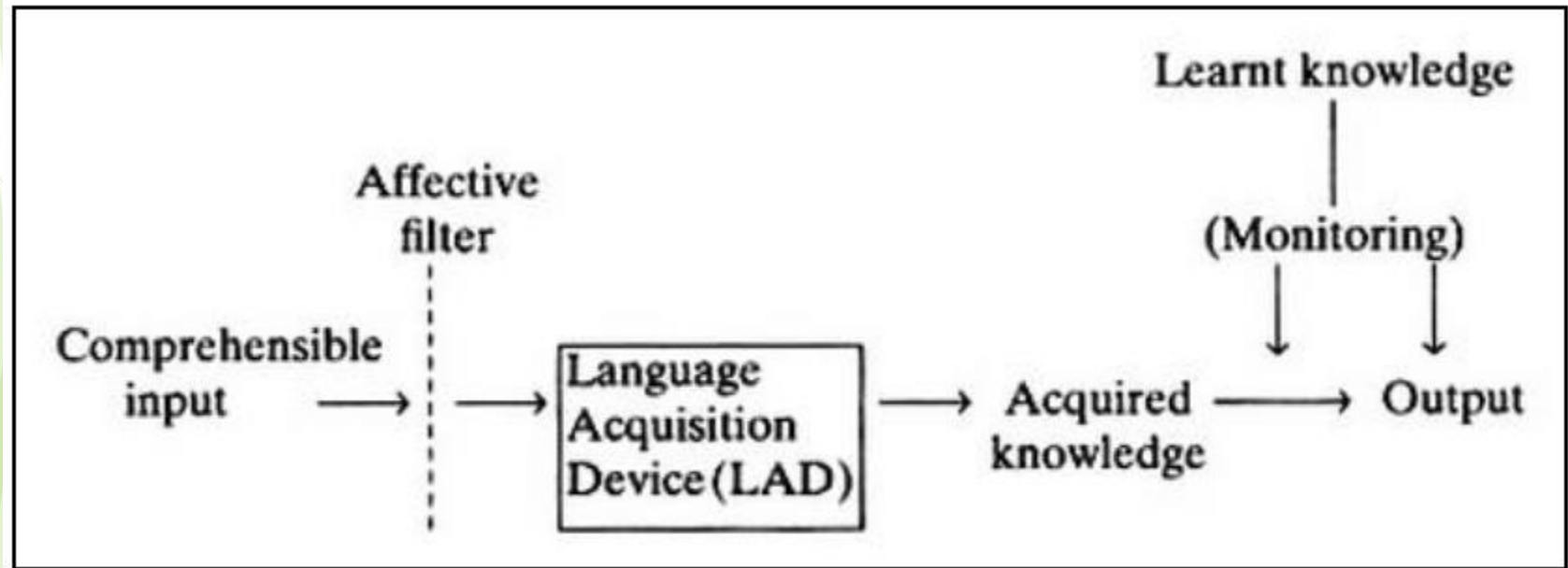
Flexible • Affordable • Accessible

Text + Sound + Phonetic Symbols + Face Gestures (TSPF)

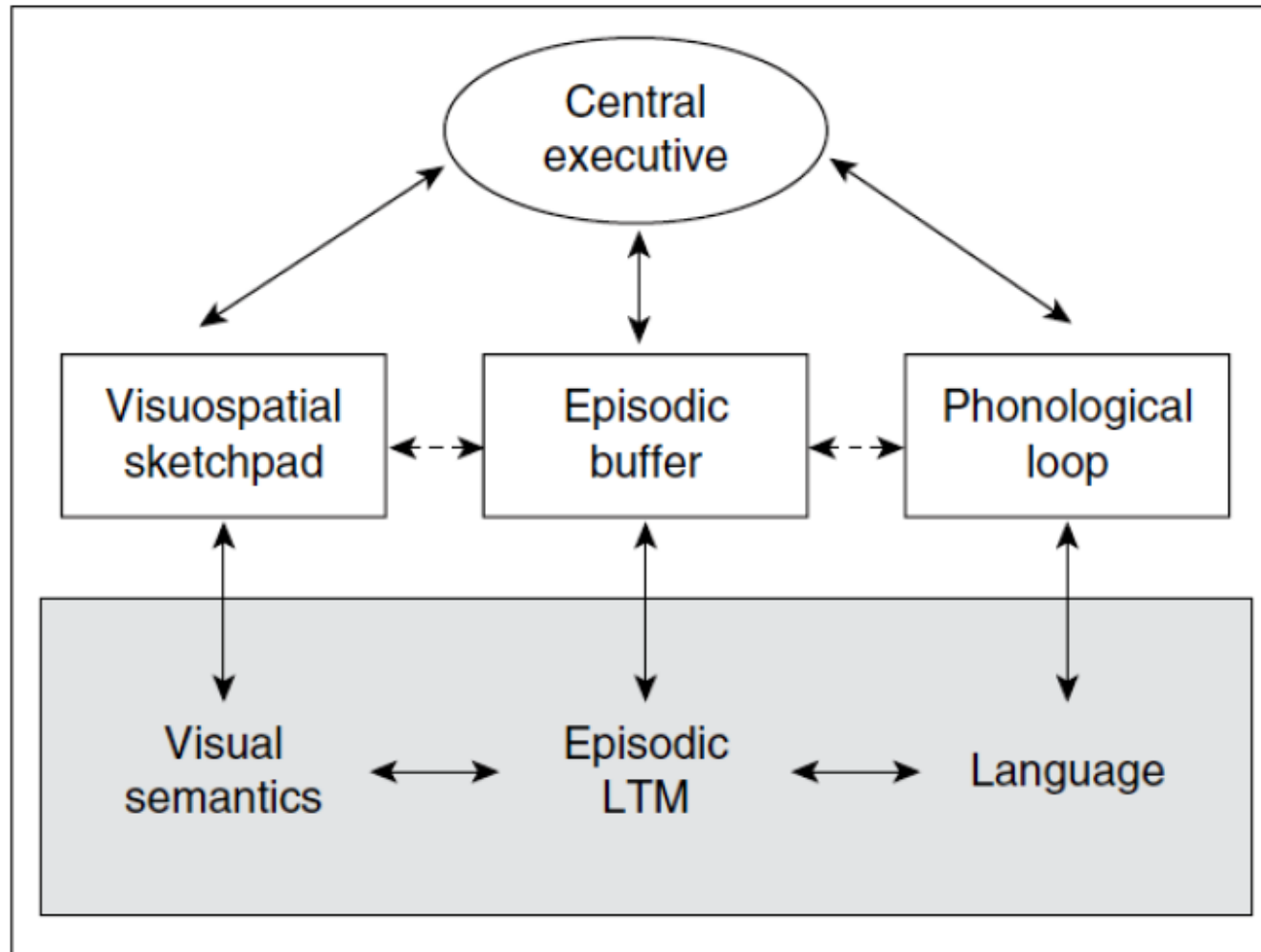
Theoretical Framework



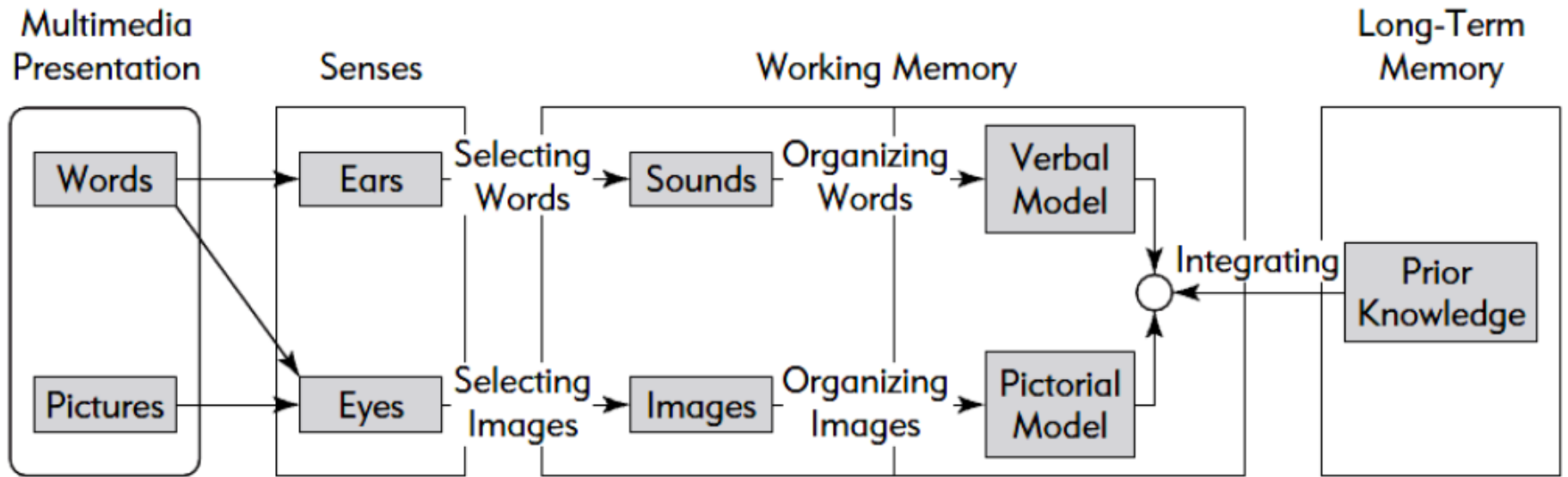
Second Language Acquisition Theory (Krashen, 1981)



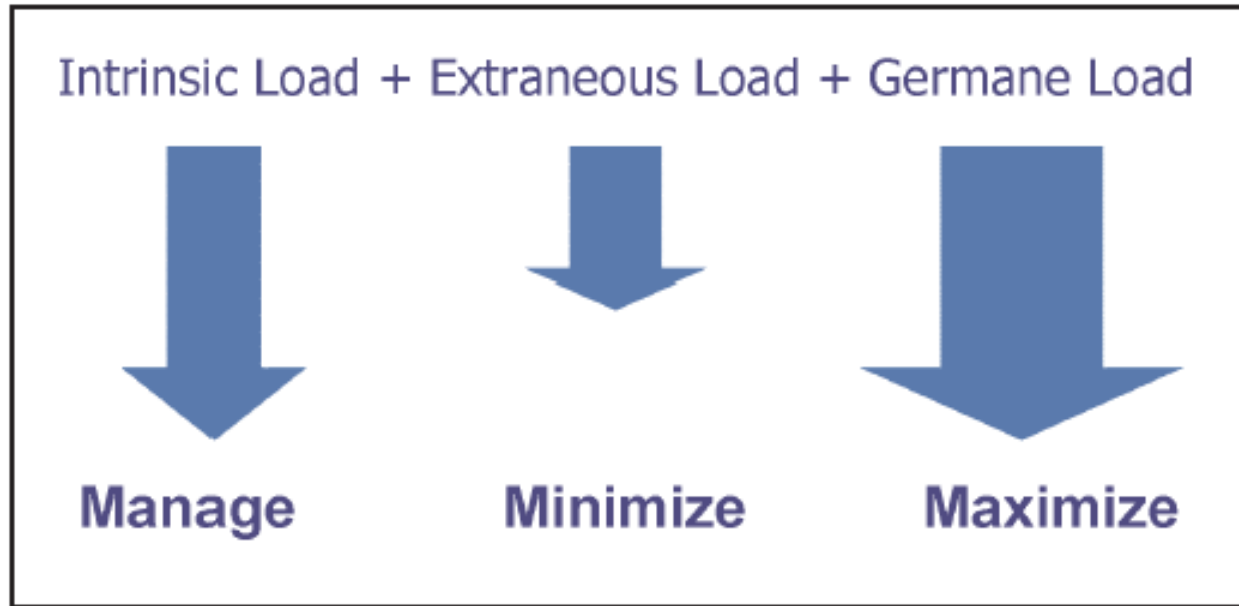
Baddeley's Model of Working Memory (Baddeley, 2000)



Cognitive Theory of Multimedia Learning (Mayer, 2001)

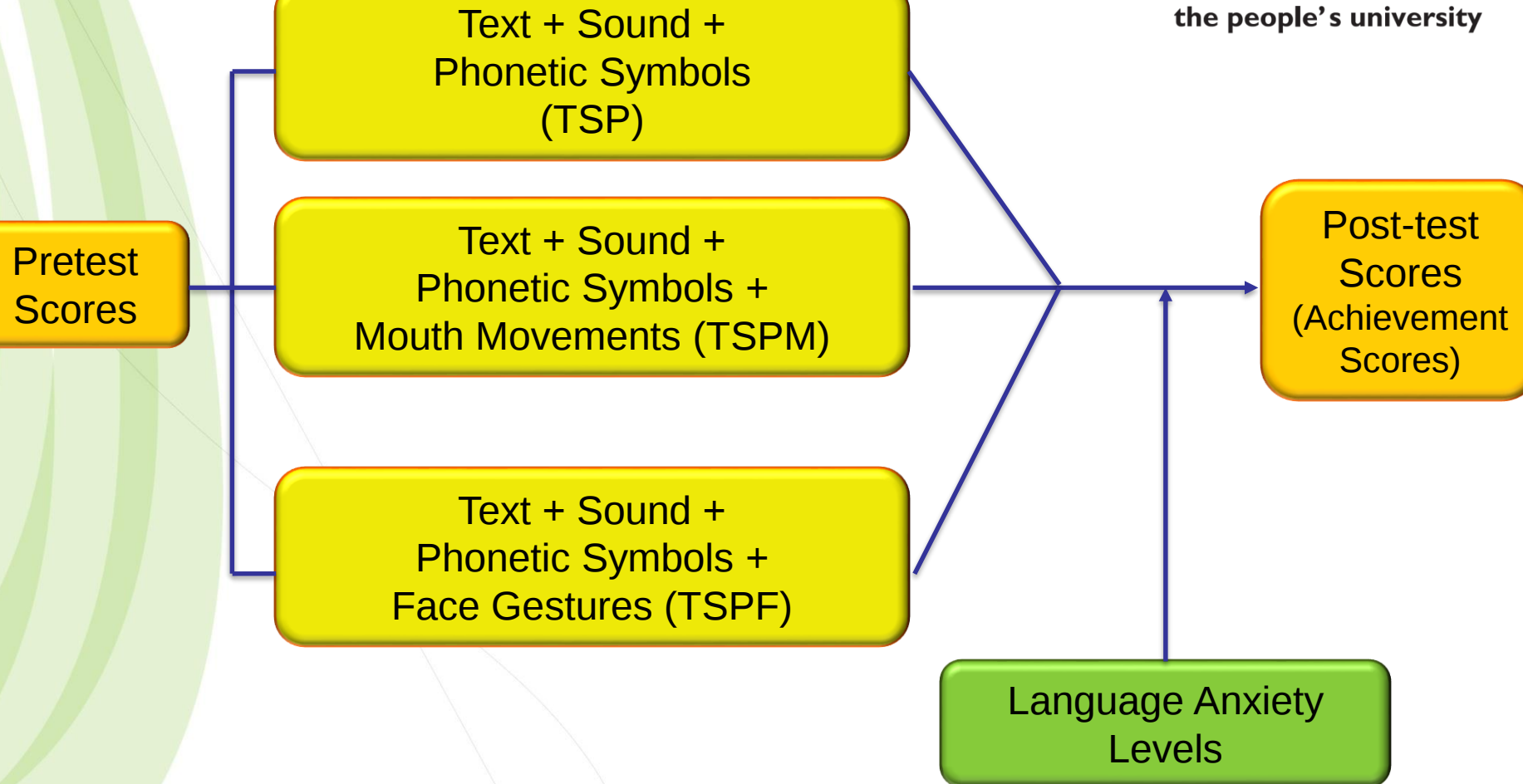


Cognitive Load Theory (Sweller, 1999)



construction of schemas - move from simple concepts to complex concepts

Conceptual Framework



Research Objectives

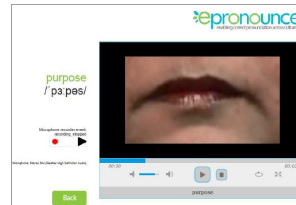
- To design and develop epronounce™ - the multimedia pronunciation learning management system incorporating phonetic symbols, mouth movements and face gestures.
- To determine whether there is any significant difference in achievement scores among students with different levels of language anxiety in using TSP, TSPM, and TSPF modes

Research Questions

By using epronounce™, will the students with different levels of language anxiety attain significantly different achievement scores in the three presentation modes?

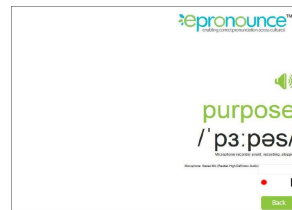
- Will students with high language anxiety (HL) using the Text + Sound + Phonetic Symbols + Face Gestures (TSPF) mode attain significantly higher achievement scores (AS) than students with high language anxiety (HL) using the Text + Sound + Phonetic Symbols + Mouth Movements (TSPM) mode?

$$AS_{HL-TSPF} > AS_{HL-TSPM}$$



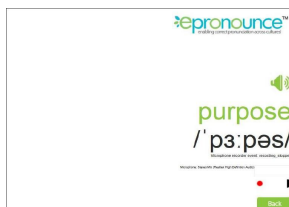
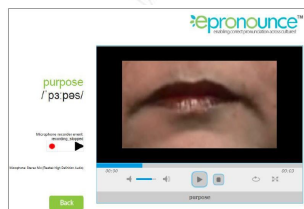
- Will students with high language anxiety (HL) using the Text + Sound + Phonetic Symbols + Face Gestures (TSPF) mode attain significantly higher achievement scores (AS) than students with high language anxiety (HL) using the Text + Sound + Phonetic Symbols (TSP) mode?

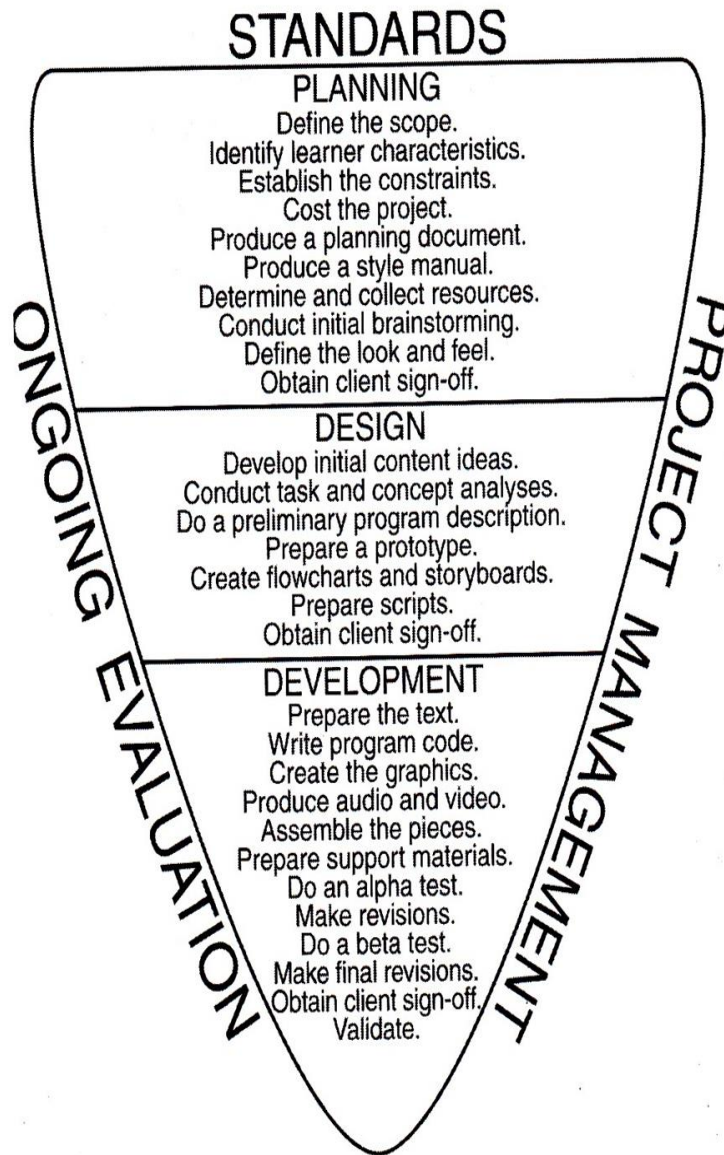
$$AS_{HL-TSPF} > AS_{HL-TSP}$$



- Will students with high language anxiety (HL) using the Text + Sound + Phonetic Symbols + Mouth Movements (TSPM) mode attain significantly higher achievement scores (AS) than students with high language anxiety (HL) using the Text + Sound + Phonetic Symbols (TSP) mode?

$AS_{HL-TSPM} > AS_{HL-TSP}$







enabling correct pronunciation across cultures!

OUR OBJECTIVES



Before you start
click here

- To design and develop an interactive multimedia student-centred learning management system to support learning of correct pronunciation among non-native English speakers
- To understand and apply phonetic symbols for correct pronunciation
- To support students and teachers as complementary learning aids
- To enable student-centred learning which is in line with the Malaysia Education Blueprint 2013-2025
- To be adopted and adapted in non-native English speaking countries, such as Thailand, Myanmar, Indonesia, China, Taiwan, Japan, and more
- To provide a Free Quality Service to the education community from the School of Educational Studies, Universiti Sains Malaysia

SIDE MENU

Home
Introduction
Background
Unique Features
Pretest
List Of Phonetic Symbols
Unit 1.0: Consonants
Unit 1.1: MCQ
Unit 1.2: True or False
Unit 1.3: Matching
Unit 2.0: Vowels
Unit 2.1: MCQ
Unit 2.2: True or False
Unit 2.3: Matching
Unit 3.0: Diphthongs
Unit 3.1: MCQ
Unit 3.2: True or False
Unit 3.3: Matching
Unit 4.0: Revision
Posttest
Contact Us

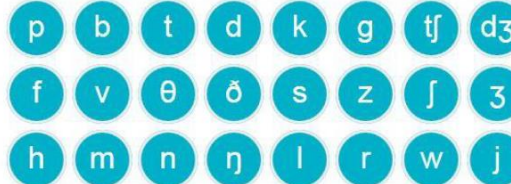
STUDENT LOGIN

Register new student, click [here](#).

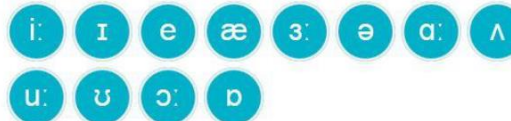
Student Email

PHONETIC SYMBOLS

CONSONANTS



VOWELS



DIPHTHONGS



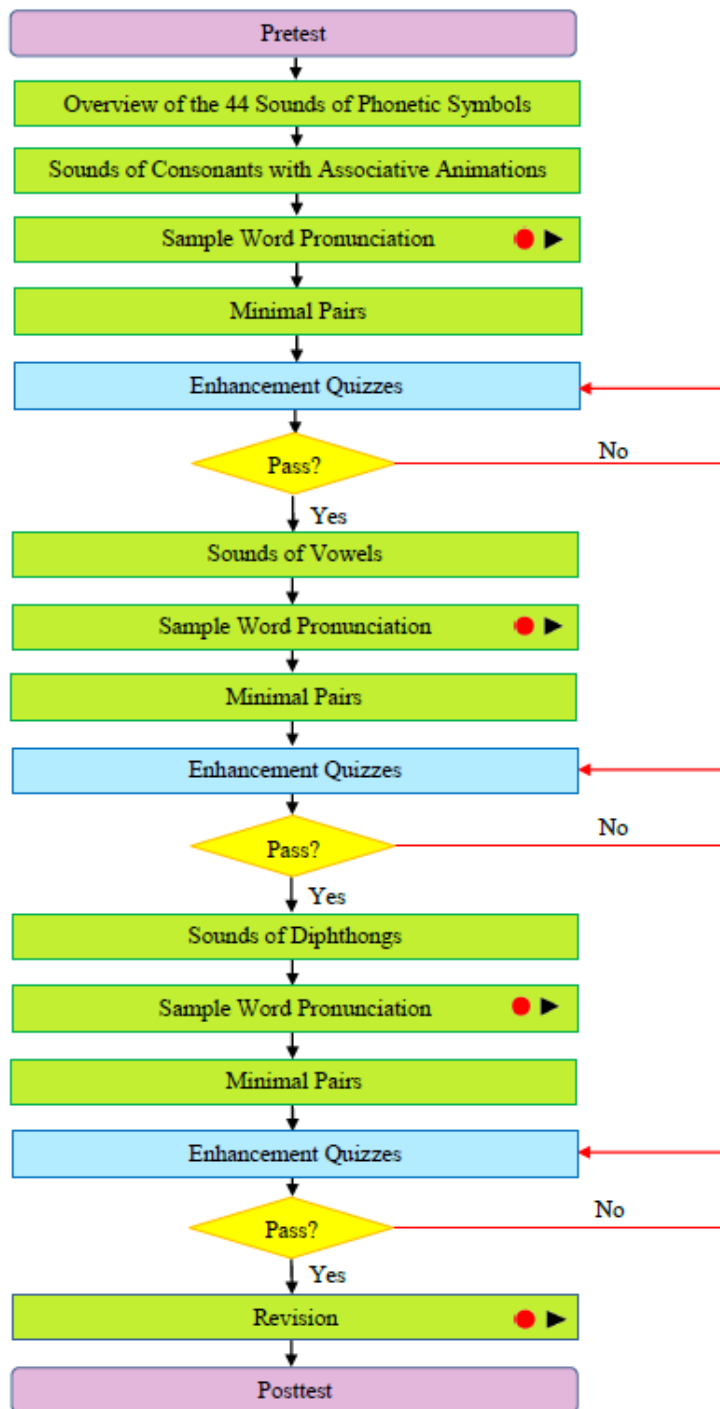
ADMINISTRATOR PANEL | Good Day, Admin!

Manage Word Pronunciation

Search Advance

Word ID	Word	Word Sound	Word Video Face	Word Video Mouth	Revise	Face	Mouth	Sound	
C01_01	pen	1_pen.mp3	1_pen_v264.mp4	1_pen_v264.mp4	☐	☐	☐	☐	Update Delete
C01_02	piece	1_piece.mp3	1_piece_v264.mp4	1_piece_v264.mp4	☐	☐	☐	☐	Update Delete
C01_03	put	1_put.mp3	1_put_v264.mp4	1_put_v264.mp4	☐	☐	☐	☐	Update Delete
C01_04	paw	1_paw.mp3	1_paw_v264.mp4	1_paw_v264.mp4	☐	☐	☐	☐	Update Delete
C01_05	people	1_people.mp3	1_people_v264.mp4	1_people_v264.mp4	☐	☐	☐	☐	Update Delete
C01_06	depot	1_depot.mp3	1_depot_v264.mp4	1_depot_v264.mp4	☐	☐	☐	☐	Update Delete
C01_07	purpose	1_purpose.mp3	1_purpose_v264.mp4	1_purpose_v264.mp4	☐	☐	☐	☐	Update Delete
C01_08	laptop	1_laptop.mp3	1_laptop_v264.mp4	1_laptop_v264_m.mp4	☐	☐	☐	☐	Update Delete
C01_09	capture	1_capture_a.mp3	1_capture_a_v264.mp4	1_capture_v264.mp4	☐	☐	☐	☐	Update Delete
C01_10	napkin	1_napkin.mp3	1_napkin_v264.mp4	1_napkin_v264.mp4	☐	☐	☐	☐	Update Delete
C01_11	reptile	1_reptile.mp3	1_reptile_v264.mp4	1_reptile_v264.mp4	☐	☐	☐	☐	Update Delete
C01_12	shop	1_shop.mp3	1_shop_v264.mp4	1_shop_v264.mp4	☐	☐	☐	☐	Update Delete

Prev | 1 | 2 | 3 | 4 | 5 | 6 | 7 | 8 | 9 | 10 | 11 | 12 | 13 | 14 | 15 | 16 | 17 | 18 | 19 | 20 | 21 | 22 | 23 | 24 | 25 | 26 | 27 | 28 | 29 | 30 | 31 | 32 | 33 | 34 | Next



UNIT 1.1: MCQ

Result: **4/5**

Listen to the pronunciation and identify the correct phonetic word.

A01		* dark /dɑ:k/	o dug /dʌg/	o tug /tʌg/	✓
A02		o say /seɪ/	* sale /seɪl/	o sad /sæd/	✗
A03		o ju /dʒu:/	o zoo /zu:/	* zoom /zu:m/	✓
A04		o deport /dɪˈpɔ:t/	* depot /ˈdeɪpəʊ/	o deco /ˈdeɪkəʊ/	✓
A05		o salad /ˈsæləd/	o shale /ʃeɪl/	* chalet /ˈʃæleɪ/	✓

UNIT 1.2: True or False

Result: **3/5**

Are the consonant sounds in the syllable initial the same? Click 'Yes' if they are same; click 'No' if they are different.

A01	zip /zɪp/		zap /zæp/		* Yes o No	✓
A02	fail /feɪl/		pail /peɪl/		* Yes o No	✗
A03	hole /həʊl/		heel /hi:l/		o Yes * No	✗
A04	neat /ni:t/		knit /nɪt/		* Yes o No	✓
A05	shun /ʃʌn/		sun /sʌn/		o Yes ☒ No	✓

UNIT 1.3: Matching

Result: **5/5**

Match the phonetic transcriptions with the words by putting the correct number.

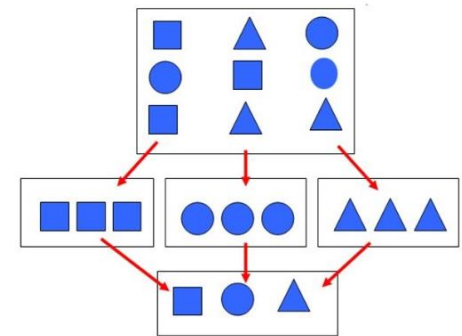
A01	thermos	/tʃaɪm/	A02	✓
A02	chime	/ˈkʌltʃə(r)/	A05	✓
A03	pirate	/jot/	A04	✓
A04	yacht	/ˈpaɪrət/	A03	✓
A05	culture	/ˈbʌsməs/	A01	✓

Samples and Sampling

- Conducted on 373 Primary Five students (aged 11), but 44 students from the overall number did not manage to complete the experiment and tests required in the study
- Final total sample size calculated for analysis : 329 students
- Stratified random sampling was employed
- Students were randomly assigned to one of the three presentation modes (TSP, TSPM and TSPF)

Instruments

- Pronunciation Competence Test (Pretest and Posttest)
- Foreign Language Classroom Anxiety Scale (FLCAS)



Results and Discussion

Table 1

Two-Way ANCOVA for Posttest Scores by Presentation Mode and Language Anxiety Level with Pretest as Covariate

Dependent Variable: Post-test

Source	Type III	Df	Mean Square	F	Sig.	Partial Eta	Observed
	Sum of Squares					Squared	Power ^b
Corrected Model	16731.667 ^a	9	1859.074	31.944	.000	.474	1.000
Intercept	13228.556	1	13228.556	227.301	.000	.416	1.000
Pretest	13010.026	1	13010.026	223.546	.000	.412	1.000
FLCAS	7715.003	2	3857.502	66.282	.000	.294	1.000
Mode	461.973	2	230.987	3.969	.020	.024	.710
FLCAS * Mode	293.571	4	73.393	1.261	.285	.016	.394
Error	18565.312	319	58.198				
Total	1392755.000	329					
Corrected Total	35296.979	328					

a. *R Squared* = .474 (*Adjusted R Squared* = .459)

b. *Computed using alpha* = .05

- Hypothesised that the learners with different levels of language anxiety will attain significantly different achievement scores in the three presentation modes
- No significant interaction effects between language anxiety levels and presentation modes of epronounce™
- Indicates that the learners of different language anxiety levels do not respond differently to epronounce™
- Their performance is at par with each other though the low and high language anxiety learners were initially expected not to perform as good as the medium language anxiety learners.
- Seemingly epronounce™ is able to bring the low and high language anxiety learners to medium language anxiety level for optimal learning under optimal learning condition as explained in the curvilinear relationship between anxiety and performance as well as the Yerkes-Dodson law (Keeley, Zayac, & Correia, 2008; Yerkes & Dodson, 1908).

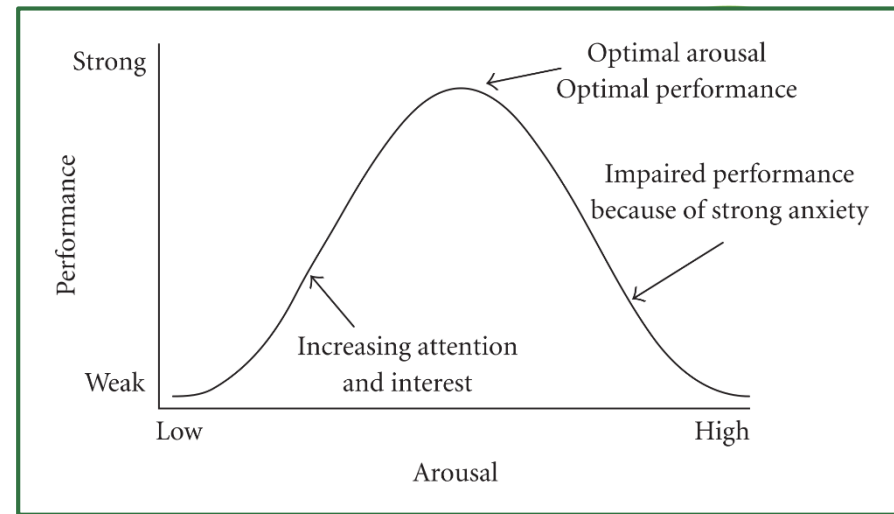


Table 2

Estimated Marginal Means by Language Anxiety Level and Presentation Mode

Dependent Variable: Post-test

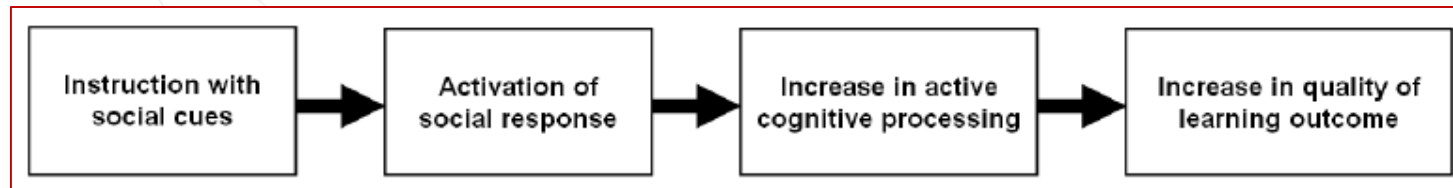
Language Anxiety Level	Presentation Mode	Mean	Std. Error	95% Confidence Interval	
				Lower Bound	Upper Bound
Low	TSP	58.832 ^a	1.805	55.280	62.384
	TSPM	60.712 ^a	1.923	56.928	64.496
	TSPF	61.454 ^a	1.688	58.134	64.775
Medium	TSP	64.820 ^a	.881	63.087	66.554
	TSPM	66.672 ^a	.852	64.996	68.348
	TSPF	71.007 ^a	.930	69.177	72.837
High	TSP	53.528 ^a	2.345	48.914	58.142
	TSPM	50.522 ^a	2.049	46.491	54.552
	TSPF	54.358 ^a	1.680	51.052	57.663

a. Covariates appearing in the model are evaluated at the following values: Pretest = 44.48.

- TSPF mode being the most effective which highlighted that learners **interacting** with a human agent in online learning environment increased in the mean achievement scores.
- Face gestures add vital information about the **intensity of the content**. When the narrator's **face gestures is integrated** into the learning process to pronounce the words, the **manner of articulation** presented by the face gestures help students to understand deeply.
- The face gestures show exactly which part of the facial movement should be moved. Therefore, students are able to follow exactly the pronunciation of words.



- The **social agency theory** that linked to Cognitive Theory of Multimedia Learning suggested people apply **social rules to media**, which in turn positively influences learning (Mayer et al., 2003).
- By using **social cues** such as face gestures, the narrator is able to **attract the students' attention** and thus help them to combine **verbal and non-verbal information** which also corresponds to Baddeley's Model of Working Memory (Baddeley, 2000).
- The **visual social cues** of the narrator, such as facial expressions, gestures, and gaze, **engage students in human-computer interaction** as a simulation of authentic human-to-human interaction.
- Social agency theory stipulated that the characteristics of a narrator prompt the **social engagement** of the students, thus allowing the students to **form a simulated human bond with the narrator** (Atkinson et al., 2005).



- The learners also feel more at ease with TSPF mode because it is a **forgiving and patient** tutor of willingly repeating the pronunciation for the learners ad infinitum with **reliable quality** in the sense of being the same every time (Pennington, 1999).

- Contrary, in traditional formal class setting, the learners experience fear when attempting to ask the **human teachers** to repeat the sounds many times because teachers may become **impatient** and other learners may also **get irritated**.
- The language anxiety of the learners is addressed as the learners get **more chance to immerse** themselves in a language learning environment **without much fear** and their pronunciation competence is enhanced.
- By increasing the **frequency of listening and practising**, the learners are **more engaged in sound discrimination and sound production** during the **information processing procedures**.

Conclusion

- The results of this study indicated epronounce™ is seemingly able to bring the learners to **medium language anxiety level** by engaging them digitally and hence optimising pronunciation learning.
- This is in line with the **Affective Filter principle of Krashen's Second Language Acquisition Theory**. Krashen claimed that the best language acquisition takes place in an environment where the **affective filter is low**.
- A low filter is associated with **relaxation, confidence to take risks** and a **pleasant learning environment**, as created by epronounce™ in this study. Krashen showed that learners who are **highly motivated** are much more likely to be successful language acquirers.
- With the innovative use of text, sound, phonetic symbols, face gestures, and interaction, the learners are more attracted to learning and therefore pay more attention to pronunciation learning. The various inputs increase learners' interest and motivation, and help establish connections between the abstract and the concrete (Boyd & Murphrey, 2002; Wald, 2008).

THANK YOU

CALL FOR PAPERS

THE 6th INTERNATIONAL SEARCH CONFERENCE 2019

27-28 June 2019

Taylor's University, Lakeside Campus, Kuala Lumpur, Malaysia

Theme: "New Media and Digital Inclusion: Embracing the 4th Industrial Revolution"

The First Industrial Revolution used water and steam power to mechanize production. The Second used electric power to create mass production. The Third used electronics and information technology to automate production. Now a Fourth Industrial Revolution is building on the Third, and more importantly, it is characterized by a fusion of technologies that is blurring the lines between the physical, digital and biological spheres.

The Fourth Industrial Revolution is evolving at an exponential rather than a linear pace, it is fundamentally changing the way we live, work and relate to one another. In order to optimize and capture the opportunities offered by new media and technologies, digital inclusion is an essential and timely topic for assessing and addressing the new revolution. The conference theme is divided into the following sub-themes but not necessarily limited to:

- *New media, digital culture and digital literacy*
- *Smart digital nation, cities, communities or generation*
- *Digital media production, distribution and consumption*
- *Digital storytelling, participatory culture, digital activism and engagement*
- *Emerging digital platforms for self-representation and identities construction*
- *Digital visuality and visibility*
- *Social connection and social networking*
- *Big data and analytics*
- *Digital transformation of media landscape, media convergence and creative industries*
- *Crowdsourcing ideas and insights*
- *Digital capabilities for public relations, advertising, marketing and branding*
- *Healthcare in digital era*
- *National and international security in digital era*
- *Policy and regulatory matters for digital accessibility, availability and affordability*
- *Privacy, consumer safety, information and network security*

Abstracts of not more than 150 words should be submitted via

[Abstract Submission Link](#)

and by email to Dr Nurzihan Hassim

search@taylors.edu.my

by **15 December 2018**.

Abstracts should include the title, research objective, methods, key findings, conclusion and maximum five keywords. Kindly include the author's title, name, affiliation, email address and a short biographical note (about 50 words in length) after the abstract.

IMPORTANT DATES

Abstract submission deadline	: 15 December 2018
Full paper submission deadline	: 28 February 2019
Early bird deadline	: 28 February 2019
Registration deadline	: 20 April 2019
Pre-conference workshop	: 26 June 2019
Postgraduate Students Colloquium	: 26 June 2019
Conference date	: 27 & 28 June 2019