



Iranian Journal of Applied Linguistics (IJAL)

Vol. 22, No. 2, September 2019, 36-70

A Persian-English Cross-Linguistic Dataset for Research on the Visual Processing of Cognates and Noncognates

Zahra Fotovatnia*, Jeffery A. Jones,

*Center for Cognitive Neuroscience and Psychology Department, Wilfrid
Laurier University, Waterloo, Canada*

Nichole E. Scheerer

University of Western Ontario, London, Canada

Abstract

Finding out which lexico-semantic features of cognates are critical in crosslanguage studies and comparing these features with noncognates helps researchers to decide which features to control in studies with cognates. Normative databases provide necessary information for this purpose. Such resources are lacking in the Persian language. We created a dataset and determined norms for the essential lexicosemantic features of 288 cognates and noncognates and matched them across conditions. Furthermore, we examined the relationship between these features and the response time (RT) and accuracy of responses in a masked-priming lexical decision task. This task was performed in English by Persian-English speakers in conditions where the prime and target words were related or unrelated in terms of meaning and/or form. Overall, familiarity with English words and English frequency were the best predictors of RT in related and unrelated priming conditions. Pronunciation similarity also predicted RT in the related condition for cognates, while the number of phonemes in the prime predicted RT for the unrelated condition. For both related and unrelated conditions, English frequency was the best predictor for noncognates. This bilingual dataset can be used in bilingual word processing and recognition studies of cognates and noncognates.

Keywords: Persian-English dataset; Cross-language studies, Bilingual word recognition; Cognates and noncognates; Lexical decision task; Priming

Article Information:

Received: 2 April 2019

Revised: 5 July 2019

Accepted: 17 July 2019

Corresponding author: 75 University Ave. W., Waterloo, ON N2L 3C5Email address: foto9420@mylaurier.ca; z.fotovatnia@iaun.ac.ir**1. introduction** *1.1. Background*

An important issue for language researchers is how formal and semantic properties of words in two languages are represented and recognized by bilinguals. Of particular interest is whether the words in each language are stored in a language-specific lexicon or a shared lexicon, and whether a relationship between phonology, orthography, and meaning (the three main features of a word) in the first (L1) and second (L2) language exists. Manipulation of phonological and semantic word features in languages with similar scripts (for example, French and English) and phonological, semantic, and orthographic word features in languages with different scripts (for example, Persian and English) provide an ideal situation for investigating these topics.

In such studies, cognates, and noncognates play an important role and need to be matched on important lexico-semantic features such as frequency, length, phonological neighborhood density, orthographic neighborhood density, imageability, concreteness, and familiarity. Cognates are translation equivalents that have an overlap between formal and semantic features across two languages (Kondrak, Marcu & Knight, 2003). Noncognates, on the other hand, are translation equivalents that display semantic but not a formal overlap. Furthermore, different degrees of overlap in terms of phonology and/or orthography result in two kinds of cognates, based on a maximum to minimum formal overlap, identical and close, respectively (Bultena, Dijkstra & van Hell, 2014). Despite the importance of such lexico-semantic features in bilingual language processing research, few studies have collected bilingual measures and made them available in Persian as the native, and English as the nonnative, language.

Research studies have investigated bilingual word representation and processing for languages with the same- (Grainger & Frenck-Mestre, 1998; Van Hell & Dijkstra, 2002) or different-scripts (Gollan, Forster & Frost, 1997). As most of these studies are conducted to compare bilinguals

with English as an L2, or less frequently as an L1, and another similar script language such as French, German, Spanish, and Dutch, online databases are available to use as resources for measuring the essential lexical and semantic features of words and creating nonwords to be used as stimuli. Examples are the MRC Psycholinguistic Database (http://websites.psychology.uwa.edu.au/school/MRCDatabase/uwa_mrc.htm) and the CLEARPOND Database (Marian, Bartolotti, Chabal, Shook, 2012) in English. In a relatively few studies, different-script languages such as Hebrew (Gollan, et al., 1997), Japanese (Nakayama, Verdonschot, Sears & Lupker, 2014), and Chinese (Zhou, Chen, Yang & Dunlap, 2010) were used. Few attempts have been made to pursue these topics with Persian-English bilinguals (Fotovatnia & Taleb, 2012, 2013). This issue could explain the lack of resources available for determining the lexicosemantic features of words in Persian. Given the importance of this type of research for determining how bilinguals represent and process words in various languages, we have created a dataset for the bilingual PersianEnglish word research and teaching. Investigating bilinguals with languages that have been less studied should increase the generalizability of the findings on the processing and representation of words in a bilingual brain.

1.2. Lexico-semantic features

A review of the literature on visual word recognition research shows that researchers control variables such as word frequency, word length, age of acquisition (AoA), neighborhood density, familiarity, imageability, and concreteness (Balota, Cortese, Sergent-Marshall, Spieler & Yap, 2004) for their potentially contaminating effects on the main variables of the study. These effects differ based on the specific requirements of the task, the language used for the task, the orthographic depth of the L1 and L2, and the participants' language proficiency in either language (De Groot, Borgwaldt, Bos & Van Den Eijnden, 2002). The existence of an interaction between the languages the participants know and the type of task researchers use to collect data justify further research on the topic. Few studies have focused on Persian-English speakers to find out whether similar lexico-semantic features of the stimuli should be controlled in cross-language studies. Creating a dataset and establishing the normative bases of the stimuli will provide a foundation for any research on word

representations and word processing of Persian-English speakers, a neglected population.

Word frequency is defined as the number of times a word appears in a corpus. In general, low-frequency words are processed more slowly than high-frequency words in both L1 (Brysbaert, Lagrou & Stevens, 2017) and L2 (Peeters, Dijkstra & Grainger, 2013), although the slower processing is more crucial in the latter. This pattern is the result of the level of exposure to words in a language (the lexical entrenchment hypothesis, Diependaele, Lemhöfer & Brysbaert, 2013). The more one is exposed to words, the richer their vocabulary knowledge becomes, and the more proficient that person will be in that language.

Words differ in the number of letters or phonemes they include. These two variables are highly related. However, each feature seems to interact with the task type, task language, and word frequency in either L1 or L2. For example, the number of phonemes affected reading latencies in Tunisian Arabic (Boukadi, Zouaidi & Wilson, 2016). However, when both the number of phonemes and letters were entered into a regression (other variables controlled), orthographic length remained a better predictor of word-naming performance. Word length and the task language showed an interaction in a study where the number of letters affected the performance of Dutch-English speakers in lexical decision and word naming in English, but not word naming in Dutch. Furthermore, an interaction between the orthographic length and the frequency of words was reported (Bakhtiar and Weekes, 2015). Bakhtiar and Weekes found that orthographic length was a better predictor of word naming performance for low-frequency words than for highfrequency words. Overall, investigating the interactions between word length in terms of the number of phonemes and letters, and task type with different language speakers should further contribute to the existing literature and help to control the lexico-semantic features of the stimuli required to collect data.

AoA is defined as the age at which a word is first learned in L1. Research shows that words learned early in life are processed and remembered more efficiently than words learned later (Brysbaert et al. 2017). AoA is not a strong predictor of lexical processing performance in L2, due to learning the language at different ages in life, and mainly after mastering L1. Furthermore, AoA measures correlate with a range of other

word features such as concreteness ($r = -.50$), imageability ($r = .72$), rated familiarity ($r = -.72$), as shown by Gilhooly and Logie (1980), and word frequency (Juhasz, 2005). Therefore, it seems reasonable to use these measures instead of AoA.

A word's orthographic neighbors are defined as a set of words that exists in L1 or L2 and that differ from the target word in one letter position. Neighbors can be created by changing one letter of the word while preserving letter positions. For example, the words *pike*, *pine*, *pole*, and *tile* are all orthographic neighbors of the word *pile* (Coltheart, Davelaar, Jonasson & Besner, 1977). Theoretically, a word with a large number of orthographic neighbors would activate a large search set (search models of word recognition, Forster, 1976), or would create more within-level inhibition resulting from greater orthographic overlap with its neighbors (interactive activation models of word recognition, McClelland & Rumelhart, 1981). Therefore, subjects should take more time and have more difficulty searching through the set, or suppressing the inhibition to perform a task. The empirical results, however, are not straightforward. For example, Grainger and Jacobs (1996) reported inhibitory influences of neighborhood frequency i.e., a word with higher frequency neighbors produced slower lexical decision times. Other studies reported no main effect of neighborhood density, but an interaction between the neighborhood density and word frequency (Balota et al., 2004, Sears, Hino, & Lupker, 1995). Neighborhood density facilitated the processing of low-frequency words, but inhibited that of high-frequency words. Familiarity is a subjective measure based on the number of times individuals have experienced a word. Therefore, it is highly related to culturally specific experiences, which vary from one language community to another (Boukadi et al., 2016). This is called subjective frequency and is highly related to the objective frequency, which reflects the number of times a word occurs in a language corpus. However, these two measures are independent of each other (Connine, Mullennix, Shernoff, & Yelen, 1990; Kreuz, 1987). The reason seems to be that familiarity affects the level of semantic activation, whereas frequency affects the level of phonological encoding in word naming and lexical decision tasks (Boukadi et al., 2016). Connine et al. (1990) reported faster reaction times for high-familiarity words in visual and auditory lexical decision tasks, as well as for word naming. Familiarity effects are also observed for pictures in picture naming studies, with faster naming for familiar objects and slower naming for

uncommon objects (Cuetos, Ellis & Alvarez, 1999). Familiarity is considered an important possible predictor of naming latencies when conducting these studies (Boukadi et al., 2016). Subjective frequency estimates were found to be a better predictor of object frequency counts in some visual and auditory word processing studies (Connine et al., 1990).

Imageability and concreteness are examples of the semantic variables of a word. Imageability is defined as the ease with which a word arouses a sensory mental image of something, while concreteness involves the extent to which a word is experienced by the senses. Imageability is confounded with concreteness in the sense that words high in imageability are more concrete than words low in imageability (the so-called abstract words). However, these two features seem to exploit partially different components. Imageability is related to the number of semantic features that develop a concept. Consequently, the concepts of highly imageability words are connected to many more semantic features than the concepts of low-imageability words (Plaut & Shallice, 1993). An imageability rating is based on the graded amounts of sensory (mainly visual) information associated with words. The concreteness rating is spatiotemporally based, concrete words are more spatiotemporally based than abstract words (Kousta, Vigliocco, Vinson, Andrews & Del Campo, 2011), meaning that they have spatial and temporal qualities that are absent in abstract words.

Most studies have shown faster responses to concrete words than abstract words (Kanske & Kotz, 2007; Kounios & Holcomb, 1994). The reason is that concrete words are represented in a verbal as well as a nonverbal code, while abstract words are represented only in a verbal code (dual-coding theory, Paivio, 1986). This means that concrete words activate verbal and image-based systems through referential connections to these systems, while abstract words activate representations in the verbal or linguistic semantic system. Furthermore, concrete words have stronger and denser associative links than abstract words (the contextavailability hypothesis, Schwanenflugel, 1991).

To further determine which lexico-semantic features of the stimuli to focus on for visual speech processing in lexical decision and word naming tasks, De Groot et al. (2002) examined the effects of 18 variables on Dutch (L1) and English (L2) lexical decision and word naming performance of 3, 4, and 5 letter words (Table 1).

Table 1. *Lexico-semantic variables affecting lexical decision and word naming in Dutch (L1) and English (L2)*

Variables	Dutch English			
	LD	WN	LD	WN
Frequency	major	+	Major	+
Orthographic length		+	+	+
Semantic variables	+		+	+
Onset variables		+		+
Neighborhood words		+		+
Cognate effect		+	+	+
Reaction time	Flow	Fast	Long	Long

As Table 1 shows, frequency mainly affected Dutch and English word recognition. The effect of frequency, however, was found to be exaggerated in the lexical decision and word naming tasks, as frequency showed interactions with familiarity in the former, and with the length of words in the latter task, in English. Considering the RT, lexical decision in Dutch was slower than word naming in the same language but took the same amount of time as word naming in English. De Groot et al. (2002) used the differences in the grapheme-to-phoneme relationship in each language to interpret their findings, which is consistent with the orthographic depth hypothesis (Katz & Frost, 1992) and the dual-route model of reading (Coltheart, Rastle, Perry, Langdon & Ziegler, 2001). Dutch is a shallower language than English. Therefore, word access uses the more effective indirect grapheme-to-phoneme route, and not the direct lexical route.

1.3. This study

The purpose of the current study is twofold. On the one hand, we created a dataset including cognates and noncognates in Persian and English and determined the lexico-semantic features that have been shown to be important in priming studies. This is an important step because no such databases are available for cross-language studies including these two

languages. On the other hand, we investigated the relationship between the lexico-semantic features of Persian primes and English targets with RT and accuracy of responses in a cross-linguistic masked lexical decision task on cognates and noncognates. The prime was masked to prevent it from reaching conscious attention (Forster & Davis, 1984). To our knowledge, no such study has focused on cognates and noncognates to determine which features of the stimuli are relevant to the RT and accuracy of responses and which features could best predict these measures in a masked-priming lexical decision task. Learning about such a relationship could prevent the confounding effects of the main variables in crosslanguage studies.

2. Method

2.1. Participants

Three groups of Persian-English speakers were recruited for this study. Two groups were involved in determining the lexico-semantic features of the words, while the third group participated in a lexical decision task. The first group completed questions about the phonological similarity of cognates in Persian and English, and familiarity with the English words. The second group included students, either graduated or studying, in the English program at the Islamic Azad University, Najafabad Branch, Iran. This group completed questions about the concreteness and imageability of the Persian words. Table 1 provides detailed information about these two groups. The third group included 32 Persian-English speakers residing in Waterloo and London, Canada, either graduated or studying at university. They had received a minimum of eight years of formal English instruction in Iran before immigration to Canada, and they all obtained a score of 5.5 - 8 ($M=6.96$, $SD=.69$) on the IELTS Academic module. Furthermore, six more proficient speakers different from those who received the questionnaires translated the words from Persian to English, and vice versa.

Table 2. Detailed information about group 1 and group 2

First Group											
Gender			Age range		Degree			Proficiency			total
			20-39	40-50	BA	MA	PhD	LI	UI	Adv.	55
N	16	38	48	5	7	34	11	2	19	30	
Second Group											

M	F	20-39	40-50	BA	MA	PhD	LI	UI	Adv.	Total
29	41	3	4	36	0	6	17	18	46	N
										1

2.2. Materials

The stimuli included only nouns, as some studies have shown larger cognate effects for nouns than verbs (Bultena et al., 2014). The stimuli were prepared following the steps below. Appendix 1 includes word targets and their matched related and unrelated primes.

1. A random list (N=150) of cognates (e.g., taxi, star) and noncognates (e.g., sparrow, ring) was created in English to use as targets in the lexical decision task in the L1 (prime) - L2 (target) direction. The words were then translated into Persian to use as related primes. Another list was created to use as unrelated primes. The related and unrelated primes were matched for the number of phonemes and letters in each word. Different from most studies, where the number of letters was primarily used for matching the related to unrelated primes, we used the number of letters and phonemes due to the specific characteristics of the Persian script. In Persian, only consonants are represented by letters in the written form. This selection allowed us to see whether word length, defined in terms of phonemes versus letters, would have a different effect. Cognate primes (e.g., تیفس /ti:'fu:s/) and their corresponding targets (e.g., typhus) shared semantic and phonological similarity, while noncognate primes and their corresponding targets shared only semantic features, (e.g., خطا /khæ'ta/, meaning error). Unrelated primes did not have any phonological or semantic relationship with the targets. Related and unrelated prime words were utterly tried to be selected from the same (e.g., living/nonliving things, animals, food, events).
2. To ensure that the translation equivalents in both languages had the same meaning, the words were presented on two random lists. The lists were given to six proficient Persian-English speakers to translate from Persian into English and vice versa. Only the words that were translated

similarly in both directions by all people were included in the list, and others were removed. For example, the word *scholar* was replaced with the word *researcher*, as the former was translated as محقق/mŪhæ'ghəgh / in Persian, but the same Persian word (محقق) was translated as *researcher* in English. Similarly, the word مسافر/mŪsa'fer/ was not included, because it was translated as *traveler* by one translator, and as *passenger* by another translator. Likewise, the word زمین/zæ'min/ was not included, because it was translated as *earth* by one, and *land* by another translator. Related primes and targets had to be translation equivalents. In fact, each prime had to activate the corresponding target and no other word. Furthermore, unrelated primes and their corresponding targets did not share any onset phoneme or letter.

3. The number of letters and phonemes, word frequency, and orthographic neighborhood density of the English stimuli were determined using the CLEARPOND database (Marian et al., Shook, 2012).
4. Nonwords were generated using the English Project Website (at alexicon.wustl.edu) and were matched with the corresponding words for the number of letters and neighbourhood density (n=230).
5. Frequency of the Persian primes was manually determined using the MAHAK (means “measure” in English) corpus (Sheykh Esmaili, et al., 2007). MAHAK is the largest Farsi test collection containing 3007 documents and 216 queries on various topics.
5. Cognate phonological similarity, familiarity with English words, and concreteness and imageability of the Persian words were determined through two 5-point Likert scale questionnaires (<http://www.qualtrics.com>). The first questionnaire included two sections: (a) 125 English and Persian cognate pairs selected in the previous stage and a few noncognates to use as fillers (n=17), and (b) all cognates and noncognates (N=250). The former section elicited the degree of similarity in the pronunciation of cognates and the latter asked for the level of familiarity with the English words. Each cognate in Persian was presented with its English equivalent and participants were asked to pronounce both words, determine how similar in the pronunciation they were, and then select one point along a 5-point Likert scale ranging from *completely*

different to completely similar. For familiarity, the question was “how frequently do you encounter these words in listening, reading, speaking, and writing?” (Bakhtiar & Weekes, 2015). The scales ranged from *never* (completely unfamiliar) to *daily* (completely familiar). The second questionnaire included only Persian words presented in two sections to be rated for concreteness and imageability. For concreteness, participants selected one point on the scale *completely abstract* to *completely concrete*. For imageability, participants determined how easily a word provoked a mental image in the form of a picture, sound, taste, or smell (*very difficult* to *very easy*). Instructions and examples were provided in Persian to ensure that participants would understand the task.

6. The neighborhood density of Persian words was determined based on the MAHAK corpus using calculations made in Microsoft Office Excel 2013 (Bakhtiar & Weeks, 2015).

7. Forty-eight words with similar features to the main stimuli were used as fillers in the lexical decision task.

2.3. Procedure

Participants received the invitation to the survey, a URL to access the questionnaire, and general instructions via e-mail. After they had access to the questionnaire, participants read specific instructions for the first feature to rate. Each page included 25 words, and each word was immediately followed by a 5-point scale. After they rated all the words for one feature, they received instructions for the next feature rating. The order of features and words in each list varied randomly for each participant. All participants signed an online consent form to participate in the study, which was approved by the Wilfrid Laurier University research ethics board (approval number, 4585).

To collect the RT and accuracy of responses to the stimuli, participants in the third group were seated in front of a computer and pressed different keys for words and nonwords they saw on the screen as quickly and as accurately as possible. The stimuli ($n=144$ target words, $n=230$ nonwords) were presented in black at the center of a white background on a 16-inch screen, through two lists counterbalanced across participants, using STIM2 software (Compumedics, NeuroScan, Charlotte, NC) in the following order: a fixation sign (+) for 500 ms, number signs to cover the prime (#####) for 500 ms, a Persian prime in 14 pt Nazanin font for 50

ms, and an English target word/nonword in 16pt New Times Roman font in lower case letters, which remained on the screen until a response was recorded. Each participant performed a 30item practice block that was similar to the main task. The experiment was conducted in the Center for Cognitive Neuroscience at Wilfrid Laurier University.

3. Results

Only the questionnaires that were more than 75% filled out were analyzed. The words with no response were treated as missing data. Then the mean of each feature was calculated by averaging the ratings across all participants. The data from the lexical decision task was analyzed by deleting incorrect answers and keeping only the RTs that were between 300-1800 ms and within 2 standard deviations of the mean. The data of one participant was removed due to having more than 25% incorrect responses. Then RTs and correct responses to each word were averaged. The first section includes the analysis of all words, and the second section shows the analysis of cognates and noncognates separately.

3.1. Analysis of all words

To understand which word features were correlated with one another, and with the RT and accuracy of responses, a Pearson product-moment correlation was run. A significant relationship was observed between the number of phonemes and the number of letters, $r(288) = .79, p < .001$, concreteness and familiarity with the English words, $r(240) = .216, p = .001$, imageability and English frequency, $r(247) = .147, p = .021$, concreteness and Persian frequency, $r(249) = -.228, p < .001$, concreteness and imageability, $r(249) = .855, p < .001$, and familiarity and English frequency, $r(248) = .402, p < .001$. Correlation coefficients for RT and accuracy and the lexico-semantic features of the stimuli are shown in Table 3.

Table 3. Correlation Coefficients Showing the Relationship Between the Lexico-Semantic Features of Words and RT and Accuracy of Responses

	English phoneme	English letters	English Frequency	Familiarity with English words	English Neighborhood Density	Farsi phoneme	Farsi letters	Farsi neighborhood density	Farsi frequency	Image ability	Concreteness
RT	.360**	.416**	-.453**	-.506**	-.354**	.200**	.224**	-.235**	-.232**	-.109	.002
Sig	.000	.000	.000	.000	.000	.001	.000	.000	.000	.087	.974
Number	286	286	286	240	262	288	288	249	250	249	249
Accuracy	-.020	-.650	.257**	.512**	.195**	.003	-.066	.010	.110	.009	-.038
Sig	.739	.273	.000	.000	.001	.953	.261	.870	.083	.883	.550
Number	286	286	286	240	262	288	288	249	250	249	249

*Correlation is significant at the 0.05 level (2-tailed).

**Correlation is significant at the 0.01 level (2-tailed).

As shown in Table 3, a larger number of phonemes and letters in English and Persian words was associated with slower RTs. On the other hand, higher word frequency and word neighborhood density in English and Persian, and familiarity with English words were associated with faster responses. For accuracy, an increased English frequency, English neighborhood density and familiarity with English words was associated with a higher number of correct responses.

To find the best predictors of RT and accuracy, multiple regression analyses were run, but before the analyses, the assumptions of normality, linearity, multicollinearity, and homoscedasticity were tested. For RT, the total variance explained by the model as a whole was 46.5 %, $F(6,151) = 21.84$, $p < .001$. The strongest significant unique contribution to the model was $\beta = .42$, which was for the familiarity of English words. The next strongest contribution was $\beta = .41$, which was for the number of letters of English targets. The third-largest contribution was $\beta = -.19$, which was for English frequency. Pronunciation similarity and neighborhood density did not significantly contribute to the model.

For accuracy, the total variance explained by the model as a whole was 28.4 %, $F(6,151) = 9.98$, $p < .001$. The only significant unique

contribution to the model was $\beta = -6.1$, which was for familiarity with English words. No other variables contributed significantly to the model.

To determine how much variance in the RT and accuracy was explained by the number of phonemes, the number of letters, frequency per million, and concreteness and imageability of the Farsi primes, a multiple regression analysis was run. For RT, the total variance explained by the model as a whole was 15.3%, $F(6,236) = 7.09$, $p < .001$. The strongest significant unique contribution to the model was $\beta = -.2$, which was for frequency of Farsi primes. The next strongest contribution was $\beta = -.3$, which was for imageability of Farsi primes. The third-largest contribution was $\beta = -.2$, which was for neighborhood density. The smallest significant contribution was $\beta = -.13$, which was for the number of phonemes.

The model was not significant for accuracy. None of the features of Farsi primes contributed significantly to the model.

3.2. Analysis of cognates and noncognates

Similar steps were followed to investigate the relationship between length, frequency, familiarity, pronunciation similarity, and the RT and accuracy of related and unrelated cognates and noncognates (Table 4).

Table 4. *Correlation coefficients showing the relationship between length, frequency, familiarity, pronunciation similarity, and RT and accuracy of responses for related and unrelated cognates and noncognates*

Frequency		correlation	.464* *		.495**			.419* *	
IJAL, Vol. 22, No. 2, September 2019									
(2tailed)			.000	.017	.000	.009	.000	.018	.001
			.142	.52					
		N	90	85	58	58	85	85	58
English Familiarity	Pearson correlation		-.556	.457	-	.352**	.466**	-.604**	.583**
	Sig. (2tailed)		.000	.000	.019	.001	.000	.000	.083
		N	81	81	44	44	81	81	44
Persian Frequency	Pearson correlation		-.257*	.054	-.275*	.360**	-.123	.055	-.325*
	Sig. (2tailed)		.018	.626	.044	.008	.271	.624	.016
		N	84	84	54	54	82	82	54
Pronunciation Similarity		correlation Pearson	-.140	-	-.098	.112	-	.221*	
		Sig. (2tailed)	.029		.215	-	-	.386	.323
		N	80	80		-	80	80	-
Persian	Pearson correlation		-.068	-.082	-.040	.064	.322**	.240*	.112
			.532	.452					-.214
Persian Letter	Pearson correlation		.163	-.007	*	-.396**		.237*	-.293*
	Sig. (2tailed)		.135	.952	.005	.002	.015	.028	.200
		N	86	86	58	58	86	86	58
Phoneme		n							
		Sig. (2-tailed)	86	86	.767	.635	.003	.026	.401
		N			.362*	58	.263*	86	.171
								58	58

*Correlation is significant at the 0.05 level (2-tailed).

**Correlation is significant at the 0.01 level (2-tailed).

Table 4 illustrates the results of running Pearson product-moment correlations on the data. As shown, the length of the related and unrelated targets correlated positively with RT in all four conditions. Longer words were processed more slowly. For accuracy, the length of targets was not an important factor, as it did not correlate significantly with accuracy in almost any of the conditions. Nevertheless, the number of letters of cognates correlated negatively with accuracy in the related prime condition; longer cognates elicited more incorrect responses. More frequent English targets elicited faster responses for both related and unrelated cognates and noncognates and more correct responses in all conditions, except in the unrelated noncognate condition. Familiarity with English words was positively related to RT and accuracy in all conditions, except for RT for unrelated noncognates. Familiar English words were processed faster and more accurately. Persian frequency was negatively related to RT for cognates and noncognates and positively related to accuracy in all conditions, except for RT in the unrelated cognate and the related noncognate conditions; more frequent primes were processed faster and more accurately. Besides, pronunciation similarity was negatively related to the RT for the related cognates. Cognates that were similar in pronunciation in Persian and English were processed faster. All in all, the frequency of the prime correlated negatively with RT in almost any of the conditions, and the length of the prime correlated positively with RT in only two conditions. Also, similar features of the targets (the number of phonemes and letters, frequency, and familiarity) significantly correlated with RT in nearly any of the conditions. On the other hand, fewer features of the prime correlated with accuracy than RT.

Word features that were correlated with RT and accuracy were entered into a multiple regression analysis for each condition. For related cognates, English and Persian frequency, pronunciation similarity, and Familiarity with English words were considered. The total variance in RT explained by the model as a whole was 63.4%, $F(3,76) = 12.6, p < .001$. The strongest significant unique contribution to the model was $\beta = -.41$, which was for Familiarity with English words. The next significant contribution to the model was $\beta = -.28$, which was for English frequency. The last significant contribution to the model was $\beta = -.19$, which was for pronunciation similarity. For related cognates, the total variance of accuracy explained by the model as a whole was 46.1%, $F(2,77) = 10.36$,

$p < .001$. The strongest significant unique contribution to the model was $\beta = .43$, which was for Familiarity with English words.

For unrelated cognates, English frequency, Familiarity with English words, and the number of Persian phonemes and letters were considered. The total variance in RT explained by the model as a whole was 72%, $F(4,75) = 20.16$, $p < .001$. The strongest significant unique contribution to the model was $\beta = -.53$, which was for Familiarity with English words. The next significant contribution to the model was $\beta = .43$, which was for Persian phonemes. The last significant contribution to the model was $\beta = -.21$, which was for English frequency. For unrelated cognates, the same four variables as those for RT were entered into the model. The total variance in accuracy explained by the model as a whole was 63.9%, $F(4,75) = 12.93$, $p < .001$. The only significant unique contribution to the model was $\beta = .57$, which was for Familiarity with English words.

For related noncognates, English and Persian frequency, Familiarity with English words, and the number of Persian letters were entered into the model. The total variance of RT explained by the model as a whole was 59.3%, $F(4,39) = 5.3$, $p = .002$. The strongest significant unique contribution to the model was $\beta = -.36$, which was for English frequency. For accuracy, English and Persian frequency, and the number of Persian letters were used. The total variance in accuracy explained by the model as a whole was 61.4%, $F(4,39) = 5.9$, $p = .001$. The strongest significant unique contribution to the model was $\beta = .38$, which was for Familiarity with English words. The next significant unique contribution to the model was $\beta = -.31$, which was for the number of Persian letters.

For unrelated noncognates, English and Persian frequency were entered into the model. The total variance in RT explained by the model as a whole was 48.3%, $F(2, 51) = 7.76$, $p = .001$. The strongest significant unique contribution to the model was $\beta = -.37$, which was for English frequency. For accuracy, Familiarity with English words, and the number of Persian letters were used. The total variance in accuracy explained by the model as a whole was 52%, $F(2,41) = 7.59$, $p = .002$. The only significant unique contribution to the model was $\beta = .44$, which was for Familiarity with English words.

Comparing the semantic features of cognates with noncognates (Table 5) showed that cognate means were larger than the corresponding noncognate means for imageability, concreteness, and familiarity. Independent samples *t*-tests showed that cognates were rated as more imageable and concrete than noncognates, $t(243) = 1.989$, $p = .048$, and $t(243) = 2.95$, $p = .003$, respectively.

Table 5. *Descriptive statistics of imageability, and concreteness of Persian words and familiarity of English words*

	Type	N	Mean	Std. Deviation	STD. Error Mean
Imageability	Cognate	89	3.90	.72	.08
	Noncognate	156	3.71	.67	.05
Concreteness	Cognate	89	3.77	.75	.08
	Noncognate	156	3.45	.85	.07
Familiarity	Cognate	78	2.68	.80	.09
	Noncognate	40	2.61	.74	.12

4. Discussion

Models of visual word recognition have primarily been developed using the outcomes of experiments with English as the L1 or L2. The existence of dependencies between orthography and word-recognition procedures (e.g., Katz & Feldman, 1983; Katz & Frost, 1992) makes it appropriate to use languages with different phonological and/or orthographic features to evaluate the validity of the findings observed in previous studies. We created a dataset in Persian and investigated the relationship between word lexico-semantic features that are found critical in studies on visual word recognition in a masked-priming lexical decision task. We further investigated the relationship between these features and the RT and accuracy of responses to cognates and noncognates to determine which features are essential to control for each word type and whether to generalize the findings to Persian-English speakers, a less studied population.

Concerning the word features investigated in this study, concreteness was found to be positively correlated with familiarity with English words and imageability of Persian words, and negatively correlated with Persian frequency. The relationship between concreteness, imageability, and familiarity supports the idea that concrete items are more imageable, and these two features together create a feeling of familiarity for the L2 words. The positive relationship between concreteness and imageability supports the dual-code theory (Paivio, 1986), which attributes the concreteness effect of words to qualitative differences between concrete and abstract words; Concrete words are more imageable than abstract words. This might be an indication that when participants were evaluating their familiarity with the English words, they could not ignore word concreteness, which was measured using Persian words. The negative relationship between concreteness and Persian frequency may be the results of the Persian language frequency being calculated using the MAHAK database, which is based on written sources. The relationship might have been different if the frequency had been determined using a speech-based database. Unfortunately, there is no speech-based database available in Persian. This finding might further confirm that concreteness and frequency are two different but related components. Interestingly, cognates were rated as more imageable and concrete than noncognates. Cognates and noncognates were randomly selected in this study. However, this finding might confirm that cognates have a special status for bilinguals, as they are encountered in both languages. Engaging with cognates in either language might result in more familiarity with this type of word.

RT was related to the word length, frequency, and neighborhood density of both primes and targets, and familiarity with the English targets. However, the best predictors of RT (for the English targets) were familiarity with English words, number of letters, and frequency of English targets, while the best predictors of RT (for the Persian primes) were Persian frequency, imageability, neighborhood density, and the number of phonemes. We found that the number of phonemes and letters of the English targets were highly correlated with each other, as well as with RT, while the number of letters was a better predictor of RT. This is similar to what Boukadi et al., (2016) found in a word-naming task in Tunisian Arabic. On the other hand, the number of phonemes of the Persian primes and not the number of letters was a significant predictor of RT. This finding might be related to a specific feature of the Persian language, as the written script does not display

all vowels in Persian. It remains to investigate whether the number of phonemes could be a better predictor of RT with Persian words as targets in a lexical decision task.

Another interesting finding is the negative relationship that was found between the neighborhood density of primes and RT. Almost all studies using a priming paradigm controlled for only the neighborhood density of the targets. The results of the present study show that the neighborhood density of primes should also be controlled, as the prime ultimately influences the processing of target words. Furthermore, given the importance of the number of phonemes and not the number of letters of the Persian primes, researchers might consider controlling the phonological neighborhood density of the primes in addition to their orthographic neighborhood density. The phonological neighborhood density of primes might be a better predictor of RT in Persian. More research on this topic is necessary.

Another issue we could infer from the findings is the relationship between English frequency and familiarity with English words. The results clearly show that these two components are highly related but different, and Familiarity with English words was a more influential feature than English frequency. Familiarity is presumably a more realistic measure of frequency for L2 speakers than the mere frequency measured based on the L1 databases. For accuracy, Familiarity with English words was a critical feature. No features of the Persian primes predicted the accuracy of responses. Once more, this finding emphasizes the importance of familiarity with L2 words.

The present study further investigated the relationship between features of cognate and noncognate targets and the RT and accuracy of responses when the targets were preceded by the related and unrelated primes. RT was positively related to the length of the English targets (the number of phonemes and letters) in all four conditions, which was expected. Common to all conditions was that familiarity with English words was the best predictor of RT. The only exceptions were the situations where the related and unrelated noncognates were concerned. For this group, English frequency was the best predictor of RT, while English frequency was the second-best predictor of RT for related and unrelated cognates. However, the third significant predictor was the pronunciation similarity for the related cognates, while it was the number of Persian phonemes for the unrelated cognates. This finding shows that cognates are special words because they share not only meaning but also a

degree of formal features in both languages. This complies with the phonological account of cognate advantage (Voga & Grainger, 2007) reported in many experiments, where related primes were processed faster and more accurately than unrelated primes. Persian frequency was not a significant predictor of RT and accuracy, although it correlated negatively with it. This finding, therefore, casts doubt on the idea that frequency in L1 results in a cognate advantage (Peeters et al., 2013) at least for languages with different scripts. Peeters and colleagues (2013) reported that the best situation for observing the cognate advantage is when frequency in L1 and L2 is high. Such findings support the idea that cognate and noncognate representations are quantitatively different, meaning that positive cognate effects result from differing exposure to cognates and noncognates. Cognates are available in both languages. Therefore, they have an advantage over noncognates that are merely available in one language.

The results of this study are consistent with those reported by De Groot et al. (2002), in that familiarity with English words was found to be the best predictor of RT and accuracy in most conditions mentioned before. In other words, participants seemed to use their familiarity with the English words, especially familiarity with cognates, to perform the lexical decision. This can further be supported by the fact that increasing the neighborhood density of targets helped participants make faster decisions.

These findings have implications for L2 teaching. Cognates were shown to be more familiar than noncognates in this study. Even this feature was more important than the frequency of English words. Undoubtedly, cognates are a part of the knowledge of L1 that is transferable to L2, and so this transferability gives cognates an advantage over noncognates in vocabulary learning. Furthermore, learners can rely on this knowledge when they have limited L2 vocabulary knowledge (Vandergrift, 1997). Thus, it is suggested that language teachers use such an advantage and include cognates on their to-do-list, especially at the early stages of L2 learning. Teachers, however, should raise the learners' awareness in L1 and L2 (Otwinowska-Kasztelanic, 2009). Raising awareness is essential, as the proportion of the similarity of form in Persian and English influenced RT in this study.

5. Conclusion

The present study produced a dataset that can be used by the researchers who would like to conduct cross-language studies in Persian and English. This dataset provides 288 cognates and noncognates whose features such as frequency, orthographic and phonological length, familiarity, orthographic neighborhood size, imageability, and concreteness were determined and matched across experimental conditions. Nevertheless, as item-selection is observed to bring about inconsistencies in the literature, more comprehensive datasets in Persian-English are needed to provide generalizability of findings observed in studies on languages other than Persian.

6. References

- Balota, D. A., Cortese, M. J., Sergent-Marshall, S. D., Spieler, D. H., & Yap, M. J. (2004). Visual word recognition of single-syllable words. *Journal of Experimental Psychology: General*, 133, 283–316.
- Bakhtiar, M., & Weekes, B. (2015). Lexico-semantic effects on word naming in Persian: Does age of acquisition have an effect? *Memory & Cognition*, 43, 298-313.
- Barber, H. A., Otten, L. J., Kousta, S., & Vigliocco, G. (2013). Concreteness in word processing: ERP and behavioral effects in a lexical decision task. *Brain and Language*, 125, 47-53.
- Boukadi, M., Zouaidi, C., Wilson, M.A., (2016). Norms for name agreement, familiarity, subjective frequency, and imageability for 348 object names in Tunisian Arabic. *Behavior Research Methods*, 48, 585-599.
- Brysbaert, M., Lagrou, E., & Stevens, M. (2017). Visual word recognition in a second language: A test of the lexical entrenchment hypothesis with lexical decision times. *Bilingualism: Language and Cognition*, 20, 530-548.
- Bultena, S., Dijkstra, T., & Van Hell, J. G. (2014). Cognate effects in sentence context depend on word class, L2 proficiency, and task. *The Quarterly Journal of experimental Psychology*, 67, 1214-1241.
- Coltheart, M., Davelaar, E., Jonasson, J. F., & Besner, D. (1977). Access to the internal lexicon. In S. Dornic (ed.), *Attention and performance VI* (pp. 535-555). Hillsdale, NJ: Erlbaum.

Coltheart, M., Rastle, K., Perry, C., Langdon, R., Ziegler, J. (2001). DRC: A dual route cascaded model of visual word recognition and reading aloud. *Psychological Review*, 108, 204-256.

Connine, C., Mullennix, J., Shernoff, E., & Yelen, J. (1990). Word familiarity and frequency in visual and auditory word recognition. *Journal of Experimental Psychology: Learning, Memory and Cognition*, 16, 1084-1096.

Cuetos, F., Ellis, A.W., Alvarez, B. (1999) Naming times for the Snodgrass and Vanderwart pictures in Spanish. *Behavior Research Methods, Instruments, and Computers*, 31, 650-658.

De Groot, A. M. B., Borgwaldt, S., Bos, M., & van den Eijnden, E. (2002). Lexical decision and word naming in bilinguals: Language effects and task effects. *Journal of Memory and Language*, 47, 91-124.

Diependaele, K., Lemhöfer, K., & Brysbaert, M. (2013). The word frequency effect in first- and second-language word recognition: A lexical entrenchment account. *The Quarterly Journal of Experimental Psychology*, 66, 843-863.

Forster, K. & Davis, C. (1984). Repetition priming and frequency attenuation in lexical access. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, (4)10, 680-698.
Doi:10.1037/02787393.10.4.680

Forster, K. I. (1976). Accessing the mental lexicon. In R. J. Wales & E. Walker (eds.), *New approaches to language mechanisms*, pp. 257-287. Amsterdam: North-Holland.

Fotovatnia, Z., & Taleb, F. (2012). Mental Representation of Cognates/Noncognates in Persian-Speaking EFL Learners. *Journal of Teaching Language Skills*, 31, 25-48.

Fotovatnia, Z. & Taleb, F. (2013). Testing conceptual routes in elementary/highly proficient Persian speaking EFL learners. *Iranian Journal of Research in English Language Teaching*, 1, 48-63.

- Gilhooly, K. J., Logie, R.H., (1980). Age-of-acquisition, imagery, concreteness, familiarity, and ambiguity measures for 1,944 words. *Behavior Research Methods & Instrumentation* 12, 395-427.
- Gollan, T. H., Forster, K. I., & Frost, R. (1997). Translation priming with different scripts: Masked priming with cognates and noncognates in HebrewEnglish bilinguals. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 23, 1122-1139.
- Grainger, J., & Jacobs, A. M. (1996). Orthographic processing in visual word recognition: A multiple read-out model. *Psychological Review*, 103, 518-565.
- Grainger, J. & Frenck-Mestre, C. (1998). Masked priming translation equivalents in proficient bilinguals. *Language and Cognitive Processes*, 13, 601- 623.
- Juhasz, B. J. (2005). Age-of-acquisition effects in word and picture identification. *Psychological Bulletin*, 131, 684–712.
- Kanske, P., & Kotz, S. A. (2007). Concreteness in emotional words: ERP evidence from a hemifield study. *Brain Research*, 1148, 138-148.
- Katz, L., & Feldman, L. B. (1983). Relation between pronunciation and recognition of printed words in deep and shallow orthographies. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 9, 157-166.
- Katz, L., Frost, R. (1992). The reading process is different for different orthographies: The orthographic depth hypothesis in L. Katz & R. Frost (eds.), *Orthography, phonology, morphology and meaning*, pp. 67-84. Amsterdam: Elsevier North Holland Press.
- Kondrak, G., Marcu, D., & Knight, K. (2003). Cognates can improve statistical translation models. In proceedings of the 2003 Conference of the North American Chapter of the Association for Computational Linguistics on Human Language Technology (NAACL 2003), Edmonton, Alberta, Canada.
- Kounios, J., & Holcomb, P. J. (1994). Concreteness effects in semantic processing: ERP evidence supporting dual-coding theory. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 20, 804-823.

Kousta, S., Vigliocco, G., Vinson, D. P., Andrews, M., Del Campo, E. (2011). The representation of abstract words: Why emotion matters. *Journal of Experimental Psychology: General*, 140, 14-34.

Kreuz, R. J. (1987). The subjective familiarity of English homophones. *Memory and Cognition*, 15, 154-168.

Marian, V., Bartolotti, J., Chabal, S., Shook, A. (2012). CLEARPOND: Cross-linguistic easy-access resource for phonological and orthographic neighborhood densities. *PloS One*, 7, e43230. <http://clearpond.northwestern.edu>

McClelland, J. L., & Rumelhart, D. E. (1981). An interactive activation model of context effects in letter perception: I. an account of basic findings. *Psychological Review*, 88, 375-407.

Nakayama, M., Verdonschot, R. G., Sears, C. R., & Lupker, S. J. (2014). The masked cognate translation priming effect for different-script bilinguals is modulated by the phonological similarity of cognate words: Further support for the phonological account. *Journal of Cognitive Psychology*, 26, 714-724.

Otwinowska-Kasztelanic, A. (2009). Raising Awareness of Cognate Vocabulary as a Strategy in Teaching English to Polish Adults. *International Journal of Innovation in Language Learning and Teaching*, 3, 131-147. Doi: 10.1080/17501220802283186

Paivio, A. (1986). *Mental representations*. New York: Oxford University Press.

Peeters, D., Dijkstra, T., & Grainger, J. (2013). The representation and processing of identical cognates by late bilinguals: RT and ERP effects. *Journal of Memory and Language*, 68, 315-332.

Plaut, D. C. and Shallice, T. (1993). Perseverative and semantic influences on visual object naming errors in optic aphasia: A connectionist account. *Journal of Cognitive Neuroscience*, 5, 89-117.

Schwanenflugel, P. (1991). Why are abstract concepts hard to understand? In P. Schwanenflugel (ed.), *The psychology of word meanings*, pp 223-250. Hillsdale, NJ, Lawrence Erlbaum Associates.

Sears, C. R., Hino, Y. & Lupker, S. J. (1995). Neighborhood size and neighborhood frequency effects in word recognition. *Journal of Experimental Psychology: Human Perception and Performance*, 21, 876-900.

Sheykh Esmaili, K., Abolhassani, H., Neshati, M., Behrangi, Rostami, & Mohammadi-Nasirie, M. (2007). Mahak: A Test Collection for Evaluation of Farsi Information Retrieval Systems. Presented at the International Conference of Computer Systems and Applications, AICCSA '07. IEEE/ACS.

Van Hell, J. G., & Dijkstra, A. (2002). Foreign language knowledge can influence native language performance in exclusively native contexts. *Psychonomic Bulletin and Review*, 9, 780-789.

Vandergrift, L. (1997). The Comprehension Strategies of Second Language French) Listeners: A Descriptive Study. *Foriegn Language Annals*, 30(3), 387-409.

Voga M, & Grainger J. (2007). Cognate status and cross-script translation priming. *Memory and Cognition*, 35, 938-952.

Zhou, H., Chen, B., Yang, M., & Dunlap, S. (2010). Language non-selective access to phonological representations: Evidence from Chinese-English bilinguals. *The Quarterly Journal of Experimental Psychology*, 63, 20512066.

Notes on Contributors:

Zahra Fotovatnia is an assistant professor of TEFL in the Department of English at the Islamic Azad University, Njafabad Branch, Iran, and a Ph.D. candidate in the Center of Cognitive Neuroscience at Wilfrid Laurier University, Canada. She is interested in studying the neural mechanisms involved in language processing in a bilingual brain.

Jeffery A. Jones is a professor of psychology in the Center of Cognitive Neuroscience at Wilfrid Laurier University, Canada. He investigates the psychological and neurological foundations of human communication. A

primary goal of his research is to understand the neural mechanisms involved in processing the multisensory information that results from producing and perceiving verbal and nonverbal communication.

Nichole E. Scheerer is a post-doctoral fellow in the Brain and Mind Institute at Western University, Canada. Nichole investigates the role of auditory feedback in the development of speech motor control. Nichole is also interested in how differences in sensory processing influence higher order processes like social communication.