



Universiteit
Leiden
The Netherlands

Word processing in languages using non-alphabetic scripts: The cases of Japanese and Chinese

Verdonschot, R.

Citation

Verdonschot, R. (2011, May 12). *Word processing in languages using non-alphabetic scripts: The cases of Japanese and Chinese*. LOT dissertation series. LOT, Utrecht. Retrieved from <https://hdl.handle.net/1887/17637>

Version: Not Applicable (or Unknown)

License: [Licence agreement concerning inclusion of doctoral thesis in the Institutional Repository of the University of Leiden](#)

Downloaded from: <https://hdl.handle.net/1887/17637>

Note: To cite this publication please use the final published version (if applicable).

Word processing in languages
using non-alphabetic scripts.

The cases of Japanese and Chinese

Published by
LOT
Trans 10
3512 JK Utrecht
The Netherlands

Phone: +31 30 253 6006

e-mail: lot@uu.nl
<http://www.lotschool.nl>

Cover photo by Rinus Verdonschot
ISBN: 978-94-6093-059-1
NUR 616

Copyright © 2011: Rinus Verdonschot. All rights reserved.

Word processing in languages using non-alphabetic scripts.

The cases of Japanese and Chinese

PROEFSCHRIFT

Ter verkrijging van
de graad van Doctor aan de Universiteit Leiden
op gezag van Rector Magnificus prof. mr. P.F. van der Heijden,
volgens besluit van het College voor Promoties
te verdedigen op donderdag 12 mei 2011
klokke 13.45 uur

door

Rinus Verdonschot

Geboren te Geldrop
in 1977

Promotie commissie

Promotor: Prof. Dr. N.O. Schiller

Co-Promotor: Dr. W. La Heij

Overige Leden: Prof. Dr. B. Hommel
Prof. Dr. L.L. Cheng
Prof. Dr. K. Tamaoka (Nagoya University)
Dr. C.C. Levelt
Dr. P.A. Starreveld (Universiteit van
Amsterdam)

Acknowledgements

I am indebted to my supervisor (Prof. Niels Schiller) and co-supervisor (dr. Wido La Heij) for allowing me to explore my fascination for the processing of logographic scripts. I have learned a lot from both of you and our discussions were always held in a constructive and pleasant way. I am also very grateful to the LUCL and the Cognitive Psychology department for supporting me when I was collecting data in Japan and China. In this respect I want to especially mention the involvement of Albertien Olthoff and Gea Hakker. Without you I could not have efficiently planned and carried out the work necessary for this dissertation. I also want to thank Prof. Katsuo Tamaoka, Prof. Qingfang Zhang, dr. Sachiko Kiyama, dr. John Phillips, Nanae Murata, Xuebing Zhu, Haoran Zhao and Chikako Ara for their invaluable assistance in recruiting and testing participants in Japan and China. During my PhD period I was lucky to have met so many nice, smart and crazy colleagues here in the Netherlands as well as in Japan and China, it would unfortunately take up too much space to personally acknowledge everybody I have a good connection with but I do want to mention my two office-mates (Marieke and Jun) for making our time together in (and outside) the office so enjoyable. Last but not least I am grateful to my friends, my two ‘paranimfen’ (Johan and Vincent) and above all to my family for their constant support!

ありがとうございました！

Contents

Acknowledgements	
Chapter 1: General Introduction.....	7
Chapter 2: Japanese kanji primes facilitate naming of multiple katakana targets	25
<i>Experiment 1: Reading aloud katakana targets preceded by kanji primes with Equal Reading Preference.....</i>	33
<i>Experiment 2: Reading aloud katakana strings preceded by kanji primes with Equal and high-ON/high-KUN readings.....</i>	38
Chapter 3: Semantic context effects when naming Japanese kanji, but not Chinese hànzi	53
<i>Experiment: Semantic context effects in Japanese and Chinese word reading</i>	59
Chapter 4: Context effects when naming Japanese (but not Chinese), and degraded Dutch nouns: evidence for processing costs?	73
<i>Experiment 1: Naming Japanese kanji with homophonic distractor pictures</i>	81
<i>Experiment 2: Naming Chinese hànzi with homophonic pictures </i>	86
<i>Experiment 3: Dutch bare noun, det+N, and degraded-word naming</i>	89
Chapter 5: The functional unit of Japanese word naming: evidence from masked priming.	105
<i>Experiment 1: Segment versus mora priming</i>	115
<i>Experiment 2: The effects of kana and romaji scripts on masked onset priming.....</i>	118
<i>Experiment 3: Effect of hiragana primes on romaji targets</i>	121
<i>Experiment 4a: Distinguishing the CV from the syllable: MOPE group.</i>	123
<i>Experiment 4b: Distinguishing the CV from the syllable: MORA group.</i>	125
Chapter 6: Summary and Conclusion.....	149
Samenvatting in het Nederlands.....	163
Curriculum Vitae.....	169

Chapter 1: General Introduction

General introduction.

This thesis is concerned with reading aloud words written in a logographic script such as Japanese kanji and Chinese hànzi which are amongst the most complex writing systems in the world. Imagine therefore, for a second, that you have to read aloud a word presented on a screen. How long would that take you? The answer is: not very long, generally around half a second. The ability and speed with which we convert visually presented material (such as alphabetic words, logographic words, symbols and pictures) to intelligible speech output is a quite impressive cognitive achievement. However, the exact amount of time it takes to name aloud visually presented words or pictures depends on a large number of factors, such as the language or script used, specific properties of a word or picture and the setting (or experimental condition) in which a target word or picture has to be named. Word properties that affect naming times are e.g. language, frequency and length, but also regularity. Glushko (1979) found, for instance, that irregular words (i.e. words that do not have a consistent spelling-to-sound correspondence) are named slower compared to regular words; however, others (e.g. Jared, McRae & Seidenberg, 1990) found this to hold only for low-frequency items (such as the first syllable /die/ in DIESEL and DIET which is pronounced /di/ or /daI/). Picture naming latencies are also found to be dependent on properties such as frequency. Oldfield and Wingfield (1964) for instance, found a linear correlation between picture name frequency and naming latency (i.e. higher frequency leads to faster RTs). For many decades now scientists have designed chronometric experiments and investigated speech errors to better understand which components of information processing are involved and which routes are taken when producing words.

Models of language production

Models of language (and word) production usually discriminate between different processing levels which may or may not (depending on the model) include: conceptualization, assignment of syntactic features, phonological word-form encoding, and ultimately articulation (for overviews see e.g., Caramazza, 1997; Dell, 1986; Levelt, Roelofs, & Meyer, 1999; Coltheart, Rastle, Perry, Langdon & Ziegler, 2001). In this thesis, I do refer to other models to account for my findings, but my primary focus lies on the most

detailed model of language production called “Word Encoding by Activation and VERification” or WEAVER++ (which is largely based on evidence from chronometric experiments) developed by Roelofs (1992) and Levelt et al. (1999). According to this model, speech production involves a specific number of levels (see Figure 1).

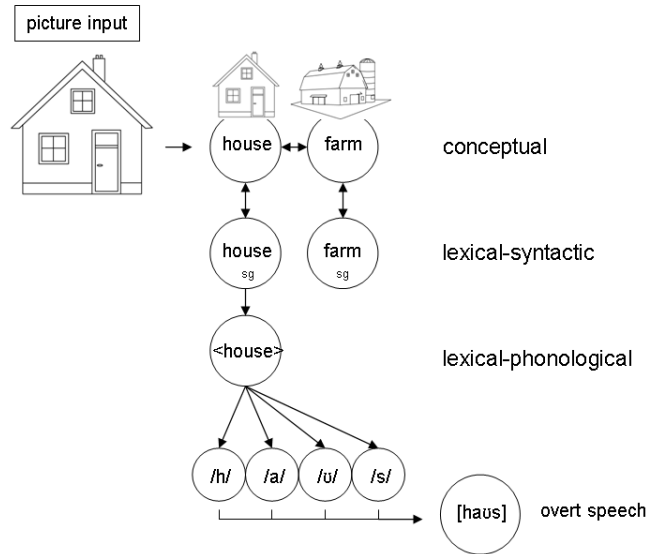


Figure 1. The route a to-be-named picture follows in naming in the WEAVER++ model of language production (Roelofs, 1992; Levelt et al., 1999; sg = singular)

First of all, at the conceptual level, contents of the utterance related to the communicative situation or intention become activated. For instance, if one is instructed to name a picture (e.g., of a “house”), the presented object is visually processed and activates the appropriate concept. It is commonly assumed that not only the target concept “house” but also semantically related concepts (such as “building” or “farm”) receive activation (Collins & Loftus, 1975; Levelt et al., 1999; Glaser & Döngelhoff, 1984; Starreveld & La Heij, 1995).

Subsequently, the activated conceptual nodes will spread activation to the level of lexical representations containing the word’s lexical-syntactic (or lemma) representation. This involves accessing parameters concerning morpho-syntactic make-up of the words. For

instance, if one has been presented with multiple houses (“house_{pl}”), then at this level the parameter for plural (+s) for this specific word will be set. Obviously, syntactic information for irregular words, such as “fish” having an irregular plural form (e.g. “fish”), is also stored and incorporated at this level. The time it takes to select the target lexical-syntactic representation depends on the number of co-activated representations (and their respective activation strengths) that are assumed to compete for selection (but see Mahon, Costa, Peterson, Vargas, & Caramazza, 2007 for an alternative proposal). Competition at the lexical-syntactic level results from cascading activation between the conceptual and lexical-syntactic levels. According to Levelt et al. (1999; see also Levelt et al., 1991), one lexical-syntactic representation is ultimately selected and only this representation activates a representation at the subsequent level of word form encoding (but see Peterson & Savoy, 1998; Roelofs, 2008).

At the level of phonological word-form encoding first of all the phonological make-up of the morphemes that constitute the word (or utterance) has to be retrieved from the mental lexicon. This entails accessing the phonological codes for all included morphemes, e.g. in the case of a picture of multiple houses accessing the free morpheme /house/ and the bound morpheme /s/. Subsequently the metrical and segmental properties of these morphemes will be processed. This involves incremental clustering of phonological segments (e.g. /h/ /a/ /o /and /s/) into a syllabic pattern. Metrical encoding at this level involves determining e.g. the amount of syllables and stress placement in a word.

The final part of the speech production process is turning these syllabified patterns into motor instructions that can be produced by the articulatory system resulting in overt speech.

In this thesis, I will be mostly concerned with the reading aloud (e.g., overt production) of visually presented words. In WEAVER++ (Roelofs, 1992; Levelt et al., 1999), a word can enter the production system in three different ways (see Figure 2).

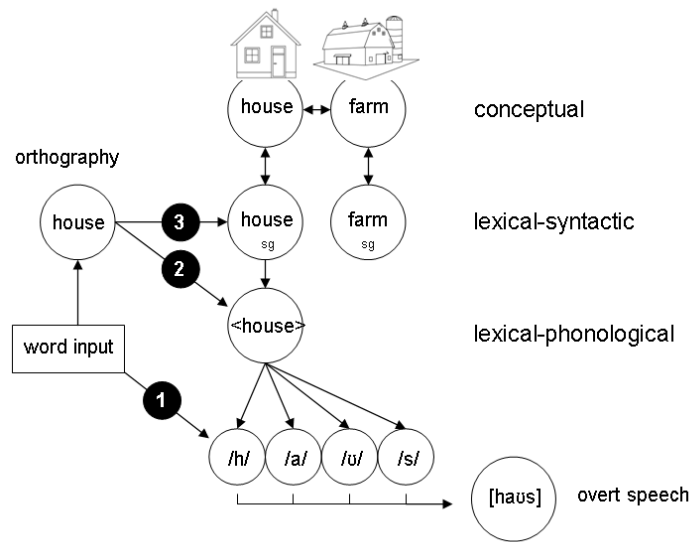


Figure 2. The route a to-be-named word can follow in the WEAVERT+ model of language production (Roelofs, 1992; Levelt et al., 1999)

To begin with, phonemes can be (non-lexically) assembled from a presented string of letters through a rule-based conversion route (Route 1 in Figure 2), also called the grapheme-phoneme rule system (e.g. Coltheart et al., 2001). The ability to pronounce non-words (which by definition do not have an entry in the lexicon), combined with the finding that non-words can facilitate production when they are phonologically to the picture's name (e.g. Lupker, 1982) and, for instance, form-priming effects which were obtained in naming non-words (Horemans & Schiller, 2004), demonstrate the involvement of this sub-lexical route during production. Secondly, words (which have an entry in the lexicon) follow a route from an orthographic to a phonological word-form representation (Route 2 in Figure 1, see also Roelofs, 2003, 2006). This can be seen, for instance, in the long-lag priming paradigm (i.e. many intervening trials between prime and target words) where morphemically (applecore [prime] →

apple [target]) related prime words speed up picture naming (Koester & Schiller, 2008; Zwitserlood, Bölte, & Dohmes, 2000). Finally, besides taking Route 2 (orthography to phonology), words also travel from orthography to their lexical-syntactic representations. This is not to say that reading aloud *necessarily* involves lexical-syntactic representations when the task is to read aloud a word. However, semantic interference effects (e.g., the distractor FARM induces more interference in naming a picture of a house than an unrelated word; Glaser & Dünghoff, 1984; Schriefers et al., 1990) and gender/determiner congruency effects (in languages with a gender system, a gender incongruent distractor word induces interference in naming a picture using a determiner + noun phrase than a gender congruent distractor word) are commonly localized (and indicate involvement of) visually presented distractor words at this specific level (Schriefers, 1993; but see Schiller & Caramazza, 2003).

The majority of word reading and picture naming research has been carried out using languages that employ alphabetic scripts (e.g., English, Dutch, etc.). However, more and more studies are focusing on whether language production theories also provide an useful account of reading aloud in non-alphabetic scripts. In this thesis, I am concerned with the production of phonology from orthography in languages that use logographic scripts (to which we report results in Japanese and Chinese). In the following I will first provide a brief overview of the essentials of logographic characters in Chinese and Japanese. Next, important differences between these two logographic scripts will be discussed. Finally, we present an overview of the main questions addressed in this thesis and how they are dealt with in the various chapters.

The term “logographic characters” usually refers to graphemes that represent a complete word or morpheme (e.g. 海 represents the word or free morpheme for “sea”). In this thesis, we concentrate on logographs which originated in China and which are currently used in both Chinese and Japanese. These specific graphemes can be divided into three categories: pictograms, ideograms and semantic-phonetic compounds (which usually consist of a phonetic and semantic radical, which is a subpart of the whole character indicating semantic membership, e.g. “animal” or a clue about pronunciation).

Table 1. Different types of logograms in Chinese and Japanese (number denotes one of four possible tones in Mandarin Chinese pronunciation, e.g. 1 = high pitch, 2 = rising, 3 =falling then rising, 4 = falling; KUN or ON denotes the origin of the pronunciation)

Type of character	Form	Chinese pronunciation	Japanese pronunciation	meaning	additional information
pictographic	山	/shan1/	/yama _{kun} / or /san _{on} /	mountain	
Ideographic (simple)	石	/shi2/	/ishi _{kun} /seki _{on} /	stone	radical for stone
ideographic compound	岩	/yan2/	/iwa _{kun} /gan _{on} /	rock	山 + 石
semantic - phonetic compound	硅	/gui1/	/kei _{on} /	silicon	石 + 圭

A pictogram is a character denoting something (for instance an object) by representing its shape. In the last column of Table 1, it can be seen how, over time, a drawing from a mountain turned into the current logograph (山). The second category, ideographs or ideographic compounds, are characters which convey an idea. This can be realized either by combining existing pictographs (e.g. combining the radical for stone 石 with mountain 山 to create rock → 岩) or by the introduction of new ideographs which reflect ideas or concepts (such as 上 for “up” or 下 for “down”).

It is important to notice that pictographs and ideographs only make up a small portion of the whole logograph inventory. The most commonly found structure for a logographic character is called the semantic-phonetic compound. Typically, semantic-phonetic compounds are made up of two parts, (1) a semantic part indicating the general meaning or category of the character and (2) a phonetic part indicating an approximate pronunciation. Considering the last example from Table 1, 硅 “silicon”. The left part denotes the radical for stone (石 or the *something-to-do-with stone or minerals* group) and the right part gives a clue to the overall pronunciation for that character, e.g. 圭 (/gui1/ for Chinese or /kei_{on}/ for Japanese). Another clear example is the character 蚂 (“ant”; /ma3/). In this character, the left radical 虫 (“insect”; /chong2/) indicates the group “insects” and the right radical 马 (“horse”; /ma3/) indicates the pronunciation for the character (for Chinese, not Japanese where this character does not exist). In other words, one could etymologically interpret such structure as “the insect which sounds like /ma3/”. Although the predictions of the phonetic radical are in many cases correct, e.g., 蟬 (“cicada”; /chan2/) with 單 (“single”; /chan2/), this is not always the case, e.g., 蛾 (“moth”; /e2/) with 我 (“I”; /wo3/). However, 我 used as a phonetic radical is still consistent in predicting /e/ without the tone such as in 俄 (“Russia or suddenly”; /e2/) or 饿 (“hungry”; /e4/).

Typically, Chinese characters have just a single pronunciation, however, there are also some characters which can carry more pronunciations (e.g. 行 /xing2/ or /hang2/, meaning ‘to go’). This is completely different for Japanese as the *majority* of Japanese kanji has more pronunciations for a single character. This issue will be

discussed in more detail shortly after a brief introduction on Japanese scripts.

Modern Japanese employs a combination of no less than three scripts (not counting letters/numbers), namely logographic kanji and hiragana and katakana. Logographic kanji are generally used to denote parts of the language which carry meaning such as nouns and verb- or adjective stems. Hiragana is usually used to write verb and adjective inflections, grammatical particles (e.g. を object marker “o”), and some native Japanese words. Katakana is mainly used to represent non-Japanese loanwords and onomatopoeia (e.g. チョキチョキ, “choki choki”, i.e. the cutting sound of scissors).

The route from orthography to phonology in Japanese kanji.

As mentioned above, Chinese usually, but not always, has a single pronunciation for a character, e.g. 虫 “insect”; /chong2/). In Japanese, however, 虫 could be pronounced /mushi_{kun}/ or /chuu_{on}/ depending on the combination it forms with other characters (e.g. intra-word context). The etymology of this complex pronunciation system lies in the fact that Japan had no written language when trade and cultural exchange with China started. As a result, over time, Chinese characters were incorporated into the Japanese language. However, in many cases, not only the character was imported, but also the Chinese pronunciation. Consider, for instance, the character for water 水 which in Chinese is pronounced /shui3/. In Japanese it can be pronounced as the original Japanese word for water e.g. /mizu_{kun}/ (also called KUN, or literally, “meaning” way) or as the incorporated Chinese word for water, e.g. /sui_{on}/ (also called ON, or literally, “sound” way). Usually, when a character stands alone, and it has a KUN reading (which not all characters have), it will be pronounced using that reading, meaning that if 水 stands alone, one would say /mizu_{kun}/ and not /sui_{on}/. However, if the character is part of a compound, usually it will take the ON reading. The fact that such rules are not invariable can be seen from examples such as 海水 (‘seawater’), in which the 水 is pronounced as /sui_{on}/ and 雨水 (‘rainwater’) in which the 水 is pronounced as /mizu_{kun}/.

The main issue this thesis tries to shed more light on is how reading aloud takes place in languages using logographic scripts. In particular whether in reading aloud words the Japanese language

production system (by means of its complicated pronunciation system) differs from Chinese.

At this point, it is useful to return briefly to the introduction mentioning that irregular words are found to be named slower than regular words (Glushko, 1979) particularly when they are low frequent (Jared et al., 1990). In addition, it has also often been reported that heterophonic homographs, i.e., words that can take more pronunciations depending on the context they appear in, also show prolonged naming latencies. The word “read”, for instance, would be pronounced differently depending on its tense (future or past tense, i.e., “I’ll read [/rid/] this book” vs. “I’ve read [/rɛ d/] this book”).

There is ample evidence that such words are in general read aloud more slowly than matched controls (Seidenberg, Waters, Barnes, & Tanenhaus, 1984; Kawamoto & Zemblidge, 1992; Folk & Morris, 1995; Gottlob, Goldinger, Stone, & Van Orden, 1999). In light of these results it has been theorized that pronunciation of such words might be complex, due to processing cost reflecting the time to select between simultaneously activated pronunciations.

As over 60% of all kanji have such multiple pronunciations, what would this imply for Japanese words? For instance, if 水 is presented in isolation, perhaps only /mizu_{kun}/ and not /sui_{on}/ will receive activation or when the compound 海水 /kai_{on}-sui_{on}/ “seawater” is presented, perhaps /mizu_{kun}/ does not receive activation, however, both may still receive activation, but only the one which reaches the highest activation (or a certain threshold) will be produced.

In an influential study, Wydell, Butterworth, and Patterson (1995) explored this issue. These authors investigated whether consistency effects in terms of competing ON and KUN readings could be found in Japanese. Words including a character like 水 were termed “inconsistent” since the pronunciation of 水 varies across orthographic neighborhood. Consistent words were defined as words with a single pronunciation for each character at a given position in a compound. Wydell and colleagues (1995) did not find (in)consistency effects in five experiments with two-kanji words and one experiment with single-kanji words. Regarding the two-kanji experimental results, Wydell et al. (1995) concluded that the computation of phonology from orthography is mainly situated at the level of the whole word (e.g. the compound) and not at its subcomponents (e.g. pronunciation

for the individual members of the compound). In other words, for 海水 /kai_{on}-sui_{on}/ “seawater” it does not matter that 海 “sea” can also be pronounced as /umi_{kun}/ and 水 “water” can also be pronounced as /mizu_{kun}/. In the stand alone single-kanji experiment, Wydell et al. (1995) found that when a kanji presented in isolation had to be named, it did not matter whether that kanji had two or three possible alternative pronunciations (in compounds), i.e. they found no consistency effects in bare kanji naming.

However, subsequent research by Fushimi, Ijuin, Patterson, and Tatsumi (1999) did find significant RT differences between consistent and inconsistent multiple reading kanji when ‘typicality’ (i.e. when a certain pronunciation was typical or not) was introduced as a factor in naming compound and non-words in Japanese. These results led to their interpretation (contrasting Wydell et al., 1995) that computation can be affected by the character-sound correspondence at the subcomponent level (individual kanji) and not only at the whole-word level. Another study by Kayamoto, Yamada, and Takashima (1998) reported contrasting results for single kanji naming depending on relative frequencies of the alternative readings. In their first experiment, participants were instructed to name kanji having only a single reading, i.e. 肉 /niku_{on}/ “meat” versus kanji having multiple reading, i.e. 数 /kazu_{kun}/ or /suu_{on}/ “number”. Kayamoto and colleagues employed high-frequency and mid-frequency kanji and found for both frequency ranges a significant increase in naming latencies for multiple reading kanji compared to single reading kanji. Interestingly, participants even produced occasional blending of both readings, i.e. 数 might have been blended to /kasuu/ (kazu_{kun}+suu_{on}). Kayamoto et al. (1998) argued that the longer RTs observed for the multiple reading Kanji might be due to the fact that in the stimulus materials used the alternative readings (i.e. ON readings) had relatively high language frequencies. Therefore, in a second experiment, the authors employed kanji with a single reading, e.g., again 肉 /niku_{on}/ “meat”, versus kanji which had a subordinate (weaker) alternative and hence a stronger dominant reading, e.g., mado in 窓 /mado_{kun}/ and /sou_{on}/ “window”. In this second experiment, the naming latencies of single- and multiple reading kanji were similar in size. Kayamoto et al. (1998) therefore concluded that competition induced by a strong alternative (ON) reading caused the

processing cost in their first experiment. The absence of an effect in their second experiment was proposed to be due to an insufficient strength of the alternative reading to be an adequate competitor to the dominant reading.

In sum, there is evidence for the activation of multiple readings when processing Japanese kanji (Kayamoto et al., 1998; Fushimi et al., 1999) as well as evidence against it (Wydell et al., 1995). Thus, the body of experimental evidence regarding phonological activation of multiple readings in kanji processing is not entirely conclusive yet. In addition, the aforementioned experimental evidence was always indirectly acquired, e.g. interpretations always necessitated the *assumption* that the activation of multiple readings caused naming latency differences. This thesis endeavors to establish whether presentation of a single kanji activates multiple pronunciations (Chapter 2; by means of assessing *directly* whether multiple readings can be primed from a single kanji prime) and subsequently aims to ascertain whether and under which circumstances the activation of multiple pronunciations leads to increased processing costs (Chapter 3 and 4; using reading aloud tasks with picture distractors).

Units of speech production in Japanese

The focus of Chapter 3-4 concerns the reading aloud of logographic words, especially whether multiple phonological representations of Japanese kanji show different latency patterns from reading aloud Chinese hànzì when presented in semantic or phonological context. Specific attention in Chapter 1 is paid to establishing whether multiple phonological representations become active in Japanese. In addition to this, Chapter 5 specifically zooms in on the chunk size of the activated phonological units of speech production in Japanese.

Theories of language production generally describe the phonemic segment (e.g. /r/) to be the basic unit in phonological encoding. However, there is also evidence that such a unit might be language-specific. In Mandarin Chinese, for instance, speakers are shown to profit from preparation of the first syllable but not from the first phonemic segment of a word (Chen, Chen, & Dell, 2002). Such findings are inconsistent with results obtained using English, Dutch, and other Germanic languages. Certain aspects usually found in

theories of word production (such as the segment being the basic phonological unit size) might not apply to all languages. Therefore, to augment the understanding of Japanese phonological encoding, it is important to also establish the size of the basic processing unit in Japanese. This might differ considerably from Chinese and Germanic languages such as English and Dutch, as Japanese is often argued to be a mora-based¹ language. We address this question in Chapter 5, where several masked priming experiments distinguish the mora from segment and syllable effects during word naming.

Overview of the experimental chapters

Chapter 2 of this thesis aims to obtain direct evidence regarding the activation of multiple readings for a single kanji character. More specifically, I investigate whether activation of an alternative reading can be detected even when the alternative reading is weak, or to put it in other words, whether activation of alternative readings only occurs under special circumstances, e.g. in case of a high frequency competitor. In this chapter, we report the results of two masked-priming experiments using kanji primes and their KUN and ON transcriptions in katakana (a syllabic Japanese script mostly used to write loan words) that show that presentation of a single kanji prime indeed activates multiple readings. In Chapter 3, participants are asked to read aloud Japanese and Chinese target words superimposed on semantically related and unrelated picture distractors. We show that in Japanese (but not Chinese) semantically related pictures speed up naming latencies of the target words. In Chapter 4, we show the same pattern of results, this time using phonologically related distractor pictures (e.g., homophones). A subsequent control experiment in Dutch confirms that the observed facilitation in Japanese (but not Chinese) is likely the result of a processing delay at the lexical-phonological level. Although this may seem counterintuitive, I propose that the processing cost allows phonologically related pictures (compared to phonologically unrelated pictures) to exert a facilitatory influence on naming latencies. In the General Discussion of Chapter 4, we conclude that the research reported in Chapters 2 to 4 shows that multiple readings can be activated when processing Japanese kanji.

¹ A mora is considered to be an independent rhythmical structure within a syllable. For instance, the well-known Japanese name “Honda” consists of two syllables /hoN/ and /da/ (N = nasal coda) and three morae, /ho/ /N/ /da/ which last approximately equally long.

The activation of multiple representations might come at a processing cost (perhaps due to competition or shared activation) which makes such stimuli susceptible to semantic and homophonic context effects from distractor pictures. Previous work (see Kinoshita, 2003, for a review) has consistently shown that when a target word is preceded by a briefly presented masked prime word sharing the onset (i.e. the first letter) with the target, reading aloud the target is facilitated compared to an unrelated prime (e.g. Forster & Davis, 1991; Schiller, 2004). This masked priming paradigm was used in four experiments described in Chapter 5 to establish whether or not onset effects could be obtained in Japanese. Throughout the experiments, we manipulated the degree of overlap between a prime word and a target word from one consonantal segment to a whole mora (CV). The first three experiments in Chapter 5 show that onset effects are not present in Japanese word reading even when allowing for a script which favors segmental processing (e.g., romaji). The fourth experiment distinguishes between the mora and the syllable, and indicates that the mora is indeed the elementary (or proximate) unit during phonological encoding in Japanese language production. The last chapter provides an overview of the results (and conclusions on) the experimental work performed in this thesis.

The following references correspond to the chapters in this thesis.

Chapter 2: Verdonschot, R. G., La Heij, W., Poppe, C., Tamaoka, K., & Schiller, N. O. (submitted). Japanese kanji primes facilitate naming of multiple katakana targets.

Chapter 3: Verdonschot, R. G., La Heij, W., Schiller, N.O. (2010). Semantic context effects when naming Japanese kanji, but not Chinese hànzi, *Cognition*, 115, 512-518.

Chapter 4: Verdonschot, R. G., Paolieri, D., Kiyama, S., Zhang, Q. F., La Heij, W., & Schiller, N. O. (submitted). Context effects when naming Japanese (but not Chinese), and degraded Dutch nouns: evidence for processing costs?

Chapter 5: Verdonschot, R. G., La Heij, W., Kiyama, S., Tamaoka, K., Kinoshita, S., & Schiller, N. O. (submitted). The functional unit of Japanese word naming: evidence from masked priming.

References

- Caramazza, A. (1997). How many levels of processing are there in lexical access? *Cognitive Neuropsychology*, 14, 177-208.
- Chen, J.-Y., Chen, T.-M., & Dell, G. S. (2002). Word-form encoding in Mandarin Chinese as assessed by the implicit priming task. *Journal of Memory and Language*, 46, 751-781.
- Collins, A. M., & Loftus, E. F., (1975). A spreading activation theory of semantic processing. *Psychological Review*, 82, 407-28.
- 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.
- Dell, G. S. (1986). A spreading-activation theory of retrieval in sentence production. *Psychological Review*, 93, 283-321.
- Folk, J. R., & Morris, R. K. (1995). The use of multiple lexical codes in reading: Evidence from eye movements, naming time, and oral reading. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 21, 1412-1429.
- Forster, K. I., & Davis, C. (1991). The density constraint on form-priming in the naming task: Interference effects from a masked prime. *Journal of Memory and Language*, 30, 1-25.
- Fushimi, T., Ijuin, M., Patterson, K., & Tatsumi, I. F. (1999). Consistency, frequency, and lexicality effects in naming Japanese Kanji. *Journal of Experimental Psychology: Human Perception and Performance*, 25, 382-407.
- Glaser, W. R., & Dünghoff, F. J. (1984). The time course of picture-word interference. *Journal of Experimental Psychology: Human Perception and Performance*, 10, 640-654.
- Glushko, R. J. (1979). The organization and activation of orthographic knowledge in reading aloud. *Journal of Experimental Psychology: Human Perception and Performance*, 5, 674-691.
- Gottlob, L. R., Goldinger, S. D., Stone, G. O., & Van Orden, G. C. (1999). Reading homographs: Orthographic, Phonologic, and Semantic Dynamics. *Journal of Experimental Psychology: Human Perception and Performance*, 25, 561-574.
- Horemans, I., & Schiller, N. O. (2004). Form-priming effects in non-word naming. *Brain and Language*, 90, 465-469.
- Jared, D., McRae, K., & Seidenberg, M. S. (1990). The basis of consistency effects in word naming. *Journal of Memory and Language*, 29, 687-715.

Kawamoto, A. H., & Zemblidge, J. (1992). Pronunciation of homographs. *Journal of Memory and Language*, 31, 349-374.

Kayamoto, Y., Yamada, J., & Takashima, H. (1998). The consistency of multiple-pronunciation effects in reading: The case of Japanese logographs. *Journal of Psycholinguistic Research*, 27, 619-637.

Kinoshita, S. (2003). The nature of masked onset priming effects in naming: A review. In S. Kinoshita & S. J. Lupker (eds.), *Masked priming. The state of the art* (pp. 223-238). Hove: Psychology Press.

Koester, D., & Schiller, N. O. (2008). Morphological priming in overt language production: Electrophysiological evidence from Dutch. *NeuroImage*, 42, 1622-1630.

Levelt, W. J. M., Schriefers, H., Vorberg, D., Meyer, A. S., Pechmann, T., & Havinga, J. (1991). The time course of lexical access in speech production: A study of picture naming. *Psychological Review*, 98:122-42.

Levelt, W. J. M., Roelofs, A., & Meyer, A. S. (1999). A theory of lexical access in speech production. *Behavioral and Brain Sciences*, 22, 1-75.

Lupker, S. J. (1982). The role of phonetic and orthographic similarity in picture-word interference. *Canadian Journal of Psychology*, 36, 349-367.

Mahon, B. Z., Costa, A., Peterson, R., Vargas, K. A., & Caramazza, A. (2007). Lexical Selection is not by competition: a reinterpretation of semantic interference and facilitation effects in the picture-word interference paradigm. *Journal of Experimental psychology, Learning, Memory and Cognition*, 33, 503-535.

Oldfield, R. C., & Wingfield, A. (1964). The time it takes to name an object. *Nature*, 202, 1031-1032.

Peterson, R. R., & Savoy, P. (1998). Lexical selection and phonological encoding during language production: Evidence for cascaded processing. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 24, 539-557.

Roelofs, A. (1992). A spreading-activation theory of lemma retrieval in speaking. *Cognition*, 42, 107-142.

Roelofs, A. (2003). Goal-referenced selection of verbal action: Modeling attentional control in the Stroop task. *Psychological Review*, 110, 88-125.

- Roelofs, A. (2006). Context effects of pictures and words in naming objects, reading words, and generating simple phrases. *Quarterly Journal of Experimental Psychology*, 59, 1764–1784.
- Roelofs, A. (2008). Tracing attention and the activation flow in spoken word planning using eye movements. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 34, 353–368.
- Schiller, N. O. (2004). The onset effect in word naming. *Journal of Memory & Language*, 50, 477–490.
- Schiller, N. O., & Caramazza, A. (2003). Grammatical feature selection in noun phrase production: Evidence from German and Dutch. *Journal of Memory and Language*, 48, 169–194.
- Schriefers, H. (1993). Syntactic processes in the production of noun phrases. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 19, 841–850.
- Schriefers, H., Meyer, A. S., & Levelt, W. J. M. (1990). Exploring the time course of lexical access in speech production: Picture-word interference studies. *Journal of Memory and Language*, 29, 86–102.
- Seidenberg, M. S., Waters, G. D., Barnes, M. A., & Tanenhaus, M. K. (1984). When does irregular spelling or pronunciation influence word recognition? *Journal of Verbal Learning and Verbal Behavior*, 23, 383–404.
- Starreveld, P. A., & La Heij, W. (1995). Semantic interference, orthographic facilitation and their interaction in naming tasks. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 21, 686–698.
- Wydell, T. N., Butterworth, B. L., & Patterson, K. E. (1995). The inconsistency of consistency effects in reading: Are there consistency effects in Kanji? *Journal of Experimental Psychology: Learning, Memory and Cognition*, 21, 1155–1168.
- Zwitzerlood, P., Bölte, J., & Dohmes, P. (2000). Morphological effects on speech production: evidence from picture naming. *Language and Cognitive Processes*, 15, 563–591.

Chapter 2: Japanese kanji primes facilitate naming of multiple katakana targets

This chapter is based on: Verdonschot, R. G., La Heij, W., Poppe, C., Tamaoka, K., & Schiller, N. O. (submitted). Japanese kanji primes facilitate naming of multiple katakana targets.

Abstract

Some types of words such as words with inconsistent grapheme-to-phoneme conversion or homographs are read aloud more slowly than matched controls, presumably due to competition processes. The majority of Japanese kanji have multiple readings for the same orthographic unit, i.e. the native Japanese reading (called KUN reading) and the Sino-Japanese reading (called ON reading). This has led to the question whether the presence of multiple readings may lead to processing costs for the selection of the correct pronunciation. Studies examining this issue have provided mixed evidence. Wydell and colleagues (1995) did not obtain processing costs for naming single kanji which had more readings compared to kanji which had fewer readings. In contrast, Kayamoto and colleagues (1998) did find such costs (longer RTs) when dominant readings were relatively weak. However, as their participants produced the dominant reading of target words, the presence or absence of processing costs in naming kanji having multiple readings only provides indirect evidence on the issue whether or not alternative readings were also active. The current study aims to obtain direct evidence on whether or not multiple readings become active and whether biasing the reading preference towards the KUN and ON reading leads to a different activation spreading pattern. The results of two priming experiments showed indeed that the reading of multiple katakana targets was facilitated by the same kanji prime indicating that multiple readings were activated. In addition, when Kanji characters are presented in isolation, phonological activation was constantly stronger for KUN-readings compared to ON-readings, even when the kanji is biased towards the ON reading.

Japanese kanji primes facilitate reading aloud of multiple targets.

Reading aloud words is a task that can be carried out easily and rapidly. However, the underlying mechanisms are complex. For instance, exception words, which do not consistently follow orthography-to-phonology conversion (OPC) rules, i.e. STEAK, are named slower compared to consistent words, i.e. BEAK. Presumably, this is caused by a conflict between stored knowledge of the word's pronunciation and OPC regularities for similar orthographic patterns (Glushko, 1979; Stanovich & Bower, 1978; see also Seidenberg, Waters, Barnes & Tanenhaus, 1984). Waters and Seidenberg (1985) replicated these findings, but only found this effect for low-frequency words. Presently, three categories have been established to describe the print-to-sound consistency of a word. The first category is termed regular-consistent, e.g. HAZE, because the so-called word bodies (e.g. -AZE) of all other words across its neighborhood match. These words are also called "friends" (e.g. MAZE, GAZE, DAZE, etc.). The second category contains regular-inconsistent which entails words such as WAVE for which some words, but not all (i.e. "friends" would be e.g. GAVE, CAVE, SAVE), in its neighborhood deviate from the print-to-sound pattern. These so-called "enemies" or exception words, in turn, form the third category (e.g. HAVE).

Glushko (1979) demonstrated that regular-inconsistent words took longer to be read aloud than regular-consistent words, and exception words usually take longer than regular-inconsistent words. However, these findings are not always obtained (see e.g. Brown, 1987).

In 1990, Jared and colleagues proposed that size of the consistency effect on reading aloud latencies might be influenced by properties of a particular word's neighborhood, specifically the relative frequencies of friends and enemies. An inconsistent word which has many friends and one or few enemies might experience less processing costs when computing its pronunciation than words having many enemies. However, a word which has many friends but also many enemies might experience a greater cost.

In 1995, Wydell, Butterworth, and Patterson extended the discussion to Japanese in an effort to find out whether consistency effects could also be found using Japanese kanji which have distinctly different properties compared to alphabetic scripts.

Modern Japanese uses two scripts, syllabic kana (consisting of hiragana and katakana) and logographic kanji. Lexical morphemes, such as nouns and the roots of verbs and adjectives, are usually written in kanji, whereas grammatical morphemes (e.g. inflections and function words) are written in hiragana; loan words are usually written in katakana. Many kanji have the unique property that one visual form can be pronounced in many different ways depending on the character(s) it combines with. This stems from the fact that around 350 AD Japanese began to adopt the Chinese script which also included taking up many Chinese words and pronunciations. Kanji pronunciation is currently divided into two types, i.e. ON-readings, derived from the Chinese pronunciation, and KUN-readings, originating from the native Japanese pronunciation. Usually, there is a tendency that when a kanji character stands alone, the KUN pronunciation is used (e.g. 水 /mizu_{kun}/ “water”), and when a kanji is part of a compound, the ON reading is used. However, many deviations from this rule exist (e.g. the kanji “水” in 海水, /kai_{on}-sui_{on}/ ‘seawater’ and 雨水 /ama_{kun}-mizu_{kun}/ ‘rainwater’).

Wydell and colleagues (1995) described consistency in terms of competing ON and KUN readings. For instance, words including the character 水 would be termed “inconsistent” as the pronunciation of 水 varies across orthographic neighborhood. Consistent words are words for which there would be, for example, just one pronunciation for each character, e.g. 駅員 /eki_{on}in_{on}/ “station attendant”. Wydell and colleagues (1995) did not find consistency effects across five experiments with two-kanji words and one experiment with single-kanji words. They therefore concluded that orthography to phonology is mainly computed at the whole word level and not at its subcomponents (e.g. individual kanji). However, Fushimi, Ijuin, Patterson, and Tatsumi (1999) using more controlled stimuli did find significant consistency effects when naming compound and non-words in Japanese, leading to a contrasting interpretation that computation can be affected by the character-sound correspondence at the subcomponent level (individual kanji) and not only at the whole-word level.

It is important to notice that a major difference between the previous consistency experiments (Gluhsko, 1979; Jared et al., 1990) and the Japanese experiments (Wydell et al., 1995; Fushimi et al.,

1999) relates to the fact that most Japanese stimuli used are in fact compounds. As Wydell et al. put it (p. 1157) comparing “HASTE vs. CASTE” would in Japanese actually be similar to comparing compounds containing heterophonic homographs, e.g. “BOW-TIE” vs. “BOW-WOW” or “LEAD-IN” vs. “LEAD-FREE”. It is important to realize there is an ongoing debate about whether compounds are stored in a decomposed way, i.e. stored in terms of their constituent morphemes at the lexical-phonological (or lexeme) level (Levelt et al., 1999), or in their full form (Caramazza, 1997). Bien, Levelt, and Baayen (2005) using a speech preparation task found that Dutch participants are sensitive to morpheme frequency, indicating decomposed storage. However, a recent study by Janssen, Bi, and Caramazza (2008) used picture naming in Chinese and English to establish whether manipulating the frequency of a compound’s constituents influences RTs. In both Chinese and English they found no evidence of storage in terms of their constituents, supporting the full-form storage account (e.g. Caramazza, 1997).

Although the debate is not fully settled, it is important to consider, as Janssen and colleagues (2008) proposed, that when the morphological complexity changes for a language the relative importance for representations at that level might also change. When one hypothesizes that for any consistency effect to arise the contrasting pronunciations of the constituents should necessarily be active at some point to cause a processing delay, it seems sensible to first compare single kanji with one or more different pronunciations to obtain a measure of consistency and subsequently assess the influences from a representation of that same kanji at a compound word level.

This is reasonable if one considers there is ample evidence that monomorphemic heterographic homographs in English (i.e., “dove” in “a dove [dʌ v] is a bird” and “he dove [dəʊ v] into the sea”) are named slower than their matched controls (Seidenberg et al., 1984; Kawamoto & Zemblidge, 1992; Folk & Morris, 1995; Gottlob, Goldinger, Stone, & Van Orden, 1999). It has been proposed that such processing delay is due to selection costs between multiple simultaneously activated pronunciations.

Bearing this in mind, one may speculate whether and how such results would generalize to the naming of single Japanese kanji. This is a relevant question because approximately 60% of all Japanese

kanji are in fact heterophonic homographs. For many kanji the meaning also changes with the pronunciation but there are also sufficient examples of kanji for which the meaning does not necessarily change for the various pronunciations (as many Sino-Japanese and original Japanese pronunciations essentially carry the same meaning). Bearing this in mind, the question arises whether kanji which have multiple readings may be named slower than kanji which have fewer or just a single reading. However, as multiple pronunciations are the rule rather than the exception in Japanese, it may be that native Japanese can efficiently address such conditions and not show any processing differences between kanji symbols with one and multiple readings. It might also be that when presented in isolation, only the KUN pronunciation, which is usually the preferred one in such case, will be most highly activated, leaving the alternative (e.g. ON) pronunciations with much less activation.

Indeed, data from Wydell et al., (1995) seem to validate the latter hypothesis. For instance, in their fifth experiment they used single kanji which were usually read using their KUN reading. These kanji could either have a single alternative (ON) reading or two alternative (ON) readings. Naming latencies did not differ between these two kanji categories. In their study, however, it was not pointed out what the ON-ratios (see Nomura, 1978 and Tamaoka & Makioka, 2004 for details) for alternative readings were (e.g. were the alternative readings frequent?). Such information is important because Kayamoto, Yamada, and Takashima (1998) reported mixed results for single kanji naming depending on relative frequencies of the alternative readings. In their first experiment, participants were asked to name kanji which had only one reading, i.e. 駅 /eki_{on}/ 'station' versus kanji which had multiple readings, i.e. 雄 /osu_{kun}/ or /yuu_{on}/ 'male'. Kayamoto et al. used high-frequency and mid-frequency kanji and found in both cases a large increase in naming latency for kanji with multiple readings compared to kanji with single readings (63 ms and 88 ms, respectively). Additionally, in some cases participants even produced blending of both readings, i.e. 雄 might be blended to /oyuu/ (osu+yuu). However, the alternative readings (ON) had a high frequency of occurrence which might have lead to the longer RTs. This led Kayamoto and co-workers (1998) to question whether the presence of alternative readings would always exert an effect or

merely in strong alternative reading kanji (e.g. high ON frequency). In a second experiment, they employed kanji having single readings versus kanji which had a subordinate (weaker) alternative and hence a stronger dominant reading. With such stimuli the competition effect disappeared, i.e. kanji with multiple readings were named non-significantly (10 ms) slower than single pronunciation kanji. Kayamoto et al. (1998) concluded that competition between readings created the processing cost in their first experiment, and the absence of competition in their second experiment was due to insufficient strength (low ON ratio) of the rival reading to be a competitor to the dominant reading. Although this conclusion seems reasonable, it is still an indirect measure and, critically, has to assume activation of multiple word readings. In addition, it has not been thoroughly fleshed out how and at what level such activation spreading and competition works.

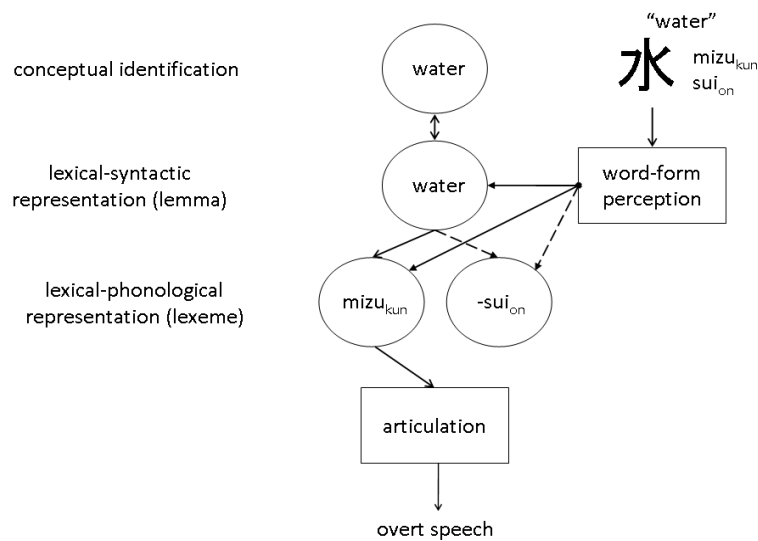


Figure 1. Routes to name a kanji character according to the Levelt et al. (1999) model. Straight lines indicate routes of activation. Dashed lines indicate possible routes of activation for bound morphemes (ON-reading).

In the model proposed by Levelt and colleagues (1999; see Figure 1) printed words activate lexical-syntactic and lexical-phonological representation in parallel. It is possible that when encountering 水 both lexical-phonological representations attached to this character (KUN and ON) are activated. However, as the ON reading is in most cases a bound morpheme (part of a compound), it is also conceivable that, when standing alone, only the Japanese (KUN) pronunciation is activated. There is evidence that a direct route from orthography to phonology to all pronunciations of a kanji (or subcomponent) is indeed activated when naming words in Japanese. For instance, Fushimi et al. (2003) found in a surface dyslexic patient (TI) who suffered from poor word comprehension. For instance, when naming kanji, the amount of consistency affected his performance on both words and non-words in a parametric way. TI showed worst performance for inconsistent-atypical words, especially when they were low-frequency, intermediate performance for inconsistent-typical, and best for consistent words. Important is that, although his semantic system had been impaired, e.g. he was unaware of the meaning of a word; he could still respond with the correct reading but more importantly, also occasionally with a non-suitable but still legitimate reading for a word (or component), e.g. naming 海水 as /umi_{kun}.mizu_{kun}/ (instead of /kai_{on}.sui_{on}/). This indicates that the links from orthography to phonology for that respective word or character had all remained intact although he could not access the meaning.

The current study aims to obtain direct evidence on whether multiple readings of a kanji character become active, e.g. even when the alternative reading is weak, or whether activation of multiple readings only occurs under special circumstances, e.g. in case of a balanced reading ratio between multiple possible pronunciations. We report the results of two masked priming experiments with kanji primes and their KUN and ON transcriptions in katakana (a syllabic Japanese script mostly used to write loan words). The first experiment employs kanji with two readings which have equal pronunciation occurrence in compounds (taken from the 2004 database by Tamaoka and Makioka which can be downloaded from <http://www.lang.nagoya-u.ac.jp/~ktamaoka/>). Such an experiment is comparable to Kayamoto et al.'s (1998) first experiment, as dominant and alternative readings were selected to be comparable in frequency

of occurrence. The second experiment employs kanji which are biased to one of the pronunciations, i.e. the dominant reading is much more frequent compared to the alternative reading. In that sense, it is comparable to the second experiment of Kayamoto et al. (1998).

Our hypotheses are straightforward: if native Japanese speakers prepare multiple pronunciations when processing a masked kanji prime with two alternative readings, then the naming of the transcriptions of those two readings in katakana when preceded by the congruent kanji prime (which is identical for both transcriptions) should show facilitation compared to control primes. Furthermore, if the activation of non-dominant readings is limited or absent, we expect the priming effect for the non-dominant katakana transcription to be smaller compared to the effect for the dominant katakana transcription or even entirely absent. Such a finding would provide a stronger empirical basis for cause of the absence of processing costs in the parts mentioned earlier of the studies by Wydell et al. (1995) and Kayamoto et al. (1998). Furthermore, because the task includes speech production components (reading aloud katakana strings) it will also provide insights into the production processes involved when Japanese participants are primed with a kanji character.

Experiment 1: Reading aloud katakana targets preceded by kanji primes with Equal Reading Preference

We employed a katakana word reading task with masked kanji primes. Masked priming is assumed to prevent possible strategic influences which may hinder or bias the results (Forster & Davis, 1984). This first experiment was performed using kanji which, in compounds, have approximately equal (or balanced) frequencies of occurrence in one or the other reading. If there is indeed simultaneous activation of multiple readings, then kanji having approximately balanced readings are the most likely candidates for inducing activation in multiple readings.

Method.

Participants. Forty undergraduate students from Hiroshima University (23 female, 17 male; average age: 20.5 years; SD = 1.8) took part in this experiment in exchange for financial compensation. All participants were native speakers (and fluent readers) of Japanese and had normal or corrected-to-normal vision.

Stimuli. Twenty-nine kanji primes with two possible pronunciations were selected which adhered to an equal ON-reading ratio (henceforth called Reading Preference or RP) of between 40% and 60% (mean summed average 50.9%), i.e. a kanji character is pronounced approximately equally often with the Sino-Japanese (ON) and the native Japanese (KUN) reading in compound words (for a detailed description of how this statistic is obtained see Tamaoka, Kirsner, Yanase, Miyaoka, & Kawakami, 2002; p. 272). Although there is a bias towards the KUN-reading for stand-alone kanji, we still expected that kanji with equal RP, i.e. which have less of a bias towards one of the multiple readings when occurring in compounds, do show activation spreading to multiple readings, even when presented in isolation (as shown by Kayamoto et al., 1998). The RPs were taken from a database by Tamaoka and Makioka (2004). Target strings were katakana transcriptions of the KUN and ON reading for a kanji character, for instance, 町 meaning city or block was transcribed as マチ /machi_{kun}/ (KUN target) and チョウ /chou_{on}/ (ON target).

We chose to transcribe words into katakana instead of hiragana to avoid tapping into lexical processes as transcriptions into katakana are less familiar to participants compared to hiragana. Furthermore, Hino, Lupker, Ogawa and Sears (2003) showed that masked repetition priming effects with kanji pronunciations transcribed as katakana can be obtained reliably. Twenty-nine kanji control primes were selected which had only one possible reading and were phonologically and semantically unrelated to the targets. The summed kanji frequencies (Yokoyama, Sasahara, Nozaki, & Long, 1998) of repetition and control primes were equated as much as possible, with a mean summed character frequency of 631.4 for repetition primes and 598.4 for control primes, $t(28) = 1.1$, ns, as were the summed stroke complexities, with a mean of 9.8 for repetition primes and 10.6 for control primes, $t(28) < 1$. All kanji were taken from the set of 1,945 commonly-used Japanese kanji as published by the Japanese government in 1981 (for detailed information, see Tamaoka & Makioka, 2004; Yasunaga, 1981). See Appendix A for an overview of the materials used in this experiment.

Design. A 2 (Prime Duration, i.e. backward mask present or absent) x 2 (Target Type, i.e. KUN and ON katakana target) x 3 (Prime Type, i.e. Congruent, Control, or Neutral) within participants

factorial design was implemented. Each participant was subjected to 348 trials (2 x 2 x 3 x 29). Four pseudo-random lists were constructed such that phonologically or semantically related primes or targets had at least a distance of two trials to avoid unintended priming effects. Within participants, the order of lists was counterbalanced. Half of the participants for a particular list saw the backward masking condition first and the other half saw the condition without backward masking first.

Procedure. In all reported experiments the software package E-prime 2.0 combined with a voice key was used for stimulus presentation and data acquisition. Participants were seated approximately 60 cm from a 17 inch LCD computer screen (Eizo Flexscan P1700 with a screen cycle refresh rate of 60 Hz) in a quiet room at Hiroshima University and tested individually. Two trial types were included, e.g. trials with or without a backward mask consisting of three hash marks (###) and lasting for 50 ms. This backward mask was introduced because besides complete masking it can be hypothesized that prolonging activation spreading without extending the actual prime exposure duration might allow alternative pronunciations to build up more activation. A trial comprised the presentation of a fixation cross (1,000 ms) followed by a forward mask – identical to the backward mask – for 500 ms, and subsequently a kanji prime (50 ms) replaced either immediately by the katakana target word (maximally three characters long), which disappeared when the participant responded or after maximally 2,000 ms, or replaced by the backward mask (50 ms), which in turn was replaced by the katakana target word. In between trials, a random inter-stimulus-interval of 400 – 800 ms was introduced to avoid expectancy effects. Naming latencies were measured from target onset. Participants were specifically instructed to respond as fast as possible while avoiding errors. They were not informed about the presence of the prime. After the experiment, as in earlier studies, informal interviewing showed that participants were found to be mostly unable to recognize the primes under the masking conditions used in this study (see also prime visibility tests reported by Schiller, 1998, 2004 under analogous conditions). Participants were furthermore presented with a questionnaire containing the kanji used in the experiment and were asked to write down their preferred pronunciation of the kanji in a script of their choice (hiragana, katakana, or romaji).

Results and Discussion. Naming latencies exceeding 2.5 standard deviations per participant per prime duration were counted as outliers (comprising 2.5% of the data). Separate analyses were carried out with participants (F1) and items (F2) as random variables. In the F2 analysis, Target Type was considered to be a between-item variable. In Table 1, mean RTs and error rates per condition are reported. Since there were overall only 0.93% errors overall distributed approximately equally across conditions in Experiment 1, errors were not analyzed.

Katakana Target	Prime Condition	Mean RT (SD)	%E
KUN reading (e.g. マチ – ‘machi’)	Congruent (町)	456 (42)	0.1
	Control (式)	474 (43)	0.0
	Congruent-Control	–18 (9)	0.1
ON reading (e.g. チョウ – ‘chou’)	Congruent (町)	462 (44)	0.0
	Control (式)	472 (42)	0.1
	Congruent-Control	–10 (8)	–0.1

To rule out that there was already a difference in bare pronunciation times between targets, we performed a 2 (Prime Duration) x 2 (Target Type) Repeated Measures ANOVA for the neutral (or no-)prime condition for both prime durations. We found that bare naming latencies did not depend on whether the target was transcribed in the KUN or ON reading (all $F_s < 1$), there was no interaction between Prime Duration and Prime Type (all $F_s < 1$). As RTs are indistinguishable when there is a neutral prime preceding the targets we decided to perform all subsequent analyses without the neutral condition. There was a main effect of Prime Duration. Introducing a backward mask of 50 ms resulted in response latencies 19 ms faster overall, $F(1,39) = 23.37$, $MSe = 1280,14$, $p < .001$;

$F(1,56) = 323.53$, $MSe = 88.12$, $p < .001$. However, there was no interaction between Prime Duration and any of the other variables (all $F_s < 1$ except the F-value in the 3-way interaction in the participant analysis between Prime Duration, Target Type and Prime Type, $F(1,39) = 2.52$, ns; $F(2,56) < 1$). As Prime Duration did not interact with Target Type or Prime Type, we collapsed the data over the two prime durations. On these collapsed data we performed a 2 (Target Type: KUN/ON) by 2 (Prime Type: Repetition, Control) Repeated Measures ANOVA. There was no effect of Target Type (KUN or ON), $F(1,39) = 3.120$, ns; $F(2,56) < 1$. However, there was a significant effect of Prime Type, $F(1,39) = 220.84$, $MSe = 35.37$, $p < .001$; $F(2,56) = 48.1$, $MSe = 132.94$, $p < .001$ and a significant interaction between Target Type and Prime Type in the participant analysis, $F(1,39) = 22.6$, $MSe = 35.0$, $p < .001$; $F(2,56) = 2.41$, ns. Planned comparisons show that KUN targets were named 18 ms faster when preceded by a Related Prime compared to a Control Prime, $t(39) = 13.5$, $SD = 8.6$, $p < .001$; $t(29) = 5.9$, $SD = 16.6$, $p < .001$ and ON targets were named 10 ms faster when preceded by a Related Prime compared to a Control Prime, $t(39) = 7.4$, $SD = 8.2$, $p < .001$; $t(29) = 3.9$, $SD = 16.0$, $p < .001$.

In this experiment we obtained evidence that a single kanji prime can cause facilitation in multiple katakana targets as e.g. both マチ /machikun/ and チ ヨ ウ /chouon/ show faster RTs when preceded by 町 compared to 式. It seems therefore that the prime 町, although presented briefly (50 ms), has activated both its pronunciations. There is a stronger facilitation effect for the KUN reading which is likely due to the fact that the kanji primes were presented in isolation, which usually entails using the KUN reading.

One may, however, argue that because the kanji primes were selected such that they adhered to an approximately balanced RP, half of the participants may have preferred one specific reading and the other half the other reading, and hence we obtained priming effects for both targets. However, it turns out that this neither complies with the experimental item analysis nor with the results of the questionnaire. If we look at an item-by-item basis, then for stimuli such as マチ /machikun/ (primes: 町 vs. 式) only one out of forty participants did not show a priming effect, and for チ ヨ ウ /chouon/ (with the same prime pairs) this number was three out of forty. Moreover, the

questionnaire demonstrated that there was strong consensus (> 90%) about the stand-alone pronunciation of the kanji (i.e. participants transcribed 町 as /machi_{kun}/ and not as /chou_{on}/), indicating that the priming effect for /chou_{on}/ is not due to the fact that this was the preferred stand-alone reading for participants.

It is conceivable that the obtained facilitation effect for both targets is only present in kanji which do not have a strong primary reading, as suggested by the data of the Kayamoto et al. (1998) study. Therefore, in Experiment 2 we examine whether or not evidence for activation of multiple pronunciations can also be obtained for kanji which have a bias towards one of the readings. Furthermore, control primes in Experiment 1 typically took only one reading (e.g. 式), which might have been responsible for the obtained priming effect. This was resolved in Experiment 2 where both congruent and control primes take multiple readings.

Experiment 2: Reading aloud katakana strings preceded by kanji primes with Equal and high-ON/high-KUN readings

In order to ascertain that the findings of Experiment 1 could be replicated, Experiment 2 also employed repetition primes which have an equal RP as well as stimuli which have a bias towards ON- or KUN reading. This offers the advantage of a comparison between these three sets of kanji within the same group of participants. Experiment 2 sought also to resolve some important issues which can be raised concerning Experiment 1, such as (1) avoiding excessive repetitions and (2) using control primes which also take more readings. As Experiment 1 did not show any interaction between Prime Duration and another variable, the backward mask was left out to avoid unnecessary repetitions. We hypothesize that when kanji primes are selected which have a bias towards a specific reading (KUN or ON), the prime will spread more activation to that reading compared to the other (unbiased) reading which will result in more facilitation by that prime for its biased transcription target.

Participants. Twenty-eight undergraduate students from Nagoya University (12 female; mean age: 19 years; SD = 3.6) took part in Experiment 2 in exchange for financial compensation. All

participants were native speakers of Japanese and had normal or corrected-to-normal vision.

Stimuli. Three Prime Bias groups of kanji were created. Sixteen kanji primes with two possible pronunciations were selected which adhered to a Reading Preference (RP; Tamaoka & Makioka, 2004) of between 40% and 60% (Equal-kanji; mean sum: 53.6%). Furthermore, sixteen kanji primes were selected which had a RP for the KUN reading between 70% and 90% (high-KUN kanji; mean sum: 78.0%) and also 16 kanji primes which had a RP between for the ON reading between 70% and 90% (high-ON kanji; mean sum 79.8%). There were no differences between the groups in their respective repetition and control kanji primes regarding summed mean frequency, $F(4,90) = 1.1$, ns and summed mean stroke complexity, $F(4,90) = 1.4$, ns. In general, the average ON-transcribed katakana target words' frequency of occurrence was higher compared to KUN target words which is due to kanji homophony being much higher for ON compared to KUN readings (e.g. the word /shin_{on}/ occurs more frequently than the string /mori_{kun}/; Tamaoka, 2005). However, the mean summed homophone count was statistically not different between the ON target items in the Equal group (13.2), high-ON group (11.8) and high-KUN group (17.6), $F(2,45) = 1.2$, ns. Appendix B provides an overview of all stimuli used in Experiment 2. Design. A 3 (Prime Bias, i.e. Equal, high-ON and high-KUN) x 2 (Target Type, i.e. KUN and ON transcribed target) x 2 (Prime Type, i.e. Congruent and Control) within participants factorial design was implemented. Each participant was subjected to half of the experimental trials (equaling 96 trials) showing each target (KUN and ON) only once to avoid unnecessary repetition. For each participant individually a pseudo-randomized list was created such that phonologically or semantically related primes or targets had at least a distance of two trials to avoid unintended priming effects and conditions appeared equally often. Across participants, the design was complete.

Procedure. In all reported experiments the software package E-prime 2.0 combined with a voice key was used for stimulus presentation and data acquisition. Participants were seated approximately 60 cm from a 17 inch LCD computer screen (Eizo Flexscan P1700; 60Hz) in a quiet room at Nagoya University and tested individually. A trial comprised the presentation of a fixation

cross for 1,000 ms followed by a forward mask for 500 ms, and subsequently a kanji prime of 50 ms replaced immediately by the katakana target word (maximally three characters long), which disappeared when the participant responded or after maximally 2,000 ms. In between trials, a random inter-stimulus-interval of 400 – 800 ms was introduced to avoid expectancy effects. Naming latencies were measured from target onset. Participants were specifically instructed to respond as fast as possible while avoiding errors. They were not informed about the presence of the prime. After the experiment informal interviewing showed that participants were found to be unable to recognize the primes under the masking conditions used in this experiment.

Results. Naming latencies exceeding 2.5 standard deviations per participant per prime duration were counted as outliers (comprising 1.2% of the data). Error rates in Experiment 2 were low (0.8%), and again distributed approximately equally across conditions; therefore they were not analyzed. In the F2 analysis, Prime Bias and Target Type were considered to be between-item variables. In Table 2, mean RTs and error rates per condition are reported.

Prime Bias	Prime Type	Target Type			
		KUN		ON	
		RT (SD)	%E	RT (SD)	%E
Equal	Congruent (e.g. 町)	449 (59)	0.0	450 (62)	0.0
	(マチ/チョウ) Control (e.g. 沢)	468 (73)	0.1	464 (63)	0.0
	Congruent-Control	-19 (28)	-0.1	-14 (23)	0.0
KUN Biased	Congruent (e.g. 森)	453 (62)	0.0	456 (60)	0.0
	(モリ/シン) Control (e.g. 鉄)	466 (62)	0.1	460 (62)	0.2
	Congruent-Control	-13 (26)	-0.1	-4 (26)	-0.2
ON Biased	Congruent (e.g. 罪)	456 (54)	0.1	460 (58)	0.1
	(ツミ/ザイ) Control (e.g. 則)	479 (65)	0.0	472 (62)	0.1
	Congruent-Control	-23 (22)	-0.1	-12 (22)	0.0

Mean RTs were submitted to a 3 (Prime Bias: Equal, high-KUN, high-ON) x 2 (Target Type: KUN or ON) x 2 (Prime Type: Congruent vs. Control) Repeated Measures ANOVA. There was a main effect of Prime Type, $F(1,27) = 35.7$, $MSe = 499.9$, $p < .001$; $F(1,90) = 40.5$, $MSe = 236.3$, $p < .001$, reflecting the fact that congruent primes induced a 14 ms facilitation effect compared to the control primes, and there was a main effect of Prime Bias, $F(2,54) = 10.2$, $MSe = 270.6$, $p < .001$; $F(2,90) = 4.1$, $MSe = 419.9$, $p < .05$, reflecting the fact that targets preceded by high-ON primes were named about 9 ms slower compared to the other two prime bias conditions.

There was a significant interaction between Prime Bias and Prime Type in the participant analysis, $F(2,54) = 3.8$, $MSe = 188.5$, $p < .05$, but not the item analysis, $F(2,90) = 1.4$, $MSe = 236.3$, ns,

reflecting the fact that there was less priming in the high-KUN prime condition compared to the other prime-bias conditions, and there was a marginal interaction between Prime Type and Target Type, $F(1,27) = 4.1$, $MSe = 341.3$, $p = .052$; $F(1,90) = 3.1$, $MSe = 236.3$, $p = .078$, showing that KUN targets obtained more priming compared to ON targets. All other interactions, including the three-way interaction between Prime Bias, Target Type, and Prime Type did not reach significance, all $F_s < 1$.

In the equal prime bias group we replicated the main finding of Experiment 1 in that congruent primes facilitated the naming of both the KUN and ON readings. For example, マチ /machi_{kun}/ and チ ヨ ウ /chou_{on}/ showed significant facilitation when preceded by a congruent prime in comparison to a control prime (e.g. 町 compared to 式). The lack of a Prime Bias x Prime Type x Target Type and of a Prime Bias x Target Type interaction indicates that both KUN and ON readings are facilitated by a congruent compared to control prime irrespective of the prime bias condition. However, the Target Type and Prime Type interaction demonstrates that the priming was always greater for KUN targets. A plausible explanation of this finding is that kanji primes when presented in isolation usually take the KUN reading. Given this characteristic, an overall bias towards the KUN reading – and hence more priming of the KUN reading – was to be expected. The important finding is, of course, that despite of this KUN bias in reading isolated kanji, ON readings were significantly facilitated as well.

The significant interaction between Prime Bias and Prime Type shows that there was less priming in the high-KUN Prime Bias condition. A likely account of this finding is that as primes were standing alone (hence favoring the KUN reading) the priming of the KUN reading had already reached its maximum (and is for that reason not larger than in the other two Prime Bias conditions), whereas the bias towards KUN did cause less spreading of activation towards its ON reading.

General Discussion

Two experiments were performed in order to establish whether multiple pronunciations for Japanese kanji become active during reading aloud. We employed two masked priming paradigms using katakana transcriptions as targets for each of the possible readings of

the congruent kanji prime. We hypothesized that if multiple readings for kanji are activated simultaneously, it should be possible to use these primes to facilitate multiple targets (by comparing congruent to control primes for each target). An additional hypothesis was that if the alternative readings were not dominant, priming may be less or even absent. The results of both experiments show that indeed multiple katakana targets are facilitated by the same kanji prime indicating that multiple readings were activated.

Katakana targets for all KUN transcribed readings were read significantly faster when preceded by a congruent prime and this priming effect was greater than that for ON transcribed readings in both experiments. Therefore, it seems warranted to conclude that when Kanji characters are presented in isolation, phonological activation is stronger for KUN-readings than for ON-readings. When turning to the ON readings, we see that katakana targets for all but one ON-transcribed reading were read significantly faster when preceded by a congruent prime. The only exception was the condition in Experiment 2 in which the ON reading was preceded by a kanji with a high KUN bias. This finding indicates that the amount of activation received by the ON reading of a kanji character is modulated by the strength of the link between a specific kanji character and its corresponding ON reading.

Our findings are important because they show that kanji characters activate multiple phonological representations. This leads to the question whether or not the two representations compete for selection. In answering that question it is important to note that the priming effects we reported surfaced in experiments in which there was no direct task on the kanji (as they were masked primes). Instead, in our task the activation spreading from a kanji prime to the target reading (KUN or ON) only adds to the large amount of activation that this target reading (KUN or ON) receives from the katakana target. Given these conditions, the activation of the competing reading (the reading that does not correspond to the katakana target) is most likely insufficient to delay the selection of the correct reading. This situation differs from the one in the Kayamoto et al. (1998) study, discussed in the introduction, in which participants were asked to name the kanji targets directly. In this situation, the two readings of the target only receive activation from the target, which makes it more likely that the activation of the incorrect, competing reading delays the selection of

the correct one. A recent study by Verdonschot, La Heij, and Schiller (2010) showed that when kanji characters need to be named, superimposed, semantically related context pictures facilitated naming in comparison to unrelated pictures. In line with the conclusion from Kayamoto and colleagues, these authors proposed a selection cost for multiple reading kanji (in contrast to Chinese hanzi) especially when the RP is close to Equal (50%) thereby allowing context pictures to exert an effect.

Our current findings need to be further extended by examining the processing of compound kanji where ON-reading priming can be studied without a stand-alone KUN influence. Furthermore, parametric mapping of all specific readings for kanji characters such as 上 /ue_{kun}/ (which combined with other kanji or hiragana can be pronounced in no less than six different ways including: /ue/, /uwa/, /jyou/, /nobo/, /a/ or /kami/) can lead to more detailed insights into how pronunciations are accessed when there are many alternatives.

It is furthermore important to notice that studies on the multiple readings of Japanese kanji can possibly inform the debate about decomposed vs. full-form storage. It has been shown before that cross-linguistic differences lead to contrasting results in this debate, e.g. presence and absence of compounds' constituent frequency effects (e.g. Bien et al., 2005; Chen & Chen, 2006; Alario, Perre, Castel, & Ziegler, 2007; Janssen et al., 2008). The likely importance of including Japanese kanji into these cross-linguistic comparisons, besides the different consistency effects found in Japanese (Wydell et al., 1995; Kayamoto et al., 1999; Fushimi et al., 1999), is supported by findings by Tamaoka and Hatsuzuka (1995) who also manipulated the frequency of individual kanji while controlling for overall compound frequency. In their experiments, compound words were presented in four conditions, i.e. (1) left and right kanji were both high frequency, e.g. 運送 /un_{on}.sou_{on}/ "transport"; (2) left was low and right was high frequency, e.g. 儀式 /gi_{on}.shiki_{on}/ "ceremony"; (3) left was high and right was low, e.g. 軍曹 /gun_{on}.sou_{on}/; or (4) both were low frequency, e.g. 我慢 /ga_{on}.man_{on}/. They found that reading aloud compound words with a high first constituent frequency was faster (30 ms) compared to low frequency, while this effect was absent for the second constituent in reading aloud (but present in a lexical decision task). However, Tamaoka and Hatsuzuka did not control for RP and

total number of readings for each constituent whereas our findings indicate this might be an important additional factor.

Overall, we conclude that when Japanese native speakers encounter a kanji character, multiple readings receive activation. This multiple activation can induce priming, as in our present study, or interference (Kayamoto et al. 1999; Verdonchot et al., 2010) depending on the relative strengths of alternative readings and the task at hand.

References

- Alario, F.-X., Perre, L., Castel, C., & Ziegler, J. C. (2007). The role of orthography in speech production revisited. *Cognition*, 102, 464-475.
- Bien, H., Levelt, W. J. M., & Baayen, H. (2005). Frequency effects in compound production. *Proceedings of the National Academy of Sciences*, 102, 17876-17881.
- Brown, G. D. A. (1987). Resolving inconsistency: A computational model of word naming. *Journal of Memory and Language*, 26, 1-23.
- Caramazza, A. (1997). How many levels of processing are there in lexical access? *Cognitive Neuropsychology*, 14, 177-208.
- Chen, T-M., & Chen, J-Y. (2006). Morphological encoding in the production of compound words in Mandarin Chinese. *Journal of Memory and Language*, 54, 491-514.
- Folk, J. R., & Morris, R. K. (1995). The use of multiple lexical codes in reading: Evidence from eye movements, naming time, and oral reading. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 21, 1412-1429.
- Forster, K. I., & Davis, C. (1984). Repetition priming and frequency attenuation in lexical access. *Journal of Experimental Psychology : Learning, Memory, and Cognition*, 10, 680-698.
- Fushimi, T., Ijuin, M., Patterson, K., & Tatsumi, I. F. (1999). Consistency, frequency, and lexicality effects in naming Japanese Kanji. *Journal of Experimental Psychology: Human Perception and Performance*, 25, 382-407.
- Fushimi, T., Komori, K., Ikeda, M., Patterson, K., Ijuin, M., & Tanabe, H. (2003). Surface dyslexia in a Japanese patient with semantic dementia: Evidence for similarity-based orthography-to-phonology translation. *Neuropsychologia*, 41, 1644-1658.
- Glushko, R. J. (1979). The organization and activation of orthographic knowledge in reading aloud. *Journal of Experimental Psychology: Human Perception and Performance*, 5, 674-691.
- Gottlob, L.R., Goldinger, S.D., Stone, G.O., & Van Orden, G.C. (1999). Reading homographs: Orthographic, Phonologic, and Semantic Dynamics. *Journal of Experimental Psychology: Human Perception and Performance*, 25, 561-574.

- Hino, Y., Lupker, S. J., Ogawa, T., & Sears, C. R. (2003). Masked repetition priming and word frequency effects across different types of Japanese scripts: An examination of the lexical activation account. *Journal of Memory and Language*, 48, 33-66.
- Janssen, N., Bi, Y., & Caramazza, A. (2008). A tale of two frequencies: Determining the speed of lexical access in Mandarin Chinese and English compounds. *Language and Cognitive Processes*, 23(7), 1191-1223
- Kawamoto, A. H., & Zemplide, J. (1992). Pronunciation of homographs. *Journal of Memory and Language*, 31, 349-374.
- Kayamoto, Y., Yamada, J., & Takashima, H. (1998). The consistency of multiple-pronunciation effects in reading: The case of Japanese logographs. *Journal of Psycholinguistic Research*, 27, 619-637.
- Levelt, W. J. M., Roelofs, A., & Meyer, A. S. (1999). A theory of lexical access in speech production. *Behavioral and Brain Sciences*, 22, 1-75.
- Nomura, Y. (1978). Kanji no johoo shori: Ondoku-kundoku to imi no fuyo [The information processing of Chinese characters (kanji): Chinese reading, Japanese reading and the attachment of meaning]. *Japanese Journal of Psychology*, 49, 190-197.
- Schiller, N. O. (1998). The effect of visually masked syllable primes on the naming latencies of words and pictures. *Journal of Memory and Language*, 39, 484-507.
- Schiller, N. O. (2004). The onset effect in word naming. *Journal of Memory and Language*, 50, 477-490.
- Seidenberg, M. S., Waters, G. D., Barnes, M. A., & Tanenhaus, M. K. (1984). When does irregular spelling or pronunciation influence word recognition? *Journal of Verbal Learning and Verbal Behavior*, 23, 383-404.
- Stanovich, K. E., & Bauer, D. (1978). Experiments on the spelling to sound regularity effect in word recognition. *Memory & Cognition*, 6, 410-415.
- Tamaoka, K., & Hatsuzuka, M. (1995). Kanji niji jukuo no shori ni okeru kanji shiyō hindo no eikyō [The effects of kanji printed-frequency on processing Japanese two-morpheme compound words], *The Science of Reading*, 39: 121-137 (in Japanese).

Tamaoka, K., Kirsner, K., Yanase, Y., Miyaoka, Y., & Kawakami, M. (2002). A Web-accessible database of characteristics of the 1,945 basic Japanese kanji. *Behavior Research Methods, Instruments & Computers*, 34, 260-275.

Tamaoka, K., & Makioka, S. (2004). New figures for a Web-accessible database of the 1,945 basic Japanese kanji, fourth edition. *Behavior Research Methods, Instruments & Computers*, 36, 548-558.

Tamaoka, K. (2005). The effect of morphemic homophony on the processing of Japanese two-kanji compound words. *Reading and Writing*, 18, 281-302.

Yasunaga, M. (1981). Jooyoo kanji-hyoo ga umareru-made [A background history of the Jooyoo kanji list]. *Gengo seikatsu* [Language Life], 355, 24-31.

Yokoyama, S., Sasahara, H., Nozaki, H., & Long, E. (1998). Shinbun denshi media no kanji: Asahi shimbun no CD-ROM-ni yoru kanji hindo hyoo. Tokyo: Sanseido [Japanese kanji in the newspaper media: Kanji frequency index from the Asahi newspaper on CD-ROM].

Verdonschot, R.G., La Heij, W., & Schiller, N.O. (2010). Semantic context effects when naming Japanese kanji, but not Chinese hànzi. *Cognition*, 115, 512-518.

Waters, G. S., & Seidenberg, M. S. (1985). Spelling-sound effects in reading: Time-course and decision criteria. *Memory & Cognition*, 13, 557-572.

Wydell, T. N., Butterworth, B. L., & Patterson, K. E. (1995). The inconsistency of consistency effects in reading: Are there consistency effects in Kanji? *Journal of Experimental Psychology: Language, Memory and Cognition*, 21, 1155-1168.

Appendix A
Stimulus Materials from Experiment 1

KUN Target	ON Target	Congruent Prime	Control Prime
サカナ ('sakana')	ギョ ('gyo')	魚	軸
タマ ('tama')	ギョク ('gyoku')	玉	僚
カガミ ('kagami')	キョウ ('kyoo')	鏡	菌
ノキ ('noki')	ケン ('ken')	軒	脈
ネ ('ne')	コン ('kon')	根	録
ヤマ ('yama')	サン ('san')	山	理
フダ ('fuda')	サツ ('satsu')	札	王
バ ('ba')	ジョウ ('joo')	場	議
クラ ('kura')	ゾウ ('zoo')	蔵	券
ミ ('mi')	シン ('shin')	身	銀
ムラ ('mura')	ソン ('son')	村	論
ムシ ('mushi')	チュウ ('chuu')	虫	卓
マチ ('machi')	チョウ ('choo')	町	式
トリ ('tori')	チョウ ('choo')	鳥	託
タ ('ta')	デン ('den')	田	百
シマ ('shima')	トウ ('too')	島	官
ケ ('ke')	モウ ('moo')	毛	到
カベ ('kabe')	ヘキ ('heki')	壁	郡
ワ ('wa')	リン ('rin')	輪	順
カワ ('kawa')	ヒ ('hi')	皮	誕
ハル ('haru')	シュン ('shun')	春	材
カミ ('kami')	ハツ ('hatsu')	髪	晩
ハタ ('hata')	キ ('ki')	旗	詩
ヒル ('hiru')	チュウ ('chuu')	昼	舍
ソコ ('soko')	テイ ('tei')	底	刊
ハラ ('hara')	フク ('fuku')	腹	毒
サマ ('sama')	ヨウ ('yoo')	様	課
タケ ('take')	ガク ('gaku')	岳	糖
フエ ('fue')	テキ ('teki')	笛	冗

Appendix B
Stimulus Materials from Experiment 2

	KUN target	ON target	Congruent Prime	Control Prime
Equal	ネ 'ne'	コン 'kon'	根	右
	タマ 'tama'	ギョク 'gyoku'	玉	油
	シマ 'shima'	トウ 'too'	島	先
	カベ 'kabe'	ヘキ 'heki'	壁	晴
	ワ 'wa'	リン 'rin'	輪	豊
	ハタ 'hata'	キ 'ki'	旗	乳
	ソコ 'soko'	テイ 'tei'	底	患
	サマ 'sama'	ヨウ 'yoo'	様	評
	マチ 'machi'	チョウ 'choo'	町	沢
	フダ 'fuda'	サツ 'satsu'	札	薬
	ヤマ 'yama'	サン 'san'	山	小
	ミ 'mi'	シン 'shin'	身	優
	ムラ 'mura'	ソン 'son'	村	計
	クラ 'kura'	ゾウ 'zoo'	蔵	林
	ハル 'haru'	シュン 'shun'	春	鮮
	カワ 'kawa'	ヒ 'hi'	皮	勉
KUN Biased	モリ 'mori'	シン 'shin'	森	鉄
	カワ 'kawa'	セン 'sen'	川	要
	ユメ 'yume'	ム 'mu'	夢	預
	ハシラ 'hashira'	チュウ 'chuu'	柱	誠
	ハシ 'hashi'	キョウ 'kyou'	橋	歩
	サカ 'saka'	ハン 'han'	坂	壊
	マド 'mado'	ソウ 'sou'	窓	斉
	ハナ 'hana'	カ 'ka'	花	音
	スジ 'suji'	キン 'kin'	筋	契
	イト 'ito'	シ 'shin'	糸	怖
	ハナ 'hana'	ビ 'bi'	鼻	呉
	キリ 'kiri'	ム 'mu'	霧	潤
	シオ 'sio'	エン 'en'	塩	捨
	ユキ 'yuki'	セツ 'setsu'	雪	範
	スミ 'sumi'	ボク 'boku'	墨	穩
	ウデ 'ude'	ワン 'wan'	腕	昼

Appendix B (continued)
Stimulus Materials from Experiment 2

	KUN target	ON target	Congruent Prime	Control Prime
ON Biased	タケ 'take'	チク 'chiku'	竹	異
	ツミ 'tsumi'	ザイ 'zai'	罪	則
	テラ 'tera'	ジ 'ji'	寺	削
	クルマ 'kuruma'	シャ 'sha'	車	西
	タビ 'tabi'	リョ 'ryo'	旅	婚
	ミズ 'mizu'	スイ 'sui'	水	投
	メシ 'meshi'	ハン 'han'	飯	括
	ヌノ 'nuno'	フ 'fu'	布	透
	アサ 'asa'	マ 'ma'	麻	昔
	ワケ 'wake'	ヤク 'yaku'	訳	絡
	ウシ 'ushi'	ギユウ 'gyuu'	牛	童
	キミ 'kimi'	クン 'kun'	君	酒
	ニワ 'niwa'	テイ 'tei'	庭	貴
	モモ 'momo'	トウ 'too'	桃	囚
	タテ 'tate'	ジュウ 'juu'	縦	桜
	チ 'chi'	ケツ 'ketsu'	血	希

Chapter 3: Semantic context effects when naming Japanese kanji, but not Chinese hànzì

This chapter is based on: Verdonschot, R. G., La Heij, W., Schiller, N.O. (2010). Semantic context effects when naming Japanese kanji, but not Chinese hànzì, *Cognition*, 115, 512-518.

Abstract

The process of reading aloud bare nouns in alphabetic languages is immune to semantic context effects from pictures. This is accounted for by assuming that words in alphabetic languages can be read aloud relatively fast through a sub-lexical grapheme-phoneme conversion (GPC) route or by a direct route from orthography to word form. We examined semantic context effects in a word-naming task in two languages with logographic scripts for which GPC cannot be applied: Japanese kanji and Chinese hànzi. We showed that reading aloud bare nouns is sensitive to semantically related context pictures in Japanese, but not in Chinese. The difference between these two languages is attributed to processing costs caused by multiple pronunciations for Japanese kanji.

Semantic context effects when naming Japanese kanji, but not Chinese hànzi

Models of word production distinguish various processing levels in object naming including conceptualization, retrieval of syntactic features, word-form encoding, and articulation (for overviews see e.g., Caramazza, 1997; Dell, 1986; Levelt, Roelofs, & Meyer, 1999; Roelofs, 1992, 2008). Some of these models are based on results obtained with both the picture-word interference paradigm (PWI) paradigm (e.g. Caramazza & Costa, 2000; Schriefers, Meyer, & Levelt, 1990; Starreveld & La Heij, 1995), in which pictures have to be named in the context of distractor words, and the “reversed” PWI paradigm in which words have to be read in the context of distractor pictures (e.g., Roelofs, 1992; 2006). An influential model of context effects in word production and word reading, WEAVER++ (Roelofs, 1992; Roelofs, Meyer, & Levelt, 1996; Levelt et al., 1999; Indefrey & Levelt, 2004; Roelofs, 2006), assumes that in alphabetic languages a visually presented word can be processed along three different routes, depicted in Figure 1 (adapted from Roelofs, Meyer, & Levelt, 1996).

(1) A sublexical route from a graphemic representation to a phonemic representation (GPC), evidenced by non-words inducing reliable phonological facilitation in picture naming (Lupker, 1982).

(2) A route from orthographic to phonological word-form representations, evidenced by form-related distractor words speeding up picture naming (Koester & Schiller, 2008; Zwitserlood, Bölte, & Dohmes, 2000).

(3) A route from orthographic word representations to a word’s lexical-syntactic representation, supported by semantic interference effects (e.g., Schriefers et al., 1990) and gender/determiner congruency effects (e.g., Schiller & Caramazza, 2003) in picture naming.

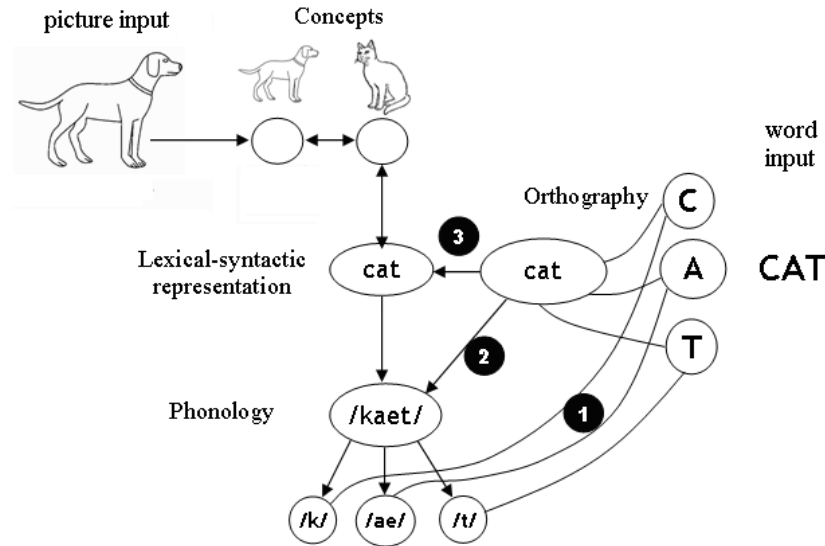


Figure 1. Three processing routes of a visually presented word.
Adapted from Roelofs, Meyer, & Levelt (1996).

Although all three routes may be used, word-naming latencies are assumed to be determined by the fastest route, which is assumed to be Route 2. This assumption is based on the observation that bare word reading is immune to semantic context effects (Glaser & Dünghoff, 1984; La Heij, Happel, & Mulder, 1990; Roelofs, 2006; but see Roelofs, 2003). Semantic context effects in reading words are only observed when information at the lemma level is required. Roelofs (2006), for instance, reported that context pictures only induced semantic facilitation in a word reading task when Dutch participants were asked to respond to a single target word using a determiner-noun phrase (e.g., “de kat”, the cat), requiring syntactic information at the lemma level (Route 3). This finding can be accounted for by assuming spreading of activation from the picture concept (DOG) to the word concept (CAT) and from there to the word’s lexical-syntactic representation (see Figure 1).

Most research on context effects in picture and word processing has been performed in alphabetic languages with scripts including sets of symbols (letters), which approximate sounds (and phonemes). However, other scripts such as logographic Chinese hànzi

and Japanese kanji characters represent words or morphemes rather than individual sounds or phonemes. These different properties may lead to different results in PWI tasks and reversed PWI tasks.

Modern Chinese employs a logographic script called hànzi. Although some characters are iconic (like 山 /shan1/, 'mountain'), most characters are ideograms representing monosyllabic units with an elementary meaning (Chen, 1992; Taft & Zhu, 1995). Many, but not all characters contain elements called radicals to indicate semantic group membership, mostly presented on the left side of the character (such as 木 'tree' in 松 'pine tree'), and radicals which are cues to the pronunciation of the character, mostly presented on the right side (such as 半 'half' in 伴 'partner' both pronounced /ban4/). Generally, Chinese words consist of two, but sometimes more characters that usually have a single pronunciation.

Modern Japanese, in contrast, employs three scripts, i.e. kanji, hiragana, and katakana. Kana characters, i.e. hiragana and katakana, were adapted from kanji to provide a means of representing native Japanese vocabulary, loanwords, proper names, and affixes. Historically, kanji were logographic characters imported from Chinese. In modern Japanese, they are used for representing words borrowed from Chinese, compounds of these words, and native Japanese vocabulary. There are two types of kanji pronunciations, i.e. ON-readings, derived from the original Chinese pronunciation, and KUN-readings originating from the Japanese pronunciation. For instance, the kanji character 上 is without context pronounced as /ue/ but has different pronunciations when it occurs as the stem of a verb (上る /nobo.ru/, 'to climb' or 上げる /a.geru/, 'to give'); when it occurs together with other kanji or kana nouns or adjectives, it can be pronounced as /ue/, /uwa/, /jyou/, or /kami/. More than 60% of the 1,945 basic kanji characters have different pronunciations in different contexts. This property of kanji might have processing consequences as shown by Kayamoto, Yamada and Takashima (1998) who compared reading aloud latencies of high frequency kanji with only one reading, such as 脳 /nou_{on}/ 'brain' with high frequency kanji which have multiple readings, such as 街 /machi_{kun}/ or /gai_{on}/ 'town'. Kanji with multiple readings were named slower, indicating that some processing cost was incurred compared to single reading kanji (Experiment 1). Such cost also emerged when mid-frequency kanji

were used, except when alternative readings were weak (Experiment 2) which is in line with findings by Wydell, Butterworth and Patterson (1995).

To recapitulate, Chinese and Japanese logographs differ in that Japanese kanji often have more than one pronunciation whereas Chinese hànzì generally have a single pronunciation. This difference may be reflected in processing differences during reading aloud. Being logographic languages, both Chinese hànzì and Japanese kanji cannot be processed via a GPC route (Siok, Perfetti, Jin, & Tan, 2004; Wydell et al., 1995). Within the model depicted in Figure 1, this excludes Route 1. If words in both languages are read via Route 3 (involving lexical-syntactic representations), we may expect semantic facilitation effects of context pictures similar to those observed by Roelofs (2006) in determiner-noun phrase production in Dutch. If, however, Chinese hànzì and Japanese kanji are read via a direct route to the word-form level (Route 2), predictions are dependent on model-specific assumptions. In a discrete model like WEAVER++, in which activated lexical-syntactic representations do not automatically spread activation to word forms, reading via the word form level should be unaffected by context pictures (Levelt et al., 1999). In contrast, in models assuming cascading of activation from lexical-syntactic to word form representations (e.g. Roelofs, 2008), context pictures may induce facilitation effects, provided that the activation from the conceptual level has enough time to affect processing at the word-form level. Given the processing costs due to resolving the correct pronunciation in Japanese kanji with multiple pronunciations (Kayamoto et al., 1998), discussed above, semantic facilitation may be larger when reading aloud Japanese kanji relative to Chinese hànzì.

To test these predictions we carried out a series of word-picture naming (i.e. reversed PWI) experiments using Japanese and Chinese stimuli in which to-be-named logograms were superimposed on semantically related or unrelated context pictures. To demonstrate that the potential absence of semantic effects in reading aloud is not due to the stimulus materials used, for both languages a standard picture-word interference (PWI) task (SOA 0 ms), using the same Japanese and Chinese materials, were administered after the reading aloud task to the same participants. In both PWI tasks semantic interference effects are expected.

Experiment: Semantic context effects in Japanese and Chinese word reading

Method.

Japanese participants. Twenty-four native Japanese speakers (14 female; mean age: 31.5 years; SD = 8.1) living in the Leiden and Amsterdam residential area in the Netherlands took part in the experiment. They had been living in the Netherlands on average for 4.5 years (SD = 5.2). The majority of the participants worked in a Japanese business-related environment and all participants reported to use kanji on a daily basis.

Chinese participants. Eighteen male college students from Dalian Maritime University (China) and two male non-university volunteers residing in Dalian took part in this experiment totalling 20 participants (mean age: 24 years; SD = 3.5).

Japanese Stimuli. Twenty target kanji having two or more readings were selected. Each kanji was paired with a semantically related and an unrelated context picture. Kanji – picture pairs were created such that semantically related and unrelated pictures occupied approximately the same screen area. It should be noted that the names of some of the context pictures are usually written in katakana (e.g. “spoon”), but since pictures enter the production system through the conceptual system (Roelofs, 1992; 2006; Caramazza, 1997; Levelt et al., 1999; see also Figure 1), this is irrelevant for our current issue. See Appendix and Figure 2 for an overview of the Japanese and Chinese stimuli.

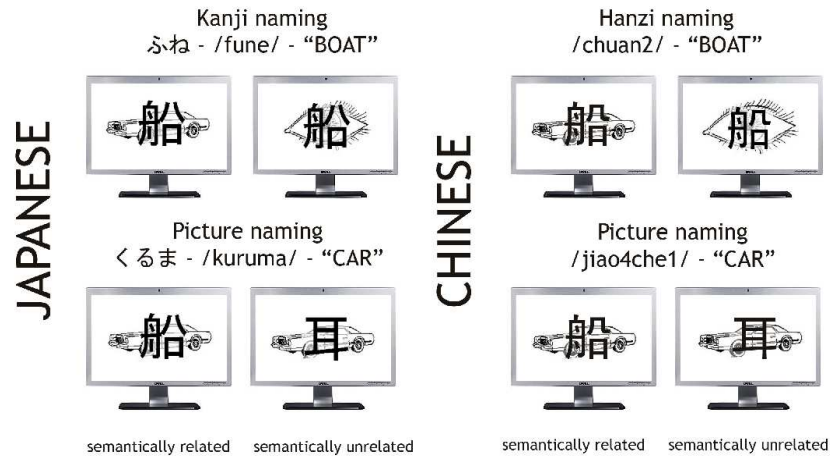


Figure 2. Examples of experimental stimuli.

Chinese Stimuli. Twenty Chinese target hànzi stimuli were selected. All but two hànzi matched the Japanese in the sense that they also consisted of a single character. No target-context pairs were phonologically overlapping. There was no significant difference in mean target frequency (per million) between Japanese (366) and Chinese stimuli (201), $t(19) = 1.72$, ns (taken from Yokoyama, Sasahara, Nozaki & Long, 1998, and Da, 2004, respectively). Reading aloud design. For both the Japanese and Chinese part a 2x2 within-subjects factorial design was implemented, with the factors SOA (0 ms, i.e. picture and word presented simultaneously, or -150 ms, i.e. picture first) and Relatedness (semantically related or unrelated context picture). Each participant was subjected to 80 naming trials presented in two blocks (one block per SOA). For each participant, pseudo-random lists were constructed per block such that there were at least two intervening trials between phonologically or semantically related characters or pictures. Across participants, the order of blocks was counterbalanced. Each block was preceded by two warm-up trials (not included in the analysis).

Reading aloud procedure. Participants were seated approximately 60 cm from a 17 inch CRT computer screen in a quiet room at Leiden University (Japanese participants) or Dalian Maritime University (Chinese participants) and tested individually. Trials consisted of a fixation point (1,000 ms) which was replaced by the

kanji/hànzì – picture pair until participants responded or after maximally 2,500 ms. The experimenter recorded whether or not a response was accurate, followed by an inter trial interval of 500 ms before the next trial started. Naming latencies were measured from target onset using a voice-key. Participants were instructed to respond as fast as possible while avoiding errors.

Picture naming procedure. For the picture naming experiment, participants first saw the to-be-named pictures on the screen with the corresponding names printed underneath. Subsequently, pictures appeared without any distractor and participants were asked to name the pictures to verify whether they used the intended name. Then the experiment proper started. Trials consisted of a fixation point (1,000 ms) followed and replaced by the picture – kanji/hànzì pair, which disappeared when participants responded or after maximally 2,500 ms. Following a response, the experimenter recorded whether or not the response was accurate before the next trial started.

Results picture naming. All analyses reported were carried out with participants (F1) and items (F2) as random variables. Naming latencies faster than 300 ms or exceeding 1,500 ms were treated as outliers (Japanese: 2.7%, Chinese: 1.8% of the data). An overview of the mean RTs and error rates is given in Table 1. A repeated measures analysis with one within-subjects factor (Relatedness) and one between-subjects factor (Language) was conducted. There was a main effect of Language, $F(1,42) = 11.64$, $MSe = 11,460.04$, $p < .001$; $F(1,38) = 21.21$, $MSe = 6,134.74$, $p < .001$, and Relatedness, $F(1,42) = 18.58$, $MSe = 893.84$, $p < .001$; $F(1,38) = 18.99$, $MSe = 890.77$, $p < .001$, but no interaction between Language and Relatedness, $F(1,42) = 1.02$, $MSe = 893.84$, ns; $F(1,38) = 1.02$, $MSe = 890.77$, ns, reflecting similar semantic interference effects for both languages. Planned comparisons demonstrated a 34 ms semantic interference effect for Chinese, $t(19) = 3.70$, $SD = 41.10$, $p < .01$; $t(19) = 3.67$, $SD = 43.70$, $p < .01$, and a 22 ms semantic interference effect for Japanese, $t(23) = 2.39$, $SD = 43.30$, $p < .05$; $t(19) = 2.46$, $SD = 40.70$, $p < .05$. These results confirm that in both languages our selected stimuli were able to elicit semantic interference.

Table 1
Mean Naming Latencies (in Milliseconds) and Error Rates (in %) in the Picture Naming Task as a Function of Language and Semantic Relatedness.

	Japanese		Chinese	
	RTs (SD)	%E	RTs (SD)	%E
Related	700 (70.4)	2.4	784 (88.1)	2.9
Unrelated	678 (64.4)	3.4	750 (92.3)	2.5
Effect (Related – Unrelated)	22 (43.3)	–1.0	34 (41.1)	0.4

Interestingly, this well-known effect of distractor words on picture naming in alphabetic languages has not been investigated in depth with Japanese kanji and Chinese hànzi (but for Chinese see Zhang & Weekes, 2009). Using a PWI task, Ishio (1990) failed to find semantic interference in the Japanese language. More recently, however, Iwasaki, Vinson, Vigliocco, Watanabe, and Arciuli (2008) reported robust semantic interference when naming actions in Japanese, and we demonstrate here that this is also the case for object naming in Japanese (and Chinese) using a standard PWI task.

Results reading aloud. RTs faster than 300 ms or exceeding 1,500 ms were treated as outliers (Japanese: 1.5%, Chinese: 1.1% of the data). An overview of the mean RTs and error rates is given in Table 2. One Japanese kanji, i.e. 肺 /hai/ ‘lung’ turned out to have only one pronunciation and was excluded from further analyses. A combined analysis of the two data sets revealed no main effect of Language, $F(1,42) < 1$; $F(1,37) = 1.75$, $MSe = 3,860.51$, ns. The main effect of SOA was significant in the items analysis, but not in the participants analysis, $F(1,42) < 1$; $F(1,37) = 4.74$, $MSe = 491.27$, $p < .05$, reflecting (by items) that at SOA –150 words were named 8 ms slower than at SOA 0 ms. There was a significant 11 ms facilitation effect of Relatedness, $F(1,42) = 23.20$, $MSe = 239.25$, $p < .001$; $F(1,37) = 11.43$, $MSe = 424.64$, $p < .001$. Two interactions yielded significant effects, i.e. Language x SOA in the items analysis, $F(1,42) = 1.59$, $MSe = 2,544.97$, ns; $F(1,37) = 8.33$, $MSe = 491.27$, $p < .01$, and Language x Relatedness, $F(1,42) = 10.30$, $MSe = 239.25$, $p < .01$; $F(1,37) = 5.51$, $MSe = 424.64$, $p < .05$. To

investigate these interactions in more detail, individual analyses were performed for both languages.

Japanese Results. Mean RTs were submitted to a 2 x 2 repeated measures ANOVA with SOA and Relatedness as within-subject factors. There was no interaction between SOA and Relatedness. There was a main effect of SOA in the analysis by items but not by participants, $F(1,23) = 2.85$, $MSe = 2,392.36$, ns; $F(1,18) = 12.61$, $MSe = 487.02$, $p < .01$, reflecting (by items) that at SOA –150 words were named 18 ms slower than at SOA 0 ms. Furthermore, there was a main effect of Relatedness, $F(1,23) = 26.94$, $MSe = 314.55$, $p < .001$; $F(1,18) = 22.84$, $MSe = 297.43$, $p < .001$. Kanji were named 19 ms faster in the context of a semantically related as compared to an unrelated picture. At SOA 0 ms, there was a significant 14 ms semantic facilitation effect, $t(23) = 2.75$, $SD = 24.53$, $p < .02$; $t(18) = 2.82$, $SD = 21.50$, $p < .02$, and at SOA –150 ms, this effect was 24 ms, $t(23) = 4.67$, $SD = 24.96$, $p < .001$; $t(18) = 3.44$, $SD = 30.30$, $p < .01$. As error rates were very low (overall < 1.1%), no error analysis was performed.

Table 2

Mean Naming Latencies (in Milliseconds) and Error Rates (in %) in the Reading Aloud Task as a Function of SOA and Semantic Relatedness of the Distractor Picture.

Japanese Kanji Naming				
	SOA –150		SOA 0	
	RTs (SD)	%E	RTs (SD)	%E
Related	531 (51.1)	1.0	519 (65.7)	0.8
Unrelated	555 (60.2)	1.3	533 (74.9)	1.3
Effect	–24 (25.0)	–0.3	–14 (24.5)	–0.5
Chinese Hànzì Naming				
	SOA –150		SOA 0	
	RTs (SD)	%E	RTs (SD)	%E
Related	519 (59.2)	1.5	525 (61.4)	3.0
Unrelated	526 (60.3)	2.5	525 (67.3)	1.3
Effect	–7 (18.9)	–1.0	0 (14.5)	1.7

Chinese Results. The analysis was identical to the one performed on the Japanese data. The main effects of SOA (both $F_s < 1$) and Relatedness, $F(1,19) = 1.92$, $MSe = 148.09$, ns; $F(1,19) < 1$, were not significant, and there was no interaction between SOA and Relatedness either, $F(1,19) = 1.61$, $MSe = 135.87$, ns; $F(1,19) < 1$. Error rates were very low (overall $< 2.1\%$) and therefore not analyzed.

Discussion

We reported four important empirical results in this study. First, a significant semantic interference effect induced by visually presented distractor words in Japanese picture naming. Second, an analogous semantic interference effect in Chinese picture naming. Third, a significant semantic facilitation effect induced by context pictures in Japanese word (kanji) naming at two SOAs (-150 ms and 0 ms). Fourth, the absence of such effects in Chinese word (hànzì) naming at the same SOAs.

The observation that naming Chinese hànzì does not show a semantic facilitation effect, despite the presence of a semantic interference effect in the corresponding picture-naming task, seems to rule out the hypothesis that words in this language are named via their lexical-syntactic representations. That is, Chinese hànzì are likely read via the direct route from orthography to phonology (Route 2). The presence of a semantic facilitation effect when reading Japanese kanji may be taken to suggest that our kanji characters are read via the lexical-semantic representation (Route 3). However, this interpretation is hard to reconcile with neuropsychological evidence indicating the use of a direct orthography to phonology route in reading kanji. Sasanuma, Sakuma and Kitano (1992) and Nakamura et al. (1998) showed that the ability of patients with Alzheimer's dementia to comprehend kanji deteriorated over time, while their ability to read kanji aloud was retained. More recently, Fushimi et al. (2003) reported a Japanese surface-dyslexic patient whose reading performance is best explained by assuming (a) an intact orthography-to-phonology route (Route 2 in Figure 1) and (b) a reduction of activation arriving from semantics.

Given these considerations, the most parsimonious interpretation of our findings is that the kanji characters used in our experiment activate multiple phonological representations (via Route 2 in Figure 1), which induces a processing delay that allows activation from the semantic system to affect response latencies. In contrast, in

reading Chinese hànzi only one phonological representation is activated and selected, allowing no time for activation from semantics to speed up this process. The observation that Japanese kanji words do not take more reading aloud time than Chinese hànzi cannot be taken as evidence against this proposal, as the kanji and hànzi words used in our experiments differed both in form and pronunciation.

In conclusion, although we cannot fully exclude the possibility that kanji characters are read via a lexical-syntactic route, our data are most parsimoniously accounted for by assuming that (a) logographic scripts, like alphabetic scripts, are read via a direct route from orthography to word-form representations and (b) this route is susceptible to semantic context effects when multiple mappings between orthography and word-form are possible.

References

- Caramazza, A. (1997). How many levels of processing are there in lexical access? *Cognitive Neuropsychology*, 14, 177-208.
- Caramazza, A., & Costa, A. (2000). The semantic interference effect in the picture-word interference paradigm: Does the response set matter? *Cognition*, 75, B51-B64.
- Chen, H. C. (1992). Reading Comprehension in Chinese: implication for character reading times. In H. C. Chen & O. J. L. Tzeng (Eds.), *Language processing in Chinese* (pp. 175-205). Amsterdam: North-Holland.
- Da, J. (2004). A corpus-based study of character and bigram frequencies in Chinese e-texts and its implications for Chinese language instruction. The studies on the theory and methodology of the digitalized Chinese teaching to foreigners. In Pu, Z., Xie, T., & Xu, J. (eds.), *Proceedings of the Fourth International Conference on New Technologies in Teaching and Learning Chinese* (pp. 501-511). Beijing: Tsinghua University Press.
- Dell, G. S. (1986). A spreading-activation theory of retrieval in sentence production. *Psychological Review*, 93, 283-321.
- Fushimi, T., Komori, K., Ikeda, M., Patterson, K., Ijuin, M., & Tanabe, H. (2003). Surface dyslexia in a Japanese patient with semantic dementia: Evidence for similarity-based orthography-to-phonology translation, *Neuropsychologia*, 41, 1644-1658.
- Glaser, W. R., & Dünghoff, F. J. (1984). The time course of picture-word interference. *Journal of Experimental Psychology: Human Perception and Performance*, 10, 640-654.
- Indefrey, P., & Levelt, W. J. M. (2004). The spatial and temporal signatures of word production components. *Cognition*, 92, 101-144.
- Ishio, A. (1990). Auditory-visual Stroop interference in picture-word processing. *The Japanese Journal of Psychology*, 61, 329-335.
- Iwasaki, N., Vinson, D. P., Vigliocco, G., Watanabe, M., & Arciuli, J. (2008). Naming actions in Japanese: Effects of semantic similarity and grammatical class. *Language and Cognitive Processes*, 23, 889-930.

Kayamoto, Y., Yamada, J., & Takashima, H. (1998). The Consistency of Multiple-Pronunciation Effects in Reading: The Case of Japanese Logographs. *Journal of Psycholinguistic Research*, 27, 619-637.

Koester, D., & Schiller, N. O. (2008). Morphological priming in overt language production: Electrophysiological evidence from Dutch. *NeuroImage*, 42, 1622-1630.

La Heij, W., Happel, B., & Mulder, M. (1990). Components of Stroop-like interference in word reading. *Acta Psychologica*, 73, 115-129.

Levelt, W. J. M., Roelofs, A., & Meyer, A. S. (1999). A theory of lexical access in speech production. *Behavioral and Brain Sciences*, 22, 1-75.

Lupker, S. J. (1982). The role of phonetic and orthographic similarity in picture-word interference. *Canadian Journal of Psychology*, 36, 349-367.

Nakamura, K., Meguro, K., Yamazaki, H., Ishizaki, J., Saito, H., Saito, N., Shimada, M., Yamaguchi, S., Shimada, Y., & Yamadori, A. (1998). Kanji predominant alexia in advanced Alzheimer's disease. *Acta Neurologica Scandinavica*, 97, 237-243.

Roelofs, A. (1992). A spreading-activation theory of lemma retrieval in speaking. *Cognition*, 42, 107-142.

Roelofs, A., Meyer, A. S., & Levelt, W. J. M. (1996). Interaction between semantic and orthographic factors in conceptually driven naming: Comment on Starreveld and La Heij (1995). *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 22, 246-251.

Roelofs, A. (2003). Goal-referenced selection of verbal action: Modeling attentional control in the Stroop task. *Psychological Review*, 110, 88-125.

Roelofs, A. (2006). Context effects of pictures and words in naming objects, reading words, and generating simple phrases. *Quarterly Journal of Experimental Psychology*, 59, 1764-1784.

Roelofs, A. (2008). Tracing attention and the activation flow in spoken word planning using eye movements. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 34, 353-368.

Sasanuma, S., Sakuma, N., & Kitano, K. (1992). Reading kanji without semantics: Evidence from a longitudinal study of dementia. *Cognitive Neuropsychology*, 9, 465-486.

Schiller, N. O., & Caramazza, A. (2003). Grammatical feature selection in noun phrase production: Evidence from German and Dutch. *Journal of Memory and Language*, 48, 169-194.

Schriefers, H., Meyer, A. S., & Levelt, W. J. M. (1990). Exploring the time course of lexical access in speech production: Picture-word interference studies. *Journal of Memory and Language*, 29, 86-102.

Siok, W. T., Perfetti, C. A., Jin, Z., & Tan, L. H. (2004). Biological abnormality of impaired reading is constrained by culture. *Nature*, 431, 71-76.

Starreveld, P. A., & La Heij, W. (1995). Semantic interference, orthographic facilitation and their interaction in naming tasks. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 21, 686-698.

Taft, M., & Zhu, X. (1995). The representation of bound morphemes in the lexicon: A Chinese study. In L. Feldman (Ed.), *Morphological aspects of language processing* (pp. 293-315). Hillsdale, NJ: Lawrence Erlbaum Associates.

Wydell, T. N., Butterworth, B. L., & Patterson, K. E. (1995). The inconsistency of consistency effects in reading: Are there consistency effects in Kanji? *Journal of Experimental Psychology: Language, Memory and Cognition*, 21, 1155-1168.

Yokoyama, S., Sasahara, H., Nozaki, H., & Long, E. (1998). Shinbun denshi media no kanji: Asahi shimbun no CD-ROM-ni yoru kanji hindo hyoo. [Japanese kanji in the newspaper media: Kanji frequency index from the Asahi newspaper on CD-ROM]. Tokyo: Sanseido.

Zhang, Q. F., & Weekes, B.S. (2009). Orthographic facilitation effects on spoken word production: Evidence from Chinese. *Language and Cognitive Processes*, 24, 1082-1096.

Zwitserslood, P., Bólte, J., & Dohmes, P. (2000). Morphological effects on speech production: evidence from picture naming. *Language and Cognitive Processes*, 15, 563-591.

Appendix

Japanese Kanji Target	Chinese Hànzì Target	Japanese Related Picture	Chinese Related Picture	Japanese Unrelated Picture	Chinese Unrelated Picture
犬 (dog; inu/ken)	狗 (dog; gou3)	cat (neko)	horse (ma3)	harp (haapu)	harp (shu4qin2)
牛 (cow; ushi/gyuu)	牛 (cow; niu2)	sheep (hitsuji)	sheep (yang2)	bed (beddo)	axe (fu3zi)
雲 (cloud; kumo/un)	云 (cloud; yun2)	sun (taiyou)	sun (tai4yang2)	chicken (niwatori)	chicken (ji1)
足 (leg; ashi/soku)	脚 (leg; jiao3)	arm (ude)	arm (shou3bi4)	cup (koppu)	cup (bei1)
窓 (window; mado/sou)	窗 (window; chuang1)	door (doa)	door (men2)	trousers (zubon)	trousers (ku4zi)
木 (tree; ki/moku)	树 (tree; shu4)	flower (hana)	flower (hua1)	brain (nou)	brain (nao3)
耳 (ear; mimi/ji)	耳 (ear; er3)	eye (me)	eye (yan3jing1)	car (kuruma)	car (jiao4che1)
弓 (bow; yumi/kyuu)	弓 (bow; gong1)	axe (ono)	axe (fu3zi)	spoon (supuun)	spoon (shao2)
箸 (chopsticks; hashi/cho)	刀 (knife; dao1)	spoon (supuun)	spoon (shao2)	axe (ono)	bed (chuang2)

豚 (pig; buta/ton)	猪 (pig; zhu1)	chicken (niwatori)	chicken (ji1)	sun (taiyou)	sun (tai4yang2)
海 (sea; umi/kai)	海 (sea; hai3)	mountain (yama)	mountain (shan1)	church (kyoukai)	church (jiao4tang2)
皿 (plate; sara/bei)	盘 (plate; pan2)	cup (koppu)	cup (bei1)	arm (ude)	arm (shou3bi4)
船 (boat; fune/sen)	船 (boat; chuan2)	car (kuruma)	car (jiao4che1)	eye (me)	eye (yan3jing1)
剣 (sword; ken/tsurugi)	剑 (sword; jian4)	pistol (pisutoru)	pistol (shou3qiang1)	bolt (boruto)	hammer (chui2zi)
机 (desk; tsukue/ki)	桌 (desk; zhuo1)	bed (beddo)	bed (chuang2)	sheep (hitsuji)	sheep (yang2)
笛 (flute; fue/teki)	笛 (flute; di2)	harp (haapu)	harp (shu4qin2)	cat (neko)	horse (ma3)
家 (house; ie/ka)	房 (house; fang2)	church (kyoukai)	church (jiao4tang2)	mountain (yama)	mountain (shan1)
肺 (lung; hai)	肺 (lung; fei4)	brain (nou)	brain (nao3)	flower (hana)	flower (hua1)
釘 (nail; kugi/tei)	锯 (saw; ju4)	bolt (boruto)	hammer (chui2zi)	pistol (pisutoru)	pistol (shou3qiang 1)
靴 (shoes; kutsu/ka)	鞋 (shoes; xie2)	trousers (zubon)	trousers (ku4zi)	door (doa)	door (men2)

Note. For the Japanese kanji targets the underlined pronunciation is the pronunciation for the kanji when standing alone, and all subjects in our experiments used these pronunciations without exception. Only the kanji 剣 can standalone be pronounced both as /ken_{on}/ or /tsurugi_{kun}/, however, our participants were consistent in naming this kanji /ken_{on}/ as this is the modern term. It was furthermore checked whether excluding from the analyses four Japanese kanji which turned out to have rather infrequent alternative readings (i.e., 箸, 皿, 釘 and 靴) would yield different results; this turned out not to be the case. The Chinese targets matched the Japanese targets with two exceptions: (1) the symbol for ‘chopsticks’ (筷 /kuai4/ or in Japanese 箸 /hashi/) which was substituted with ‘knife’ (刀 /dao1/) as Chinese readers would find it quite unusual to pronounce 筷 without the nominal suffix 子 (筷子 /kuai4zi/) and (2) the Japanese symbol for ‘nail’ 釘 (/kugi/) which was replaced in Chinese by ‘saw’ (锯 /ju4/) for the same reason. Although the pictures enter the production process through the conceptual system (Roelofs, 1992; 2006; Caramazza, 1997; Levelt et al., 1999), we nevertheless decided to change the semantically related picture ‘cat’ (猫) into ‘horse’ (马) since in Chinese (not Japanese 犬) its semantic radical would have overlapped with ‘dog’ (狗). Finally, for the Chinese pictures ‘bolt’ was replaced by ‘hammer’ to yield a categorically related context picture for the target ‘saw’.

Chapter 4: Context effects when naming Japanese (but not Chinese), and degraded Dutch nouns: evidence for processing costs?

This chapter is based on: Verdonschot, R. G., Paolieri, D., Kiyama, S., Zhang, Q. F., La Heij, W., & Schiller, N. O. (submitted). Context effects when naming Japanese (but not Chinese), and degraded Dutch nouns: evidence for processing costs?

Abstract

Reading bare words in alphabetic languages has been shown to be rather immune to effects of context stimuli, even when these stimuli are presented in advance of the target word (e.g. Glaser & Döngelhoff, 1984; Roelofs, 2003, 2006). However, recently, semantic context effects of distractor pictures on the naming latencies of Japanese kanji (but not Chinese hànzi) words have been observed (Verdonschot, La Heij & Schiller, 2010). In the present study, we further investigated this issue using phonologically related (i.e. homophonic) context pictures when naming target words in either Chinese or Japanese. We found that pronouncing bare nouns in Japanese is sensitive to phonologically related context pictures, whereas this is not the case in Chinese. The difference between these two languages is attributed to processing costs caused by multiple pronunciations for Japanese kanji. A subsequent experiment using Dutch degraded stimuli words demonstrated that context effects could arise even in bare noun naming using an alphabetic language when stimulus characteristics (i.e. visual degradation) induce a processing cost. A possible way to model these findings is discussed.

Context effects when naming Japanese (but not Chinese), and degraded Dutch nouns: evidence for processing costs?

Word naming (i.e. reading aloud words) has been intensively studied in recent years and several models have emerged to explain how word naming is accomplished. The influential Dual-Route Cascading model (or DRC; Colheart, Rastle, Perry, Langdon and Ziegler, 2001), depicted on the left-hand side of Figure 1, assumes that there are two routes through which a word can be read aloud: the lexical and non-lexical route. The lexical route can be further divided into two parts: the lexical non-semantic route entails the involvement of the components of the mental lexicon that contains the correct pronunciation of a specific word (Route 2 in Figure 1). The lexical-semantic route (Route 1 in Figure 1) within the DRC involving the word's semantic representation has not been implemented (Colheart et al., 2001; p. 217). The non-lexical route converts orthographic information ("graphemes") into pronounceable output by means of orthography-to-phonology conversion rules (OPC; Route 3 in Figure 1). The existence of the OPC route is evidenced by the fact that we can name non-words such as "DELK" which, by virtue of being a non-word, do not have entries in the mental lexicon. In contrast, words with an "irregular" pronunciation, such as "TWO" /tu/, would have to be looked up in the mental lexicon, as simple conversion would produce overgeneralization errors, i.e. /two/.

An influential word-production model that also simulates word naming is WEAVER++ (Roelofs, 1992; Roelofs, Meyer, & Levelt, 1996; Levelt, Roelofs, & Meyer, 1999; Indefrey & Levelt, 2004; Roelofs, 2006). Regarding the naming of objects, this model distinguishes a number of processing levels including conceptualization, retrieval of syntactic features, phonological word-form encoding, and ultimately articulation. In addition to an OPC route, word naming is assumed to involve the same stages (see the right-hand side of Figure 1). As is evident from Figure 1, the three routes mentioned above, the lexical-syntactic (lexical-semantic in DRC terminology) route, the lexical-phonological or direct route (lexical non-semantic in DRC terminology), and the OPC route are present in both models.

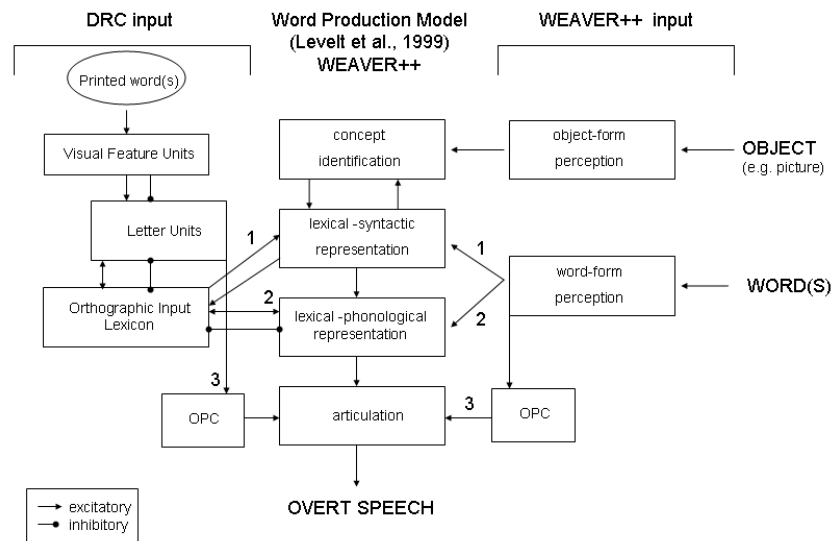


Figure 1. Input from DRC (Coltheart et al., 2001) and WEVER++ into the Word production model of Levelt et al. (1999). This figure is partly adapted from Coltheart et al. (2001) and Roelofs (2006).

The DRC model states that the visual characteristics of a to-be-named word first activate the letter units of a word. These letters in turn activate the word's entry in the orthographic input lexicon, which subsequently activates the corresponding entry in the phonological lexicon (the phonological word-form), which in turn activates the word's phonemes in parallel. The WEVER++ model (e.g. Roelofs, 1992, 2006) assumes that to-be-named words automatically activate the lexical-syntactic (Route 1) and lexical-phonological (Route 2) routes in parallel. If the task does not require information at the lexical-syntactic level, the fastest route will determine the reading latencies, i.e. Route 2. This entails phonological word form retrieval, syllabification, and ultimately turning syllables into motor action instructions (e.g. overt articulation). However, Route 1 determines reading latencies if the task requires information stored at the lexical-syntactic level. Support for the usage of a direct Route 2 without involvement of Route 1 comes from an observation by Glaser and Döngelhoff (1984). These authors found that semantically related distractor words slowed down picture naming compared to unrelated distractor words, but that the reverse effect was not found:

semantically related distractor pictures did not affect the naming of single words. A simple horse-race explanation for this asymmetry was rejected on the basis of the finding that context pictures did not even affect word naming when presented 400 ms before the target word. This finding suggests that words can be named via a fast route that bypasses the lexical semantic/syntactic level.

Roelofs (2003, 2006) investigated semantic context effects in naming pictures or words as well. When participants were required to name visually presented nouns (N) or det (determiner)+N phrases, e.g. “HOND” [dog] or “DE HOND” [the dog], respectively, context pictures did not induce semantic or grammatical gender effects, suggesting the use of Route 2. If the task was to generate a det+N phrase given a single noun (e.g. responding with “de hond” [the dog] to the stimulus “hond”), context pictures did induce semantic context effects. Semantically related context pictures (e.g. CAT) now facilitated the production of the det+N phrase (“de hond”) compared to unrelated context pictures. The author proposed that this finding is due to the fact that to generate the correct gender-marked determiner lemma access (via Route 1) is required.

Recently, Verdonchot, La Heij, and Schiller (2010) investigated semantic context effects of pictures on naming Japanese kanji and Chinese hànzi words. Japanese kanji form a unique set of words in that over 60% of them are homographic heterophones, meaning that most kanji have at least two different pronunciations. This contrasts with most alphabetic languages (and Chinese hànzi) in which the majority of visually presented words only have one pronunciation. The etymology of these multiple readings of Japanese kanji lies in the fact that they were originally imported from China. In those days not only the script itself was imported but also the Chinese pronunciation of the characters. For instance, the original name for “water” in Japanese is /mizu/ (called the KUN-reading), and the Chinese name for “water” is /shui3/. Over time the Chinese-derived ON-reading in Japanese changed to some extent (e.g. /sui/), but the character for “water” 水 still has two potential readings in modern Japanese, i.e. /mizu_{kun}/ and /sui_{on}/, depending on whether character it combines with (e.g. 海水 /kai_{on}.sui_{on}/ “seawater” and 雨水 /ama_{kun}.mizu_{kun}/ “rainwater”).

In their study, Verdonschot et al. (2010) combined kanji targets with semantically related and unrelated context pictures and found that at two stimulus-onset asynchronies (SOAs) of 0 ms (simultaneous presentation) and –150 ms (context picture first) semantically related distractor pictures sped up word-naming latencies. This result is at variance with both the lack of a picture-context effect in reading Chinese characters and the lack of picture-context effects in naming words in alphabetic languages discussed above (Glaser & Dünghoff, 1984; Roelofs, 2003, 2006). Verdonschot et al. suggested two possible accounts of their finding with Japanese kanji: (a) naming kanji requires lexical-syntactic information to determine which pronunciation is the correct one (Route 1) and (b) naming kanji faces a processing cost at the lexical-phonological level (due to the necessity of pronunciation selection), which provides the opportunity for context pictures to exert an effect on naming latencies. Although their data did not exclude the involvement of lexical-syntactic representation, the authors opted for the latter, more parsimonious, alternative (which was furthermore supported by neuropsychological evidence indicating the use of a direct orthography to phonology route in reading kanji, e.g. Sasanuma, Sakuma, & Kitano, 1992; Nakamura et al., 1998; Fushimi et al., 2003).

As noted above, kanji are unique because over 60% are homographic heterophones. Although much smaller in number, homographic heterophones are also present in alphabetic languages, like the word “read” in English (i.e. “I’ll read [rɪd/] this book” vs. “I’ve read [rɛ d/] this book”). There is evidence that such words show longer naming latencies compared to matched controls (Seidenberg, Waters, Barnes, & Tanenhaus, 1984; Kawamoto & Zemblidge, 1992; Folk & Morris, 1995; Gottlob, Goldinger, Stone, & Van Orden, 1999). It has been proposed that this is due to the time necessary to select between two or more simultaneously activated pronunciations.

In both the WEAVER++ and DRC models there are at least two ways for a word such as “read” to activate one of its pronunciations (/rɪd/ or /rɛ d/). One option is that a single orthographic unit, i.e. “read”, activates both pronunciations and that one of these pronunciations is ultimately selected. The second option

is that such a word is read via the lexical-syntactic route resulting in the selection of a representation (for instance, on the basis of syntactic or semantic context) subsequently leading to the activation of the corresponding phonological representation(s).

It seems realistic to assume that in Japanese, a heterophonic kanji could follow the same two routes: the kanji for “water” 水, for example, could either be read via its lexical-syntactic representation or via the direct route from orthography to phonology (see Figure 2).

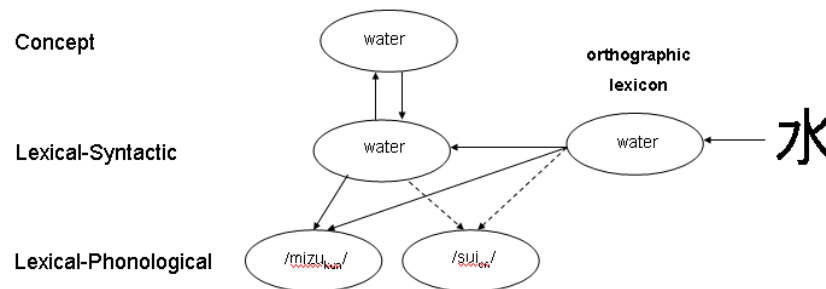


Figure 2. Activation from orthographic kanji input to its pronunciations according to WEAVER++.

Within the basic model depicted in Figure 2, the kanji symbol for water 水 will activate its representation in the orthographic lexicon and activation will spread to the phonological word form /mizu_{kun}/ as this character, when standing alone, is typically pronounced in this way. However, as argued before, it could also activate the alternative word form /sui_{on}/. Evidence for the activation of /sui_{on}/, although this pronunciation is not used when standing alone, comes from a study by Kayamoto, Yamada, and Takashima (1998). They reported that single kanji that have a frequent alternative reading when part of a compound are named slower than their matched controls (but see Wydell, Butterworth, and Patterson, 1995; Exp. 5). Furthermore, Fushimi, Ijuin, Patterson, and Tatsumi (1999) found significant consistency effects when naming compound kanji and non-words in Japanese when typicality was introduced as a factor. Pure consistent kanji were kanji compounds for which its constituents have the same pronunciation in all words containing that constituent in that position (e.g. 医 and 学 in target word 医学 /i_{on}.gaku_{on}/ “medical science”;

other words are e.g. 医者 /i_{on}.sha_{on}/ “doctor” and 科学 /ka_{on}.gaku_{on}/ “science”). Inconsistent but typical kanji are target compounds for which the constituents can take more than one pronunciation but there is a statistically common pronunciation (e.g. compounds using that kanji at that position usually take that reading). Inconsistent but atypical kanji are target words for which its constituents can have alternative pronunciations and the current reading is not typical amongst words in that same position (e.g. 人 and 間 in target 人間 /ni_{on}.gen_{on}/ “mankind”; other words are e.g. 人手 /hito_{kun}.de_{kun}/ “crowd” and 時間 /ji_{on}.kan_{on}/ “time”). Consistent words typically took less time to name compared to inconsistent words, especially when they were of low frequency. Furthermore, consistency effects between inconsistent-atypical and typical words were also observed. This shows that at a constituent level (individual kanji) character-sound correspondences exerted an effect, which suggests involvement of multiple pronunciations (e.g. /hito_{kun}/ for 人 in 人間).

Finally, a study by Verdonschot, La Heij, Poppe, Tamaoka, and Schiller (submitted) reported that a single kanji prime could facilitate its multiple readings when those readings were both transcribed in Japanese katakana script (e.g. 町 “town” which can be pronounced /machi_{kun}/ or /chou_{on}/), i.e. マチ “machi” and チョウ “chou”, compared to an unrelated prime. This indicates that multiple readings were activated during the short time span the prime was presented.

As mentioned earlier, Verdonschot et al. (2010) obtained facilitation effects from semantically related pictures compared to unrelated pictures when naming Japanese kanji but not when naming Chinese hànzi. If this effect originates from the fact that Japanese kanji is read through the direct route (Route 2 in Figure 1) and this route is susceptible to context effects when a processing cost is incurred, then also phonologically related context pictures are expected to speed up naming latencies in Japanese (but not Chinese). The current study further examines this issue by means of three experiments involving phonological (homophonic) and semantic effects of context pictures on word naming. The first two experiments employ to-be-named Japanese/Chinese logographic characters, which are superimposed on context pictures. The names of these pictures are either homophones of the correct kanji/hànzi reading, or

phonologically unrelated to the correct reading. First of all, for Chinese the predictions are straightforward, i.e. Chinese hànzi naming proceeds via the fast direct route from orthography (Route 2 in Figure 1) in line with the interpretation by Verdonschot et al. (2010). Therefore, distractor pictures with homophonic names will not facilitate Chinese hànzi naming as the fast direct route and the lack of multiple pronunciations prevents any influence from picture processing. However, for Japanese the story becomes different. In this case, we propose that naming Japanese kanji also proceeds via the direct lexical-phonological level, however, the fact that multiple pronunciations are activated (due to kanji heterophony) causes a processing cost which in turn leads to the same susceptibility to context effects as observed (for semantic context) in Verdonschot et al. (2010). Therefore, we hypothesize that introducing homophonic context pictures in our first two experiments should give rise to different effects for Japanese (Experiment 1) and Chinese (Experiment 2).

In order to further investigate the possible role of processing costs in the emergence of semantic and phonological context effects, we ran a similar experiment in Dutch (an alphabetic language). This Experiment 3 employs the naming of bare words and det+N phrase naming in Dutch (see also Roelofs, 2003) but also includes a novel degraded-word condition in which the word naming process is made more difficult, thereby artificially inducing a processing cost. We hypothesize that if processing costs at the lexical-phonological level are responsible for the observed context effects, then naming degraded bare nouns in Dutch should also become susceptible to context effects, similar to Japanese kanji.

Experiment 1: Naming Japanese kanji with homophonic distractor pictures

In this study, kanji target words are presented with distractor pictures whose name is homophonic with the dominant reading of the (standing alone) kanji character. For instance, the kanji for “white” 白 (/shiro_{kun}/ or /haku_{on}/) was superimposed on a picture of a “castle” which is also named /shiro_{kun}/ (note: the kanji for “castle” is 城 /shiro_{kun}/ or /jyou_{on}/) compared to an unrelated picture. As any semantic or orthographic relationship between picture distractor and

target word is absent in our stimuli, a possible facilitation effect of homophonic pictures is expected to be localized at the lexical-phonological level. Note that phonological facilitation of picture names has been observed in word production tasks (picture naming and color naming; Kuipers & La Heij, 2009; Morsella & Miozzo, 2002; Navarrete & Costa, 2005) indicating that, at least under some circumstances, context pictures are processed up to the level of phonological word forms (but see Jescheniak et al., 2009, Bloem & La Heij, 2003 and Bloem, van den Boogaard & La Heij, 2004).

Method

Participants. Twenty-one undergraduate students from Yamaguchi University, Japan (15 female, average age: 20.3 years; SD = 1.3) took part in the Experiment in exchange for financial compensation. All participants were native speakers (and fluent readers) of Japanese and had normal or corrected-to-normal vision. *Stimuli.* We selected 22 kanji characters for which we could also select an appropriate picture bearing the same pronunciation. For instance, the kanji 造 for “construction” which is pronounced /zou/ was superimposed on a picture of an elephant (which carries the same pronunciation, /zou/). The control picture of a tree (pronounced /ki/) does not bear any phonological relationship with the target kanji.

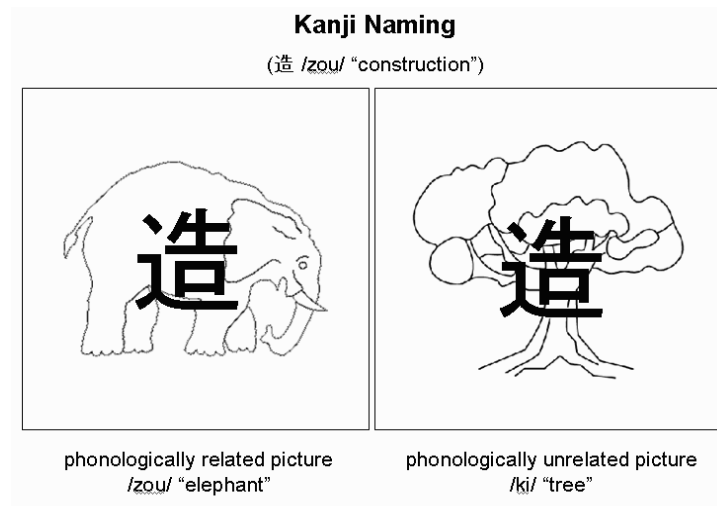


Figure 3. Examples of Japanese experimental stimuli.

To avoid effects due to the nature (e.g. visual properties) of the pictures, we balanced the distractor pictures so they made so-called equal pairs with the targets, e.g. for the target 器 (“bowl”, /ki/) the same two pictures were used as for 造 (“construction”, /zou/), only their roles were reversed in this case. Figure 3 provides examples of kanji-picture pairs and Appendix A lists all Japanese stimuli. We also selected 30 kanji characters that were paired with two unrelated pictures to act as filler items (thereby reducing the homophonic proportion to 26.8%) to reduce the likelihood that participants became aware of the homophonic relation between some of the target-picture pairs. Kanji target characters had summed average kanji-to-sound correspondence (ranging from 1 [not adequate] to 7 [very adequate]) for the KUN-reading of 5.58 (SD = 1.4) and for the ON-reading of 5.9 (SD = 0.8; Amano & Kondo, 2000).

Design. A 2x2 within-subjects factorial design was implemented, with the factors SOA (0 ms, i.e. picture and word presented simultaneously, or –150 ms, i.e. picture first) and phonological relatedness (homophonic or unrelated context picture). Each participant was subjected to 208 kanji naming (88 experimental + 120 filler) trials presented in four blocks (two blocks per SOA). For each participant, pseudo-random lists were constructed per block such that there were minimally two intervening trials between phonologically or semantically related characters or pictures. Across participants, the order of blocks was counterbalanced. Each block started with three warm-up trials (all filler trials).

Procedure. Participants were seated approximately 60 cm from a 17 inch LCD computer screen (Eizo Flexscan P1700 at 60 Hz) in a quiet room at Yamaguchi University. The E-prime 2.0 software package was used to present the stimuli and record the responses. Trials consisted of a fixation point presented for 750 ms followed and replaced by the picture – kanji pair (using the appropriate SOA for that block), which disappeared when participants responded or after maximally 2,000 ms. Following a response, the experimenter recorded whether or not the response was accurate before the next trial started. Naming latencies were measured from target onset using a voice-key. Participants were instructed to respond as fast as possible while avoiding errors.

Reaction Time Results. Naming latencies below 300 ms and above 1,500 ms and voice key errors were counted as outliers (comprising 1.5% of the data). Other errors (i.e. incorrect target names) accounted for 4.3% of the data. Table 1 shows the mean RTs and percentages of errors in the various conditions. An analysis of variance (ANOVA) with SOA (0 ms and –150 ms) and Phonological Relatedness (homophonic versus unrelated) as within-subject variables showed a marginal effect of SOA in the items (but not the subjects) analysis, $F(1,20) = 1.66$, n.s.; $F(1,21) = 4.32$, $MSe = 1329.3$, $p = .05$ and a main effect of Phonological Relatedness in the subjects (but not the items) analysis, $F(1,20) = 17.15$, $MSe = 611.6$, $p < .001$; $F(1,21) = 1.42$, n.s, reflecting in the subject analysis that overall homophonic target-distractor pairs were named faster. More importantly, there was a significant interaction between SOA and Phonological Relatedness in the subjects (not the items) analysis, $F(1,20) = 8.18$, $MSe = 549.7$, $p = .01$; $F(1,21) = 2.81$, $MSe = 3717.6$, $p = .11$. Planned t-tests show that at SOA = 0 the 8 ms facilitation effect of homophonic pictures on kanji naming latencies as compared to unrelated pictures was not significant, all t s < 1 . However, for SOA = –150 homophonic pictures sped up naming of the target kanji as compared to unrelated pictures by 37 ms, $t(20) = 5.85$, $SD = 29.00$, $p < .001$; $t(21) = 2.34$, $SD = 70.08$, $p < .05$.

Table 1
Mean Naming Latencies (in Milliseconds) and Error Rates (in %) in the Kanji Naming Task as a Function of SOA and Phonological Relatedness.

	SOA = –150	%E	SOA = 0	%E
Homophonic relation	552 (54)	3.6	587 (79)	4.2
Phonologically Unrelated	589 (64)	4.9	595 (93)	4.5
Homophonic context effect	–37 (29)	–1.3	–8 (38)	–0.3

Error results. An identical ANOVA was performed on the error percentages. This analysis showed no main effect of SOA, all F s < 1 , but there was a main effect of Phonological Relatedness, $F(1,20) = 6.2$, $MSe = 1.6$, $p < .05$; $F(1,21) = 5.6$, $MSe = 1.7$, $p < .05$ indicating that more errors were made with unrelated pictures. Furthermore, there was an interaction (marginally significant by

items) between SOA and Phonological Relatedness, $F(1,20) = 6.2$, $MSe = .69$, $p < .05$; $F(1,21) = 3.4$, $MSe = 1.2$, $p = .08$. To explore the interaction in more detail, planned comparisons were carried out, at $SOA = 0$ there was no effect of Phonological Relatedness on error rates, all t s < 1 , however at $SOA = -150$ more errors were made in the phonologically unrelated condition, $t(20) = 3.1$, $SD = 1.68$, $p < .01$; $t(21) = 2.5$, $SD = 2.07$, $p < .05$.

Discussion. Our results show that homophonic distractor pictures speed up kanji naming latencies when presented 150 ms before target onset. This is an important finding as for English and Dutch (alphabetic) words such context effects are absent (see for instance Glaser & Döngelhoff, 1984). Also it corroborates well with the findings obtained in Verdonschot et al., (2010) who found semantic context effects of pictures on the naming latencies of kanji at $SOA -150$ and $SOA 0$. Our current findings can be accounted for by assuming that the distractor pictures activate their conceptual representations and that this activation cascades to the lexical-syntactic and the lexical-phonological level and exerts an effect at the latter level. Note that the phonologically related picture name should not affect the processing of the target word at the lexical-syntactic level, as picture and word are not semantically or orthographically related. The target word is supposed to activate its representation in the orthographic lexicon and via the fast direct route (Route 2) its phonological word-form. Although Route 2 is usually fast, the results show an effect of homophonic distractor pictures when the pictures are given a 150 ms head start. This susceptibility of kanji naming to context effects stands in marked contrast to the general lack of context effects in naming single words in alphabetic languages (Glaser & Döngelhoff, 1984; La Heij, Happel, & Mulder, 1990; Roelofs, 2003).

The most parsimonious explanation for the homophonic facilitation effect is the fact that the Japanese kanji characters used have multiple readings (thereby requiring a time-consuming selection process at the word-form level). To test this hypothesis, logographic characters in Japanese should be examined that do not have multiple readings. However, as it turns out to be hard to find a set of single ON or KUN reading characters that could be equally well matched with homophonic pictures in Japanese; we decided to employ Chinese logographs in Experiment 2. Chinese hànzi characters (leaving specific grammatical differences between languages aside) are similar

to the Japanese kanji stimuli with the difference that a Chinese hànzi character usually has a single pronunciation.

Experiment 2: Naming Chinese hànzi with homophonic pictures

The setup of this experiment is identical to Experiment 1. In this experiment, word targets are again accompanied by homophonic and control distractor pictures. The issue is whether the significant facilitation effects of homophonic pictures on naming latencies of Japanese kanji can be replicated using Chinese hànzi.

Method

Participants. Twenty-four undergraduate university students (who were enlisted in a database of the psychology department of the Chinese Academy of Sciences in Beijing, China; 17 female, average age: 24.0 years; SD = 1.6) took part in the experiment in exchange for financial compensation. All participants were native speakers (and fluent readers) of Mandarin Chinese and had normal or corrected-to-normal vision.

Stimuli. As in Experiment 1, we selected 22 hànzi characters and corresponding semantically unrelated pictures with the same name. For instance, the hànzi 珠 for “pearl” which is pronounced /zhu1/ was superimposed on a picture of a pig (the Chinese name which has the same pronunciation and tone, e.g. /zhu1/). The control picture of a chicken /ji1/ does not bear any phonological relationship with the target hànzi. For target hànzi and distractor pictures tones were always kept the same. There was no significant difference in mean target frequency (per million) between Japanese (594) and Chinese stimuli (365), $t(42) = 1.20$, ns (taken from Yokoyama, Sasahara, Nozaki, & Long, 1998, and Da, 2004, respectively). Again, we created equal pairs (as in Experiment 1). Figure 4 provides examples of hànzi-picture pairs and Appendix B lists all Chinese stimuli. We also selected 30 hànzi characters paired with unrelated pictures to act as filler items, to reduce the likelihood that participants became aware of the homophonic relation between some of the target-picture pairs.

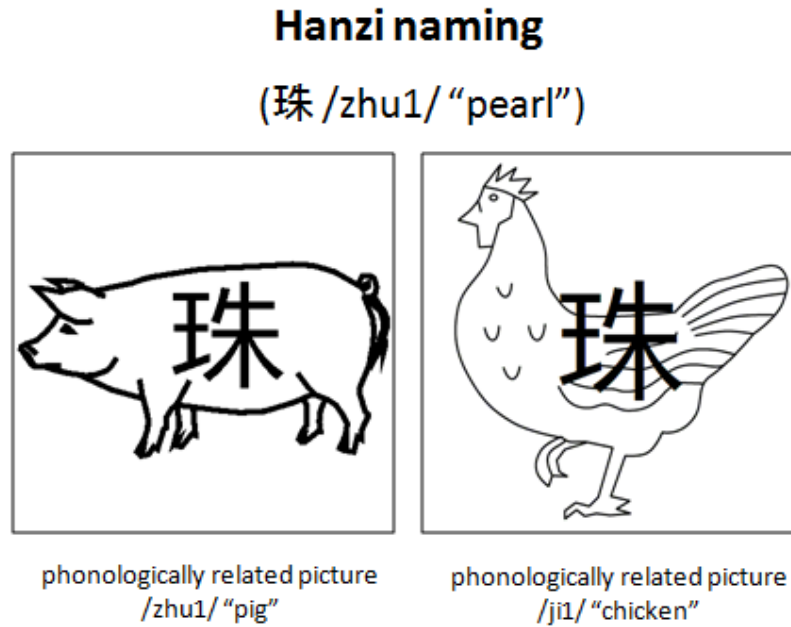


Figure 4. Examples of Chinese experimental stimuli.

Design. The design was identical to Experiment 1.

Procedure. Participants were seated approximately 60 cm from a 17 inch CRT computer screen in a quiet room at the Institute of Psychology at the Chinese Academy of Sciences. The rest of the procedure was identical to Experiment 1.

Reaction Time Results. Naming latencies below 300 ms and above 1,500 ms were counted as outliers (comprising 1.0% of the data); other errors (e.g. incorrect target names) accounted for another 1.0%. Table 2 shows the mean correct RTs in the various conditions. An ANOVA was performed with SOA (0 ms vs. –150 ms) and Phonological Relatedness (homophonic vs. unrelated) as within subject variables. The analysis showed no main effect of SOA, $F(1,23) = 1.5$, $MSe = 900.6$, n.s.; $F(1,21) = 3.4$, $MSe = 376.7$, $p = .08$ and no main effect of Phonological Relatedness, all $F_s < 1$, and there was no interaction between SOA and Phonological Relatedness, all $F_s < 1$.

Table 2. Mean Naming Latencies (in Milliseconds) and Error Rates (in %) in the Chinese hànzi Naming Task as a Function of SOA and Phonological Relatedness.

	SOA -150	%E	SOA 0	%E
Homophonic	539 (68)	0.4	531 (61)	1.2
Phonologically unrelated	538 (61)	2.0	530 (59)	0.4
Phonological context effect	1 (26)	-1.6	1 (23)	0.8

Error results. An identical ANOVA was performed on the error percentages. This analysis showed no main effect of SOA in the subjects analysis, $F(1, 23) < 1$, but it approached significance in the items analysis, $F(1, 21) = 4.1$, $MSe = .07$, $p = .06$. There was no main effect of Phonological Relatedness, $F(1, 23) = 1.0$, n.s.; $F(1, 21) = 1.3$, n.s., but there was a significant interaction between SOA and Phonological Relatedness in the subjects analysis, $F(1, 23) = 4.8$, $MSe = .22$, $p < .05$, but not the items analysis, $F(1, 21) < 1$. Planned t-tests showed that at SOA = 0 ms there was no effect of Phonological Relatedness on error rates, $t(23) = 1.1$, n.s.; $t(21) < 1$; however, it was marginally significant at SOA = -150 ms in the subjects analysis, $t(23) = 2.1$, $SD = 0.7$, $p = .05$, but not the items analysis, $t(21) = 1.3$, n.s., reflecting slightly more errors (1.3%) in the unrelated compared to the homophonic condition.

Discussion.

Our results show that Phonological Relatedness (homophony) of distractor pictures with target hànzi does not speed up naming latencies at any SOA. Mean RTs obtained with homophonic and control distractor pictures are virtually identical. Therefore, the homophonic context effect observed in naming Japanese kanji (Experiment 1) does not generalize to naming Chinese hànzi (Experiment 2). The absence of this effect in Chinese can be accounted for by assuming that the activation of the phonological representation of a Chinese word (via Route 2 in Figure 1) builds up so fast that pictures are unable to exert an effect on naming latencies. Some support for this assumption is provided by the faster overall naming latencies in Chinese than in Japanese (a difference of 46 ms), $F(1, 43) = 6.44$, $MSe = 14,658.14$, $p < .05$; $F(1, 42) = 12.76$, $MSe =$

11,310.82, $p < .001$. However, as hànzi and kanji words differ both in form and pronunciation, it is difficult to draw strong conclusions from this observation. However, the absence of context effects in Chinese word reading corroborates our hypothesis that the context effects observed in Japanese kanji reading is due to a processing cost induced by the activation of multiple word-form candidates in that language. One way to further investigate this issue is by inducing a processing delay in naming words in languages using alphabetic scripts (e.g. English or Dutch). In the introduction, we discussed a study by Roelofs (2003) in which he employed both Dutch bare-noun naming (e.g. naming the word “kat” [cat] as /kat/) and det+N phrase generation (e.g. responding with “de kat” to the word “kat”). He found semantic facilitation effects of context pictures in the latter but not in the former task. Roelofs accounted for this finding by assuming that in det+N naming the lexical-syntactic level has to be involved in order to select the correct gender-marked determiner for the utterance (see Schiller & Caramazza, 2003, 2006 for additional arguments). It is at this lexical-syntactic level that a semantically related context picture can exert its effect.

Our Experiment 3 sought to replicate Roelofs’ (2003, 2006) context effects in Dutch word reading with the addition of a degraded bare noun naming condition. This condition was introduced to induce a cost in stimulus processing without the requirement of lexical-syntactic access, as Dutch bare nouns can be named without accessing this level.

Experiment 3: Dutch bare noun, det+N, and degraded-word naming

This experiment seeks to replicate and extend Experiment 1 of Roelofs (2006), in which Dutch target words were accompanied by semantically related and unrelated context pictures and participants were asked to perform two tasks: bare noun naming and det+N naming. In the current experiment, participants performed three tasks: (1) a bare noun naming task (e.g. simply respond “hond” [dog] when presented with the word “hond”), (2) a det+N generation task (responding “de hond” [the dog; common gender] to the word “hond”), and (3) a degraded bare noun naming task (responding “hond” [dog] when presented with “\$H\$O\$N\$D\$”). If semantic context effects (as observed in Experiment 1, and in Verdonchot et al., 2010) are not due to a processing cost at the lexical-phonological

level, then we expect to find effects of semantically related context pictures only in the det+N naming task and not in the bare noun naming and degraded bare noun naming task. If processing costs do play a role in the emergence of the semantic context effect, then we expect to find semantic context effects of pictures also in the degraded bare noun naming task.

Participants. Eighteen paid participants from Leiden University (9 female; mean age: 22.8 years; SD = 4.4) took part in the Experiments. All participants were native speakers of Dutch and had normal or corrected-to-normal vision.

Materials. In order to replicate and extend Roelofs' (2003, 2006) semantic context effects, we used the same 32 objects from eight different semantic categories, with their basic level terms in Dutch (see Roelofs, 2006). Half of the objects in a particular category (e.g. ANIMALS) had names with neuter gender (e.g. "het konijn"), and the other half were picture names with common gender (e.g. "de zwaan"). In addition, eight different pictures taken from two non-included semantic categories were selected to serve as a practice items. All pictures were black line drawings on white backgrounds (for an overview of the stimuli see Appendix C or Roelofs, 2006).

Design. A 3 x 2 within-subjects factorial design was implemented, with the factors Task (bare word naming, determiner word naming and degraded word naming) and Semantic Relatedness (semantically related vs. unrelated) as within participant variables. Tasks were blocked and per block participants received 32 word-picture pairings from the same semantic category (related condition), plus 32 word-picture pairings from different semantic category (unrelated condition), yielding 64 trials for each task, and totaling 192 trials for all three tasks. Within each block, trial randomization was subjected to the following constraints: (1) items belonging to the same semantic category did not appear in consecutive trials and (2) target words are not repeated in consecutive trials. Task order was counterbalanced across participants.

In the bare noun naming condition, participants were asked to simply read aloud the word on the screen, while ignoring the picture in the background (e.g. respond "HOND" [dog]). In the det+N naming condition, participants were asked to produce not only the word but also the correct gender-marked determiner, e.g. respond "de hond" (the dog) for common gender words and e.g. "het paard" (the horse)

for neuter gender words. For the degraded word naming task, participants were again simply required to read aloud the word presented on screen. However, unlike in the bare noun naming task, the words were degraded by inserting dollar signs between each letter of the word. For example, if the to be named word was 'HOND' (dog), the stimulus was presented on screen as '\$H\$O\$N\$D\$'. Each target word was combined with a picture depicting an object from either the same semantic category (related condition) or with a random picture from another semantic category (unrelated condition). The word and the picture name always carried the same grammatical gender.

Procedure. Participants were individually tested in a quiet and dimly lit room. The stimuli were presented using E-Prime (PST Software). Participants were presented with the stimuli on a 100 Hz CRT monitor at a viewing distance of about 50 cm. RTs were measured from the onset of the stimulus to the beginning of the naming response using a voice key (SRBOX). The experimental session lasted about 30 minutes. Before the beginning of the experiment, participants were familiarized with the paradigm. After a participant had read the instructions, a block of 16 practice trials (not part of the proper experiment) was administered, which was followed by the experiment proper. Each trial contained the following sequence: a fixation point (+) was presented for 500 ms in the center of the screen. Next, the screen was cleared for 500 ms, followed by the distractor picture. One-hundred-and-fifty ms later, the target word was added to the display (i.e. SOA = -150 ms). The choice of this SOA was based on the findings of Roelofs (2003, 2006), Verdonschot et al. (2010), and of Experiment 1 in the present study. Both response speed and accuracy were emphasized. After each trial, the experimenter registered whether or not the response was accurate and whether or not the voice key malfunctioned.

Results. Naming latencies below 300 ms and above 1,500 ms were counted as outliers. Also voice key errors were excluded (in total comprising 3.2% of the data). The mean reading latencies, standard deviations, and error rates are shown in Table 3.

Table 3. Mean Reaction Times (in ms; SD between parentheses) and Error Percentages in Experiment 3 as a Function of Task and Condition.

Distractor Type	Tasks					
	Bare word reading		Word reading with determiner		Word reading with degraded stimuli	
	RT	E%	RT	E%	RT	E%
Sem. related	500 (51)	0.3	640 (91)	2.8	730 (181)	1.9
Sem. unrelated	502 (51)	0.7	657 (105)	3.5	768 (191)	2.8
Context effect	-2	-0.4	-17	-0.7	-38	-0.9

Results.

Reaction Times. The mean correct RTs were subjected to an ANOVA with Task (bare noun, det+N, and degraded noun naming) and Semantic Relatedness (semantically related vs. unrelated) as within-participant variables. This analysis showed a main effect of Task, $F(2,34) = 31.0$, $MSe = 18021.6$, $p < .001$; $F(2,62) = 223.0$, $MSe = 4290.2$, $p < .001$, and Semantic Relatedness, $F(1,17) = 22.1$, $MSe = 433.6$, $p < .001$; $F(1,31) = 17.9$, $MSe = 1123.8$, $p < .001$, as well as an interaction between these factors, $F(2,34) = 7.7$, $MSe = 395.0$, $p < .01$; $F(2,62) = 6.7$, $MSe = 968.7$, $p < .01$. Planned t-tests showed no effect of Semantic Relatedness in the bare noun naming, all $t_s < 1$; however, there was such an effect in the det+N naming, $t(17) = 2.5$, $SD = 29.6$, $p < .05$; $t(31) = 2.6$, $SD = 38.0$, $p < .05$, and the degraded noun naming task, $t(17) = 4.2$, $SD = 38.5$, $p < .001$; $t(31) = 3.6$, $SD = 65.8$, $p < .001$, reflecting the fact that responses were significantly faster when target words were presented with semantically related distractor pictures compared to unrelated distractor pictures (17 ms and 38 ms, respectively).

Error Results. The same ANOVA was performed on the error percentages. This analysis revealed a main effect of Task, $F(2,34) = 3.5$, $MSe = 1.9$, $p < .05$; $F(2,62) = 5.9$, $MSe = .6$, $p < .01$, indicating

that fewer errors were made in the bare noun naming task compared to the other two tasks. Task did not interact with Semantic Relatedness, all $F_s < 1$, and the main effect of Semantic Relatedness was not significant, either, $F(1,17) = 3.2$, $MSe = .4$, $p = .09$; $F(1,31) = 2.2$, n.s.

Discussion. Experiment 3 replicated Roelofs (2003; Experiment 3), i.e. we obtained no semantic context effect of distractor pictures in bare noun naming. Moreover, also in accordance with Roelofs (2003, 2006), we obtained a significant context effect of semantically related distractor pictures in the det+N production task. Crucially, we also found a significant semantic context effect in the degraded word naming task. The absence of semantic context effects in the bare noun naming condition indicates that the lexical-syntactic representation is not involved in producing bare nouns such as “hond” (dog). The most likely interpretation is that the word “hond” directly activated its lexical-phonological representation and could readily be pronounced. If participants, however, face a cost at some point in this process, context pictures get a chance to induce a measurable effect on the activation of the phonological word form, as evidenced by the degraded bare noun naming condition.

General Discussion

In three experiments, we investigated to which degree the naming of words in Japanese (kanji; a logographic language), Chinese (hànzì; a logographic language), and Dutch (an alphabetic language) can be influenced by context pictures. We found that homophonic context pictures induced facilitation on naming Japanese kanji characters with multiple readings (Experiment 1). However, comparable homophonic context pictures induced no such effect on naming Chinese hànzì (Experiment 2). This finding parallels results obtained in earlier work in our lab, which showed the same pattern for semantic context effects in Japanese and Chinese character naming (Verdonschot et al., 2010). In Experiment 3, using Dutch (an alphabetic language), semantically related context pictures yielded significant facilitation when naming gender-marked (det + N phrases), but not when naming bare nouns (replicating Roelofs, 2003).

Crucially, the novel condition (naming degraded bare nouns) in the experiment also showed facilitation by semantically related

distractor pictures (in comparison to unrelated pictures). In the remainder, we discuss the implications of our findings.

First of all, Chinese hànzì naming did not show a homographic facilitation effect. This finding suggests that Chinese hànzì are read via the fast, direct, route from orthography to phonology (Route 2 in Figure 1). Secondly, in contrast to the Chinese results, we found a homophonic facilitation effect in Japanese. It is unlikely that this effect arose at the lexical-syntactic level, as there was no semantic (nor orthographic) relation between target words and related context pictures. Furthermore, there is ample neuropsychological evidence showing that kanji activation spreads via orthography to phonology.

For instance, Sasanuma et al. (1992) as well as Nakamura et al. (1998) showed that patients with Alzheimer's dementia, whose comprehension of kanji was deteriorated, still maintained their ability to read kanji aloud. In addition, Fushimi et al. (2003) reported that a Japanese surface-dyslexic patient (TI) had an intact orthography-to-phonology route in combination with a decrease of activation coming from semantics. It seems as such plausible that the effect we observed in reading kanji arises at the lexical-phonological level and is due to a processing cost that results from the heterophony in Japanese kanji. This cost, a result of the necessity to select one of the activated word forms, may have allowed for the homophonic distractor pictures to exert a facilitation effect. This account is further corroborated by the findings of Experiment 3, which showed that bare noun naming in Dutch is only affected by semantically related context pictures when the processing of the target words is hampered by visual degradation.

The idea underlying our account is that the speed with which a phonological representation is activated upon the presentation of the target word determines whether a context effect (i.e. due to a context picture) surfaces. Reading words in alphabetic languages (like Dutch or German) or reading characters in Chinese may involve very strong links between orthographic and phonological representations that makes the process rather invulnerable to effects of context stimuli, even when these stimuli are presented in advance of the target word (even at SOA -400 ms in Glaser & Döngelhoff, 1984). However, when activation of the phonological representations takes more time due to, for instance, a one-to-many relation between orthography and phonology (as in many Japanese kanji; our Experiment 1) or reduced legibility (our Experiment 3), there is room for context stimuli to

affect the speed with which activation builds up and, hence, affect naming latencies. Figure 5 illustrates one way to model this proposal. When the build-up of activation of a word's phonological representation is fast (the steep increase in activation in the left panel of Figure 5), additional activation from context stimuli has little effect on the time necessary to reach a threshold value (illustrated by the dashed line in Figure 5). This might explain why there is no effect when the context picture is presented in advance of the target word (negative SOAs, see Glaser & Döngelhoff, 1984). However, when the build-up of activation is relatively slow (right panel), the effect of picture context can be much larger. This proposal naturally requires further substantiation by future experiments which could for instance address this matter by manipulating target frequency.

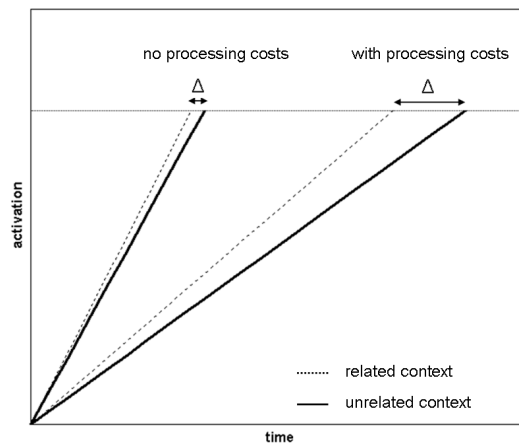


Figure 5. Hypothetical activation of the target word's phonological representation with an unrelated context picture (solid lines) and with homophonic context picture (dashed lines) under conditions of no processing cost and with processing cost.

In conclusion, we propose that kanji characters (like Chinese characters and alphabetic words) are most likely named via a direct route from orthography to phonology (Route 2 in Figure 1). In addition, context pictures can affect processing along this route when the stimulus characteristics induce a processing cost.

References

- Amano, N., & Kondo, K. (2000). *Nihongo-no goi tokusei* [Lexical properties of Japanese]. Tokyo: Sanseido.
- Bloem, I., & La Heij, W. (2003). Semantic facilitation and semantic interference in word translation: Implications for models of lexical access in language production. *Journal of Memory and Language*, 48, 468–488.
- Bloem, I., Van den Boogaard, S., & La Heij, W. (2004). Semantic facilitation and semantic interference in language production: Further evidence for the conceptual selection model of lexical access. *Journal of Memory and Language*, 51, 307–323.
- Coltheart, M., Rastle, K., Perry, C., Langdon, R., & Ziegler, J. C. (2001). DRC: A dual route cascaded model of visual word recognition and reading aloud. *Psychological Review*, 108, 204–256.
- Da, J. (2004). A corpus-based study of character and bigram frequencies in Chinese e-texts and its implications for Chinese language instruction. The studies on the theory and methodology of the digitalized Chinese teaching to foreigners. In Z. Pu, T. Xie, & J. Xu (Eds.), *Proceedings of the fourth international conference on new technologies in teaching and learning Chinese* (pp. 501–511). Beijing: Tsinghua University Press.
- Folk, J. R., & Morris, R. K. (1995). The use of multiple lexical codes in reading: Evidence from eye movements, naming time, and oral reading. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 21, 1412–1429.
- Fushimi, T., Ijuin, M., Patterson, K., & Tatsumi, I. F. (1999). Consistency, frequency, and lexicality effects in naming Japanese Kanji. *Journal of Experimental Psychology: Human Perception and Performance*, 25, 382–407.
- Fushimi, T., Komori, K., Ikeda, M., Patterson, K., Ijuin, M., & Tanabe, H. (2003). Surface dyslexia in a Japanese patient with semantic dementia: Evidence for similarity-based orthography-to-phonology translation. *Neuropsychologia*, 41, 1644–1658.
- Glaser, W. R., & Dünghoff, F. J. (1984). The time course of picture–word interference. *Journal of Experimental Psychology: Human Perception and Performance*, 10, 640–654.

Gottlob, L. R., Goldinger, S. D., Stone, G. O., & Van Orden, G. C. (1999). Reading homographs: Orthographic, Phonologic, and Semantic Dynamics. *Journal of Experimental Psychology: Human Perception and Performance*, 25, 561–574.

Indefrey, P., & Levelt, W. J. M. (2004). The spatial and temporal signatures of word production components. *Cognition*, 92, 101–144.

Jescheniak, J. D., Oppermann, F., Hantsch, A., Wagner, V., Mädebach, A., & Schriefers, H. (2009). Do perceived context pictures automatically activate their phonological code? *Experimental Psychology*, 56, 56–65.

Kawamoto, A. H., & Zemplidige, J. (1992). Pronunciation of homographs. *Journal of Memory and Language*, 31, 349–374.

Kayamoto, Y., Yamada, J., & Takashima, H. (1998). The consistency of multiple-pronunciation effects in reading: The case of Japanese logographs. *Journal of Psycholinguistic Research*, 27, 619–637.

Kuipers, J. R. & La Heij, W. (2009). The limitations of cascading in the speech production system. *Language and Cognitive Processes*, 24, 120–135.

La Heij, W., Happel, B., & Mulder, M. (1990). Components of Stroop-like interference in word reading. *Acta Psychologica*, 73, 115–129.

Levelt, W. J. M., Roelofs, A., & Meyer, A. S. (1999). A theory of lexical access in speech production. *Behavioral and Brain Sciences*, 22, 1–75.

Morsella, E., & Miozzo, M. (2002). Evidence for a cascade model of lexical access in speech production. *Journal of Experimental Psychology: Learning, Memory, & Cognition*, 28, 555–563.

Nakamura, K., Meguro, K., Yamazaki, H., Ishizaki, J., Saito, H., Saito, N., Shimada, M., Yamaguchi, S., Shimada, Y., & Yamadori, A. (1998). Kanji predominant alexia in advanced Alzheimer's disease. *Acta Neurologica Scandinavica*, 97, 237–243.

Navarrete, E., & Costa, A. (2005). Phonological activation of ignored pictures: Further evidence for a cascade model of lexical access. *Journal of Memory and Language*, 53, 359–377.

Roelofs, A. (1992). A spreading-activation theory of lemma retrieval in speaking. *Cognition*, 42, 107–142.

- Roelofs, A., Meyer, A. S., & Levelt, W. J. M. (1996). Interaction between semantic and orthographic factors in conceptually driven naming: Comment on Starreveld and La Heij (1995). *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 22, 246–251.
- Roelofs, A. (2003). Goal-referenced selection of verbal action: Modeling attentional control in the Stroop task. *Psychological Review*, 110, 88–125.
- Roelofs, A. (2006). Context effects of pictures and words in naming objects, reading words, and generating simple phrases. *Quarterly Journal of Experimental Psychology*, 59, 1764–1784.
- Sasanuma, S., Sakuma, N., & Kitano, K. (1992). Reading kanji without semantics: Evidence from a longitudinal study of dementia. *Cognitive Neuropsychology*, 9, 465–486.
- Schiller, N. O., & Caramazza, A. (2003). Grammatical feature selection in noun phrase production: Evidence from German and Dutch. *Journal of Memory and Language*, 48, 169–194.
- Schiller, N. O., & Caramazza, A. (2006). Grammatical gender selection in language production: The case of diminutives in Dutch. *Language and Cognitive Processes*, 21, 945–973.
- Seidenberg, M. S., Waters, G. D., Barnes, M. A., & Tanenhaus, M. K. (1984). When does irregular spelling or pronunciation influence word recognition? *Journal of Verbal Learning and Verbal Behavior*, 23, 383–404.
- Verdonschot, R. G., La Heij, W., Poppe, C. P., Tamaoka, K., & Schiller N. O. (submitted). Japanese kanji primes facilitate reading aloud of multiple katakana targets.
- Verdonschot, R. G., La Heij, W., & Schiller, N. O. (2010). Semantic context effects when naming Japanese kanji, but not Chinese hànzi. *Cognition*, 115, 512–518.
- Yokoyama, S., Sasahara, H., Nozaki, H., & Long, E. (1998). Shinbun denshi media no kanji: Asahi shimbun no CD-ROM-ni yoru kanji hindo hyoo. Tokyo: Sanseido [Japanese kanji in the newspaper media: Kanji frequency index from the Asahi newspaper on CD-ROM].
- Wydell, T. N., Butterworth, B. L., & Patterson, K. E. (1995). The inconsistency of consistency effects in reading: Are there consistency effects in Kanji? *Journal of Experimental Psychology: Learning, Memory and Cognition*, 21, 1155–1168.

Appendix A – Japanese stimuli
Stimulus Materials from Experiment 1

<u>Target</u>		<u>Related Picture</u>		<u>Unrelated Picture</u>	
<u>Pronunciation (kun/on)^a</u>	<u>Meaning</u>	<u>Pronunciation</u>	<u>Meaning</u>	<u>Pronunciation</u>	<u>Meaning</u>
箸 'hashi: 6.42' 'chou: 4.08'	chopsticks	hashi	bridge	me	eye
器 'utsuwa: 6.02' 'ki: 6.58'	bowl	ki	tree	zou	elephant
刃 'ha: 6.42' 'jin: 4.08'	blade	ha	leaf	su	nest
芽 'me: 6.42' 'ga: 5.42'	seedling/sprout	me	eye	hashi	bridge
緒 'o: 6.12' 'sho: 5.75'	cord/strap	o	tail/ridge	hata	flag
応 'kota: 5.92' 'ou: 6.71'	application	ou	king	nami	wave
可 'be: 3.42' 'ka: 6.58'	possible/passable	ka	mosquito	shita	tongue
雨 'ame: 6.79' 'u: 6.00'	rain	ame	candy	kutsu	shoes
端 'hashi: 6.04' 'tan: 5.96'	edge	hata	flag	o	tail
下 'shita: 6.71' 'ka: 5.79'	under	shita	tongue	ka ^b	mosquito
券 'fuda: 2.17' 'ken: 6.54'	ticket	ken	sword	shima	island
並 'nami: 6.42' 'hei: 5.75'	ordinary	nami	wave	ou	king
縞 'shima: 5.58' 'kou: 5.21'	stripe	shima	island	ken	sword
白 'shiro: 6.67' 'haku: 6.12'	white	shiro	castle	hi	fire
酢 'su: 6.58' 'saku: 4.58'	vinegar	su	nest	ha	leaf

造 'tsuku: 6.17' 'zou: 6.21'	construction	zou	elephant	ki	tree
比 'kura: 5.83' 'hi: 6.62'	comparison/ratio	hi	fire	shiro	castle
便 'tayo: 5.21' 'ben: 6.46'	mail/post/flight	bin	bottle	hon	book
屈 'kaga: 4.29' 'kutsu: 6.25'	leading/outstanding	kutsu	shoes	ame	candy
回 'mawa: 6.33' 'kai: 6.38'	counter occurrence	kai	seashell	nou	brain
翻 'hirugae: 5.5' 'hon: 6.54'	change ones mind	hon	book	bin	bottle
農 'nariwai: 1.75b' 'nou: 6.54'	farming/agriculture	nou	brain	kai	seashell

^aNumbers denote kanji-reading correspondences. These indices were taken from the NTT Japanese Word Database (Amano & Kondo, 2000). This index ranges from 1 (not adequate at all) to 7 (very adequate) judging kanji to sound correspondence.

^bThis unrelated item accidentally turned out to be a possible ON-reading for 下. A re-analysis without this distractor and the low kun-reading correspondence character 農 did not change the experimental findings and interpretation; therefore we decided to leave both in.

Appendix B – Chinese stimuli
Stimulus Materials from Experiment 2

<u>Target</u>		<u>Related Picture Distractor</u>		<u>Unrelated Picture Distractor</u>	
<u>Pronunciation</u>	<u>Meaning</u>	<u>Pronunciation</u>	<u>Meaning</u>	<u>Pronunciation</u>	<u>Meaning</u>
离 'li2'	to leave	li2	pear	wang2	king
晚 'wan3'	late	wan3	bowl	fu3	axe
掩 'yan3'	to cover	yan3	eye	tong3	bucket
亡 'wang2'	die away	wang2	king	li2	pear
螳 'tang2'	mantis	tang2	candy	qi2	flag
棋 'qi2'	chess	qi2	flag	tang2	candy
件 'jian4'	a piece	jian4	sword	bao4	leopard
播 'bo1'	to broadcast	bo1	wave	xia1	shrimp
抱 'bao4'	to hug	bao4	leopard	jian4	sword
评 'ping2'	to evaluate	ping2	bottle	xie2	shoes
斜 'xie2'	slanted	xie2	shoes	ping2	bottle
备 'bei4'	back-up	bei4	seashell(s)	ku4	trousers
瞎 'xia1'	blind	xia1	shrimp	bo1	wave
珠 'zhu1'	pearl	zhu1	pig	ji1	chicken
陵 'ling2'	tomb	ling2	bell	yun2	cloud
腐 'fu3'	rotten	fu3	axe	wan3	bowl

- 102 -

匀 'yun2'	equal	yun2	cloud	ling2	bell
公 'gong1'	male	gong1	a bow	gu3	bone
古 'gu3'	old, ancient	gu3	bone	gong1	a bow
酷 'ku4'	cool	ku4	trousers	bei4	seashell(s)
机 'ji1'	machine	ji1	chicken	zhu1	pig
统 'tong3'	to unify	tong3	bucket	yan3	eye

Appendix C – Dutch stimuli
Stimulus materials from Experiment 3

	<u>target</u>	<u>determiner</u>	<u>semantic</u>	<u>distractor</u>	
				<u>unrelated</u>	<u>identical</u>
animals	zwaan	de	schildpad	rok	zwaan
	schildpad	de	zwaan	beker	schildpad
	konijn	het	hert	paleis	konijn
	hert	het	konijn	bureau	hert
clothing	trui	de	rok	dolk	trui
	rok	de	trui	zwaan	rok
	hemd	het	vest	oor	hemd
	vest	het	hemd	kasteel	vest
transportation	fiets	de	trein	kast	fiets
	trein	de	fiets	arm	trein
	schip	het	vliegtuig	been	schip
	vliegtuig	het	schip	glas	vliegtuig
buildings	molen	de	fabriek	kom	molen
	fabriek	de	molen	neus	fabriek
	kasteel	het	paleis	vest	kasteel
	paleis	het	kasteel	konijn	paleis
weapons	dolk	de	speer	trui	dolk
	speer	de	dolk	tafel	speer
	kanon	het	pistool	bord	kanon

- 104 -

service	pistool	het	kanon	bed	pistool
	beker	de	kom	schildpad	beker
	kom	de	beker	molen	kom
	glas	het	bord	vliegtuig	glas
furniture	bord	het	glas	kanon	bord
	tafel	de	kast	speer	tafel
	kast	de	tafel	fiets	kast
	bed	het	bureau	pistool	bed
bodyparts	bureau	het	bed	hert	bureau
	arm	de	neus	trein	arm
	neus	de	arm	fabriek	neus
	been	het	oor	schip	been
	oor	het	been	hemd	oor

Chapter 5: The functional unit of Japanese word naming: evidence from masked priming.

This chapter is based on: Verdonschot, R. G., La Heij, W., Kiyama, S., Tamaoka, K., Kinoshita, S., & Schiller, N. O. (submitted). The functional unit of Japanese word naming: evidence from masked priming.

Abstract

Theories of language production generally describe the segment to be the basic unit in phonological encoding (e.g. Dell, 1988; Levelt, Roelofs, & Meyer, 1999). However, there is also evidence that such a unit might be language-specific. Chen, Chen and Dell (2002), for instance, using a preparation paradigm found no effect of single segments. To shed more light on the functional unit of phonological encoding in Japanese, a language often described as being mora-based, we report the results of four experiments using word reading tasks and masked priming. Experiment 1 using Japanese kana script demonstrates that primes, which overlapped in the whole mora with target words, sped up word reading latencies but not when just the onset overlapped. Experiments 2 and 3 investigated a possible role of script by using combinations of romaji (Romanized Japanese) and hiragana, and again found facilitation effects only when the whole mora overlapped, but not the onset segment. The fourth experiment distinguished mora priming from syllable priming and revealed that the mora priming effects obtained in the first three experiments are also obtained when a mora is part of a syllable (and again found no priming effect for single segments). Our findings suggest that the mora and not the segment (phoneme) is the basic functional phonological unit in Japanese language production planning.

The functional unit of Japanese word naming: evidence from masked priming

Despite the fact that languages throughout the world display a great deal of variation, most research on language production has focused on West Germanic and Romance languages such as English, Dutch and French. As a consequence, many theories of language production (e.g. Dell, 1988; Levelt, Roelofs, & Meyer, 1999) have proposed, in one way or the other, that word-form construction is performed by incrementally clustering phonological segments (phonemes) into syllabic patterns. For instance, for the word “Japan” the segments /dʒ / and /ə/ would be clustered into the first syllable /dʒ ə/ and the next three segments /p/ /æ/ /n/ would be clustered into /pæn/ thereby creating the phonological form /dʒ ə-pæn/. However, the functional unit size may differ substantially between languages. For instance, Meyer (1991) and Roelofs (2006), using a so-called implicit priming task, also known as the preparation paradigm (a paradigm which we will describe in more detail later), found that Dutch target words which overlapped in the first phoneme (e.g. boek, bijl, beer [book, axe, bear]) were read faster than words that differed in their first phoneme, indicating that the first phoneme of a word (a sub-syllabic segment) is a functional unit in Dutch. In contrast Chen, Chen, and Dell (2002; Exp 5.), using the same task, found that when to-be-named Mandarin Chinese target words overlapped in the first phoneme (e.g. mo, ma, mu, mi), no facilitation effect was apparent. Facilitation was found, however, when words overlapped in the complete first syllable. Chen and colleagues concluded that Mandarin Chinese does not allow planning at a sub-syllabic level and that syllables (not segments) are the functional units linked to speech production in Mandarin Chinese. This is in agreement with Chen, Lin, and Ferrand (2003) who reach a similar conclusion from data obtained in a masked priming study.

The idea of variable functional unit size between languages is furthermore in line with results obtained by Ferrand, Segui, and Grainger (1996) who investigated the role of sublexical phonological units (in particular the syllable) in French. Using a masked priming technique, in which participants were required to read aloud words (or name pictures), whilst keeping the amount of overlap (in segments) between masked prime words and targets constant, they manipulated

whether or not the prime constituted a whole syllable. For instance, the CV prime ba% % % % comprises the whole first syllable of BA.LADE but it is only part of the first syllable of BAL.CON. Furthermore, the CVC prime bal% % % transcends the first syllable of BA.LADE, whereas it makes up the whole first syllable of BAL.CON. Ferrand et al. (1996) always obtained greater priming (facilitation) when the prime equaled the syllable compared to when it did not. This effect was also found in non-word and picture naming. However, it disappeared in a lexical decision task, indicating that the effect likely finds its origin in the generation of articulatory output, which according to these results, is syllabically structured in French.

This is in contrast to results obtained by Schiller (1998) who in Dutch (using the same paradigm as the previous study, i.e. masked priming in word and picture naming) found that CVC primes always caused greater priming (as compared to CV primes). These results indicate that the syllable does not constitute a functional unit in the production of Dutch phonology as it does in French. In fact, Schiller (1998) found that the more segments overlap between prime and target, the more priming will be obtained in Dutch, leading to the segmental overlap hypothesis.

In the current study, we aim to extend the discussion of language-specific functional units to Japanese, a language that has been argued to be mora-timed in contrast to stress-timed Dutch/English and syllable-timed French/Chinese (see Port, Dalby, & O'Dell, 1987; Warner & Arai, 2001). In Japanese phonology, a distinction can be made between the syllable and the mora. For instance, a word such as Nihon (Japan) consists of two syllables (ni and hoN; N = nasal coda) but one can further divide this word into three moras (e.g. ni.ho.N). Moras are considered metrical units. They typically correspond to an equal number of kana symbols (Japanese script; e.g. にほん) and each mora is assumed to be generally constant in duration (e.g. ni, ho, and N last roughly equally long). Therefore, in Japanese, moras (such as the nasal coda) form an independent rhythmical structure within a syllable. The range of Japanese moras is quite limited with only 108 different items divided into 5 types (Otake, Hatano, Cutler, & Mehler, 1993), i.e. CV (Consonant – Vowel), CjV, V, N (Nasal Coda) and Q (Geminate). Japanese words usually involve simple moraic CV combinations (e.g. /ka/ or /mi/), which, in CV form, equal a syllable. However, there are also other

combinations possible: e.g. long vowels take two moras (e.g. [CVV] 脳 /nou/ “brain”), geminate (denoted as Q; e.g. [CVQ.CV] 切手 /kiQ.te/ “stamp”) and nasal coda (e.g. [CVN] 本 /hoN/ “book”) elements also take one mora each. Furthermore, a diphthong (VV) takes one mora per element (e.g. [VV.CV] 英語 /ei.go/ “English”). See Table 1 for an overview of a selection of some Japanese words and their properties.

Table 1. Example of Japanese words differing in structure, syllable, and number of moras.

Meaning	Kanji	Transcript	Kana	Structure*	Syll	Moras
paper	紙	/ka.mi/	かみ	CV.CV	2	2
book	本	/hoN/	ほん	CVN	1	2
stamp	切手	/kiQ.te/	きって	CVQ.CV	2	3
eigo	英語	/ei.go/	えいご	VV.CV	2	3
dragon	龍	/ryuu/	りゅう	CjVV	1	2

**Note: N – nasal coda, Q – geminate, Cj – consonant with palatal glide (C+ や ゆ よ).*

As stated before, mora structures in Japanese are well reflected in the kana script, which constitutes moraic symbols divided into hiragana and katakana. These scripts were adapted from the more complicated logographic kanji characters to allow a more precise phonological representation of native Japanese vocabulary words and their inflectional morphology. The kana scripts are phonological in nature by being based on a one-to-one correspondence of kana-to-mora, but their usage differs. Katakana is mainly used to represent loan words, typically those from languages with alphabetic writing systems, proper nouns, and proper names from foreign languages (e.g. マクドナルド /ma.ku.do.na.ru.do/ for “McDonalds”, the fast food restaurant). Hiragana, on the other hand, is used for native elements in the language as, for example, function words, inflectional affixes, onomatopoeia, as well as those borrowings that have been assimilated into the language.

Evidence demonstrating that the mora plays an important role in language production comes from speech errors. For instance,

Kubozono (1989) found that Japanese are more likely to elicit errors which respect mora boundaries, e.g. /toma(re)/ “stop!” and /(suto)Qpu/ “stop!” would be blended into /to.maQ.pu/ (thereby including the geminate (= mora) and not simply the last syllable). Furthermore, many Japanese language games also respect moraic structure. Consider a game played by children (and adults) such as “shiritori” (Katada, 1990; p. 641) in which players have to come up with a follow-up word that starts with the last element (i.e. mora) of the previously heard word, e.g. when /kao/ “face” is heard, /oN.ga.ku/ “music” would be a valid answer, indicating that not the entire diphthong, but only the last mora i.e. /o/, is important. When a player says /ka.mi/ “paper/god/hair”, a good continuation would be /mi.zu/ “water”; however, /mi.zu/ cannot be followed by /zu.boN/ “trousers” since Japanese has no word that begins with a nasal coda, and as such the player would lose the game.

Empirical evidence that the mora plays an important role in the Japanese language comes from studies on speech segmentation (e.g. Otake et al., 1993; Cutler & Otake, 1994). In these studies, participants were required to monitor Japanese words for the appearance of specific strings of segments, e.g. CV (“na”) or CVN (“naN”). They found that there was no detection advantage for CV structures whether they were part of a CV.CV.CV or CVN.CV target word, contrary to what a syllable hypothesis would predict (CVN.CV being more difficult as CV target is only part of the syllable CVN). In addition, CVN targets (e.g. naN) were much easier to detect in CVN.CV words such as “naN.ka” where the nasal coda served as a separate mora compared with “na.no.ka” where the target is only part of a mora. These patterns were not obtained when the same experiment was repeated with English participants, leading Otake et al. (1993) and Cutler and Otake (1994) to conclude that Japanese speech segmentation respects mora boundaries².

However, there could be other factors playing a role such as the type of task being used. This becomes evident as the same authors (Otake & Cutler, 2003) using a Word Reconstruction paradigm found

² It is worthwhile to mention that in the speech segmentation literature an effect of script on segmentation has been shown. Inagaki, Hatano, and Otake (2000) compared Japanese preschool children who were not yet able to read to older Japanese children who were able to read kana (moraic script). They found that the literate children (e.g. who could read kana) shifted their segmentation preference from a syllable-based representation towards a mora-based representation.

that subjects were sensitive to sub-moraic (e.g. segmental) information. In this paradigm, subjects had to judge whether auditorily presented non-words could be reconstructed into real words. Otake and Cutler showed that participants found it easier to reconstruct /ka.me.ra/ from non-words which had a partial mora preserved, such as /ki.me.ra/ and /na.me.ra/ compared to /ni.me.ra/ in which the whole mora was different. In a subsequent lexical decision task (LDT), participants heard words and non-words for which the so-called non-word uniqueness point (NUP), i.e. the point in the word when it starts differing from a real word (for example, the “i” in Japin), was manipulated. Otake and Cutler found that the sooner a word became a non-word, the faster participants could reject it. In their regression analysis, Duration and Phoneme (but not Moras) significantly accounted for a portion of the variance as independent predictors. Therefore, Otake and Cutler (2003) concluded that segmental information contributes to word recognition. This is in line with findings reported by Tamaoka and Taft (1994) who presented participants with katakana strings to which a word non-word judgment had to be made (LDT). Participants found it harder to reject /ko.me.ra/ which is one mora and one phoneme different from the real loan-word /ka.me.ra/ compared to /so.me.ra/ which is one mora and two phonemes different. Tamaoka and Taft concluded therefore that participants were sensitive to segmental information when processing Japanese kana.

A more recent study by Kureta, Fushimi, and Tatsumi (2006) directly investigated the functional encoding unit in Japanese speech production using the preparation (or implicit priming) paradigm (e.g. Meyer, 1991). In this paradigm, participants are typically required to learn small sets of semantically related word pairs (called prompt-response pairs). Participants are subsequently asked to produce the corresponding response word upon presentation of a prompt, for instance, a participant should respond by saying “ring” after seeing the prompt “marriage”. Creating small blocks differing in phonological consistency can influence the presence and absence of priming. For instance, prompts can be presented that result in phonologically congruent (or homogeneous) blocks, such as “rule, rain, ring” (all starting with the phoneme /r/), or prompts can be presented that have no phonological relationship (so-called heterogeneous blocks) such as “cloud, book, ring”. Typically, reaction times are shorter in the

homogeneous blocks compared to the heterogeneous blocks when there is phonological congruency in the first syllable (but not the rhyme; see Meyer, 1991). Kureta et al. (2006) found that for Japanese this form preparation effect only occurred when initial consonant and vowel (CV) were similar (e.g. katsura, kabuki, kaban) but not when just the consonant (C) or consonant + palatal glide (Cj) were similar (e.g. katsura, kujira, kofun and gyakuten, gyuuniku, gyousei, respectively). Kureta et al. concluded on the basis of these data that the mora plays a crucial role in the construction of the phonological form of a Japanese word. Although their study was well designed and their results and interpretation were clear-cut, there are some additional issues that, in our view, deserve further examination.

One important issue concerns the fact that although Kureta et al. (2006) used different scripts (to avoid character repetition in homogenous blocks, e.g. かつら [katsura; hiragana], 歌舞伎 [kabuki; kanji] and 鞆 [kaban; kanji]), they did not include stimuli written in romaji (Romanized Japanese, e.g. using alphabetic script). Kureta et al. (2006) did not include these words as words written in romaji have a low orthographic plausibility (i.e. a subjective rating scale concerning the preferred orthographic form a particular word is usually written in; Amano & Kondo, 1999). However, including romaji would have had the benefit to include a script that involves processing of individual phonemes (as in languages with an alphabetic script), and although the orthographic plausibility is low, it is taught at primary school, and many Japanese frequently use romaji to input Japanese text on computers, cell phones, and other electronic devices. Therefore, many Japanese are well able to read and write Japanese using romaji.

To briefly summarize: there are studies demonstrating contrasting effects (e.g. Cutler & Otake, 1994 [speech segmentation] vs. Otake & Cutler, 2003 [word reconstruction]) with respect to functional unit size. Yet, a study by Kureta et al. (2006) yielded preparation effects for the whole mora only and not the segment.

We believe that although there is a growing body of evidence for the mora as a functional unit of phonological encoding in Japanese (e.g. Kureta et al., 2006) it is not conclusive at this moment. In this paper, we are especially interested in whether the so-called Masked Onset Priming Effect (MOPE), a form priming effect due to the initial

segment only, can be found in Japanese. This, to our knowledge has not been tested so far. The MOPE refers to the finding that when a target word (e.g. HOME) is preceded by a prime word sharing the onset (e.g. hill) briefly presented under visually masked conditions, reading aloud of the target is facilitated compared to when prime (e.g. pill) and target (HOME) do not share their onset (e.g., Grainger & Ferrand, 1996; Forster & Davis, 1991; Schiller, 2004; see Kinoshita, 2003, for a review)). Forster and Davis proposed that the MOPE originates in the non-lexical route in naming, i.e., when the sub-lexical phonology from a letter string is computed from its orthography (also called orthography-to-phonology conversion or OPC). In contrast, Kinoshita (2000) proposed that the locus of the MOPE occurs later than OPC, i.e. at the level of phonological encoding of the speech response, which she termed the speech-planning account. More specifically, the phonology from the prime and the target are proposed to compete for a slot during the segment-to-frame association process in speech planning. A mismatch between the onset of the prime and target will cause an inconsistency in the speech plan and resolving this conflict comes at a processing cost and hence takes longer (Kinoshita & Woollams, 2002).

In this paper, we will examine native speakers of Japanese using a masked priming paradigm to establish whether or not a MOPE can be found in Japanese. If so, this could be taken as evidence that the phoneme can function as an independent functional planning unit. A factor of influence, as argued earlier, could be the role of script in the detection and processing of phonological information. Consider, for instance, findings from two masked priming experiments using Korean by Kim and Davis (2002) who found a (marginal) onset effect when primes and targets were presented in hangul (script favoring segments) but not when the target was presented in hanja (a script favoring syllables). These authors concluded that the underlying nature of the script may have been responsible for the discrepancy in results between the two experiments. Such findings may be interpreted in terms of the nature of target script governing the units that drive the phonological encoding process, mediated by the unit of the orthography-to-phonology mapping process. That is, when the target script is alphabetic, as in Korean Hangul, Cyrillic script and, importantly, romaji, the phonology would likely become available segment-by-segment, and the segment-to-frame association process of

the target, as the name implies, could indeed proceed on a segment-by-segment basis. A MOPE could be observed in this case because the segmental information in the prime overlapping with the target is useful because the target script allows the phonological encoding process to proceed in a segment-by-segment fashion.

In contrast, when the target script is hanja, the target phonology would become available only in the syllable-sized unit. In this case, the segment-to-frame association process of the target could not proceed segment-by-segment, but syllable-by-syllable, and hence overlapping segmental information made available by the hangul prime may not produce onset priming. Thus, manipulating the target script in Japanese (mora-based kana vs. segment-based romaji) is important for interpreting the presence of absence of a MOPE: It may be either because segments are not the functional unit of phonological encoding in Japanese, or because the unit of orthography-to-phonology mapping was not a segment. On the reasonable assumption that the nature of target script governs the unit of orthography-to-phonology mapping, a MOPE in Japanese may depend on the nature of the target script. In contrast, if the functional unit of phonological encoding in Japanese is a mora and not a segment, a MOPE should be absent even when the target script is segmental (e.g. romaji).

We present the empirical results of four experiments. Throughout the experiments, the degree of overlap was always manipulated from one consonantal segment in the MOPE condition (e.g. target: ka.ze – primes: ko.to vs. so.to) to whole mora (CV) overlap (e.g. target: ka.ze – primes: ka.mi vs. na.mi). The first experiment aimed to ascertain whether the masked priming paradigm could provide empirical results that could distinguish between onset and mora priming. The second and third experiments investigated the possible role of script, i.e. does romaji allow for onset priming instead of kana. Lastly, the fourth experiment aimed to distinguish mora from syllable priming by introducing C and CV primes in nasal coda and geminated targets (e.g. CVN.CV and CVQ.CV). We hypothesized that if the mora (and not the phoneme or syllable) reflects the functional phonological unit in Japanese, we should only obtain a priming effect when the whole mora overlaps. However, if the paradigm (masked vs. implicit priming) and/or script type (kana vs. romaji) indeed have an influence, we may also expect segmental priming effects (e.g. MOPE) to surface.

Experiment 1: Segment versus mora priming

In the first experiment, the participants' task was to name hiragana strings that were presented on the computer screen. Targets were preceded by masked primes and participants were divided in two groups. In both groups, the target was presented in hiragana (e.g. すし /su.shi/), but although the MOPE group also received the primes in hiragana (e.g. せん /seN/ vs. れん /reN/) the mora group received the primes in katakana (e.g. スミ /su.mi/ vs. グミ /gu.mi/) to avoid visual repetition. As Japanese does not have capital letters, presenting both prime and target in hiragana in the MOPE setting would amount to a complete visual repetition (e.g. すみ /su.mi/ – すし /su.shi/). If the segment functions as an independent unit in Japanese, we expect to find priming effects in both groups, with perhaps a larger effect for the mora group (as more segments are overlapping). If the mora and not the segment functions as an independent planning unit, we expect to find priming effects only in the mora group.

Method

Participants. Twenty-two undergraduate students from Yamaguchi University, Japan (15 female, 7 male; average age: 20.3 years; SD = 1.3) took part in the MOPE (segment overlap) part of this experiment and twenty undergraduate students from Reitaku University, Japan (15 female, 5 male; average age: 23.7 years; SD = 4.9) took part in the MORA (mora overlap) part. All students received financial compensation for their participation. All participants were native speakers (and fluent readers) of Japanese and had normal or corrected-to-normal vision.

MOPE Stimuli. Forty-two bi-moraic words were selected as targets and an additional forty-two bi-moraic, semantically unrelated words as primes. All words (primes and targets) were selected such that a broad variety of moras appeared at the first position (ka, ki, ku, ke, ko; sa, shi, su, se, so; etc.). Primes and targets were presented in hiragana (Japanese syllabic script). There was no visual overlap between prime and target (e.g. target かぜ /kaze/ with the onset prime こと /koto/ vs. the control prime そと /soto/). Some mora combinations appear more frequently than others, for instance, the bi-moraic combination こと /koto/ is more frequent than the bi-moraic combination カオ /kao/ (for more information, see the freely available database on bi-moraic frequencies on <http://www.lang.nagoya->

u.ac.jp/~ktamaoka/down_en.htm). Therefore, whenever possible, primes were matched to form pairs thereby avoiding bi-moraic frequency effects (e.g. the target せき /seki/ would be paired with the onset prime そと /soto/ vs. the control prime こと /koto/).

MORA Stimuli. Thirty bi-moraic words from the same corpus as the MOPE stimuli were selected as targets and an additional thirty bi-moraic, semantically unrelated words as primes. As the MORA prime and target when both presented in hiragana have complete character repetition, primes were presented in katakana and targets in hiragana. Care was taken to avoid onset moras having visually matching characteristics in both scripts, e.g. words starting with “ka” (hiragana: か and katakana: カ) were not selected. Similar to the MOPE part, whenever possible the MORA primes formed pairs avoiding any frequency effects. See Appendix A for an overview of the stimuli used in these two parts of Experiment 1.

Design. In both the MOPE and MORA parts of this experiment, targets were preceded either by an overlapping prime (first segment in the MOPE case or complete mora in the MORA case) or a control prime. Participants in the MOPE part received 84 trials and participants in the MORA part 60 trials. Pseudo-random lists were constructed for each individual participant such that phonologically or semantically related primes or targets had at least a distance of two trials to avoid unintended priming effects.

Procedure. In all reported experiments, we used the software package E-prime 2.0 combined with a voice key for stimulus presentation and data acquisition. Participants were seated approximately 60 cm from a 17 inch LCD computer screen (Eizo Flexscan P1700 with a screen cycle refresh rate of 60 Hz) in a quiet room at Yamaguchi University (MOPE part) or Reitaku University (MORA part) and were tested individually. After a short explanation of the experimental paradigm and two warm-up trials, participants started the experiment proper. A trial comprised the presentation of a fixation cross (750 ms) followed by a forward mask consisting of hash marks (##) (500 ms) and subsequently a hiragana (MOPE part) or katakana (MORA part) prime (50 ms) that was replaced immediately by the target, which disappeared when the participant responded or after maximally 3,000 ms. Masks, primes, and targets were presented using the MS Mincho font (36 pt). All items appeared as black

characters on white background. Target words consisted of two hiragana characters which covered about $1.5^\circ \times 2.7^\circ$ of visual angle. After each trial, the experimenter recorded the accuracy of the response. There was a short break halfway the experiment after which two warm-up trials preceded the continuation of the experiment. Naming latencies were measured from target onset. Participants were specifically instructed to respond as fast as possible while avoiding errors. They were not informed about the existence of the prime. After the experiment informal interviewing showed that participants were generally unaware of the presentation of the primes.

MOPE Results. Naming latencies exceeding three standard deviations per participant per condition and voice-key errors were excluded from the analysis (comprising 6.7% of the data). As errors were few (1.7% of the data) and equally distributed across conditions, an error analysis was not performed. The reaction time analysis showed that there were no significant priming effects for onset overlap compared to control primes, both $t_s < 1$.

MORA Results. Naming latencies exceeding three standard deviations per participant per condition and voice-key errors were excluded from the analysis (comprising 7.7% of the data). As errors were few (0.1% of the data) and equally distributed across conditions, an error analysis was not performed. The reaction time analysis revealed that there was a significant priming effect (15 ms) when the whole mora prime overlapped compared to control, $t_1(19) = 3.72$, $SD = 18.11$, $p < .01$; $t_2(29) = 3.45$, $SD = 24.23$, $p < .01$, see Table 2 for an overview.

Table 2.
Mean Naming Latencies (in Milliseconds) and Error Rates (in %) in Experiment 1 on the MOPE and MORA groups as a Function of Relatedness (overlap and control).

	MOPE group (N=22)		MORA group (N=20)	
	RT (SD)	%E	RT (SD)	%E
control	519 (59)	1.6	597 (26)	0.0
overlap	520 (65)	1.8	582 (27)	0.2
effect	-1	-0.2	15	-0.2

Discussion. The results of Experiment 1 show no effect of onset priming, but do show an effect of mora overlap priming. This finding suggests that the functional unit in Japanese language production is not the segment, but the mora. However, it should be noted that prime and target were presented in moraic kana (hiragana/katakana) script, which does not favor individual activation of segmental elements and as such might have led to the observed pattern of results. Furthermore, in the MORA prime group, there was a script switch between prime and target (from katakana to hiragana) that was absent in the MOPE prime group (both hiragana). This might have led to the shorter RTs observed for this group and in turn due to a floor effect might have obscured a potential priming effect. Therefore, we decided to include romaji stimuli in Experiments 2 and 3 to determine whether or not including a script favoring the processing of segments might lead to different results whereas Experiment 4 deals with the script change issue.

Experiment 2: The effects of kana and romaji scripts on masked onset priming

This experiment was designed to examine whether the findings of Experiment 1 could be replicated using an experimental situation that uses visually presented segmental information. Experiment 2 employed hiragana and romaji targets, which were all preceded by romaji primes. If the type of script used caused the absence of onset priming in Experiment 1, Experiment 2 should remedy this situation. If, however, Japanese do plan their speech in moraic units, then Experiment 2 should not show any onset priming either, despite of the romaji script advantage towards the segment.

Method

Participants. Forty undergraduate students from Nagoya University, Japan (25 female, 15 male; average age: 22.6 years; SD = 4.8) took part in this experiment in exchange for financial compensation. All participants were native speakers (and fluent readers) of Japanese and had normal or corrected-to-normal vision.

Stimuli. Forty-two bi-mora words were selected as targets and an additional forty-two bi-mora semantically unrelated words as primes (see Appendix B). Using the pairing procedure from Experiment 1, these targets were combined with the primes to form

168 target-prime pairs. Targets appeared either in hiragana or romaji (capitals). All primes were in romaji.

Design. A 2 (Target Type: hiragana or romaji) x 2 (Prime Type: MOPE or MORA) x 2 (Relatedness: overlap vs. control) within subjects factorial design was implemented. To avoid frequent repetition of targets and primes, each participant was subjected to only two repetitions instead of eight per target resulting in 84 trials per participant. All forty-two targets were presented once in romaji and once in hiragana. To ensure that each condition for each Target Type and Prime Type appeared equally often, participants were assigned to individually generated pseudo-random lists by E-prime 2.0 which were constructed such that both MORA/MOPE conditions for the whole Experiment appeared equally often per Target Type. Furthermore, in these lists phonologically or semantically related primes or targets had at least a distance of two trials to avoid unintended priming effects. Between participants, the design included all of the experimental conditions. Within participants, the order of Target Type (romaji or hiragana) was counterbalanced.

Procedure. The apparatus was the same as the first experiment. Participants were tested individually in a quiet room at Nagoya University. The trial sequence was identical to Experiment 1. Masks, primes, and targets were displayed in Courier New font (28 pt). All items appeared in the center of the screen as black characters on white background. Each upper-case romaji target word covered approximately 0.95° of visual angle from a viewing distance of 60 cm. Romaji target words were between four and five letters in length, subtending between 2.7° and 3.2° of visual angle. Hiragana targets all consisted of 2 mora characters and covered about $0.95^\circ \times 1.6^\circ$ of visual angle.

Table 3.
Mean Naming Latencies (in Milliseconds) and Error Rates (in %) in
Experiment 2 as a function of Target Type (hiragana vs. romaji),
Prime type (MOPE vs. MORA) and Relatedness (overlap vs. control).

	Hiragana Target				Romaji Target			
	MOPE		MORA		MOPE		MORA	
	RT (SD)	%E	RT (SD)	%E	RT (SD)	%E	RT (SD)	%E
C	527 (76)	0.8	521 (87)	0.0	753 (123)	3.2	765 (132)	2.4
O	530 (89)	0.8	525 (75)	0.8	737 (110)	1.6	732 (134)	0.8
E	-3	0.0	-4	-0.8	16	1.6	33	1.6

*C = Control, O = Overlap, E = effect

Results. Naming latencies exceeding three standard deviations per participant per condition and voice-key errors were excluded from the analysis (comprising 1.0% of the data). As errors were few (1.3% of the data), an error analysis was not performed. Please see Table 3 for an overview of the results. There was a main effect of Target Type (hiragana or romaji) indicating that targets in hiragana were read aloud 221 ms faster than targets in romaji, $F(1,39) = 215.10$, $MSe = 18170.98$, $p < .001$; $F(1,41) = 675.14$, $MSe = 6076.46$, $p < .001$; $\eta^2_p(1,62) = 163.1$, $p < .001$. Furthermore, there was a main effect of Relatedness in the subject analysis (overlap vs. control), $F(1,39) = 5.48$, $MSe = 3021.98$, $p < .05$; $F(1,41) = 2.03$, $MSe = 3710.03$, n.s.; $\eta^2_p(1,67) = 1.48$, $p = .23$. Target Type interacted significantly with Relatedness in the subject analysis, $F(1,39) = 6.92$, $MSe = 2307.84$, $p < .05$, and approached significance in the item analysis, $F(1,41) = 3.61$, $MSe = 4258.12$, $p = .07$; $\eta^2_p(1,74) = 2.37$, $p = .13$. Given this interaction, paired t-tests were conducted for each Target Type, Prime Type, and Trial Type. These t-tests showed that there were no significant effects when Target Types were hiragana, all $t_s < 1$. When the Target Type was romaji, there was a significant priming effect (33 ms) in case the whole mora overlapped, $t(39) = 2.64$, $SD = 79.22$, $p < .05$; $t(41) = 2.02$, $SD = 95.23$, $p < .05$, but not when only the first segment (16 ms) overlapped, $t(39) = 1.26$, $SD = 82.71$, n.s.; $t_2 < 1$.

Discussion. Reading romaji compared to hiragana takes significantly more time (221 ms). This could account for the absence

of priming effects from romaji primes during hiragana target naming. The prime might simply be too late to exert an influence. However, when the target is also in romaji (and also takes more time to read), then both prime and target are delayed, thereby again allowing for priming effects to surface, as evidenced by our results. In that latter case, it is again found that only when the whole mora overlaps, facilitation from priming becomes significant. However we feel that, although not significant, the observation of the sizable 16 ms difference between overlap and control prime in the MOPE part induces the necessity to test the romaji targets again, though, this time with hiragana primes (thereby eliminating any processing cost due to script unfamiliarity).

Experiment 3: Effect of hiragana primes on romaji targets

Experiment 2 showed that romaji targets take longer to read than hiragana targets. Combined with romaji primes (which also took longer), no effect was found for any hiragana target, but again a MORA and not MOPE effect was found when romaji targets were used. This third experiment aims to replicate and extend these findings by again using romaji targets, thus again allowing the response to dictate the unit of speech planning (phonemes). However, current primes were all in hiragana, which allows for a comparison between fast (current) and slow access (Experiment 2) to the prime.

Method

Participants. Thirty-one undergraduate students from Nagoya University, Japan (16 female, average age: 22 years; SD = 3.7) took part in this experiment in exchange for financial compensation. All participants were native speakers (and fluent readers) of Japanese and had normal or corrected-to-normal vision.

Stimuli. Forty-two bi-moraic words were selected from the same pool as in Experiment 2 (see Appendix B) except targets were all in romaji (letters) and all primes were in hiragana.

Design. A 2 (Prime Type: MOPE or MORA) x 2 (Relatedness: overlap vs. control) within subjects factorial design was implemented. To avoid recurrent repetition of targets and primes, each participant received two repetitions instead of four per target equalling 84 trials per participant. Each participant received two blocks of 42 trials, encompassing all 42 targets once per block. Participants were assigned to pseudo-random lists (generated per participant) which were

constructed such that MORA/MOPE conditions appeared equally often per group and phonologically or semantically related primes or targets had at least a distance of two trials. Between participants the design included all the experimental conditions, and within participants the order of blocks was counterbalanced.

Procedure. The same procedure was used as in Experiment 2.

Results. Naming latencies exceeding three standard deviations per participant per condition and voice-key errors were excluded from the analysis (comprising 1.3% of the data). As errors were few (2.2% of the data) and evenly distributed across conditions, an error analysis was not performed. See Table 4 for an overview of the results. There was no main effect of Prime Type (MOPE or MORA), $F(1,30) = 1.24$, $MSe = 933.20$, n.s.; $F(1,41) = 1.33$, $MSe = 815.20$, n.s.; $\min F'(1,68) < 1$, but there was a main effect of Relatedness (overlap vs. control). Targets preceded by overlapping primes were read on average 11 ms faster compared to being preceded by control primes, $F(1,30) = 5.75$, $MSe = 653.30$, $p < .05$; $F(1,41) = 5.12$, $MSe = 1,439.90$, $p < .05$; $\min F'(1,70) = 2.71$, $p = .10$. There was also a significant interaction between Prime Type and Relatedness, $F(1,30) = 11.5$, $MSe = 1,360.60$, $p < .01$; $F(1,41) = 19.10$, $MSe = 1,324.00$, $p < .001$; $\min F'(1,61) = 7.18$, $p < .001$. To explore this interaction, we performed paired t-tests and found that MOPE primes again did not supply reliable priming effects, $t(30) = 1.52$, $SD = 42.00$, n.s.; $t(41) = 1.46$, $SD = 50.30$, n.s., but MORA overlapping primes yielded facilitation effects compared to their control primes (33 ms), $t(30) = 3.92$, $SD = 47.59$, $p < .001$; $t(41) = 4.48$, $SD = 54.77$, $p < .001$.

Table 4.
Mean Naming Latencies (in Milliseconds) and Error Rates (in %) in Experiment 3 on the romaji Targets as a Function of Prime Type (MOPE/MORA) and Relatedness (Overlap and Control).

	MOPE prime		MORA prime	
	RT (SD)	%E	RT (SD)	%E
control	733 (98)	2.0	761 (104)	2.4
overlap	744 (110)	1.6	728 (107)	2.8
effect	-11	0.4	33	-0.4

Discussion. This experiment yielded essentially the same results as Experiment 2: MOPE primes do not induce a significant facilitation effect. The results from Experiments 2 and 3 indicate that even when romaji script is used to present the prime (thereby favoring segmental processing), the functional unit of phonological encoding in Japanese is likely to be the mora, as evidenced by the significant CV but not C priming effects reported in the previous three experiments. However, two main issues need to be addressed: (a) the script change issue of Experiment 1 (MORA: katakana primes and hiragana targets) and (b) the fact that in most stimuli of Experiments 1-3 moras were indistinguishable from syllables (e.g. *こと* /ko.to/).

Experiment 4a: Distinguishing the CV from the syllable: MOPE group.

A possible reason for the discrepancy between the MOPE and MORA group in Experiment 1 might be that in the MORA group there was a script change (from katakana to hiragana), whereas this change was absent in the MOPE group. We remedied this in Experiments 4a and 4b by using katakana primes and hiragana targets in both the MOPE and MORA groups. However, a more important issue may be the fact that in most stimuli we used in the previous experiments the mora was indistinguishable from the syllable. For instance, the moras /ko/ and /to/ in *こと* /ko.to/ constitute two moras but also two syllables. Therefore, in Experiment 4, we included three groups of target words. In the first group, bi-moraic words (where mora equals syllable) were used. In the second and third group, the first CV does not constitute the whole syllable, i.e. nasal coda words (CVN; e.g. /hoN.da/ “Honda” [CVN.CV]; group 2) and geminate obstruents, (CVQ; e.g. /saQ.ka/ “soccer” [CVQ.CV]; group 3). As previous experiments never showed any sign of onset priming in Japanese, we expect that Experiment 4a (which only employs onset overlap primes) would not show priming in any of the groups. Furthermore, regarding Experiment 4b, we predict that if the functional unit in Japanese is the mora, we will obtain CV priming effects in all three groups, without an interaction between Group and Relatedness, as the priming effect should be similar between groups. In contrast, when the previously observed priming effects were in fact

due to syllabic overlap, we expect CV priming only to occur in the bi-moraic group and not the nasal coda and geminate groups.

Method

Participants. Twenty-seven undergraduate students from Yamaguchi University, Japan (21 female, average age: 22 years; SD = 7.5) took part in this experiment in exchange for financial compensation. All participants were native speakers of Japanese and had normal or corrected-to-normal vision.

Stimuli. Forty-two bi-mora words were selected for each group, totalling 126 targets (see Appendix C). Targets were all in hiragana and primes in katakana. Congruent prime-target pairs were combined such that they overlapped in the onset (MOPE).

Design. A 3 (Group: bi-moraic, nasal coda, and geminate) x 2 (Relatedness: overlap vs. control) within subjects factorial design was implemented. To avoid recurrent repetition of targets and primes, each participant was subjected to one repetition instead of two per target equalling 126 trials per participant. Each participant was subjected to two blocks of 63 trials (groups were mixed in blocks). Participants were assigned to pseudo-random lists (generated per participant) which were constructed such that groups/conditions appeared approximately equally often per block and phonologically or semantically related primes or targets had at least a distance of two trials. Between participants, the design included all the groups and experimental conditions, and within participants the order of blocks was counterbalanced.

Procedure. The same procedure was used as in Experiment 2. Masks, primes, and targets were displayed in Courier New font (36 pt). All items appeared as black characters on white background. Each target word covered approximately 1.3° of visual angle from a viewing distance of 60 cm. Target words were between two and five hiragana characters in length, subtending between 1.9° and 4.7° of visual angle.

Results and Discussion. Table 5 provides an overview of the results. Naming latencies exceeding three standard deviations per participant per condition and voice-key errors were excluded from the analysis (comprising 5.1% of the data). As errors were few (0.5% of the data), an error analysis was not performed. There was a significant main effect of group, $F(2,52) = 7.50$, $MSe = 1178.16$, $p < .001$, $F(2, 123) = 4.60$, $MSe = 2653.40$, $p < .05$; $\eta^2(2,169) = 2.85$, $p = .06$,

reflecting the fact that the bi-moraic targets were named faster than the nasal coda and geminate targets, which did not differ from each other. There was no main effect of Relatedness in the subjects analysis, $F(1,26) = 1.70$, $MSe = 678.80$, n.s., but it was significant in the items analysis, $F(1,123) = 6.60$, $MSe = 468.20$, $p < .05$; $\eta^2(1,41) = 1.35$, $p = .25$. The interaction between Group and Relatedness was not significant, both $F_s < 1$. Planned t-tests showed that there was no reliable effect of overlap for any group: bi-mora, all $t_s < 1$; nasal coda, $t_1 < 1$, $t_2(29) = 1.4$; $SD = 30.14$, n.s.; and geminate, $t_1 < 1$; $t_2(29) = 1.97$, $SD = 34.23$, $p = .056$. Overall, there seems to be no reliable effect of segment overlap on group level when only the segment overlaps.

Table 5.
Mean Naming Latencies (in Milliseconds) and Error Rates (in %) in Experiment 4a as a Function of Group (bi-moraic, nasal coda and geminate) and Relatedness (overlap and control).

	bi-moraic		nasal coda		geminate	
	RT (SD)	%E	RT (SD)	%E	RT (SD)	%E
control	527 (69)	0.6	552 (92)	0.6	545 (78)	0.0
overlap	522 (66)	0.0	547 (93)	1.2	539 (84)	0.6
effect	5	0.6	5	-0.6	6	-0.6

Experiment 4b: Distinguishing the CV from the syllable: MORA group.

Method.

Participants. Twenty-eight undergraduate students from Nagoya University, Japan (12 female, average age: 19 years; $SD = 3.3$) took part in this experiment in exchange for financial compensation. All participants were native speakers of Japanese and had normal or corrected-to-normal vision.

Stimuli. Thirty-two bi-moraic words were selected for each group, totalling 90 targets (see Appendix C). Targets were all in hiragana and primes in katakana. Congruent prime-target pairs were combined such that they had CV (mora) overlap.

Design. A 3 (Group: bi-moraic, nasal coda, and geminate) $\times$ 2 (Relatedness: overlap vs. control) within subjects factorial design was

implemented. To avoid recurrent repetition of targets and primes, each participant was subjected to one repetition instead of two per target equalling 90 trials per participant. Each participant was subjected to two blocks of 45 trials (groups were mixed in a block). Participants were assigned to pseudo-random lists (generated per participant) which were constructed such that groups/conditions appeared approximately equally often per block and phonologically or semantically related primes or targets had at least a distance of two trials. Between participants, the design included all the groups and experimental conditions, and within participants, the order of blocks was counterbalanced.

Procedure. Same procedure was used as in Experiment 4a. Results and Discussion. Table 6 gives an overview of the results. Naming latencies exceeding three standard deviations per participant per condition and voice-key errors were excluded from the analysis (comprising 2.2% of the data). As errors were few (0.8 % of the data), an error analysis was not performed. There was a main effect of Group, $F(2,54) = 38.50$, $MSe = 365.46$, $p < .001$; $F(2,87) = 12.20$, $MSe = 1212.64$, $p < .001$; $\min F' (2,130) = 9.26$, $p < .001$, which reflects significant differences in the overall means of the three groups. Furthermore, there was a main effect of Relatedness (overlap vs. control), $F(1,27) = 24.10$, $MSe = 279.70$, $p < .001$; $F(1,87) = 17.20$, $MSe = 379.24$, $p < .001$; $\min F' (1,97) = 10.04$, $p < .001$, reflecting the fact that on average targets which overlapped with the prime were named 13 ms faster. Most importantly, there was no interaction between Group and Relatedness, all $F_s < 1$ indicating that the effect of Relatedness holds for all Target Groups. This was further confirmed by planned t-tests which showed that the effect of overlap held for each group: bi-mora, $t(27) = 3.05$, $SD = 25.93$, $p < .01$; $t(29) = 2.86$; $SD = 27.45$, $p < .01$, nasal coda, $t(27) = 2.54$, $SD = 22.72$, $p < .05$; $t(29) = 2.30$; $SD = 23.34$, $p < .05$; and geminate, $t(27) = 3.07$, $SD = 20.95$, $p < .01$; $t(29) = 2.10$, $SD = 31.26$, $p < .05$. A point of concern (which was also raised in the introduction and for Experiment 1) is that the moraic nature of the kana script used in Experiments 4a and 4b might favor the processing from orthography-to-phonology. Although we cannot unequivocally refute the claim that the use of kana script has been responsible for the obtained effect, we do believe that our Experiments 2 and 3 reasonably demonstrated that

when romanized Japanese (romaji) targets are employed, no reliable priming occurs for the segment.

Table 6.
Mean Naming Latencies (in Milliseconds) and Error Rates (in %) in Experiment 4b as a Function of Group (bi-moraic, nasal coda, and geminate) and Relatedness (overlap vs. control).

	bi-moraic		nasal coda		geminate	
	RT (SD)	%E	RT (SD)	%E	RT (SD)	%E
control	471 (61)	0.0	501 (74)	2.4	489 (69)	0.6
overlap	456 (68)	0.0	490 (77)	1.8	477 (72)	0.0
effect	15	0.0	11	0.6	12	0.6

General Discussion

Many studies using different paradigms have shown that the segment (e.g. phoneme) can be conceived as an independent functional element of speech production planning in many alphabetic languages (e.g. implicit priming/preparation effect: Meyer, 1991; masked onset priming effect: Forster & Davis, 1991; Schiller, 2004; Kinoshita, 2003). Studies having investigated other non-alphabetic languages, e.g. Chinese (Chen, Chen and Dell, 2002) found that this unit might be language specific, as segment preparation in Mandarin Chinese was absent (but syllable priming was present). In our target language Japanese contrasting patterns were reported in comprehension tasks (Cutler & Otake, 1994; Otake & Cutler, 2003) and currently there is only one empirical paper to-be-found regarding production tasks (Kureta et al., 2006). The aim of the present research was to further examine a possible role for the segment as independent functional unit in Japanese speech production by making use of the well-established masked priming paradigm.

The outcomes of four experiments support previous findings obtained by Kureta et al. (2006) who proposed that the mora and not the segment is the functional unit of phonological encoding in Japanese. Introducing romaji stimuli, which by their visual characteristics were assumed to favor segmental processing (our Experiments 2 and 3) also did not lead to onset priming. This indicates that the observed absence of onset priming (e.g. our findings and those

of Kureta et al., 2006) cannot have been due to the fact that kana and/or kanji scripts were used. Importantly, we ruled out the possibility that the absence of MOPE was due to the nature of target script governing the unit of phonological encoding, as presenting target in romaji (in which the orthography-to-phonology mapping process proceeds by letter-to-segment) did not produce MOPE.

Before accepting the conclusion that the mora and not the phoneme is the functional unit of language production in Japanese, we consider two points. The first point is that the conclusion that there is no MOPE in Japanese corresponds to the null hypothesis. Of course, conventional significance testing cannot state the evidence for a null hypothesis: Non-significant p-values indicate “failures to reject” the null hypothesis, not the support for it. However, recent developments in the Bayesian data analysis technique (e.g., Dienes, 2008; Rouder, Speckman, Sun, & Morey, 2009) have made this possible in the form of computation of the Bayes factor. Bayes factor is generally interpreted as the weight of evidence provided by the data, and is essentially an odds ratio, ranging between 0 to infinity. Bayes factor of 1 means that the data provide equal evidence for the null and alternative hypotheses, and values greater than 1 favor the null hypothesis, and values less than 1 favor the alternative hypothesis. According to Jeffreys’ (1961) classification scheme, Bayes factor of 1-3 would be considered anecdotal evidence (“worth no more than a bare mention”) for the null hypothesis, 3-10 would be substantial evidence, 10-30 would be strong evidence, and so on. We computed the Bayes factor for the null hypothesis (that there is no MOPE) for each experiment, using the Bayes factor calculator provided by Rouder et al. (2009, available at <http://pcl.missouri.edu/bayesfactor>). Bayes factors based on the Jeffrey-Zellner-Siow (JZS) priors (recommended by Rouder et al. as a “default Bayes factor” requiring minimal assumptions) are presented in Table 7.

Table 7.
Bayes factor values for the null hypothesis for MOPE

Experiment and condition	Subjects	Items
Experiment 1	5.92	7.12
Experiment 2: Hiragana target	7.26	8.15
Experiment 2: Romaji target	3.79	6.29
Experiment 3	2.43	3.00
Experiment 4: bi-mora	4.30	5.20
Experiment 4: nasal coda	5.35	3.25
Experiment 4: geminate	5.35	1.35

It can be seen that in all cases, the Bayes factor favored the null hypothesis, with the evidence being “substantial” in most cases. We take these Bayes factor values to argue that our data do indicate that there is no MOPE in Japanese.

The second point to note is that each mora contains more phonological content than a phoneme. Consider, for instance, results by Schiller (2004), who found that when two segments overlap, more priming is obtained. For instance, ba% % % % produced more priming than b% % % % when naming the word “BALLET” (‘ballet’) but, importantly, br% % % % also produced more priming than b% % % % for a complex onset word such as “BROEDER” (‘brother’)³. That is, the finding of MORA priming but not MOPE may simply reflect a greater amount of phonological overlap, not the special status of mora as a functional unit of language production. To test this, an experiment is needed in which the amount of initial phonemic overlap is increased without increasing the number of mora. Unfortunately, this is not easy to realize in Japanese, due the absence of consonant clusters. The closest match to such an experiment would be the inclusion of the palatal glide (denoted j) which is occasionally situated between particular initial consonants and their subsequent vowels. In that case, the initial Cj cluster, for instance, /ky/ of a word such as /kyu.u/ [CjV.R] “sudden, haste” (two moras, namely: kyu [CjV] and a long vowel [R]) would contain more phonological content than the phoneme /k/ in a CV mora structure such as /ku/ does. This has thus

³ Note however that the increase in priming due to an additional overlapping phoneme beyond the onset (e.g., sif-SIB vs. suf-SIB) is small and may not be significant (e.g., Kinoshita, 2000; Mousikou, Coltheart, & Saunders, 2010)

far not been tested using the masked priming paradigm, however, Kureta and colleagues (2006; Experiment 2) specifically tested for this possibility using the preparation paradigm. They did not find significant preparation effects even when an additional segment (i.e. palatal glide) was added to the prevocalic consonant, which favors the hypothesis that the mora is the functional unit in Japanese speech production. Still, additional masked onset priming experiments using Cj initial overlap might be warranted to provide information as to whether the absence of the greater overlap (as found by Kureta et al., 2006) also generalizes to masked priming.

A recent proposal considering functional unit size (O'Seaghdha, Chen, & Chen, 2010) puts forward a new term called the proximate units principle. In short, these authors propose that, "proximate units (PU) are the first selectable phonological units below the level of the word/morpheme" (p. 285). They additionally put forward that languages may differ in their specific proximate unit sizes. In Dutch and English the PU would be the segment, in Chinese the syllable (Chen, Chen, & Dell, 2002), whereas in Japanese the PU would be the mora (Kureta et al., 2006, and our current findings). Interestingly, one can infer from this proposal that across-language comparisons might also show different types of speech errors. For instance, Japanese speakers may be less prone to make phoneme exchange errors (e.g. spoonerisms), as their proximate unit is larger (e.g. a mora). Indeed, evidence provided by Kubozono (1989) shows that in Japanese most speech errors occur at mora, and not at segment or syllable boundaries. An interesting question (for further examination) then arises regarding how bilingual speakers of languages differing in proximate unit (such as Japanese and English) encode the functional unit of their L2. Would for instance Japanese/English vs. English/Japanese bilinguals show diverging results on masked and implicit priming tasks, or could the proximate unit be switched depending on the language in use? Yet, another possibility might be even the transfer of the proximate unit from e.g. the L1 to the L2 depending on the task at hand.

In conclusion, in line with Kureta et al. (2006), our results corroborate the view that in Japanese moras and not segments or syllables are the functional (or proximate) units, i.e. the units, which play a crucial role in the construction of the phonological form of a word in speech production.

References

- Amano, S., & Kondo, T. (1999). NTT deetaa beesu siriizu: Nihongo no goi tokusei [NTT database series: Lexical properties in Japanese]. Tokyo: Sanseido.
- Chen, J.-Y., Chen, T.-M., & Dell, G. S. (2002). Word-form encoding in Mandarin Chinese as assessed by the implicit priming task. *Journal of Memory and Language*, 46, 751–781.
- Chen, J.-Y., Lin, W.-C., & Ferrand, L. (2003). Masked priming of the syllable in Mandarin Chinese speech production. *Chinese Journal of Psychology*, 45, 107-120.
- Cutler, A., & Otake, T. (1994). Mora or phonemes? Further evidence for language-specific listening. *Journal of Memory and Language*, 33, 824–844.
- Dell, G. S. (1988). The retrieval of phonological forms in production: Tests of predictions from a connectionist model. *Journal of Memory and Language*, 27, 124–142.
- Dienes, Z. (2008). Understanding psychology as a science: An introduction to scientific and statistical inference. New York: Palgrave MacMillan.
- Ferrand, L., Segui, J., & Grainger, J. (1996). Masked priming of word and picture naming: The role of syllabic units. *Journal of Memory and Language*, 35, 708–723.
- Forster, K. I., & Davis, C. (1991). The density constraint on form-priming in the naming task: Interference effects from a masked prime. *Journal of Memory and Language*, 30, 1–25.
- Grainger, J., & Ferrand, L. (1996). Masked orthographic and phonological priming in visual word recognition and meaning: Cross-task comparisons. *Journal of Memory and Language*, 35, 623–647.
- Inagaki, K., Hatano, G., & Otake, T. (2000). The effect of kana literacy acquisition on the speech segmentation unit used by Japanese young children. *Journal of Experimental Child Psychology*, 75, 70–91.
- Jeffreys, H. (1961). Theory of probability. Oxford: UK: Oxford University Press.
- Kim, J., & Davis, C. (2002). Using Korean to investigate phonological priming effects without the influence of orthography. *Language and Cognitive Processes*, 17, 569–591.
- Kinoshita, S. (2000). The left-to-right nature of the masked onset effect in naming. *Psychonomic Bulletin & Review*, 7, 133–141.

Kinoshita, S. (2003). The nature of masked onset priming effects in naming: A review. In S. Kinoshita & S. J. Lupker (eds.), *Masked priming. The state of the art* (pp. 223–238). Hove: Psychology Press.

Kinoshita, S., & Woollams, A. (2002). The masked onset priming effect in naming: Computation of phonology or speech-planning. *Memory & Cognition*, 30, 237–245.

Kubozono, H. (1989). The mora and syllable structure in Japanese: Evidence from speech errors. *Language and Speech*, 32, 249–278.

Kureta, Y., Fushimi, T., & Tatsumi, I. F. (2006). The functional unit of phonological encoding: Evidence for moraic representation in native Japanese speakers. *Journal of Experimental Psychology: Learning, Memory and Cognition*, 32, 1102–1119.

Levelt, W.J.M., Roelofs, A., & Meyer, A.S. (1999). A theory of lexical access in speech production. *Behavioral and Brain Sciences*, 22, 1–75.

Meyer, A. S. (1991). The time course of phonological encoding in language production: Phonological encoding inside a syllable. *Journal of Memory and Language*, 30, 69–89.

Mousikou, P., Coltheart, M., & Saunders, S. (2010). Computational modeling of the masked onset priming effect in reading aloud. *Quarterly Journal of Experimental Psychology*, 22, 725–763.

Norris, D. (2006). The Bayesian reader: Explaining word recognition as an optimal Bayesian decision process. *Psychological Review*, 113, 327–357.

Norris, D., & Kinoshita, S. (2008). Perception as evidence accumulation and Bayesian inference: Insights from masked priming. *Journal of Experimental Psychology: General*, 137, 434–455.

O'Seaghdha, P.G., Chen, J.-Y., Chen, T.-M. (2010). Proximate units in word production: Phonological encoding begins with syllables in Mandarin Chinese but with segments in English. *Cognition*, 115, 282–302.

Otake, T., Hatano, G., Cutler, A., & Mehler, J. (1993). Mora or syllable? Speech segmentation in Japanese. *Journal of Memory and Language*, 32, 258–278.

Otake, T., & Cutler, A. (2003). Evidence against “Units of Perception”. In Shohov, S.P. (Ed.), *Advances in Psychology Research: Volume 24*, Nova Science Publishers, New York, pp. 59–84.

Port, R. F., Dalby, J., & O'Dell, M. (1987). Evidence for mora-timing in Japanese. *Journal of the Acoustical Society of America*, 81, 1574–1585.

Roelofs, A. (2006). The influence of spelling on phonological encoding in word reading, object naming, and word generation. *Psychonomic Bulletin & Review*, 13, 33–37.

Rouder, J.N., Speckman, P.L., Sun, D., & Morey, R. (2009). Bayesian t tests for accepting and rejecting the null hypothesis. *Psychonomic Bulletin & Review*, 16, 225–237.

Schiller, N. O. (1998). The Effect of Visually Masked Syllable Primes on the Naming Latencies of Words and Pictures. *Journal of Memory and Language*, 39, 484–507.

Schiller, N. O. (2004). The onset effect in word naming. *Journal of Memory & Language*, 50, 477–490.

Tamaoka, K., Taft, M. (1994). Hakuha on-in shori no saisho tan-i to nariuruka: Giji gairaigo no goi seigo handan karano kousatsu. (in Japanese) [Is the smallest unit in phonological processing equivalent to the smallest unit in orthographic processing? Lexical judgements of Katakana non-words]. *The Japanese Journal of Psychology*, 65, 377–382

Warner, N., & Arai, T. (2001). Japanese mora-timing: A review. *Phonetica*, 58, 1–25.

Appendix A
Stimulus Materials from Experiment 1

target	MOPE part onset prime	control prime	MORA part overlap prime	control prime
かぜ 'kaze'	こと 'koto'	そと 'soto'	—	—
きく 'kiku'	くら 'kura'	むら 'mura'	—	—
くに 'kuni'	かみ 'kami'	なみ 'nami'	クイ 'kui'	ルイ 'rui'
けさ 'kesa'	きか 'kika'	びか 'bika'	ケツ 'ketsu'	ネツ 'netsu'
こま 'koma'	けつ 'ketsu'	ねつ 'netsu'	コト 'koto'	ソト 'soto'
さら 'sara'	しか 'shika'	ちか 'chika'	サメ 'same'	マメ 'mame'
しき 'shiki'	さめ 'same'	まめ 'mame'	シカ 'shika'	チカ 'chika'
すし 'sushi'	せん 'sen'	れん 'ren'	スミ 'sumi'	グミ 'gumi'
せき 'seki'	そと 'soto'	こと 'koto'	—	—
そふ 'sofu'	すみ 'sumi'	ぐみ 'gumi'	ソト 'soto'	コト 'koto'
まど 'mado'	めい 'mei'	べい 'bei'	マネ 'mame'	サメ 'same'
みみ 'mimi'	まめ 'mame'	さめ 'same'	ミコ 'miko'	ヒコ 'hiko'
むし 'mushi'	もり 'mori'	のり 'nori'	—	—
めし 'meshi'	むら 'mura'	くら 'kura'	メイ 'mei'	ゲイ 'gei'
もち 'mochi'	みこ 'miko'	ひこ 'hiko'	—	—
なつ 'natsu'	のり 'nori'	もり 'mori'	ナミ 'nami'	タミ 'tami'

にわ 'niwa'	ぬか 'nuka'	ぶか 'buka'	クニ 'niku'	リク 'riku'
ぬま 'numa'	ねつ 'netsu'	けつ 'ketsu'	ヌカ 'nuka'	ブカ 'buka'
ねこ 'neko'	なみ 'nami'	かみ 'kami'	ネツ 'netsu'	ケツ 'ketsu'
のど 'nodo'	にく 'niku'	りく 'riku'	ノリ 'nori'	モリ 'mori'
らく 'raku'	るり 'ruri'	くり 'kuri'	—	—
りか 'rika'	れん 'ren'	せん 'sen'	—	—
るす 'rusu'	りく 'riku'	にく 'niku'	ルイ 'rui'	クイ 'kui'
れつ 'retsu'	ろく 'roku'	ほく 'hoku'	—	—
ろじ 'roji'	らち 'rachi'	さち 'sachi'	ロク 'roku'	ホク 'hoku'
たけ 'take'	てき 'teki'	へき 'heki'	タミ 'tami'	ナミ 'nami'
てつ 'tetsu'	とみ 'tomi'	ごみ 'gomi'	テキ 'teki'	ヘキ 'heki'
とり 'tori'	たく 'taku'	がく 'gaku'	トミ 'tomi'	ゴミ 'gomi'
はな 'hana'	ほく 'hoku'	ろく 'roku'	ハラ 'hara'	バラ 'bara'
ひふ 'hifu'	はら 'hara'	ばら 'bara'	ヒコ 'hiko'	ミコ 'miko'
へび 'hebi'	ひこ 'hiko'	みこ 'miko'	—	—
ほり 'hori'	へい 'hei'	げい 'gei'	ホク 'hoku'	ロク 'roku'
ぼつ 'batsu'	ぼし 'boshi'	とし 'toshi'	バラ 'bara'	ハラ 'hara'
びわ 'biwa'	ばら 'bara'	はら 'hara'	ビカ 'bika'	キカ 'kika'
ぶた 'buta'	びか 'bika'	きか 'kika'	ブカ 'buka'	ヌカ 'nuka'
べつ 'betsu'	ぶか 'buka'	ぬか 'nuka'	—	—
ぼく 'boku'	べい 'bei'	めい 'mei'	ボシ 'boshi'	トシ 'toshi'

- 136 -

が か 'gaka'	ぎ む 'gimu'	じ む 'jimu'	—	—
ぎ し 'gishi'	ぐ み 'gumi'	す み 'sumi'	—	—
ぐ ち 'guchi'	が く 'gaku'	た く 'taku'	グ ミ 'gumi'	ス ミ 'sumi'
げ き 'geki'	ご み 'gomi'	と み 'tomi'	ゲ イ 'gei'	メ イ 'mei'
ご ご 'gogo'	げ い 'gei'	へ い 'hei'	ゴ ミ 'gomi'	ト ミ 'tomi'

Appendix B
Stimulus Materials from Experiments 2 and 3

target	MOPE part onset prime	control prime	MORA part overlap prime	control prime
かぜ 'kaze'	こと 'koto'	そと 'soto'	かみ 'kami'	なみ 'nami'
きく 'kiku'	くら 'kura'	むら 'mura'	きか 'kika'	びか 'bika'
くに 'kuni'	かみ 'kami'	なみ 'nami'	くら 'kura'	むら 'mura'
けさ 'kesa'	きか 'kika'	びか 'bika'	けつ 'ketsu'	ねつ 'netsu'
こま 'koma'	けつ 'ketsu'	ねつ 'netsu'	こと 'koto'	そと 'soto'
さら 'sara'	しか 'shika'	ちか 'chika'	さめ 'same'	まめ 'mame'
しき 'shiki'	さめ 'same'	まめ 'mame'	しか 'shika'	ちか 'chika'
すし 'sushi'	せん 'sen'	れん 'ren'	すみ 'sumi'	ぐみ 'gumi'
せき 'seki'	そと 'soto'	こと 'koto'	せん 'sen'	れん 'ren'
そふ 'sofu'	すみ 'sumi'	ぐみ 'gumi'	そと 'soto'	こと 'koto'
まど 'mado'	めい 'mei'	べい 'bei'	まめ 'mame'	さめ 'same'
みみ 'mimi'	まめ 'mame'	さめ 'same'	みこ 'miko'	ひこ 'hiko'
むし 'mushi'	もり 'mori'	のり 'nori'	むら 'mura'	くら 'kura'
めし 'meshi'	むら 'mura'	くら 'kura'	めい 'mei'	べい 'bei'
もち 'mochi'	みこ 'miko'	ひこ 'hiko'	もり 'mori'	のり 'nori'
なつ 'natsu'	のり 'nori'	もり 'mori'	なみ 'nami'	かみ 'kami'

にわ 'niwa'	ぬか 'nuka'	ぶか 'buka'	にく 'niku'	りく 'riku'
ぬま 'numa'	ねつ 'netsu'	けつ 'ketsu'	ぬか 'nuka'	ぶか 'buka'
ねこ 'neko'	なみ 'nami'	かみ 'kami'	ねつ 'netsu'	けつ 'ketsu'
のど 'nodo'	にく 'niku'	りく 'riku'	のり 'nori'	もり 'mori'
らく 'raku'	るり 'ruri'	くり 'kuri'	らち 'rachi'	さち 'sachi'
りか 'rika'	れん 'ren'	せん 'sen'	りく 'riku'	にく 'niku'
るす 'rusu'	りく 'riku'	にく 'niku'	るり 'ruri'	くり 'kuri'
れつ 'retsu'	ろく 'roku'	ほく 'hoku'	れん 'ren'	せん 'sen'
ろじ 'roji'	らち 'rachi'	さち 'sachi'	ろく 'roku'	ほく 'hoku'
たけ 'take'	てき 'teki'	へき 'heki'	たく 'taku'	がく 'gaku'
てつ 'tetsu'	とみ 'tomi'	ごみ 'gomi'	てき 'teki'	へき 'heki'
とり 'tori'	たく 'taku'	がく 'gaku'	とみ 'tomi'	ごみ 'gomi'
はな 'hana'	ほく 'hoku'	ろく 'roku'	はら 'hara'	ばら 'bara'
ひふ 'hifu'	はら 'hara'	ばら 'bara'	ひこ 'hiko'	みこ 'miko'
へび 'hebi'	ひこ 'hiko'	みこ 'miko'	へい 'hei'	げい 'gei'
ほり 'hori'	へい 'hei'	げい 'gei'	ほく 'hoku'	ろく 'roku'
ぼつ 'batsu'	ぼし 'boshi'	とし 'toshi'	ばら 'bara'	はら 'hara'
びわ 'biwa'	ばら 'bara'	はら 'hara'	びか 'bika'	きか 'kika'
ぶた 'buta'	びか 'bika'	きか 'kika'	ぶか 'buka'	ぬか 'nuka'
べつ 'betsu'	ぶか 'buka'	ぬか 'nuka'	べい 'bei'	めい 'mei'
ぼく 'boku'	べい 'bei'	めい 'mei'	ぼし 'boshi'	とし 'toshi'

が か 'gaka'	ぎ む 'gimu'	じ む 'jimu'	が く 'gaku'	た く 'taku'
ぎ し 'gishi'	ぐ み 'gumi'	す み 'sumi'	ぎ む 'gimu'	じ む 'jimu'
ぐ ち 'guchi'	が く 'gaku'	た く 'taku'	ぐ み 'gumi'	す み 'sumi'
げ き 'geki'	ご み 'gomi'	と み 'tomi'	げ い 'gei'	へ い 'hei'
ご ご 'gogo'	げ い 'gei'	へ い 'hei'	ご み 'gomi'	と み 'tomi'

Appendix C
Stimulus Materials from Experiments 4a and 4b

target	MOPE part onset prime	control prime	MORA part overlap prime	control prime
かぜ 'kaze'	コト 'koto'	ソト 'soto'	—	—
きく 'kiku'	クラ 'kura'	ムラ 'mura'	—	—
くに 'kuni'	カミ 'kami'	ナミ 'nami'	クイ 'kui'	ルイ 'rui'
けさ 'kesa'	キカ 'kika'	ビカ 'bika'	ケツ 'ketsu'	ネツ 'netsu'
こま 'koma'	ケツ 'ketsu'	ネツ 'netsu'	コト 'koto'	ソト 'soto'
さら 'sara'	シカ 'shika'	チカ 'chika'	サメ 'same'	マメ 'mame'
しき 'shiki'	サメ 'same'	マメ 'mame'	シカ 'shika'	チカ 'chika'
すし 'sushi'	セン 'sen'	レン 'ren'	スミ 'sumi'	グミ 'gumi'
せき 'seki'	ソト 'soto'	コト 'koto'	—	—
そふ 'sofu'	スミ 'sumi'	グミ 'gumi'	ソト 'soto'	コト 'koto'
まど 'mado'	メイ 'mei'	ベイ 'bei'	マメ 'mame'	サメ 'same'
みみ 'mimi'	マメ 'mame'	サメ 'same'	ミコ 'miko'	ヒコ 'hiko'
むし 'mushi'	モリ 'mori'	ノリ 'nori'	—	—
めし 'meshi'	ムラ 'mura'	クラ 'kura'	メイ 'mei'	ゲイ 'gei'
もち 'mochi'	ミコ 'miko'	ヒコ 'hiko'	—	—
なつ 'natsu'	ノリ 'nori'	モリ 'mori'	ナミ 'nami'	タミ 'tami'
にわ 'niwa'	ヌカ 'nuka'	ブカ 'buka'	ニク 'niku'	リク 'riku'

ぬま 'numa'	ネツ 'netsu'	ケツ 'ketsu'	ヌカ 'nuka'	ブカ 'buka'
ねこ 'neko'	ナミ 'nami'	カミ 'kami'	ネツ 'netsu'	ケツ 'ketsu'
のど 'nodo'	ニク 'niku'	リク 'riku'	ノリ 'nori'	モリ 'mori'
らく 'raku'	ルリ 'ruri'	クリ 'kuri'	—	—
りか 'rika'	レン 'ren'	セン 'sen'	—	—
るす 'rusu'	リク 'riku'	ニク 'niku'	ルイ 'rui'	クイ 'kui'
れつ 'retsuo'	ロク 'roku'	ホク 'hoku'	—	—
ろじ 'roji'	ラチ 'rachi'	サチ 'sachi'	ロク 'roku'	ホク 'hoku'
たけ 'take'	テキ 'teki'	ヘキ 'heki'	タミ 'tami'	ナミ 'nami'
てつ 'tetsu'	トミ 'tomi'	ゴミ 'gomi'	テキ 'teki'	ヘキ 'heki'
とり 'tori'	タク 'taku'	ガク 'gaku'	トミ 'tomi'	ゴミ 'gomi'
はな 'hana'	ホク 'hoku'	ロク 'roku'	ハラ 'hara'	バラ 'bara'
ひふ 'hifu'	ハラ 'hara'	バラ 'bara'	ヒコ 'hiko'	ミコ 'miko'
へび 'hebi'	ヒコ 'hiko'	ミコ 'miko'	—	—
ほり 'hori'	ヘイ 'hei'	ゲイ 'gei'	ホク 'hoku'	ロク 'roku'
ばつ 'batsu'	ボシ 'boshi'	トシ 'toshi'	バラ 'bara'	ハラ 'hara'
びわ 'biwa'	バラ 'bara'	ハラ 'hara'	ビカ 'bika'	キカ 'kika'
ぶた 'buta'	ビカ 'bika'	キカ 'kika'	ブカ 'buka'	ヌカ 'nuka'
べつ 'betsu'	ブカ 'buka'	ヌカ 'nuka'	—	—
ぼく 'boku'	ベイ 'bei'	メイ 'mei'	ボシ 'boshi'	トシ 'toshi'
がが 'gaka'	ギム 'gimu'	ジム 'jimu'	—	—

ぎし 'gishi'	グミ 'gumi'	スミ 'sumi'	—	—
ぐち 'guchi'	ガク 'gaku'	タク 'taku'	グミ 'gumi'	スミ 'sumi'
げき 'geki'	ゴミ 'gomi'	トミ 'tomi'	ゲイ 'gei'	メイ 'mei'
ごご 'gogo'	ゲイ 'gei'	ヘイ 'hei'	ゴミ 'gomi'	トミ 'tomi'
かんじ 'kanji'	コト 'koto'	ソト 'soto'	—	—
きんし 'kinshi'	クラ 'kura'	ムラ 'mura'	—	—
くんし 'kunshi'	カミ 'kami'	ナミ 'nami'	クイ 'kui'	ルイ 'rui'
けんか 'kenka'	キカ 'kika'	ビカ 'bika'	ケツ 'ketsu'	ネツ 'netsu'
こんど 'kondo'	ケツ 'ketsu'	ネツ 'netsu'	コト 'koto'	ソト 'soto'
さんそ 'sanso'	シカ 'shika'	チカ 'chika'	サメ 'same'	マメ 'mame'
しんか 'shinka'	サメ 'same'	マメ 'mame'	シカ 'shika'	チカ 'chika'
すんか 'sunka'	セン 'sen'	レン 'ren'	スミ 'sumi'	グミ 'gumi'
せんし 'senshi'	ソト 'soto'	コト 'koto'	—	—
そんけい 'sonkei'	スミ 'sumi'	グミ 'gumi'	ソト 'soto'	コト 'koto'
まんが 'manga'	メイ 'mei'	ベイ 'bei'	マメ 'mame'	サメ 'same'
みんわ 'minwa'	マメ 'mame'	サメ 'same'	ミコ 'miko'	ヒコ 'hiko'
むんず 'munzu'	モリ 'mori'	ノリ 'nori'	—	—
めんきょ 'menkyo'	ムラ 'mura'	クラ 'kura'	メイ 'mei'	ゲイ 'gei'
もんく 'monku'	ミコ 'miko'	ヒコ 'hiko'	—	—

なんばん 'nanban'	ノリ 'nori'	モリ 'mori'	ナミ 'nami'	タミ 'tami'
にんき 'ninki'	ヌカ 'nuka'	ブカ 'buka'	ニク 'niku'	リク 'riku'
ぬんちゃく 'nunchaku'	ネツ 'netsu'	ケツ 'ketsu'	ヌカ 'nuka'	ブカ 'buka'
ねんど 'nendo'	ナミ 'nami'	カミ 'kami'	ネツ 'netsu'	ケツ 'ketsu'
のんき 'nonki'	ニク 'niku'	リク 'riku'	ノリ 'nori'	モリ 'mori'
らんし 'ranshi'	ルリ 'ruri'	クリ 'kuri'	—	—
りんじ 'rinji'	レン 'ren'	セン 'sen'	—	—
るんば 'runba'	リク 'riku'	ニク 'niku'	ルイ 'rui'	クイ 'kui'
れんが 'renga'	ロク 'roku'	ホク 'hoku'	—	—
ろんぎ 'rongi'	ラチ 'rachi'	サチ 'sachi'	ロク 'roku'	ホク 'hoku'
たんご 'tango'	テキ 'teki'	ヘキ 'heki'	タミ 'tami'	ナミ 'nami'
てんし 'tenshi'	トミ 'tomi'	ゴミ 'gomi'	テキ 'teki'	ヘキ 'heki'
とんぼ 'tonbo'	タク 'taku'	ガク 'gaku'	トミ 'tomi'	ゴミ 'gomi'
はんこ 'hanko'	ホク 'hoku'	ロク 'roku'	ハラ 'hara'	バラ 'bara'
ひんど 'hindo'	ハラ 'hara'	バラ 'bara'	ヒコ 'hiko'	ミコ 'miko'
へんか 'henka'	ヒコ 'hiko'	ミコ 'miko'	—	—
ほにゃ 'honya'	ヘイ 'hei'	ゲイ 'gei'	ホク 'hoku'	ロク 'roku'
ばんち 'banchi'	ボシ 'boshi'	トシ 'toshi'	バラ 'bara'	ハラ 'hara'
びんせん	バラ 'bara'	ハラ 'hara'	ビカ 'bika'	キカ 'kika'

'binsen'

ぶにゃ 'bunya'	ビカ 'bika'	キカ 'kika'	ブカ 'buka'	ヌカ 'nuka'
べんり 'benri'	ブカ 'buka'	ヌカ 'nuka'	—	—
ぼんち 'bonchi'	ベイ 'bei'	メイ 'mei'	ボシ 'boshi'	トシ 'toshi'
がんそ 'ganso'	ギム 'gimu'	ジム 'jimu'	—	—
ぎんか 'ginka'	グミ 'gumi'	スミ 'sumi'	—	—
ぐんじ 'gunji'	ガク 'gaku'	タク 'taku'	グミ 'gumi'	スミ 'sumi'
げんし 'genshi'	ゴミ 'gomi'	トミ 'tomi'	ゲイ 'gei'	メイ 'mei'
ごんげ 'gonge'	ゲイ 'gei'	ヘイ 'hei'	ゴミ 'gomi'	トミ 'tomi'
かつき 'kakki'	コト 'koto'	ソト 'soto'	—	—
きつぷ 'kippu'	クラ 'kura'	ムラ 'mura'	—	—
くっしん	カミ 'kami'	ナミ 'nami'	クイ 'kui'	ルイ 'rui'

'kusshin'

けっか 'kekka'	キカ 'kika'	ビカ 'bika'	ケツ 'ketsu'	ネツ 'netsu'
こっき 'kokki'	ケツ 'ketsu'	ネツ 'netsu'	コト 'koto'	ソト 'soto'
さっか 'sakka'	シカ 'shika'	チカ 'chika'	サメ 'same'	マメ 'mame'
しっぽ 'shippo'	サメ 'same'	マメ 'mame'	シカ 'shika'	チカ 'chika'
すっきり	セン 'sen'	レン 'ren'	スミ 'sumi'	グミ 'gumi'

'sukkiri'

せっし 'sesshi'	ソト 'soto'	コト 'koto'	—	—
そっき 'sokki'	スミ 'sumi'	グミ 'gumi'	ソト 'soto'	コト 'koto'

まっちゃ 'maccha'	メイ 'mei'	ベイ 'bei'	マメ 'mame'	サメ 'same'
みつつ 'mittsu'	マメ 'mame'	サメ 'same'	ミコ 'miko'	ヒコ 'hiko'
むつつ 'muttsu'	モリ 'mori'	ノリ 'nori'	—	—
めつき 'mekki'	ムラ 'mura'	クラ 'kura'	メイ 'mei'	ゲイ 'gei'
もっか 'mokka'	ミコ 'miko'	ヒコ 'hiko'	—	—
なっとう 'nattou'	ノリ 'nori'	モリ 'mori'	ナミ 'nami'	タミ 'tami'
につき 'nikki'	ヌカ 'nuka'	ブカ 'buka'	ニク 'niku'	リク 'riku'
ぬっと 'nutto'	ネツ 'netsu'	ケツ 'ketsu'	ヌカ 'nuka'	ブカ 'buka'
ねつき 'nekki'	ナミ 'nami'	カミ 'kami'	ネツ 'netsu'	ケツ 'ketsu'
のっとり 'nottori'	ニク 'niku'	リク 'riku'	ノリ 'nori'	モリ 'mori'
らっか 'rakka'	ルリ 'ruri'	クリ 'kuri'	—	—
りっぱ 'rippa'	レン 'ren'	セン 'sen'	—	—
るっくす 'rukksu'	リク 'riku'	ニク 'niku'	ルイ 'rui'	クイ 'kui'
れっとう 'rettou'	ロク 'roku'	ホク 'hoku'	—	—
ろっぷ 'roppu'	ラチ 'rachi'	サチ 'sachi'	ロク 'roku'	ホク 'hoku'
たった 'tatta'	テキ 'teki'	ヘキ 'heki'	タミ 'tami'	ナミ 'nami'
てっぽう	トミ 'tomi'	ゴミ 'gomi'	テキ 'teki'	ヘキ 'heki'

'teppou'				
とって 'totte'	タク 'taku'	ガク 'gaku'	トミ 'tomi'	ゴミ 'gomi'
はっぱ 'happa'	ホク 'hoku'	ロク 'roku'	ハラ 'hara'	バラ 'bara'
ひっこし 'hikkoshi'	ハラ 'hara'	バラ 'bara'	ヒコ 'hiko'	ミコ 'miko'
へっつい 'hettsui'	ヒコ 'hiko'	ミコ 'miko'	—	—
ほっと 'hotto'	ヘイ 'hei'	ゲイ 'gei'	ホク 'hoku'	ロク 'roku'
ばった 'batta'	ボシ 'boshi'	トシ 'toshi'	バラ 'bara'	ハラ 'hara'
びっくり 'bikkuri'	バラ 'bara'	ハラ 'hara'	ビカ 'bika'	キカ 'kika'
ぶっし 'busshi'	ビカ 'bika'	キカ 'kika'	ブカ 'buka'	ヌカ 'nuka'
べっそう 'bessou'	ブカ 'buka'	ヌカ 'nuka'	—	—
ぼっちゃん 'bocchan'	ベイ 'bei'	メイ 'mei'	ボシ 'boshi'	トシ 'toshi'
がっき 'gakki'	ギム 'gimu'	ジム 'jimu'	—	—
ぎっしり 'gisshiri'	グミ 'gumi'	スミ 'sumi'	—	—
ぐったり 'guttari'	ガク 'gaku'	タク 'taku'	グミ 'gumi'	スミ 'sumi'
げっぷ 'geppu'	ゴミ 'gomi'	トミ 'tomi'	ゲイ 'gei'	メイ 'mei'

- 147 -

ごっかん
'gokkan'

ゲイ 'gei'

ヘイ 'hei'

ゴミ 'gomi'

トミ 'tomi'

Chapter 6: Summary and Conclusion

Summary and conclusion.

In this thesis, I specifically reported studies on word processing in languages which employ non-alphabetic scripts with an emphasis on logographic scripts such as Chinese and particularly Japanese. Detailed focus was put on the underlying principles of the unique heterophonic nature of kanji (i.e. about 60% of all kanji have more than one pronunciation which depends for selection on the intra-word context). I sought to reveal whether activation in the production network spreads to multiple pronunciations or whether it is restricted to a single (preferred) pronunciation (see Chapter 2). In addition, I attempted to gain an understanding regarding which pronunciations of a kanji character receive activation by means of its sensitivity to semantic and homophonic context effects (see Chapters 3 and 4). Finally, I intended to clarify whether or not the nature (or size) of the fundamental phonological unit in Japanese was comparable to that of West-European languages on which theories of language production are usually built (see Chapter 5).

Chapter 1

In Chapter 1, a brief overview was given of the most influential models of language production called WEAVER++ (Roelofs, 1992; Levelt, Roelofs, & Meyer, 1999). In this model, a to-be conveyed message first enters the conceptual level, after which it spreads to the lexical-syntactic level where syntactic features are integrated and lexical selection takes place (typically by competition; but see Mahon et al., 2007). Subsequently, phonological word-form encoding will take place, including the computation of metrical information and syllabification (for overviews see Schiller, 2000; Schiller et al., 2006). Finally, the utterance is overtly produced through the articulatory system. It was pointed out that pictures enter this system exclusively through the conceptual system but that written words can enter the production system in three ways. First, through grapheme to phoneme conversion evidenced by the fact we can name non-words, which by definition have no lexical representation.

Second, via the lexical-phonological route which is inferred from the fact that prime words morphologically related to picture names speed up naming in the long-lag priming paradigm as prime and target share a morpheme at this level (Koester & Schiller, 2008, in press; Zwitserlood, Bölte, & Dohmes, 2000). Lastly, by activating

their lexical-syntactic (lemma) representations from orthography, which is often inferred from semantic interference in a picture-word interference task (e.g. Glaser & Döngelhoff, 1984; see also MacLeod, 1991, for a review).

Most models of language production (including WEAVER++) are built on empirical data obtained with Indo-European languages typically employing alphabetic scripts. However, Chapter 1 discussed evidence that caution must be taken when assuming that all parameters of language production models would automatically apply to languages around the world. Japanese, for instance, has very different linguistic properties compared to Dutch, and therefore might also require a different interpretation in certain aspects of the production model, such as the absence of orthography to phonology conversion (Siok, Perfetti, Jin, & Tan, 2004). It was put forward that irregular words and heterophonic homographs in alphabetic languages are frequently found to be named slower, presumably due to competing pronunciations (e.g. Glushko, 1979; Kawamoto & Zemplidze, 1992). In most Indo-European languages, heterophonic homographs occur infrequently, however, due to specific reasons having to do with language history and especially the development of pronunciation (i.e. the assimilation of traditional Chinese pronunciations vs. the already existing Japanese pronunciations for words when importing Chinese characters) for kanji in Japanese this is the rule rather than the exception. Therefore, Japanese might have developed different mechanisms to more efficiently handle pronunciation selection for heterophonic words.

Chapter 2

Japanese kanji and Chinese hànzi differ in the sense that Japanese kanji usually have more ways to pronounce the same logographic symbol; kanji is therefore less consistent in the orthography to phonology route than Chinese. Previous studies such as Wydell, Butterworth, and Patterson (1995) explored whether kanji having more alternative pronunciations would show a processing difference compared to kanji, which have fewer pronunciations (indicated by longer RTs, for instance). However, results in the literature regarding this matter are not entirely consistent. Wydell and colleagues (1995) did not find consistency effects in Japanese, however, Fushimi, Ijuin, Patterson, and Tatsumi (1999) did. Also

Kayamoto, Yamada, and Takashima (1998) found consistency effects in a single kanji naming study, but only when the alternative pronunciation was of high frequency. What all these studies have in common is that they do not directly assess whether multiple pronunciations become active. Rather, they assume that multiple pronunciations receive activation to account for differences in naming latency.

This chapter presents the results from two masked priming studies which directly assess activation spreading in kanji having multiple pronunciations. The pronunciations from a kanji prime (e.g. 水 /mizu_{kun}/ /sui_{on}/) were transcribed into katakana (i.e. a syllabic script mostly used for loanwords) e.g. ミズ (mizu) and スイ (sui). If both katakana transcriptions benefit from the 水 prime (compared to an unrelated prime), then activation would have spread to both pronunciations. In addition, the frequency of both pronunciations was manipulated (as Kayamoto et al., 1999, showed that this might influence RTs). The results of two priming experiments presented in this chapter show that the reading of multiple katakana targets is facilitated by the same kanji prime indicating that multiple readings are activated. In addition, when kanji characters are presented in isolation, phonological activation was always stronger for KUN-readings compared to ON-readings, even when the kanji was biased towards the ON-reading. Therefore, the results from the experiments presented in this chapter provide the first direct empirical evidence of multiple activation spreading in kanji processing.

Chapter 3

As established in Chapter 2, Japanese kanji spread activation to multiple pronunciations. This, according to the data from e.g. Kayamoto et al. (1999) and Fushimi et al. (1999), leads to longer RTs due to activation of multiple readings (possibly leading to competition). In the current chapter, we are interested which route Japanese kanji and Chinese hànzì words take when entering the production system. It is quite firmly established that alphabetic words enter the production system via three different routes: (1) orthography-to-phonology conversion; (2) orthography-to-lexical phonological representation; or (3) orthography-to-lexical syntactic representation.

The first route, e.g. orthography-to-phonology (or grapheme-phoneme conversion), which exists for most alphabetic languages, is not possible in Chinese and Japanese as there are generally no parts of a character uniquely mapping to a phoneme (Siok et al., 2004; Wydell et al., 1995). This leaves the system with the orthography-to-lexical phonological and lexical-syntactic mappings. In a picture-word interference (PWI) study, Zhang and Weekes (2009) demonstrated that semantically related Chinese word distractors induce interference (compared to semantically unrelated distractors), indicating lexical-syntactic involvement. However, to date, a comparable PWI experiment has not been reported for Japanese. Furthermore, the effect of context pictures on reading aloud Japanese kanji and Chinese hànzi targets has not been examined, either. In alphabetic languages, such word naming experiments usually show that there is no effect of semantically related context pictures as compared to unrelated context pictures when naming bare nouns (Glaser & Dünghoff, 1984; La Heij, Happel, & Mulder, 1990). This is assumed to be due to the use of the direct route from orthography to the lexical-phonological level, i.e. a fast process that prevents context pictures to exert an influence. However, when the lexical-syntactic level is required for naming, as for instance in determiner noun phrase generation, semantic context effects of pictures are shown to emerge (Roelofs, 2003, 2006).

Chapter 3 presents the results of a word-naming task in Japanese and Chinese using logographic scripts. In addition, a standard PWI experiment was carried out to assert the validity of the stimuli and to obtain as yet unreported PWI results in Japanese. In this chapter, I report a number of novel empirical findings. First, a significant semantic interference effect is induced by semantically related distractor words in both Japanese and Chinese picture naming, indicating that both Japanese and Chinese characters access the lexical-syntactic level and delay the selection of the target word at that level (due to selection by competition). These findings are identical to the ones obtained with alphabetic distractor words. Second, context pictures induced a significant semantic facilitation effect in Japanese word (kanji) naming at two SOAs (–150 ms and 0 ms) but no such effect in Chinese word (hànzi) naming at the same SOAs. To account for the latter findings, I proposed that Chinese words, analogous to alphabetic words, are read via the orthography-to-phonology route that is too fast to allow for an effect of semantic context pictures to

emerge. To account for the presence of a semantic facilitation effect in Japanese kanji naming, I proposed that Japanese kanji are also named via the orthography-to-phonology but that the processing is slowed down due to the activation of multiple phonological representations (and subsequent processing delay due to competition in the selection of the appropriate pronunciation), which in turn allows for context pictures to exert an effect. I do acknowledge that it cannot be entirely ruled out that kanji are read via their lexical-syntactic representations, but this latter hypothesis seems not in line with neuropsychological evidence from Japanese patients (see: Sasanuma, Sakuma, & Kitano, 1992; Nakamura et al., 1998; Fushimi et al., 2003). Therefore, I consider it less plausible.

Chapter 4

Chapter 3 reported semantic facilitation by context pictures on the naming of Japanese (but not Chinese) target words. In Chapter 4, I aimed to extend these findings by introducing phonologically related picture distractors, more specifically pictures which are homophones to the target word. I found that Japanese kanji, but not Chinese hànzi, were susceptible to phonological context effects. This finding is completely in line with the results obtained with semantically related context pictures in Chapter 3. In a subsequent experiment, using Dutch stimuli, we replicated Roelofs' (2003) finding that semantic context effects of distractor pictures are obtained in det+N naming (which requires lexical-syntactic access), but not in bare-noun naming (which allows for the use of the direct orthography-to-phonology route). However, when bare nouns were visually degraded (e.g. \$H\$O\$N\$D\$), semantic context effects again surfaced. The idea underlying our account is that Japanese kanji and degraded words in Dutch are read via the lexical-phonological route. However, the speed with which a phonological representation is selected upon the presentation of a target word determines whether a context effect can emerge. For bare nouns in alphabetic language such as Dutch or logographic characters in Chinese (hànzi), the activation builds up too fast for any context information to affect it. In contrast, for degraded Dutch nouns or Japanese kanji, the slope of the activation buildup reaching the selection threshold is altered in such a way that context information can influence processing along this route. Subsequent

experiments (including frequency manipulation) should provide further evidence to substantiate this claim.

Chapter 5

In this chapter, I was especially concerned with the unit of phonological planning in Japanese language production. Generally, theories of language production (e.g. Dell, 1986; Levelt et al., 1999) take the basic phonological planning unit to be the segment/phoneme (as most data are acquired with Indo-European languages such as English and Dutch). However, there is also evidence that a phonological planning unit might be language-specific. For instance, in Mandarin Chinese, Chen, Chen, and Dell (2002) and O'Seaghdha, Chen, & Chen (2010) using a preparation (or implicit priming) paradigm found no effect of single segments (even in simple monosyllabic words). These authors therefore concluded that syllables and not segments are the phonological planning units linked to speech production in Chinese. In contrast to English and Dutch (segment) as well as Chinese (syllable), Japanese is often characterized as a mora-timed language (for an overview on mora and mora-timing, see Warner and Arai, 2001). In Japanese, moras form independent rhythmical structures within a syllable, for instance in the well-known brand name "honda" /hoN.da/ has two syllables (/hoN/ and /da/), but three moras (the nasal coda of the first syllable is an independent mora) /ho/, /N/ and /da/, all lasting roughly equally long.

Previous evidence for the mora as a functional unit showed mixed results. For instance, Otake et al. (1993) as well as Cutler and Otake (1994) using a speech segmentation paradigm found Japanese native speakers were just as likely to detect CV structures in CVN.CV structures where moras are part of the syllable compared to simple CV.CV structures when moras equal the syllable. Also, targets such as /nan/ (CVN) were more easily detected when the consonant N was a mora by itself (naN.ka) compared to when it was part of a mora (na.no.ka) in spite of the fact that in both cases the overlap was equal when counted in number of overlapping segments⁴). However, in contrast, using another paradigm (i.e. word reconstruction; see Otake

⁴ note: although usually onsets are more salient, also the position of /n/ *within* the syllable varied (from being the coda in /naN/ to being the onset in /no/) in this experiment.

and Cutler, 2003), the same authors found Japanese subjects to be sensitive to segmental information.

A recent production study by Kureta, Fushimi, and Tatsumi (2006) using the implicit priming (or preparation) paradigm (as did Chen et al., 2002 and O'Seaghdha et al., 2010) found Japanese native speakers to prepare their utterances in units of mora. This study (and previous studies), however, made use of kanji and kana scripts, which do not allow for processing of individual phonemes, such as romaji (alphabetic Romanized Japanese). In this chapter, I therefore included stimuli written in romaji to rule out the possibility that the previous usage of non-segmental scripts might have occluded the detection of a segmental effect in Japanese.

Chapter 5 presents the empirical results of four experiments designed to establish the nature of the fundamental phonological planning unit in Japanese. In all experiments, the participants' task was to name strings that were presented on the computer screen. Targets were preceded by masked primes which could be related in the segment or in the whole mora, e.g. /sushi/ could be primed by /sen/ (segment overlap, i.e. "s") vs. /ren/ (control) which was termed the masked onset priming (or MOPE) condition or by /sumi/ (mora overlap, i.e. "su") vs. /gumi/ (control) which was termed the MORA priming condition. In four experiments, I varied the script type of target and prime as well as the amount of overlap (segment or mora), see also Table 1.

The first experiment shows that only the group which received stimuli overlapping in the whole mora showed significant facilitation effects; i.e. the group which received the segment/onset overlap stimuli showed no significant priming. The results of Experiments 2 and 3 demonstrate that we can discard the possibility that the absence of segment priming was due to the nature of target script governing the unit of phonological encoding, as presenting targets in romaji (favoring letter-to segment mapping) did not produce a segment priming effect, either.

Table 1.
Experimental conditions and results for each Experiment in Chapter 5.

Experiment	Target	Prime	Overlap	priming
Experiment 1a (Group1)	hiragana	hiragana	segment	no
Experiment 1b (Group2)	hiragana	katakana	mora	yes
Experiment 2	hiragana	romaji	segment	no
	hiragana	romaji	mora ^a	no
	romaji	romaji	segment	no
	romaji	romaji	mora	yes
Experiment 3	romaji	hiragana	segment	no
	romaji	hiragana	mora	yes
Experiment 4a (Group 1)	hiragana	katakana	segment	no
Experiment 4b (Group 2)	hiragana	katakana	mora	yes

a) The absence of this mora priming effect was attributed to the relative processing speeds of romaji primes (slow) and hiragana targets (fast)

Finally, the fourth experiment introduced stimuli for which the mora was always part of the syllable (e.g. /suN.ka/) in order to rule out the possibility that the observed priming effect was in fact a syllable priming effect. The first group of participants receiving stimuli overlapping in the segment did not show significant priming effects, however, Group 2 receiving mora overlapping stimuli showed significant priming effects. Therefore, the results of these four experiments provide support for the claim that the mora and not the segment is the functional unit of phonological encoding in Japanese. This is in line with the before-mentioned implicit priming findings by Kureta et al. (2006) and also with the interpretation that languages may differ in their specific phonological basic unit size, i.e. the phoneme/onset in Dutch and English, the syllable in Chinese, and the mora in Japanese.

Conclusions and suggestions for future research

This thesis set out to: (1) establish in a direct way whether or not presentation of a single kanji activates its multiple pronunciations and compare this with Chinese hànzì production; (2) ascertain which

route Japanese and Chinese character naming take and whether the activation of multiple pronunciations for Japanese kanji would come at a processing cost; (3) establish the size of the functional phonological unit in Japanese language production. On the basis of the empirical data gathered in Chapters 2-5 we conclude that

When presented in isolation, Japanese kanji spread activation to multiple pronunciations, and phonological activation is generally stronger for KUN-readings compared to ON-readings.

The naming of Japanese kanji and Chinese hànzi occurs via the direct orthography to lexical-phonological representation route; however, at that level, due to the activation of multiple pronunciations, Japanese kanji incur a processing cost.

The mora and not the segment or syllable is the functional unit of phonological encoding in Japanese.

For future research, it is worthwhile to extend the findings and conclusions mentioned in (1) by including compound kanji words. In Chapter 2, it is observed that KUN priming is always stronger. This is likely due to the fact that single kanji were named (which in most cases take the KUN reading). It would be important to establish whether kanji being part of compound words, such as 海 and 水 in 海水 “seawater” /kai_{on}.sui_{on}/, would still be able to prime their respective KUN readings, namely ウミ “sea” /umi_{kun}/ or ミズ “water” (mizu_{kun}). If so, this may indicate that computation of phonology from orthography, in Japanese compound words, also does allow full phonological activation of its subcomponents. Such a finding could further provide new openings for research concerning the debate regarding decomposed versus stored representation of compound words (for more information on this debate see Chapter 2; Caramazza, 1997; Bien, Levelt, & Baayen, 2005; Janssen, Bi, & Caramazza, 2008).

With regard to findings and conclusions mentioned in (2), it is important to establish whether kanji which only have a single pronunciation, such as 式 “ceremony” /shiki_{on}/ are not sensitive to semantic context effects (as there is no processing cost due to having a single pronunciation). Likewise, it would be valuable to determine whether Chinese hànzi, which can be pronounced in multiple ways (e.g. 行 /hang2/ or /xing2/), or heterophonic homographs in English (e.g. dove or read) are sensitive to semantic and/or phonological

context effects of pictures. If so, this would further substantiate the conclusion that the generation of the pronunciation of such words comes at a processing cost.

Furthermore, regarding (3), it would be interesting to compare bilingual speakers (or learners) of languages differing in functional unit (such as Japanese, Chinese, and English). If indeed the underlying nature of phonological encoding differs for some languages, then comparing such participants might shed more light on the underlying speech production mechanisms. For instance, native Chinese/Japanese speakers talking in English (L2) might show less segment exchange errors, or perhaps might “chunk” larger elements of words together due to the underlying nature of their L1. For instance, many Japanese speakers, when speaking English, show segment insertions following a moraic pattern (as Japanese does not allow for consonant clusters; e.g. “milk” becoming /mi.ru.ku/).

To conclude, this thesis shows that using the Japanese and Chinese language and script in psycholinguistic studies will reveal potential mechanisms underlying language production that cannot easily be investigated with western, alphabetic languages. Moreover, the involvement of languages such as Chinese and Japanese leads to new findings that are likely to foster the development of language-universal accounts of language production.

References

- Bien, H., Levelt, W., & Baayen, R. H. (2005). Frequency effects in compound production. *Proceedings of the National Academy of Sciences of the USA*, 102, 17876–17881.
- Caramazza, A. (1997). How many levels of processing are there in lexical access? *Cognitive Neuropsychology*, 14, 177-208.
- Chen, J.-Y., Chen, T.-M., & Dell, G. S. (2002). Word-form encoding in Mandarin Chinese as assessed by the implicit priming task. *Journal of Memory and Language*, 46, 751–781.
- Cutler, A., & Otake, T. (1994). Mora or phonemes? Further evidence for language-specific listening. *Journal of Memory and Language*, 33, 824–844.
- Dell, G. S. (1986). A spreading-activation theory of retrieval in sentence production. *Psychological Review*, 93, 283–321.
- Fushimi, T., Ijuin, M., Patterson, K., & Tatsumi, I. F. (1999). Consistency, frequency, and lexicality effects in naming Japanese Kanji. *Journal of Experimental Psychology: Human Perception and Performance*, 25, 382–407.
- Fushimi, T., Komori, K., Ikeda, M., Patterson, K., Ijuin, M., & Tanabe, H. (2003). Surface dyslexia in a Japanese patient with semantic dementia: Evidence for similarity-based orthography-to-phonology translation. *Neuropsychologia*, 41, 1644–1658.
- Glaser, W. R., & Döngelhoff, F. J. (1984). The time course of picture–word interference. *Journal of Experimental Psychology: Human Perception and Performance*, 10, 640–654.
- Glushko, R. J. (1979). The organization and activation of orthographic knowledge in reading aloud. *Journal of Experimental Psychology: Human Perception and Performance*, 5, 674–691.
- Janssen, N., Bi, Y., & Caramazza, A. (2008). A tale of two frequencies: Determining the speed of lexical access in Mandarin Chinese and English compounds. *Language and Cognitive Processes*, 23, 1191-1223.
- Kawamoto, A. H., & Zemblidge, J. (1992). Pronunciation of homographs. *Journal of Memory and Language*, 31, 349-374.
- Kayamoto, Y., Yamada, J., & Takashima, H. (1998). The consistency of multiple-pronunciation effects in reading: The case of Japanese logographs. *Journal of Psycholinguistic Research*, 27, 619–637.

Koester, D., & Schiller, N. O. (2008). Morphological priming in overt language production: Electrophysiological evidence from Dutch. *NeuroImage*, 42, 1622-1630.

Koester, D., & Schiller, N. O. (in press). The functional neuroanatomy of morphology in language production.

Kureta, Y., Fushimi, T., & Tatsumi, I. F. (2006). The functional unit of phonological encoding: Evidence for moraic representation in native Japanese speakers. *Journal of Experimental Psychology: Learning, Memory and Cognition*, 32, 1102-1119.

La Heij, W., Happel, B., & Mulder, M. (1990). Components of Stroop-like interference in word reading. *Acta Psychologica*, 73, 115-129.

Levelt, W. J. M., Roelofs, A., & Meyer, A. S. (1999). A theory of lexical access in speech production. *Behavioral and Brain Sciences*, 22, 1-75.

Lupker, S. J. (1982). The role of phonetic and orthographic similarity in picture-word interference. *Canadian Journal of Psychology*, 36, 349-367.

Mahon, B. Z., Costa, A., Peterson, R., Vargas, K. A., & Caramazza, A. (2007). Lexical Selection is not by competition: a reinterpretation of semantic interference and facilitation effects in the picture-word interference paradigm. *Journal of Experimental psychology, Learning, Memory and Cognition*, 33, 503-535.

MacCleod, C. M. (1991). Half a century of research on the Stroop effect: An integrative review. *Psychological Bulletin*, 109, 163-203.

Nakamura, K., Meguro, K., Yamazaki, H., Ishizaki, J., Saito, H., Saito, N., Shimada, M., Yamaguchi, S., Shimada, Y., & Yamadori, A. (1998). Kanji predominant alexia in advanced Alzheimer's disease. *Acta Neurologica Scandinavica*, 97, 237-243.

O'Seaghdha, P.G., Chen, J.-Y., Chen, T.-M. (2010). Proximate units in word production: Phonological encoding begins with syllables in Mandarin Chinese but with segments in English. *Cognition*, 115, 282-302.

Otake, T., Hatano, G., Cutler, A., & Mehler, J. (1993). Mora or syllable? Speech segmentation in Japanese. *Journal of Memory and Language*, 32, 258-278.

- Otake, T., & Cutler, A. (2003). Evidence against “Units of Perception”. In Shohov, S. P. (Ed.), *Advances in Psychology Research: Volume 24* (pp. 59–84). New York: Nova Science Publishers.
- Roelofs, A. (1992). A spreading-activation theory of lemma retrieval in speaking. *Cognition*, 42, 107-142.
- Roelofs, A. (2003). Goal-referenced selection of verbal action: Modeling attentional control in the Stroop task. *Psychological Review*, 110, 88–125.
- Sasanuma, S., Sakuma, N., & Kitano, K. (1992). Reading kanji without semantics: Evidence from a longitudinal study of dementia. *Cognitive Neuropsychology*, 9, 465–486.
- Schiller, N. O. (2000). Single word production in English: The role of subsyllabic units during phonological encoding. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 26, 512-528.
- Schiller, N. O., Jansma, B. M., Peters, J., & Levelt, W. J. M. (2006). Monitoring metrical stress in polysyllabic words. *Language and Cognitive Processes*, 21, 112-140.
- Schriefers, H., Meyer, A. S., & Levelt, W. J. M. (1990). Exploring the time course of lexical access in speech production: Picture-word interference studies. *Journal of Memory and Language*, 29, 86-102.
- Siok, W. T., Perfetti, C. A., Jin, Z., & Tan, L. H. (2004). Biological abnormality of impaired reading is constrained by culture. *Nature*, 431, 71–76.
- Warner, N., & Arai, T. (2001). Japanese mora-timing: A review. *Phonetica*, 58, 1–25.
- Wydell, T. N., Butterworth, B. L., & Patterson, K. E. (1995). The inconsistency of consistency effects in reading: Are there consistency effects in Kanji? *Journal of Experimental Psychology: Learning, Memory and Cognition*, 21, 1155–1168.
- Zhang, Q. F., & Weekes, B. S. (2009). Orthographic facilitation effects on spoken word production: Evidence from Chinese. *Language and Cognitive Processes*, 24, 1082–1096.
- Zwitserslood, P., Bölte, J., & Dohmes, P. (2000). Morphological effects on speech production: evidence from picture naming. *Language and Cognitive Processes*, 15, 563-591.

Samenvatting in het Nederlands

Samenvatting in het Nederlands.

Dit proefschrift richt zich op het hardop lezen van woorden in talen die een logografisch schrift gebruiken, in het bijzonder het Japans en het Chinees. Het gebruik van schrift is een vaardigheid die evolutionair recent verworven is en is een van de meest complexe menselijke gedragingen. Echter, ondanks de complexiteit, vindt het herkennen van woorden en de afleiding van betekenis en uitspraak binnen een tijdsspanne van ongeveer een halve seconde plaats.

Talen zoals het Nederlands en het Italiaans zijn vrij regelmatig in de omzetting van letters naar spraak. Er zijn echter ook talen (zoals het Engels) waarvan de uitspraak niet altijd te voorspellen is op basis van het uiterlijk van een woord. Denk hierbij bijvoorbeeld aan het Engelse woord voor “twee” (“two”) dat als /tu/ wordt uitgesproken en niet als /two/, zoals het volgens de spelling-naar-klank regels zou moeten klinken. Talen die logografische karakters gebruiken, zoals het Chinees en het Japans hebben geen vaste regels om letters naar klanken om te zetten. Met andere woorden, er zijn geen duidelijke visuele kenmerken die omgezet kunnen worden naar een bepaalde uitspraak.

In veel Chinese karakters zijn wel zogenaamde fonetische radicalen aanwezig (onderdeel van een karakter, aan de rechterkant geplaatst) zoals in het karakter 蚂 (“mier”) wat wordt uitgesproken als /ma3/ in het Chinees. In dit geval geeft het fonetische radicaal (het rechtergedeelte van het karakter, dus 马 /ma3/) een indicatie hoe het gehele karakter zou kunnen worden uitgesproken (in dit geval dus ook /ma3/). Het fonetische radicaal voorspelt de uitspraak echter lang niet altijd.

Misschien wel het moeilijkste schrijfsysteem in de wereld is het Japans. Japanners gebruiken drie schriften door elkaar heen namelijk kanji, hiragana en katakana. Hierbij zijn Japanse woorden uitgespeld in het alfabet (ook wel ‘romaji’ genaamd) en het gebruik van cijfers niet eens meegeteld. Een van deze schriften, het logografische kanji, is overgenomen uit het Chinees. Een opmerkelijk feit en belangrijk verschil met het Chinees is, dat in het Japans naast de bestaande Japanse uitspraken voor woorden in veel gevallen ook de Chinese uitspraak wordt gebruikt. Een voorbeeld hiervan is het karakter voor ‘water’ (水) wat in het Chinees wordt uitgesproken als /shui3/. In het Japans kan dat als /mizu/ (Japanse uitspraak, ook wel

KUN lezing genoemd) of als /sui/ (afgeleide Chinese uitspraak, ook wel ON lezing genoemd) worden uitgesproken. De correcte uitspraak hangt af van de combinatie die de kanji aangaat met een ander karakter. Bijvoorbeeld, voor het Japanse woord voor ‘regenwater’ wordt de Japanse uitspraak gebruikt, 雨水 /ama_{kun}.mizu_{kun}/ en in het woord voor ‘zeewater’ de Chinese uitspraak, 海水, /kai_{on}.sui_{on}/. Het grote verschil tussen Chinese en Japans karakters is dat Chinese karakters meestal maar één enkele uitspraak hebben terwijl de meeste Japanse karakters er meerdere hebben. Ondanks het feit dat als een karakter alleen staat vaak de Japanse (KUN) uitspraak, en als het in een samengesteld (compound) woord staat de Chinese uitspraak wordt gebruikt, is dit geen universeel principe (zie bovengenoemd voorbeeld van “water”).

In Hoofdstuk 2 van dit proefschrift wordt de vraag beantwoord of er bij lezers van het Japans meerdere uitspraken (de KUN en ON lezing) worden geactiveerd bij het zien van een kanji karakter. Dit werd gedaan door in een gemaskeerd priming experiment zeer kort (zodat bewuste waarneming niet mogelijk is), een kanji op een computerbeeldscherm aan te bieden (dit wordt de “prime” genoemd; bijvoorbeeld 水 ‘water’) en die te laten volgen door een op te lezen woord (de target). Het is bekend uit de literatuur dat een targetwoord sneller wordt opgelezen wanneer heel kort daarvoor een identiek woord wordt aangeboden. De target in dit hoofdstuk, die werd aangeboden in katakana, correspondeerde met de KUN lezing ミズ (mizu) of met de ON lezing スイ (sui), van water (水). Als de prime 水 wordt aangeboden en zowel /mizu_{kun}/ als /sui_{on}/ actief worden dan zou men voor beide target transcripties sneller moeten zijn dan bij een ongerelateerde prime. Dit patroon is precies wat er in Experiment 1 van dit hoofdstuk gevonden werd. In het tweede experiment werd de voorkeur voor één van de twee uitspraken gemanipuleerd, dat wil zeggen, de gebruikte kanji prime wordt vaker met één van de twee uitspraken in verband gebracht dan met de andere. De voorkeur is hierbij gebaseerd op de uitspraak in samengestelde woorden. Bijvoorbeeld, in samengestelde woorden wordt de kanji voor ‘bos’ 森 vaker met de Japanse (KUN) uitspraak gelezen en de kanji voor ‘misdaad’ 罪 vaker met de Chinese (ON) uitspraak. Wat er gevonden werd in dit experiment is dat de transcriptie voor de Japanse lezing (KUN) altijd meer profijt had van de prime dan de Chinese lezing

(ON), zelfs al was er een voorkeur voor de Chinese uitspraak. Dit komt waarschijnlijk omdat de kanji geïsoleerd werd aangeboden, waarbij in de praktijk vaak de Japanse uitspraak wordt gebruikt.

Gegeven de bevinding dat veel Kanji meerder uitspraken blijken te activeren is een vervolgvraag hoe de lezer van het Japans de juiste uitspraak selecteert. Deze vraag werd onderzocht door het lezen van Japanse Kanji te vergelijken met het lezen van Chinese hanzi. Zoals eerder genoemd heeft het Chinees vaak maar één uitspraak voor een karakter (er zijn enkele uitzonderingen zoals 行). Waarschijnlijk zal daardoor de selectie van de juiste uitspraak in het Chinees minder ingewikkeld zijn dan in het Japans.

In de Hoofdstukken 3 en 4 wordt onderzocht of het benoemen van Chinese en Japanse karakters beïnvloed wordt door plaatjes die gelijktijdig of 150 milliseconden eerder dan het op te lezen woord worden aangeboden (zogenaamde “contexteffecten”). Uit onderzoek is bekend dat de snelheid van benoemen van plaatjes in het Nederlands of Engels (bijvoorbeeld het plaatje van een hond als /hond/) beïnvloed wordt door een eerder, of gelijktijdig aangeboden contextwoord. Als men een semantisch gerelateerd woord op een plaatje van een hond afbeeldt (bijvoorbeeld ‘kat’) wordt het plaatje langzamer benoemd (als “hond”) dan wanneer er een niet-gerelateerd woord op staat (bv. ‘tafel’). Tevens is bekend dat als het contextwoord qua uitspraak gerelateerd is aan de naam van het plaatje (bv. ‘hop’) het plaatje sneller wordt benoemd (weer in vergelijking met een ongerelateerd woord). Als men de context en target in deze taak echter omdraait, dus het plaatje als context en het woord voor te lezen target gebruikt, dan verdwijnen deze effecten. De algemeen aanvaarde verklaring hiervoor is dat het oplezen van woorden via een snelle directe route gaat, die niet of nauwelijks beïnvloed kan worden door een gerelateerde context.

In de Hoofdstukken 3 en 4 werd deze taak gebruikt om te onderzoeken of de ambiguïteit bij het lezen van Japanse Kanji een selectieprobleem veroorzaakt dat de taak gevoeliger zou kunnen maken voor context stimuli. De experimenten in beide hoofdstukken toonden aan dat het lezen van Chinese karakters even ongevoelig is voor context stimuli als het oplezen van woorden in alfabetisch schrift. Als proefpersonen Chinese karakters oplazen maakte het niet uit of er een semantisch gerelateerd, fonologisch gerelateerd, of ongerelateerd plaatje als context werd gebruikt. Echter, bij het oplezen

van Japanse karakters maakte dat wel degelijk verschil. Semantisch en fonologisch gerelateerde context plaatjes (ten opzichte van niet-gerelateerde context plaatjes) versnelden het oplezen van Japanse kanji.

Eén mogelijke verklaring voor de gevonden contexteffecten bij het lezen van Japanse Kanji is dat de verwerking van deze karakters relatief veel tijd kost, waardoor de contextplaatjes tijd krijgen om een effect uit te oefenen. Zo'n verklaring lijkt echter minder waarschijnlijk, omdat het alleen maar eerder aanbieden van contextplaatjes bij het lezen van alfabetisch schrift niet voldoende is om contexteffecten te verkrijgen. Daarom wordt in Hoofdstuk 4 voorgesteld dat de zogenaamde activatiecoëfficiënt (die de sterkte van toename van activatie van een klankrepresentatie over de tijd bepaalt) minder groot is voor Japanse kanji karakters die meerdere uitspraken hebben. Daardoor is er meer tijd nodig om een bepaalde drempelwaarde te bereiken die nodig is om een uitspraak te selecteren. Gegeven het feit dat de toename in activatie over de tijd minder sterk is, is er meer kans voor een contextplaatje om een invloed op de opleestijden uit te oefenen. Samenvattend concluderen we dat alfabetische, Chinese en Japanse woorden waarschijnlijk via dezelfde route worden gelezen, maar doordat er bij het lezen van de meeste Japanse kanji een "cost" optreedt (doordat meerdere uitspraken mogelijk zijn), kunnen context plaatjes toch een effect uitoefenen.

Het laatste experimentele hoofdstuk, Hoofdstuk 5, gaat in op de vraag welke segmenten/onderdelen van uitspraken geactiveerd kunnen worden in het Japans. Uit eerdere studies is bekend dat in talen zoals het Nederlands en het Engels woorden zijn opgebouwd uit eenheden die 'fonemen' (klanken) worden genoemd. Bijvoorbeeld het woord /kat/ wordt opgebouwd door de fonemen /k/, /a/ en /t/ incrementeel te clusteren alvorens uit te spreken. Echter, uit ander onderzoek bleek dat in het Chinees soortgelijke resultaten niet gevonden werden en dat in die taal een andere unit in de opbouw wordt gebruikt, namelijk de lettergreep (syllabe). De Japanse taal wordt vaak beschouwd als een taal die op een eenheid van tijdsduur is gebaseerd die 'mora' wordt genoemd. De meest voorkomende mora combinaties in het Japans (ongeveer 60%) bestaan uit een medeklinker en een klinker, zoals in de mora /ka/ of /ta/. Andere combinaties zijn echter mogelijk en een enkele klinker zoals /a/ kan ook een mora zijn. Een medeklinker plus 'palatale approximant',

bijvoorbeeld “kyo”, een nasal coda ‘N’ en een geminatie (of verdubbeling) kunnen allen ook een ‘mora’ duren. Een goed voorbeeld is de bekende Japanse naam “Honda”. Dit woord heeft twee lettergrepen /HoN/ en /da/ maar 3 mora’s /ho/, /N/, en /da/ die elk ongeveer even lang duren.

Eerder onderzoek naar de vraag wat de Japanse basis eenheid van fonologische planning is, maakte gebruik van het zogenaamde *implicit priming paradigm* en wees in de richting van de mora. In dit onderzoek werden echter altijd kana of kanji karakters gebruikt als stimulus materiaal. Dit is niet ideaal omdat deze scripts geen individuele fonemen (of klanken) kunnen representeren. Het onderzoek in Hoofdstuk 5 maakt gebruik van een *masked priming paradigm* gecombineerd met stimuli die zowel in kana script (hiragana/katakana) als in romaji (Japans uitgespeld in letters) gepresenteerd werden. De nadruk lag hierbij op de vraag of er een priming effect gevonden kon worden voor één enkel overlappend foneem óf dat minimaal de gehele mora moets overlappen. Dus, met andere woorden, of een prime zoals ‘sumi’ すみ (versus ‘gumi’ ぐみ) benoemtijden voor een target zoals ‘seki’ (foneem overlap) of ‘sushi’ (mora overlap) zou kunnen beïnvloeden. De resultaten tonen aan dat het laatste het geval is. Er werd geen significant priming effect van het eerste foneem gevonden (bijvoorbeeld /s/) maar wel als prime en target die overlaptten in de gehele mora (bijvoorbeeld /su/). Een controle experiment toonde aan dat deze priming effecten ook optraden wanneer de mora onderdeel was van de lettergreep (dus niet de hele lettergreep zelf).

Concluderend wijzen de onderzoeksresultaten erop dat Japanse kanji karakters en Chinese hànzi karakters beide waarschijnlijk via een directe route van orthografie (letters) naar fonologie (klanken) gelezen kunnen worden en dat veel Japanse kanji daarbij meerdere uitspraken activeren. De verwerkings “cost” die daarvan het gevolg is, leidt ertoe dat context bij het lezen van Kanji van invloed kan zijn op de snelheid waarmee een woord kan worden gelezen. Tot slot wijzen onze resultaten in Hoofdstuk 5 erop dat de verschillende uitspraken in het Japans in eenheden van mora (en niet in fonemen) worden opgebouwd.

Curriculum Vitae

Rinus Verdonshot was born in Geldrop on the 13th of March 1977. He attended the College Asten-Someren in Asten from 1989-1995. After that he studied at the Conservatory Maastricht to obtain his bachelor degree in Music (1996-2001). Subsequently, he studied Biological Psychology (BSc) from 2002-2005 and Cognitive Neuroscience (MSc, *cum laude*) from 2005-2006 at Maastricht University. From november 2006-2010 he worked on his PhD project “Word processing in languages using non-alphabetic scripts: The cases of Japanese and Chinese” under supervision of Prof. Dr. N.O. Schiller and Dr. W. La Heij at Leiden University. From 2011 on he will be a Canon Foundation Research Fellow at Nagoya University, Japan.