

日本語母語話者による韓国語習得における語彙能力 と読解の因果関係

齊藤, 信浩
九州大学留学生センター : 准教授

玉岡, 賀津雄
名古屋大学大学院国際言語文化研究科 : 教授

<https://hdl.handle.net/2324/4483179>

出版情報 : STUDIA LINGUISTICA. 28, pp.111-123, 2014-12-05. 名古屋大学言語文化研究会
バージョン :
権利関係 :



日本語母語話者による韓国語習得における 語彙能力と読解の因果関係¹

齊藤 信浩²

玉岡 賀津雄³

要約: 韓国語は日本語と統語的に類似した言語ではあるが語彙体系は全く異なっている。また表記体系も日本語との類似点が見られない。そのため、韓国語の読解には語彙的な負荷が考えられる。本稿では、韓国語を第二言語として学習する日本語母語話者 61 名を対象に品詞と語種の下位項目からなる 48 問の韓国語語彙能力テストと、韓国語能力試験 (TOPIK) とハングル能力検定試験の配当級に準拠した 12 問からなる読解テストを実施し、語彙能力と読解との因果関係を構造方程式モデリング (SEM, structural equation modeling) の手法で検討した。その結果、語彙能力テストの品詞別に検討した場合、 $\beta=0.99$ ($p<.001$) で語彙能力が読解に極めて高い影響を及ぼしていた。同様に、語彙能力テストの語種別に検討した場合においても、 $\beta=0.98$ ($p<.001$) と極めて高い影響が見られた。この結果、韓国語の読解に語彙能力が強く影響していることが観察され、語彙能力を向上することで、読解が促進されることが予想される。

キーワード: 日本人韓国語学習者 因果関係 語彙能力 読解 構造方程式モデリング

¹ English title: A causal relationship between lexical ability and reading comprehension in Korean by native Japanese speakers

² SAITO, Nobuhiro: Associate Professor, International Student Center, Kyushu University, Japan, E-mail: nsaito88@isc.kyushu-u.ac.jp

³ TAMAOKA, Katsuo: Professor, Graduate School of Languages and Cultures, Nagoya University, Japan E-mail: ktamaoka@lang.nagoya-u.ac.jp

1. はじめに

韓国語は日本語と同じ漢字圏の言語として、漢字語彙の歴史的な借用関係があり、日本語と韓国語の漢字語彙の 80～90%は同形同義であると言われる(李漢燮, 1984; 曹喜澈 1991)。しかし、漢字語が同一であるとは言え、その音自体は異なった読み方がされるため、漢字の一音一音を個別に覚えなければならない。例えば、「사회(社会)」は「회사(会社)」に造語できるとは言え、「社」を「사[sa]」と「会」を「회[hwe]」と覚えなければならないという負荷がある。日本語の漢字音から、近い漢字音を予測的に想定することは可能であっても、初級から中級の段階において、漢字語彙を覚えるという負荷はやはり大きいと言える。また、漢字語彙の体系を有していても、現代韓国語においては漢字表記が用いられておらず、日本語母語話者が韓国語の読解を行う際に、視覚的に漢字を追って理解するというストラテジーが使えない。中国などの漢字語圏の場合は漢字表示された語彙が諸概念を示してくれるので、漢字語彙の視覚情報から読解を達成する傾向が見られる。しかしながら、韓国語の文章の読解では、漢字による視覚的な補助を得ることができないため、ハングル表記から音韻を介した語彙理解が要求される。この処理プロセスにおける認知負荷は大きいと予測される。韓国語と日本語では、固有語においても大きな違いが存在する両言語は統語構造上の類似性は指摘されつつも、語彙においては全く異なる体系であり、とりわけ固有語においては、両言語の類似性がないと言ってよい。このように、韓国語の習得において語彙の習得は非常に大きな比重を占めている。

そこで、本研究では、日本語母語話者に焦点を絞って、韓国語の語彙の習得が読解にどのような因果関係を持つかを検討することにした。

2. 先行研究

文章の読解は、文字や単語の符号化から文章全体の理解に至るまでの様々な下位過程から構成されている認知的活動であり、読解を促進する下位能力として語彙能力は重要な要素である(Haberlandt, 1984)。Yamashita (2002)は重回帰分析によって日本人英語学習者を対象に読解能力を構成する下位過程の能力を調べたところ、第二言語である英語の読解に対して語彙能力の影響が強いことを示している。このような読解に対する語彙能力の影響の因果関係の検証のために、

構造方程式モデリング(SEM, structural equation modeling)の方法で検証した研究がいくつかある。玉岡・宮岡・福田・母 (2007)は、中国語母語話者の日本語文章の読解の際に、語彙能力と文法能力がどのような因果関係にあるのかを構造方程式モデリングによって分析した。その結果、読解に対して、語彙能力に強い因果関係が見られ、文法能力の影響は弱かったと報告している。更に、Tamaoka, Miyaoka, Lim, Kim & Sakai(2007)においても、日本語を学習する中国語母語話者と韓国語母語話者の大きなデータプールから、性別、日本語学習歴、年齢、読解の変数がもっとも近い数値となるペアを 80 対選んだ。そして、均質な中国語母語話者 80 名と韓国語母語話者 80 名の被験者群を比較検討した。これは、ペアマッチサンプリング(pair-matched sampling)という手法である。

このサンプルに対して、構造方程式モデリングによって、語彙能力と文法能力のどちらの能力がより談話理解(discourse comprehension; 二つ以上の文の関係の理解)に影響しているかを検証した。その結果、中国語母語話者の場合は、語彙能力がより談話理解に貢献しており、一方で、韓国語母語話者は語彙と文法の両方の能力が談話理解に影響していたと報告している。これらの研究で示されているのは、語彙能力は常に有意な因果関係を持って、読解理解に影響を与えているということである。そこで、本研究でも、日本語を母語とする韓国語学習者の韓国語習得において、読解と語彙能力の因果関係を証明するために構造方程式モデリングの手法によって検証した。

3. 語彙能力の検証

韓国語の語彙能力を測定するためのテストが齊藤・玉岡(2014)によって開発されている。この語彙能力は、48 問の設問項目からなるテストであり、宮岡・玉岡・酒井(2011)が開発した日本語語彙能力の測定テストをモデルとして、韓国語習得研究のための能力測定テストとして開発した。韓国語の語彙能力テストの設問語彙は、韓国語能力試験(TOPIK)の出題語彙および長谷川・曹・大名(2012)によって検証された教育基幹語彙を参照にして抽出した。抽出した韓国語の語彙は、TOPIK の過去に出題された配当級、教育基幹語彙の調査による初級教科書相当および中級教科書相当の分類を照らし合わせながら、初級が 12 問、初中級が 12 問、中級が 12 問、上級が 12 問の 4 レベルの合計 48 問の構成として、各レベ

ルで問題数が同じになるように調整した。さらに、語彙の下位カテゴリとして、固有語を 24 問、漢字語を 24 問の二分類とした。これにより語種の視点からの分析が可能である。また、同様に下位カテゴリ

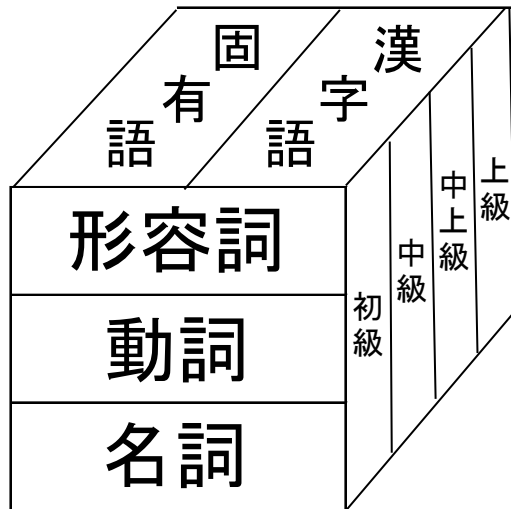


図1 韓国語語彙能力テストの構成

として、形容詞が 16 問、動詞が 16 問、名詞が 16 問による 3 分類も設けており、品詞の視点からも分析が可能のようにデザインした。

例えば、漢字語形容詞の初級問題では、「한국요리 중에서 김치는 가장 ().」とい文を与え、()に入る最適な表現を以下の 4 つから選択する四者択一形式である。具体的 な 回答 は , 「a 유명합니다」「b 조용합니다」「c 시원합니다」「d 미안합니다」の四択で、この場合の正答は「a 유명합니다」である。正答および錯乱肢は、同一のレベルの同一の語種から選択して作成した。つまり、初級配当級の中から作られた漢字語形容詞の設問文を選び、さらに錯乱肢は漢字語形容詞をレベルと語種が同じレベルになるよう統制した。これにより、レベル、語種および品詞の混在ができるだけ避けられるであろう。

本語彙テストの問題文において、四肢の選択肢が a, b, c, d の4つであり、ランダムな回答であれば 25%の確率が期待される。また、正答の位置を均等に配置するために、12 問ごとに異なる位置を正答とした。具体的には、a の正答が 12 問、c の正答が 12 問、d の正答が 12 問で、全体で 48 問の正答となる。これらを問題紙の中でランダムに配置した。

この語彙能力テストの得点は、斉藤・玉岡(印刷中)が項目応答理論(IRT, Item Response Theory)を使って、信頼性が高いことが検証されている。PASW (SPSS) Statistics Version 18.0 による下位カテゴリ間の弁別力についても、被験者の学習歴によって分けられた上位群(学習歴 24 ヶ月以上)と下位群(学習歴 24 ヶ月未満)の間での *t* 検定の結果、品詞別、語種別の全てにおいて 0.1%水準で有意差が見られ、レベル別の弁別力が明らかであった。また、学習歴による語彙能力の違いを

このテストで明瞭に弁別できることが報告されている。さらに、品詞別、語種別に弁別力を、反復のある一元配置の分散分析によって検討した結果、下位カテゴリの弁別力も確認されている。

4. 読解問題

読解問題は、韓国語能力試験(TOPIK)の上級(6級・5級)から3問、中級(4級・3級)から2問を選定し、加えて、ハングル能力検定試験の準2級(最高級は1級)から1問、3級から2問、4級から2問、5級から2問、合計12問を用意した。うち、会話形式の読解文が5問、説明文形式の読解文が7問である。会話文は対話形式によって展開するまとまったテキストの文章であり、独和形式のものも1つ含まれている(読解番号3;嬉しい日)。説明文は文語体によって内容説明が展開するまとまったテキストの文章である。

表1 韓国語読解問題の構成

読解番号	参照試験種別	配当級	種別	文体	タイトル	文字数	延べ語数	異なり語数
1	ハングル能力検定試験	5級	会話文	해요	図書館で	75	27	21
2	ハングル能力検定試験	5級	会話文	합쇼	位置を示す	67	24	20
3	ハングル能力検定試験	4級	会話文	합쇼	嬉しい日	72	23	22
4	ハングル能力検定試験	4級	会話文	해요	財布が欲しい	116	47	44
5	ハングル能力検定試験	3級	会話文	합쇼	貧乏人の夢	166	58	48
6	ハングル能力検定試験	3級	説明文	해라	ユミちゃん	145	50	44
7	ハングル能力検定試験	準2級	説明文	합쇼	商売と欲	360	125	106
8	韓国語能力試験	中級	説明文	해라	電子レンジ	111	36	30
9	韓国語能力試験	中級	説明文	해라	楽器	113	41	36
10	韓国語能力試験	上級	説明文	해라	古代の鏡	132	42	36
11	韓国語能力試験	上級	説明文	해라	詩人	111	39	35
12	韓国語能力試験	上級	説明文	해라	紙幣偽造	142	50	44

1つの読解文に対して、その内容を問う四肢選択式の問題が1つ用意されており、各読解問題は1点で、読解問題全体で満点は12点となる。読解問題の構成は以下、表1のようになっている。読解問題のタイトルは内容に合わせて本稿で名づけたものであり、原文には題名は付けられていない。文字数は初級レベルのもので70字前後、中級レベルで110字前後、上級レベルで110～140字前後であり、全体的に比較的短めの読解文を用意した。韓国語の対者敬語法の分類に従うと

(남기심・고영근, 1985; Lee, et al. 2001; 野間, 2012), 会話文は, 합체와 해체의文章であり, 説明文は, 합체와 해체의文章である。韓国語では文体という用語はあまり用いられないが, ここでは日本語の枠組みに従って, 文体という用語を用いて表に記載した。

5. 調査

韓国語を主専攻または副専攻とするコースを持つ, 福岡県, および, 山口県に所在する5箇所の大学で学ぶ, 日本語を第一言語とする韓国語学習者 61 名を被験者として調査を行った。61 名のうち, 男性は 3 人, 女性は 58 名, 調査実施日時点の年齢は, 平均で 20 歳 7 か月, 標準偏差は 13 歳と 8 ヶ月であった。2013 年 6 月～7 月に調査を実施し, 試験時間は 90 分, 試験監督をつけ, 辞書の使用を禁止して解答をさせた。

構造方程式モデリング(SEM)で, 因果関係を検討するにあたり, 読解と語彙能力を潜在変数(latent variable)として設定した。これらの潜在変数は, 品詞別にみると3種類, 語種別にみると2種類の観測変数として設定できる。以下に, その詳細を示す。

表 2 潜在変数と観測変数の満点, 平均, 標準偏差および信頼性係数

潜在変数	観測変数	満点	平均	標準偏差
語彙能力 ($\alpha=.945$)	形容詞	16	9.77	3.92
	動詞	16	10.00	4.16
	名詞	16	9.97	4.23
	固有語	24	14.21	5.40
	漢字語	24	15.11	5.99
	合計	48	29.74	11.51
読解 ($\alpha=.529$)	会話文	5	3.41	1.37
	説明文	7	2.25	1.26
	合計	12	5.66	2.14

語彙テストの被験者数 61 名の平均点は 29.74, 標準偏差は 11.51 であった。語彙能力テストのクロンバックの信頼度係数(Cronbach's reliability coefficient)は, 被験者数が 61 名であるにも拘わらず, 問題項目数が 48 問の条件で, $\alpha=.945$ ($N=48$)となり, 極めて高い信頼性を示した。下位カテゴリごとに見ても, 形容詞

($N=16$)が、 $\alpha=.860$, 動詞($N=16$)が $\alpha=.834$, 名詞($N=16$)が、 $\alpha=.875$ となり、問題数 16 問でも、どの品詞においても高い信頼度を示していた。

一方、被験者数 61 名での読解の平均は 5.66, 標準偏差は 2.14 であった。読解テストのクロンバックの信頼度係数は、被験者数 61 名で、問題数が 12 問で、 $\alpha=.529$ となり、妥当とみなされる 0.700 に届かなかった。しかし、問題数が少ない条件であることを考えると決して悪い信頼度でもないと考えられよう。読解の内容および設問は背景知識からも影響されるので、語彙能力テストのような単一の知識の検証とはなり難い。そのため、12 問の読解文のテストの信頼性係数は低くなっているのだと考えられる。下位カテゴリ別に見ると、会話文 5 問は $\alpha=.605$ で、説明文 7 問は $\alpha=.294$ で、説明文の信頼性が低かったことを報告しておく。

6. 分析

6.1. 構造方程式モデリングの適合度指標の概観

本研究の因果関係のモデルの検証のために、構造方程式モデリング(SEM)を用いて因果関係の分析を行った。その解析には、PASW Statistics Version 18.0 の AMOS の統計ソフトを使用した(Arbuckle 2009)。まず、構造方程式モデリングの適合度指標について言及する。

GFI(Goodness-of-fit index)の指標は、Marsh and Grayson(1995)によれば、1 に近いほど良い指標であり、0.95 以上であれば、良好な適合、0.90 以上であれば許容範囲内の適合を示している。自由度(df)の影響を考慮して GFI を補正した AGFI(adjusted GFI)も 1 に近いほど良い指標であり、0.90 以上で良好であるとされる。

NFI に自由度の影響を考慮した CFI(Comparative Fit Index)の指標は、0.97 以上が良好で、0.95 以上が許容範囲内の適合を示すとされている(Schermelleh-Engel et.al 2003)。

NFI (Normed Fit Index)の指標は、0.95 以上であれば良好な適合 (Kaplan, 2000), 0.90 以上で許容できる程度の適合とされる(Marsh & Grayson, 1995)。

標本数と自由度で基準化したカイ二乗統計値である RMSEA (Root Mean Square Error of Approximation)の指標は、0.05 以下であれば良好な適合度を示していると言える(Browne and Cudeck 1993)。

6.2. 構造方程式モデリングによる因果関係の検証

語彙能力を予測するために、品詞(形容詞・動詞・名詞)の 3 つの下位カテゴリから語彙能力を検証し、読解と語彙能力との因果関係を検証した。SEM の分析結果は、図2の通りである。図2のモデルはカイ二乗適合度検定の χ^2 値が有意ではなかった[$N=61$, $\chi^2(4)=6.006$, $p=.199$, ns]ので、データがモデルとよく適合しており、モデルを有効に説明していると言える。図2のモデルは、GFI が.984, AGFI が.940, CFIが.995, NFIが.986で非常に良い適合度を示した。RMSEA は.059で、良い適合指標とされる.050よりは大きい数値が出ているが、全体の適合度指標としては、非常に良好な結果であったと言える。

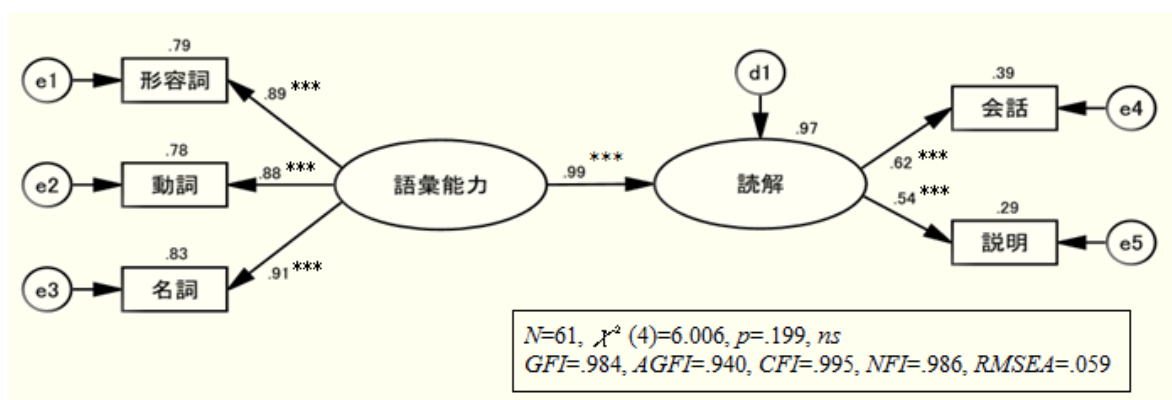


図2 品詞別に見た語彙能力と読解との因果関係

データがモデルを適切に説明していることが実証されたので、実際の因果関係を図2にしたがって試みていく。読解の下位カテゴリは会話文と説明文からなり、会話文が $\beta=0.62$ ($p<.001$), 説明文が $\beta=0.54$ ($p<.001$)という標準偏回帰係数を示した。読解という潜在変数をそれぞれに有意に構成する観測変数となっている。語彙能力を構成する下位カテゴリの品詞は、形容詞が $\beta=0.89$ ($p<.001$), 動詞が $\beta=0.88$ ($p<.001$), 名詞が $\beta=0.91$ ($p<.001$)であった。語彙能力の潜在変数を、これら3つの観測変数は、きわめて有意に貢献していた。本研究の最も重要な語彙能力から読解への因果関係は、 $\beta=0.99$ ($p<.001$)と極めて高い因果関係を示した。このことから、日本人の韓国語学習者は、語彙能力が読解の得点をほぼ決定することが分かった。

次に、語種で分けた固有語と漢字語の2つの観測変数から語彙能力を検証し、読解と語彙能力との因果関係を検証した。結果は、図3に示した通りである。カイ二乗適合度検定の χ^2 値が有意であった[$N=61$, $\chi^2(1)=5.241$, $p<.05$]。したがって、データがこのモデルと最適な適合とは言えなかった。図3のモデルは、GFI は.982で高いが、補正した AGFI は.824 で 0.90 を下回った。しかし、CFI は.984, NFI は.983 で非常に良い適合指標を示した。ただし、RMSEA は.172 で、良い適合指標とされる.050 よりも大きい数値となった。語彙能力を品詞で測定した場合に比べて、適合度が落ちた。しかし、全体の適合度指標を総括してみると、検討の余地のあるモデルとは言えよう。

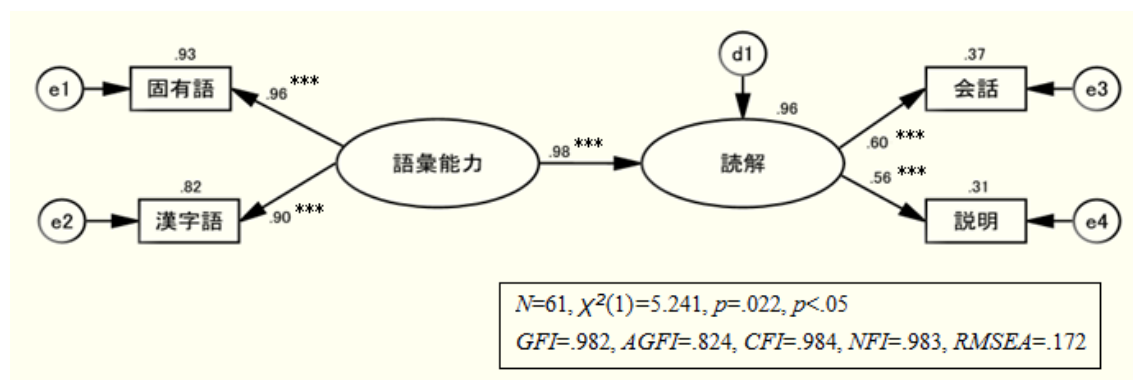


図3 語種別に見た語彙能力と読解との因果関係

図2の品詞別で語彙能力を測定した場合よりも、語種別で測定した図3の方が、会話文 $\beta=0.60$ ($p<.001$) では貢献度が落ちたが、説明文 $\beta=0.56$ ($p<.001$) では上がっている。結果的に、読解という潜在変数を構成する両観測変数の貢献度は、語彙能力を構成する観測変数の下位カテゴリーが多少変わっても大きく変化しないようである。語彙能力を語種別に測定した場合も、固有語が $\beta=0.96$ ($p<.001$)、漢字語は $\beta=0.90$ ($p<.001$) で、潜在変数の語彙能力に、両観測変数共に非常に高く貢献をしていた。語彙能力から読解への標準偏回帰係数は、 $\beta=0.98$ ($p<.001$) で、これも品詞別に見た場合と同様に、極めて高い因果関係であった。

以上のように、読解を予測するという因果関係を説明する上で、モデルとデータの適度度は、語彙能力を3つの品詞で測定した場合(図2)の方が、2つの語種で測定した場合(図3)よりも、より適合しているという結果であった。SEM のような因果

関係モデルでは、潜在変数を説明する観測変数は2つよりも3つの方がより適合すると言われている。しかし、データとモデルの適合度を置いて考えると、品詞別でも語種別でも、語彙能力が読解の理解を強く促進する因果関係が見られ、日本語を母語とする韓国語学習者にとって、語彙の獲得が読解に極めて大きく貢献することが実証された。

7. おわりに

韓国語語彙能力テストの得点は、読解テストの得点に有意な因果関係を持ち、語彙能力の読解を非常に強く促進することが分かった。日本語を母語とする韓国語学習者にとって、読解の向上には語彙能力を向上させていくことで最大の効果が期待できそうである。また、語彙能力という潜在変数を測定するにあたっては、品詞別の方が語種別よりも適合度が高くなることが分かった。因果関係の研究としては、品詞別の観測変数を使って語彙能力を定義することを勧める。

本研究の残された課題としては、まず、検証可能な文法能力テストの開発がまだであり、語彙能力と文法能力との比較ができなかった。これに加えて、日本語母語話者のみのデータの検証では、読解に語彙能力が大きな貢献しているかどうか普遍的な因果関係として考えて良いかどうか判定し難い。今後、他の言語を母語とする韓国語学習者へと拡大した調査と分析を進めて行く必要がある。

[付記]

本研究は、文部科学省科学研究費助成事業挑戦的萌芽研究(C) (研究課題番号:24652123)「日本語母語話者の韓国語習得における文法知識と語彙知識が読解に与える影響について」の成果の一部を報告したものである。

[引用文献]

- Arbuckle, J. L. (2009). *AMOS 18.0 User's Guide*. Chicago, IL: SPSS Inc.
- Browne, M. W., & Cudeck, R. (1993). Alternative ways of assessing model fit. In K. A. Bollen & J. S. Long (Eds.), *Testing structural equation models*. Newbury Park, CA: Sage, 136-162.

- Haberlandt, K. (1984). *Components of Sentence and Word Reading Times*. New Methods in Reading Comprehension Research ed. by D.E. Kieras & M.A. Just, Hillsdale, NJ: Erlbaum, 219-252.
- Iksop, Lee, & S. Robert Ramsey. (2001). *The Korean Language*, State University of New York Press
- Jöreskog, K., & Sörbom, D. (1993). *LISREL 8: Structural equation modeling with the SIMPLIS command language*. Chicago, IL: Scientific Software International.
- Kaplan, D. (2000). *Structural equation modeling: Foundation and extensions*. Thousand Oaks, CA: Sage.
- Marsh, H. W., & Grayson, D. (1995). Latent variable models of multitrait-multimethod data. In R. Hoyle (Ed.), *Structural equation modeling: Concepts, issues and applications*, Thousand Oaks, CA: Sage, 177-198.
- Schermelleh-Engel, K., Moosbrugger, H., & Müller, H. (2003). Evaluating the fit of structural equation models: Tests of significance and descriptive goodness-of-fit measures. *Methods of Psychological Research*, 8, 23-74.
- Tamaoka, Katsuo, Yayoi Miyaoka, Hyunjung Lim, Sujin Kim, and Hiromu Sakai (2007). Differences in discourse comprehension strategies for L2 (second language) Japanese as employed by pair-matched L1 (first language) Chinese and Korean Speakers. 日本言語学会第135回大会予稿集, 316-321.
- Yamashita, Junko (2002). Reading strategies in L1 and L2: Comparison of four groups of readers with different reading ability in L1 and L2. *ITL: Review of Applied Linguistics*, 135 & 136, 1-35.
- 李漢燮 (1984) 「日韓同形の漢字表記語彙」『日本語学』3月号, 明治書院
- 齊藤信浩・玉岡賀津雄・母育新 (2012) 「中国人日本語学習者の文章および文レベルの理解における語彙と文法能力の影響」, 『ことばの科学』25, 愛知: 名古屋大学言語文化学部言語文化研究会.
- 齊藤信浩・玉岡賀津雄 (印刷中) 「項目応答理論による韓国語語彙能力テストの開発」『朝鮮学報』
- 曹喜澈 (1991) 「日韓同形漢字語の語義・用法の相違」, 『日本近代語研究』, ひつ

じ書房.

玉岡賀津雄・宮岡弥生・福田倫子・母育新 (2007). 中国語を母語とする日本語学習者の語彙と文法の知識が聴解・読解および談話能力に及ぼす影響, 2007 年度日本語教育学会秋季大会予稿集, 131-136.

豊田秀樹 (1998). 共分散構造分析(入門編)—構造法的式モデリング, 東京: 朝倉書店.

長谷川由起子・曹美庚・大名力 (2010)『韓国語教材における語彙使用頻度調査研究』, 福岡: 九州大学大学院言語文化研究院.

野間秀樹 (2012)「待遇表現と待遇法を考えるために」, 『韓国語教育論講座 第2巻』, くろしお出版

宮岡弥生・酒井弘・玉岡賀津雄 (2011)「日本語語彙テストの開発と信頼性—中国語を母語とする日本語学習者によるテスト評価—」『広島経済大学研究論集』34(1), 1-17.

남기심・고영근 (1985)『표준국어문법론』, 탐출판사.

齐藤信浩 (SAITO, Nobuhiro)

nsaito88@isc.kyushu-u.ac.jp

九州大学留学生センター准教授

玉岡賀津雄 (TAMAOKA, Katsuo)

ktamaoka@lang.nagoya-u.ac.jp

名古屋大学大学院国際言語文化研究科教授

A causal relationship between lexical ability and reading comprehension in Korean by native Japanese speakers

SAITO, Nobuhiro

TAMAOKA, Katsuo

Although the Korean language is similar to Japanese in the perspective of syntactic structure, lexical items greatly differ from each other with dissimilar scripts used for each language, *kana* and *kanji* for Japanese and *hangeul* for Korean. The lexical differences between the languages impose difficulty for native Japanese speakers reading Korean texts. Therefore, by testing native Japanese speakers learning Korean, the present study investigated the causal relation of lexical knowledge to reading comprehension in Korean. Both a vocabulary test, created by Saito and Tamaoka (in press), and a reading comprehension test, questions taken from the Test of Proficiency in Korean (TOPIK) and the Korean Language Proficiency Test, were given to measure both lexical knowledge and reading comprehension respectively. Using structural equation modeling (SEM) it was revealed that the Korean lexical knowledge measured by three word class categories (i.e., noun, adjective, and verbs) significantly contributed to reading comprehension ($\beta=0.99$, $p<.001$). Likewise, the lexical knowledge of two word type categories (i.e., Korean-originated words and Chinese-originated words) also significantly contributed to reading comprehension ($\beta=0.98$, $p<.001$). Consequently, during the early stages of Korean learning, lexical knowledge is an important factor of success for Japanese speakers to be able to comprehend Korean texts.

Keywords: native Japanese speakers learning Korean, lexical knowledge, reading comprehension, casual relation, structural equation modeling (SEM)

