

기본 어휘의 선정 기준

—영어 어휘를 중심으로—

신동광*

<차 례>

- I. 서론
- II. 이론적 배경
- III. 어휘 선정 기준
- IV. 어휘 선정 기준의 적용 방안
- V. 결론 및 제언

I. 서론

최근 언어교육에 활용되고 있는 가장 흥미로운 디지털 기술 중 하나를 들자면 ‘코퍼스(corpus)’라고 말할 수 있을 것이다. 코퍼스는 원어린이 표현한 문자 혹은 음성언어를 문서화한 언어 자료를 의미하지만 최근에는 보통 컴퓨터로 분석 가능한 텍스트 파일 형태의 디지털화된 언어 자료를 통칭한다. 즉, 코퍼스는 본질적으로 컴퓨터가 읽을 수 있는 텍스트(computer-readable text)의 집합체이며, 때로는 수백만 단어로 구성되며, 언어 사용에 있어 특정 영역의 대표형(representative)으로 여겨진다. 코퍼스는 초기 사전(dictionary) 제작 또는 언어학적으로 규범문법(prescriptive grammar)에 대응한 기술문법(descriptive grammar)의 근거를 마련하는 데 주로 사용되었지만 이제

* 한국교육과정평가원(sdhera@kice.re.kr)

는 언어교육적인 측면에서 그 활용 범위를 넓혀가고 있다.

초기 코퍼스의 활용은 코퍼스로부터 유용한 표현을 추출하거나 사전의 예문을 들어주거나 주요 어휘를 기준으로 수준별 교재를 제작하는 데 주로 사용되었다. 컴퓨터를 바탕으로 한 코퍼스 기술이 활성화되면서 가능하게 된 이러한 초기 코퍼스 활용을 소위 “Behind the Scenes Approach”라고 한다. 이 접근법은 말 그대로 현장에서 한걸음 물러나 미리 준비한다는 의미로 온라인으로 코퍼스를 실시간 활용하는 것이 아니라 오프라인으로 코퍼스를 분석하여 학습내용을 선정하거나 교재를 개발하는 데 활용하는 것이다. 1990대 초반까지는 이 접근법이 주를 이루었고 현재까지도 많은 연구들이 진행되어 왔다(예, Chung, 2003; Coxhead, 2000; Laufer, Elder, Hill, & Congdon, 2004; Leech, Rayson, & Willson, 2001; Meara, 2005; Nation, 2006). 이 접근법에서는 기본적으로 가장 많이 쓰이는 표현들이 가장 유용하다는 가정 하에 빈도수 조사를 통해 학습 내용을 선정하였다. 또한 단정할 수는 없지만 실제로도 어느 정도는 빈도수가 습득 순서(acquisition order)와도 연관성을 가지고 있을 것으로 사료된다. 학습자는 자주 쓰이는 표현에 더 많이 노출될 것이고 자연스럽게 그 표현은 빠르게 기억될 것이기 때문이다.

이러한 컴퓨터 분석을 위해서는 어휘 단위의 개념을 명확히 할 필요가 있다. 현재 국내에서 개발된 어휘 목록은 어휘의 개념에 비추어 볼 때 일관성 면에서 여러 문제점을 노출하고 있다. 예를 들어 영어과 교육과정에서는 어휘의 기본형으로 기본 어휘를 제시하고 있지만 ‘act’와 ‘active’, ‘beauty’와 ‘beautiful’ 등의 단어들은 별개의 어휘로 등록되어 있다. 다음은 일반적으로 통용되는 어휘의 정의에 따른 굴절형 또는 파생형의 포함 범위를 나타낸 것이다.

- (1) 출현형(token) : 텍스트를 구성하고 있는 총 단어를 의미한다.
- (2) 낱말 유형(type) : 텍스트를 구성하고 있는 단어 중에서 중복이 없이 순수하게 한 번씩만 집계했을 때 보이는 단어의 유형으로 단어의 형태가 다르면 다른 유형으로 간주한다.
- (3) 사전 등재형(lemma) : 굴절 변이형을 포함하는 기본형으로 예를 들

어 ‘swim’, ‘swims’, ‘swam’, ‘swimming’은 모두 동사라는 품사를 그대로 유지한 채 문법적인 굴절만을 보여주고 있으므로 이 네 단어의 lemma는 기본형 ‘swim’이다. 하지만 굴절의 범위는 여러 이론에 따라 약간의 이견을 보이기도 한다.

- (4) 단어군(word family) : 가장 포괄적인 기본형으로 품사에 상관없이 굴절과 파생 변이형을 모두 포함한다.

어휘의 빈도 분석은 어느 정도의 어휘수를 알고 있으면 어느 수준까지 영어 표현을 이해하고 표현할 수 있는지 그 포괄 범위의 예측을 가능하게 해주었다. 다음의 텍스트 포괄 범위(text coverage)의 예이다.

- (1) ‘the’ : 7%
 - (2) 최상위 빈도 10 단어 : 25%(coverage)
 - (3) 최상위 빈도 100 단어 : 50%
 - (4) 최상위 빈도 1000 단어 : 75%
 - (5) 최상위 빈도 2000 단어 : 80%(문어)/90%(구어)
- * 기능어=47%

코퍼스 활용에 관한 연구가 가속화되고 웹기반 인프라가 급격히 발전함에 따라 코퍼스 활용에 대한 새로운 요구들이 생겨나기 시작했다. 이러한 접근법은 코퍼스의 샘플을 온라인으로 직접 연결하여 그 예를 학생들에게 제시하거나 코퍼스를 바탕으로 예시 문항을 현장에서 바로 제작하여 활용하는 것으로 소위 “On Stage Approach”라 한다. 이 접근법은 Johns (1986, 1988, 1991)가 제시한 “Data-Driven Learning(DDL)”으로 더 널리 알려져 있다. 말 그대로 데이터인 코퍼스에서 자료를 끌어와 교수학습에 활용하는 것으로 학습자들은 많은 샘플을 보면서 목표로 하는 의미를 찾아보거나 특정 표현의 쓰임을 직접 접해보면서 문법 또는 어휘를 귀납적으로 습득하는 발견식 학습을 그 원리로 하고 있다. DDL은 최근 들어 Cobb (2007)와 같은 학자들이 코퍼스 활용에 기초를 둔 웹사이트를 구축하여 그

활용을 극대화 하고 있는데 이러한 코퍼스를 이용한 언어교육에서의 진보의 기반이 되었던 것은 어휘 분석 및 선정에 대한 체계적인 연구에서 비롯된다.

II. 이론적 배경

단어군을 산정하는 데 있어 우선 고려되어야 하는 요인이 바로 텍스트 포괄 범위이다. 다음은 단어군 수에 따른 텍스트 포괄 범위를 보여주는 예이다.

〈표 1〉 단어군과 텍스트 포괄 범위

	단어군	텍스트 포괄범위(%)
General Service List	2,000	80~90
Academic Word List	570	2[g]/10~13[a]
Technical words	2,000	3~5[g]/95[t]
Low-frequency words	123,200	2

* g : general text, a : academic text, t : technical text

Nation(1990, 2001)에 의하면 가장 빈도수가 높은 2,000개의 word family로 구성된 West(1953)의 General Service List(GSL)는 구어(spoken English)에서 80% 그리고 문어(written English)에서는 90%를 포괄한다고 한다. 하지만 West의 어휘 목록은 그 당시의 다른 어휘 목록에 비해 많은 장점을 지닌 반면에 유용성이 낮은 어휘들, 예를 들어 ‘lump’, ‘loyal’, ‘mannerism’, ‘mere’, ‘ornament’, ‘curse’, ‘vessel’, ‘urge’ 등의 어휘가 포함되어 있었으며 Richards(1974)는 West의 어휘 목록이 1930년대 이전에 만들어진 것이기 때문에 ‘internet’, ‘television’ 등과 같이 빈도수가 높은 낱말들이 포함되어 있지 않다고 지적하였다. 결국 GSL은 신조어들이 계속해서 추가되지 않기

때문에 시간이 흐를수록 실제 텍스트 포괄 범위는 하락할 수밖에 없으며 현재 분석 시에도 Nation의 분석에 비해 다소 낮게 나타날 수 있다. 그럼에도 불구하고 GSL은 여전히 중요한 어휘 선정의 기준이 되고 있다. 이밖에도 Coxhead(1998, 2000)의 Academic Word List(AWL)는 28개 교과목, 3,600,000개의 token으로 구성된 Academic Corpus로부터 추출된 570개의 단어군으로 구성되어 있다. 이 570개의 단어군은 대학 교과목에 상관없이 공통적으로 대학의 교과목을 이수하기 위해 필요한 필수 단어들이며 일반적인 텍스트에서는 2% 내외의 단어들을 포괄하지만 학문적인 텍스트에선 10~13%의 총 어휘수를 포괄한다. 뿐만 아니라 전문적인 영역에서 집중되는 전문용어는 일반적으로 총 텍스트 구성 어휘 중 3~5%를 정도만 포괄하지만 전문 영역, 예를 들어 Ward(1999)에 따르면 기술 관련(engineering) 텍스트에서 2,000개의 전문용어(technical word)가 무려 95%의 어휘를 포괄한다는 것을 보여주었다. 따라서 이러한 어휘들은 글의 분야나 장르에 따라 큰 편차를 보인다는 것을 알 수 있다. 반면 이외의 123,200개의 빈도수가 낮은 어휘들을 저빈도 단어(low-frequency word)라 하며 이들 어휘는 많은 수에도 불구하고 일반 텍스트의 약 2% 내외의 어휘만을 포괄한다. 따라서 어휘 목록의 작성은 어휘 목록의 사용 목적에 따라 선정 기준이 달라질 수 있다. 국가수준 영어능력인증시험 개발에서는 실용적이고 일상적인 내용뿐 아니라 일부 학문적인 내용도 포괄할 수 있는, 하지만 전반적으로는 보편적이며 대중적인 어휘를 선별하는 것이 그 목표이다.

Hu와 Nation(2000)은 어떠한 글을 읽던 간에 최소한으로 알아야 할 기본 어휘량을 텍스트의 80%로 제시하고 98% 이상을 알아야 모르는 단어를 접했을 때도 외부의 어떠한 도움이 없이도 문맥에서 단어의 뜻을 유추하며 독립적으로 글을 읽을 수 있다고 주장한다. 아래 표에서 보이듯 위의 98%의 텍스트 포괄 범위를 안정적으로 확보하기 위해서는 문어에서 8,000~9,000개의 단어군이 필요하고 구어에서는 6,000~7,000개의 단어군을 요구한다.

〈표 2〉 문어와 구어에서의 텍스트 포괄 범위 비교(Nation, 2006, p. 79)

Levels	No. of Levels	Approximate Written Coverage(%)	Approximate Spoken Coverage(%)
1st 1000	1	78-81	81-84
2nd 1000	1	8-9	5-6
3rd 1000	1	3-5	2-3
4th-5th 1000	2	3	1.5-3
6th-9th 1000	4	2	0.75-1
10th-14th 1000	5	< 1	0.5
Proper nouns	1	2-4	1-1.5
not in the lists	1	1-3	1

하지만 6,000~9,000개에 육박하는 어휘수는 EFL(English as a Foreign Language) 환경에서 쉽게 수용되기에는 힘든 수치이다. 현실적이며 가장 보편적으로 추천되는 텍스트 포괄 범위는 95%이다(Read, 2004). 위의 표에 의하면 5,000단어(word family) 수준일 때 구어에서는 92~98%, 문어에서는 89.5~96%까지의 텍스트 포괄 범위를 보여주고 있어 가장 이상적인 어휘수라고 판단된다. 뿐만 아니라 Laufer(1992)에 따르면 영어로 강의가 진행되는 대학교육에서 요구되는 어휘수는 5,000개로 대학교재의 개발이나 일상적인 의사소통, 학문적인 과제를 해결하기 위해 반드시 필요하다고 보고 있다.

본 연구에서는 어휘 선정의 기준에 대한 고찰과 선정 기준이 어휘 목록에 미치는 영향 그리고 그러한 어휘 선정 기준의 적용 방안에 대하여 살펴보고자 한다.

III. 어휘 선정 기준

어휘 선정 기준에는 여러 가지가 있지만 인간의 주관을 배제하고 객

관적인 분석만을 통해 추출하기 위해서는 국제적으로 크게 빈도수(Frequency), 사용범위(Range), 사용분포(Dispersion) 세 가지 기준을 적용한다. 다음은 각 기준들에 대한 설명이다.

1. 빈도수

최근까지 만들어진 대부분의 어휘 목록은 빈도수가 높은 순으로 어휘를 추출하여 어휘 목록을 제작하였다. 가장 쉽게 어휘 목록을 제작할 수 있을 뿐만 아니라 빈도수가 높다는 것은 일상생활에서 그만큼 많이 쓰이고 유용하다고 볼 수 있기 때문에 경제성의 원리에서도 빈도수가 가지는 장점은 매우 크다.

2. 사용범위

사용범위는 하나의 어휘가 얼마나 다양한 텍스트에 사용되는가를 측정하는 것이다. 예를 들어 ‘kiwi’라는 단어는 여러 가지 뜻을 지니는데, “날지 못하는 특정 새”를 지칭하기도 하고 “특정 과일”을 지칭하기도 하며 “뉴질랜드인”을 가리키기도 한다. 따라서 뉴질랜드 코퍼스에서는 빈도수가 상당히 높게 나타나고 있다. 하지만 ‘kiwi’라는 단어는 뉴질랜드에 국한해서 빈도수가 높은 것이지 다른 지역의 영어에서는 빈도수가 극히 낮거나 아예 나타나지 않는다. 따라서 절대 빈도수만 보고 어휘를 추출하는 것은 보편성을 대표하는 데 문제가 있다고 볼 수 있으며 이러한 문제를 극복하기 위해서는 얼마나 다양한 환경에서 쓰이는지를 측정하는 사용범위가 보다 적당한 기준이라고 볼 수 있다. 실제로 Coxhead(1998, 2000)의 AWL의 경우에는 28개의 대학교과목 중 최소 15개의 교과목에서 사용되는 어휘만을 추출하였다.

3. 사용분포

사용분포는 “Spread Frequency”라고도 불리며 사용범위가 단순히 한 단어의 빈도수가 한 번이라도 상관없이 몇 개의 다른 코퍼스에서 사용했는지를 측정한 반면 사용분포는 한 코퍼스 안에서 일정한 빈도수를 유지하는 것을 측정하는 것이다. 예를 들어 AWL의 제작 시에는 4개의 대영역 즉 Law, Science, Arts, Commerce 모두에서 각각 10번이상의 빈도수를 가지는 단어들만 추출하였다. 사용분포를 보다 체계적으로 구하는 공식은 다음과 같다.

$$\text{Dispersion(사용분포)} = 100 \times [1 - (V/\text{사용한 corpus의 수} - 1)]$$

*V = 각 corpus에서 나타나는 type들이 가지는 빈도수의 표준편차/각 corpus를 구성하는 token의 평균

국가영어능력평가시험(National English Ability Test, NEAT)은 2007 ~ 2008년 사전 연구에서는 5개의 개별 검사지로 개발이 되었으나 현재는 대학수학능력시험의 외국어(영어) 영역을 대체하기 위해 성인용 1개 검사지를 제외한 2개 검사지로 개발 방향이 전환되었다. 이에 따라 당초 각 등급별 1,000 단어씩을 배정하여 총 5,000 단어군의 어휘 목록을 개발하였으나 이중 2,000 단어는 3급 검사지에 적용이 되고 있고 2급의 경우에는 1,000 단어가 추가된 3,000 단어 수준으로 검사지가 개발되고 있다.

본 연구에서는 기존에 개발된 대표적인 어휘 목록을 기준으로 3,000단어군 목록이 2010년 현재 영어과 교과서 검정을 통과한 초·중·고 영어 교과서(활동책 포함)와 해외 영어 코퍼스를 분석하였을 때 어느 정도의 텍스트 포괄 범위를 보이는지 비교·분석하였다. 비교·분석에 사용된 어휘 목록은 국가영어능력평가시험의 어휘 목록인 NEAT3000, Nation(2004)이 BNC에서 추출한 14,000개의 단어군 가운데 상위 3000개를 발췌한 소위 BNC3000, 그리고 Oxford 출판사에서 개발한 OXFORD3000이다. 다음의

표는 각 어휘 목록 개발 시 적용된 어휘 선정 기준 및 적용 순서이다.

〈표 3〉 대표적인 어휘 목록 제작 시 어휘 선정 기준 및 적용 순서

어휘 목록	어휘 선정 기준 및 적용 순서
NEAT3000	사용범위 → 사용분포 → 빈도수
BNC3000	사용범위 → 빈도수 → 사용분포
OXFORD3000	빈도수 → 사용범위 → 친숙도

대부분의 어휘 목록은 빈도수를 첫 번째 어휘 선정 기준으로 적용한 다. 그것은 빈도수가 가지는 노출의 기회와 유용성에 대한 전통적인 믿음에서 비롯되었다. 하지만 NEAT3000의 경우에는 위의 표에서 보듯 기존의 어휘 선정 기준의 적용 방식과는 차이가 있다는 것을 알 수 있다. 일반적으로 원어민의 언어자료를 바탕으로 텍스트 포괄 범위(text coverage)를 분석해 보면 최대의 텍스트 포괄 범위를 확보하기 위해서는 빈도수를 최우선 기준으로 배정해야 한다는 것이 정설이다. 하지만 국가영어능력평가시험의 경우 학생용으로 개발되는 시험이고 최고 등급을 원어민 기준이 아니라 국가교육과정을 성취수준에 맞추었기 때문에 보다 일상적이고 일반적인 어휘들의 선별에 초점을 두었다. 이러한 이유로 얼마나 다양한 텍스트에서 사용되는지를 측정하는 사용범위와 얼마나 편차 없이 고르게 사용되는지를 측정하는 사용분포를 빈도수보다 우선하는 기준으로 적용한 것이다. 그리고 사용범위나 사용분포를 우선했을 경우 총 텍스트의 포괄 범위에서는 다소 떨어지지만 보다 일상적이고 일반적인 어휘를 추출할 수 있다는 가정을 증명하기 위해 소규모의 실험 분석으로 진행하였다. 먼저 ACE(호주 문어 코퍼스), Frown(미국 문어 코퍼스), FLOB(영국 문어 코퍼스), BNC7(BNC spoken section, 영국 구어 코퍼스), WWC(뉴질랜드 문어 코퍼스), WSC(뉴질랜드 구어 코퍼스)로 구성된 12,308,634(약 1200백만)단어의 대규모 코퍼스를 소스로 하여 1) 사용범위→빈도수 순과 2) 빈도수→사용범위 순으로 나누어 3,000단어로 구성된 단어군 목록을 제작하였다. 사용분포의 경우 빈도수와의 상관관계가 높아 실험에서는 제외하였다. 분석의 편의를 위해

알파벳, 단위, 고유명사 등의 예외 항목을 두지 않고 어휘 선정 기준에 의해 산출된 수치에 따라 상위 3,000 단어만을 선별한 단어군 목록을 제작에 활용하였다. 제작된 두 개의 단어군 목록을 Range Program에 각각 탑재하여 검정된 교과서를 포함한 다양한 코퍼스를 분석하였다.

〈표 4〉 어휘 선정 기준 적용 순서에 따른 텍스트 포괄 범위 비교

	일치도	Frown	FLOB	WSC	교과서
[A형] 사용범위(Range) ↓ 빈도수(Frequency)	90% (type 기준)	87.45%	88.63%	91.24%	89.40%
[B형] 빈도수(Frequency) ↓ 사용범위(Range)		88.73%	89.53%	92.19%	89.54%

* 텍스트 포괄 범위는 출현형(token)을 기준으로 산정됨

[A형]의 경우 사용범위를 빈도수에 우선 시하여 적용하였고 [B형]은 빈도수를 사용범위에 우선 시하여 적용하였다. [A형]의 3,000 단어군은 18,082개의 낱말유형(type)을 포함하고 [B형]은 18,077개의 낱말유형(type)을 포함하였다. [A형]과 [B형]은 낱말유형을 기준으로 약 90%가 일치했고 10%는 각각 다른 낱말유형을 포함하고 있었다. 단어군으로 기준으로하면 361개의 단어군에서 서로 차이를 보여 약 12%의 차이를 보였다. 3,000개의 단어군 중 상위 185개는 어휘 선정 기준의 적용순서가 달랐음에도 불구하고 순위까지 일치하여 사용 빈도가 상위 185개 단어에 매우 집중되어 있다는 것을 알 수 있다. 텍스트 포괄 범위를 살펴보면 모든 분석 자료에서 [B형]이 더 높은 수치를 보여 빈도수가 텍스트 포괄 범위에서는 보다 중요한 기준임을 확인할 수 있다. 단어군 목록 제작에 사용된 소스가 영국 구어 자료와 뉴질랜드 자료의 비중이 컸던 이유로 영국 문어 코퍼스인 FLOB에서의 텍스트 포괄 범위가 미국 문어 코퍼스인 Frown 보다 높았고 뉴질랜드 구어 코퍼스인 WSC에서는 90%가 넘는 수치를 보였다. 검정된

교과서의 합본에서는 [B형]이 약간 더 높은 수치를 보였지만 [A형]과의 차이는 0.14%로 거의 차이를 보이지 않아 교과서는 원어민 자료와는 달리 상대적으로 통제된 어휘를 사용하고 있음을 간접적으로 확인할 수 있었다. 또한 빈도수를 우선하는 것보다 사용범위를 우선하여 선정하는 것이 일상적이고 일반적인 어휘 추출을 하는 데에는 보다 효과적이라는 가정을 확인하기 위해서 Coxhead(1998, 2000)의 AWL로 분석한 결과, [A형]에만 포함된 361개의 단어군 중 21개만이 학문적인 단어인 반면 [B형]은 361개 중 100개의 단어가 학문적인 단어인 것으로 나타났다. 즉, 본 연구에서 가정한 대로 텍스트 포괄범위는 다소 낮아지더라도 일반적이고 일상적인 어휘들을 우선적으로 추출하기 위해서는 사용범위가 빈도수에 우선하여 적용되어야 한다는 점이 어느 정도 입증되었다. 또한 NEAT3000의 개발 취지와 어휘 선정 기준 적용이 적절했다는 것 또한 이를 통해 확인할 수 있었다.

어휘 선정 기준과 관련하여 또 하나의 이슈는 친숙도(Familiarity)의 적용 여부이다. 위에서 제시한 3개의 어휘 목록 중 OXFORD3000의 경우는 친숙도를 어휘 선정 기준의 하나로 적용하고 있다는 점이 다른 2개의 어휘 목록과 차별화된 부분이다. 하지만 친숙도는 다른 어휘 선정 기준인 빈도수, 사용범위, 사용분포와는 달리 객관적인 수치로 나타내는데 어려움이 있고 주관성이 개입된다는 점에서 문제점이 제기되곤 한다(신통광과 주현우, 2008). 하지만 영어를 배우고 사용하는 환경에 따라 유용성과 필요성 면에서 원어민의 언어 환경과는 다르게 사용되는 어휘들이 분명 존재하는 것은 사실이다. 예를 들어 ‘classroom’, ‘textbook’, ‘livingroom’, ‘watermelon’, ‘cute’, ‘song’ 등 교실영어나 우리나라 환경에서 친숙한 과일이나 음식명 또는 한국어에서 자주 사용되는 단어들로서 영어로 번역되었을 때 대응되는 영어 단어들이다. 이러한 단어들은 객관적인 수치 면에서는 상위권에 포함되지 않지만 한국 학생들에게는 친숙하여 교재나 기본 어휘에 포함되거나 포함되어야 한다고 요구되는 단어들이기도 하다. 그러나 이미 언급한 바와 같이 소수의 주관적 판단에 따라 기본 어휘에 이러한 단어들을 포함시키는 것은 신뢰성에 문제가 있어 친숙도를 측정하는 보다 타당하고

신뢰성 있는 방법이 필요하다. OXFORD3000의 제작 매뉴얼에서는 영어 교사들 대상으로 친숙도를 측정하였다고만 언급을 하고 있어 어떠한 방식으로 친숙도를 측정했는지는 정확히 알 수 없지만 친숙도를 측정하는 방법은 기존에도 어느 정도 연구가 된 바 있다. Zimmerman(1997)은 어휘지식을 측정하기 위해 다음과 같은 4단계의 척도를 이용하였다.

〈표 5〉 Zimmerman(1997)의 어휘 지식 측정 척도(Knowledge Scale)

등급	친숙도 정도
1 단계	단어를 알지 못한다.
2 단계	단어를 전에 본 적은 있으나 의미를 확실히 알지는 못한다.
3 단계	단어를 문장에서 듣거나 보고 의미를 이해할 수는 있으나 실제 말하기, 쓰기에서 사용하지는 못한다.
4 단계	단어를 문장에 포함하여 사용할 수 있다.

Zimmerman의 모델에서는 이해어휘 지식, 표현어휘 지식과 더불어 문장에 포함하여 사용할 수 있는 능력으로 단어 간의 관계(연어)나 파생과 굴절과 같은 기본형을 넘어선 다양한 낱말유형까지도 적절하게 실제 문맥에서 사용이 가능한 지를 최종 단계에서 묶어서 측정하고자 하였다.

Waring(2000)은 측정요인을 이해어휘 지식과 표현어휘 지식으로 단순화하고 척도를 한 단계 더 늘려 5개의 척도로 구성된 어휘 지식 측정 모델을 제안하였다.

〈표 6〉 Waring(2000)의 어휘 지식 측정 척도

등급	친숙도 정도
1 단계	단어를 알지 못한다.
2 단계	단어의 의미를 이해한다고 생각하지만 실제 사용하지는 못한다.
3 단계	단어의 의미를 이해하지만 실제 사용하지는 못한다.
4 단계	단어의 의미를 이해한다고 생각하고 실제 사용할 수도 있다.
5 단계	단어의 의미를 이해하고 실제 사용할 수도 있다.

Zimmerman(1997)이나 Waring(2000)의 경우에는 모델을 처음 고안할 때 문자로 제시하여 측정하는 방법을 전제로 하고 있고 현재 대부분의 어휘 평가가 그렇듯 한 단어의 다양한 의미를 측정하지는 못하고 있으며 그와 마찬가지로 의미를 측정하는 데 있어서도 하나의 낱말유형의 의미를 측정하는 데 머무르고 있다. 하지만 영어가 모국어나 공용어가 아닌 우리나라와 같은 언어 환경에서는 같은 수준의 이해어휘 지식을 갖고 있는 경우라도 듣고 단어의 의미를 파악할 수 있지만 문자로 제시되었을 때에는 그 단어의 철자를 몰라 의미를 파악하지 못하는 경우도 있고 그와는 반대로 단어를 본적이 있어 보고 의미를 파악할 수 있으면서도 들어 본적은 없어 음성으로 제시되었을 때는 의미 파악이 불가능 경우도 있을 수 있다. 따라서 듣고 이해하는 능력과 보고 이해하는 능력이 같은 능력으로 보고 묶어서 측정한다고 한다면 음성과 문자로 함께 제시해야 할 것이다. Zimmerman(1997)이나 Waring(2000)이 연구를 진행할 당시에는 두 가지 형태의 제시방법이 불가능했을 수도 있지만 이제는 대부분의 온라인 무료 사전에서도 Text-To-Speech(TTS) 프로그램을 이용하여 음성정보와 문자정보를 동시에 제공 받을 수도 있기 때문에 제시 방법에도 세심한 고려가 필요할 것으로 보인다.

본 연구에서는 단어의 음성정보와 문자정보를 동시에 제공하는 것을 전제로 하여 어휘의 측정 요인도 세분화 그리고 명시화하여 제시하였다. 측정 요인으로는 먼저 이해어휘 지식과 표현어휘 지식으로 구분하여 반영하였고 실제 사용 능력도 의미의 다양성과 굴절과 파생변이형의 적절한 사용 능력까지 명시적으로 제시하였다. 또한 척도의 수도 6개로 늘려 변별력을 높였다. 다음의 표는 위에서 언급한 기존의 어휘 친숙도 측정 척도의 문제점을 보완하여 고안된 새로운 어휘 친숙도 측정 척도이다.

〈표 7〉 어휘 친숙도 측정 척도

등급	친숙도 정도
1 단계	단어를 듣거나 본적이 없어 의미 또한 알지 못한다.
2 단계	단어를 듣거나 본 적은 있으나 의미를 알지 못한다.

등급	친숙도 정도
3 단계	단어를 듣거나 보고 의미를 파악할 수 있지만 말하거나 쓰기에서 사용할 수는 없다.
4 단계	단어를 듣거나 보고 의미를 파악하여 실제 쓰임에 한정된 형태로는 말하기에서 사용할 수는 있지만 쓰기에서 사용할 없다.
5 단계	단어를 듣거나 보고 의미를 파악하여 실제 쓰임에 한정된 형태로 말하기, 쓰기에서 사용할 수 있다.
6 단계	단어를 듣거나 보고 다양한 의미를 파악하여 실제 쓰임에 적절한 형태로 말하기, 쓰기에서 쓸 수 있다.

어휘 친숙도 측정 척도를 사용하면 어휘들의 친숙도를 수치화하여 어휘 선정 기준으로 반영할 수 있을 것으로 기대된다. 하지만 친숙도가 어휘 선정 결과에 어떠한 영향을 미치는지는 논란이 많은 이슈이니만큼 별도의 대규모 연구를 통해 입증할 필요가 있다. 본 연구에서 친숙도가 미치는 영향을 직접 측정할 수는 없지만 친숙도가 반영된 OXFORD3000을 다른 어휘 목록과 비교·분석하여 대략적인 추정을 시도하였다.

다음은 NEAT3000, BNC3000, OXFORD3000으로 약 600만 단어로 구성된 초·중·고등학교 교과서 합본(활동책 포함)을 분석하여 각각의 텍스트 포괄 범위를 제시한 것이다.

〈표 8〉 어휘 목록별 초·중·고 교과서 합본 텍스트 포괄 범위

WORD LIST	출현형/%	낱말유형/%	단어군
NEAT3000	4,994,865/86.00	10,112/32.76	2,897
BNC3000	5,294,379/89.37	10,113/32.76	2,892
OXFORD3000	5,060,675/85.42	9,163/29.68	2,519
Proper etc(공통)	361,741/6.11	3,214/10.41	3,214

* 초·중·고등학교 교과서 합본은 5,924,461단어로 구성, 약 6백만 단어

초·중·고등학교 교과서 합본을 분석한 결과는 NEAT3000은 텍스트 포괄 범위에서 86%의 수치를 보여주었고, BNC3000은 예상대로 NEAT

3000 보다 약간 높은 89%의 수치를 보여주었다. OXFORD3000은 85%로 3개의 어휘 목록 가운데 가장 낮은 수치를 나타냈다. 고유명사가 공통적으로 6%를 차지하고 있는 것으로 감안하면 3개 어휘 목록의 실질적인 텍스트 포괄 범위는 모두 90%를 넘어서고 있기 때문에 공통적으로 높은 수준의 유용성을 보이고 있다고 판단할 수 있다. 또한 단어군 기준으로 볼 때는 NEAT3000의 3,000단어군 가운데 검정 교과서에서는 2,892개의 단어군이 사용되고 있어 BNC3000 보다도 약간 높은 수치를 보이고 있으며 고유명사의 비율이 일반 코퍼스보다 높은 것은 교육용으로 제작된 교과서의 특징으로 다양한 인명과 지명 등이 포함되었기 때문이다.

교과서 자료가 일반적으로 접하게 되는 원어민의 자료와는 차이가 있기 때문에 실제 원어민 코퍼스를 바탕으로 3개의 어휘 목록이 차지하는 텍스트 포괄 범위를 추가로 분석해 보았다. 다음은 100만단어로 구성된 미국의 문어 코퍼스인 Frown 코퍼스를 분석한 결과이다.

〈표 9〉 어휘 목록별 Frown 텍스트 포괄 범위

WORD LIST	출현형/%	낱말유형/%	단어군
NEAT3000	892,615/87.13	12,263/28.2	2,992
BNC3000	893,130/87.19	10,113/32.76	2,892
OXFORD3000	868,705/84.80	10,930/25.13	2,521
Proper etc(공통)	25,339/2.47	3,873/8.91	3,873

* Frown은 총 1,024,395단어로 구성된 미국 문어 코퍼스

NEAT3000의 텍스트 포괄범위는 약 87%이고 BNC3000도 87%로 매우 유사했으나 OXFORD3000의 경우는 약 85%로 교과서 분석과 마찬가지로 상대적으로 다소 낮게 나타났다. Frown 코퍼스에서 사용된 단어군의 종류는 NEAT3000의 경우 총 3,000개의 단어군 중 2,992개가 사용되어 교과서 분석 때보다도 높은 수치를 보였고 BNC3000 보다도 100개, OXFORD 3000보다도 471개의 단어군이 더 다양하게 사용되어 보다 보편적인 어휘로 구성되어 있다는 것을 확인할 수 있었다. 다음은 100만단어로 구성된

영국의 문어 코퍼스인 FLOB 코퍼스를 분석한 결과이다.

〈표 10〉 어휘 목록별 FLOB 텍스트 포괄 범위

WORD LIST	출현형/%	낱말유형/%	단어군
NEAT3000	901,069/88.22	12,410/28.54	2,999
BNC3000	903,262/88.44	12,381/28.48	2,993
OXFORD3000	877,638/85.93	11,052/25.42	2,518
Proper etc(공통)	26,170/2.56	4,453/10.24	4,453

* FLOB은 총 1,021,310단어로 구성된 영국 문어 코퍼스

영국 FLOB 코퍼스의 분석결과를 보면 NEAT3000과 BNC3000의 경우 텍스트 포괄 범위의 수치가 더 높아지는 것을 볼 수 있고 그 수치도 서로 매우 유사하는 것을 알 수 있다. 그 이유는 어휘 목록 제작에 사용된 코퍼스에 BNC 구어 자료가 공통적으로 많은 부분을 차지하고 있기 때문이다. 하지만 모든 분석결과에서 OXFORD3000은 가장 낮은 수치를 보이고 있다. 이는 어휘 목록 제작 시 어휘 선정 기준에서 차이가 있었기 때문이라고 추정된다. OXFORD3000 또한 어휘 목록의 제작 시 Oxford 출판사에서 자체 수집한 Oxford Corpus Collection이라는 영국 코퍼스를 사용했기 때문에 큰 차이를 보이는 점은 오직 영국의 교사들이 가지는 단어에 대한 친숙도 수치를 반영한 것 외에는 그 차이를 찾아 볼 수 없기 때문이다. 물론 영국 교사들이 친숙하게 생각하는 단어들과 한국 교사나 학생들이 생각하는 친숙한 단어에는 분명 차이가 있을 것이다. 그럼에도 불구하고 위의 분석 자료를 본다면 친숙도가 가지는 주관성이 텍스트 포괄 범위와 같은 유용성의 수치를 떨어뜨리는 요인일 가능성은 여전히 높다고 판단된다. 다음은 위의 3개 어휘 목록 간의 일치도를 비교한 것이다.

〈표 11〉 어휘 목록별 일치도

WORD LIST	낱말유형/%	단어군
NEAT3000 ⇔ OXFORD3000	12,850/82.74	2,245
BNC3000 ⇔ NEAT3000	14,891/90.47	2,665
OXFORD3000 ⇔ BNC3000	12,737/77.38	2,290

먼저 NEAT3000과 BNC3000은 낱말유형 기준으로 90%의 일치도를 보였고 단어군 기준으로는 3,000 단어 가운데 2,665개가 일치하였다. 일부 차이는 NEAT3000이 알파벳, 로마자, 단위, 나라이름, 국적, 언어이름과 같은 예외 항목을 두어 3,000 단어에서 포함하지 않고 구성한데 반해 BNC3000은 인명, 지명 등의 고유명사와 로마자 정도만 예외 항목으로 제외한데 기인했다고 판단된다. 또 다른 이유는 어휘 선정 기준의 적용순서와 두 어휘 목록 제작을 위해 사용한 코퍼스가 둘 다 영국 영어이기만 하지만 BNC3000이 영국 구어 코퍼스인 BNC 구어 자료만을 자료로 활용한 반면 NEAT3,000은 다양한 영어의 코퍼스를 활용하였고 어휘 선정 기준에서도 대중성과 보편성을 더 고려했다는 데 있다. 결과적으로 NEAT3000은 BNC3000과 비교하여 영국적인 색채를 띠는 어휘를 배제하는 데 성공하여 보편성면에서는 BNC3000 보다 더 뛰어난 것으로 보인다. NEAT3000과 OXFORD3000은 약 83%의 일치도를 보였고 OXFORD3000과 BNC3000은 77% 정도의 일치도만 보여 유사성에서 가장 낮게 나타났다.

IV. 어휘 선정 기준의 적용방안

국가영어능력평가시험의 어휘 목록인 NEAT3000은 국내에서는 최초로 코퍼스를 자료로 하여 사용범위, 사용분포, 빈도수와 같은 객관화할 수 있는 어휘 선정 기준으로 적용하여 제작한 사례이다. NEAT3000의 제작 절차는 한국어 어휘 목록 제작에도 충분한 시사점을 줄 수 있는 새로운

시도였다는 점에서 그 개발 과정을 살펴볼 필요가 있다.

1. 어휘 추출을 위한 분석 자료

국가영어능력평가시험은 대학수학능력시험 외국어(영어)영역 대체는 물론 국제적인 통용성을 염두에 두고 개발되고 있기 때문에 어휘 추출을 위한 분석 자료에도 가능한 다양한 영어 자료를 사용하였다. 먼저 상대적으로 최근에 개발된 코퍼스 위주로 구성하였고 미국, 영국, 뉴질랜드의 영어를 포함하였으며 구어와 문어 코퍼스를 모두 포함하되 구어를 주로 포함하였다.

〈표 12〉 국가수준 영어능력인증시험 어휘 목록 개발을 위해 사용된 코퍼스

Freiburg-Brown (Frown)	Freiburg-LOB (FLOB)	Wellington Written Corpus(WWC)	Wellington Spoken Corpus(WSC)	Seven Spoken Sections of the British National Corpus (BNC)
1991-1996	1991-1996	1986-1992	1986-1992	1991-1994
미국 문어	영국 문어	뉴질랜드 문어	뉴질랜드 구어	영어 구어
1,000,000 단어	1,000,000 단어	1,000,000 단어	1,000,000 단어	10,000,000 단어

위의 표처럼 BNC의 구어의 7개영역을 하나의 개별 코퍼스로 본다고 한다면 총 11개의 corpus, 즉 14,000,000개의 어휘로 구성된 대규모 코퍼스(mega-corpus)를 어휘 추출용 자료로 사용하였다. 코퍼스에 대한 연구가 대부분은 영국권에서 이루어지고 코퍼스도 영국 영어가 대부분이라 미국 영어의 비율이 상대적으로 낮았으며 구어 자료의 경우에는 영국 영어만 포함되었다는 제한점을 가지기는 하지만 교육용 어휘 목록 개발을 위해서 대표성 있는 어휘를 추출하기 위해서는 자료에만 의존하는 것이 아니라 선정 기준을 다시 적용하기 때문에 이러한 문제를 대부분 해소할 수 있었다. 또한 어휘 선정 기준에는 위에서 소개한 바와 같이 사용범위, 사용분

포, 빈도수 순으로 세 가지 기준을 적용하였다.

2. 국가영어능력평가시험에서 사용된 단어군의 정의 및 범위

Bauer와 Nation(1993)은 1,000,000 단어로 이뤄진 Lancaster-Oslo-Bergen (LOB) 코퍼스를 분석하여 파생어의 빈도를 조사하고, 그 결과를 다음과 같이 8가지의 굴절어와 파생어 분류 기준에 따라 7 단계로 분류했다.

- 빈도(Frequency)
- 생산성(Productivity)
- 예측성(Predictability)
- 문자로 된 기본형의 규칙성(Regularity of the written form of the base)
- 구어로 된 기본형의 규칙성(Regularity of the spoken form of the base)
- 접사 철자의 규칙성(Regularity of the spelling of the affix)
- 접사 발음의 규칙성(Regularity of the spoken form of the affix)
- 기능의 규칙성(Regularity of function)

위의 8가지 기준을 통해 분석 결과로 얻은 각 단계별 굴절형과 파생 접사는 다음과 같다.

〈표 13〉 굴절 및 파생 접사의 분류표(Bauer와 Nation, 1993)

1 단계	book, books와 같이 기본형이 같더라도 낱말들의 형태가 다르면 모두 다른 낱말로 간주하는 단계이다.
2 단계	모든 굴절형을 기본형과 함께 묶어서 동일한 낱말로 간주한다. 여기에 해당하는 굴절형은 복수형, 3인칭 단수 현재형, 과거형, 과거분사형, -ing형, 비교급, 최상급과 소유격이 있다.
3 단계	-빈도가 높고, 생산적인 편이며, 종종 어근에 변화를 주는 접사 위에서 제시한 8 가지 기준이 모두 적용되는 단계이다. 철자상 -y가 -i로 바뀌는 것 외에는 규칙적이다. 이에 해당하는 접사로는 -able, -er, -ish, -less, -ly, -ness, -th, -y, un-, non-이 있다.

4 단계	-빈도가 낮고, 생산적이지 않으며 어근에 변화가 적은 접사 -al, -ation, -ess, -ful, -ism, -ist, -ity, -ize, -ment, -ous, in-의 11개 접사를 포함한다.
5 단계	-빈도는 높으나 철자법이 규칙적이지 않은 접사 -age (leakage), -al (arrival), -ally (idiotically), -an (American), -ance (clearance), -ant (consultant), -ary (revolutionary), -atory (confirmatory), -dom (kingdom: officialdom), -eer (black marketeer), -en (wooden), -en (widen), -ence (emergence), -ent (absorbent), -ery (bakery: trickery), -ese (Japanese: officialese), -esque (picturesque), -ette (usherette: roomette), -hood (childhood), -i (Israeli), -ian (phonetician: Johnsonian), -ite (Paisleyite: also chemical meaning), -let (coverlet), -ling (duckling), -ly (leisurely), -most (topmost), -ory (contradictory), -ship (studentship), -ward (homeward), -ways (crossways), -wise (endwise: discussion-wise), anti- (anti-inflation), ante- (anteroom), arch- (archbishop), bi- (biplane), circum- (circumnavigate), counter- (counter-attack), en- (encage: enslave), ex- (ex-president), fore- (forename), hyper- (hyperactive), inter- (inter-African, interweave), mid- (mid-week), mis- (misfit), neo- (neo-colonialism), post- (post-date), pro- (pro-British), semi- (semi-automatic), sub- (subclassify: subterranean), un- (untie: unburden의 50개 접사를 포함한다.
6 단계	-빈도는 높으나 철자법이 규칙적이지 않은 접사 -able, -ee, -ic, -ify, -ion, -ist, ition, -ive, -th, -y, pre-, re-의 12개 접사를 포함한다.
7 단계	-고전어 어근 및 접사(classical roots and affixes) -ar(circular), -ate(compassionate, captivate, electorate), -et (packet, casket), -some(troublesome), -ure (departure, exposure), ab-, ad-, com-, de-, dis-, ex-('out'), in-('in'), ob-, per-, pro-('in front of'), trans-

이와 같은 절차에 따라 분석한 접사의 단계별 예시는 다음과 같다. 1 단계는 모든 단어를 형태 별로 모두 다르게 보는 단계이므로 제외했다.

〈표 14〉 Bauer와 Nation(1993)과 Cobb(2007)의 단어군 범위 비교

Level	Bauer와 Nation(1993)	Cobb(2007)
2	develop develops developed developing	develop develops developed developing
3	developable undevelopable developer(s) underdeveloped	developer(s) undeveloped underdeveloped underdevelopment
4	development(s) developmental developmentally	development(s) developmental developmentally
5	developmentwise semideveloped anti-development	
6	redevelop predevelopment	redevelop redevelops redeveloped redevelopment redevelopments

Bauer와 Nation(1993)의 기준과 Cobb(2007)의 기준을 ‘develop’이란 단어를 기준으로 비교해 보면 전반적으로 일치하고 있으나 5단계와 6단계에서 차이를 보이고 있다. 특히나 5단계에서는 Cobb의 단어군 범위 안에는 Bauer와 Nation이 포함하고 있는 단어를 전혀 포함하고 있지 않다는 것을 볼 수 있다. 이것은 Cobb(2007)가 기본적으로는 Bauer와 Nation의 기준을 따랐지만 현장 교사들의 의견을 수렴하여 명백히 같은 단어군에 포함되는 단어들이라고 판단되는 것은 포함시켰고 지나치게 확장되었다고 판단되거나 빈도수가 낮은 단어들은 그 범위에서 탄력적으로 제외시켰기 때문이다. 위에서 제시한 어휘 목록의 단어군 범위는 Cobb의 기준이 보다 현실적이라고 판단되어 Cobb의 기준에 따라 그 범위를 설정하였다. 그리고 두 기준 모두 미국식 철자법과 영국식 철자 모두 포괄하고 있어 두 가지의

철자법(예, realize, realise)을 모두 반영하였다. Cobb는 www.lex tutor.ca에서 Familizer라는 프로그램을 무료로 제공하고 있으며 이 프로그램을 바탕으로 단어군 목록이 제작되었다. 물론 Cobb가 제공하는 Familizer¹⁾가 프로토타입(prototype)이라 많은 오류를 포함하고 있었지만 수차례의 검토와 수정을 통하여 단어군 목록을 완성할 수 있었다.

V. 결론 및 제언

본 연구에서는 영어권에서 사용하는 영어 기본 어휘를 선정하는 객관적인 기준을 소개하고 그러한 기준들을 실제로 적용한 국가영어능력평가 시험의 개발 사례 그리고 어휘 선정의 기준과 기준의 적용 순서가 어휘 목록에 미치는 영향 등을 살펴보았다. 한국어는 교착어(glued language)로 조사가 단어에 붙어있는 형태로 모두 개별 단어로 떨어져 있는 영어와는 달리 어휘 선정 시 사용되는 코퍼스 분석이 용이하지 않다. 그러한 이유로 대부분의 한국어 어휘 연구가 질적 연구에 집중되고 있다. 하지만 심층적으로 어휘습득 및 어휘의 사용을 살펴보는 것도 있지만 대표성이 있는 데이터 수집을 통해 객관화할 수는 시사점을 이끌어내는 것도 또한 중요하다. 아무리 문법학자들이 특정 표현이 문법적이고 원어민은 이렇게 어휘를 사용한다고 규정한다고 할지라도 실생활에서 원어민 그렇게 예측한대로 어휘를 구사한다고 장담할 수는 없다. 언어사용의 실제성(authenticity)이 강조 되는 외국어 교육의 추세를 비추어 본다면 한국어 교육에서도 이제 질적 연구에 더하여 양적 연구에도 관심을 쏟아야 할 시점이 아닌가 생각해 본다. 한국어 어휘 연구에서의 양적 연구를 위해서는 먼저 선행되어야 두 가지 조건이 있다. 첫째, 위에서 살펴본 <표 14>와 마찬가지로

1) 현재는 Lexipia(www.cafe.daum.net/sdhera)에서 제공되는 Familizer를 통해 보다 쉽고 간편하게 단어군 목록을 제작할 수 있다.

한국어에서도 기본형에 굴절·파생 변이형을 모두 명시된 체계화된 단어군 목록이 개발되어야 한다. 즉, 가능한 큰 규모의 단어군 목록 은행을 개발할 필요가 있다. 둘째, 한국어 단어군 목록이 개발되면 단어군 목록을 탑재하여 한국어 코퍼스에서 코퍼스를 구성하는 단어들의 빈도수는 물론 사용범위를 측정할 수 있는 프로그램이 개발되어야 한다. 이러한 두 가지 조건이 전제되면 한국어 교육에서도 수준별 어휘 목록 개발이 가능하며 다시 이러한 수준별 어휘 목록을 분석 프로그램을 탑재하면 어휘수가 체계적으로 통제된 수준별 교재 개발도 가능하리라 생각된다. 해외의 한국어 능력시험의 응시자가 매년 증가하고 있고 국내에서도 아동들을 위한 수준별 교재 시장이 서서히 기지개를 펴고 있다는 점을 생각해 본다면 한국어 교육에서 본 연구와 같은 양적 연구에 대한 필요성은 더 절실히 보인다.*

* 본 논문은 2011. 2. 18. 투고되었으며, 2011. 3. 5. 심사가 시작되어 2011. 3. 30. 심사가 종료되었음.

▣ 참고문헌

- 신동광, 주현우(2008), “영어과 교과서 김정용 어휘 목록 개발 : 어휘 검색 프로그램 및 기본 어휘 목록에 대한 고찰”, 『멀티미디어언어교육』 11(3), 93-111.
- Bauer, L., & Nation, I. S. P.(1993), Families. *International Journal of Lexicography*, 6, 253-79.
- Chung, T. M.(2003), A corpus comparison approach for terminology extraction. *Terminology*, 9(2), 221-245.
- Coxhead, A.(1998), *The development and evaluation of an academic word list*. Unpublished master's thesis. Victoria University of Wellington, New Zealand.
- Coxhead, A.(2000), A new academic word list. *TESOL Quarterly*, 34(2), 213-238.
- Cobb, T.(2007), *Familizer : A program for making word families* [software]. Available at <http://www.lexutor.ca/familizer/>
- Hu, M. H-C., & Nation, I. S. P. (2000), Unknown vocabulary density and reading comprehension. *Reading in a Foreign Language*, 13(1), 403-430.
- Johns, T.(1986), Micro-concord : A language learner's research tool. *System*, 14(2), 151-162.
- Johns, T.(1988), Whence and whither classroom concordancing? In T. Bongaerts, P. de Haan, S. Lobbe & H. Wekker (Eds.), *Computer applications in language learning* (pp.9-27), London : Foris.
- Johns, T.(1991), Should you be persuaded - two examples of data-driven learning. In T. Johns & P. King (Eds.), *Classroom concordancing : English Language Research Journal*, 4(pp.1-13), University of Birmingham : Centre for English Language Studies.
- Laufer, B.(1992), How much lexis is necessary for reading comprehension? In Arnaud, P. J. L. and Bejlint, H. (Eds.), *Vocabulary and Applied Linguistics*(pp.126-132), London : Macmillan.
- Laufer, B., Elder, C., Hill, K., & Congdon, P.(2004), Size and strength : Do we need both to measure vocabulary knowledge? *Language Testing*, 21, 202-226.
- Leech, G., Rayson, P., & Wilson, A.(2001), *Word frequencies in written and spoken English : Based on the British National Corpus*. London : Longman.
- Meara, P.(2005), Lexical frequency profiles : A Monte Carlo analysis. *Applied Linguistics*, 26(1), 32-47.
- Nation, I. S. P.(1990), *Teaching and learning vocabulary*. New York : Newbury House.
- Nation, I. S. P.(2001), *Learning vocabulary in another language*. Cambridge : Cambridge

University Press.

- Nation, I. S. P.(2004), A study of the most frequent word families in the British National Corpus. In P. Bogaards & B. Laufer (Eds.), *Vocabulary in a second language : Selection, acquisition, and testing*(pp.3-13), Amsterdam : John Benjamins.
- Nation, I. S. P.(2006) How large a vocabulary is needed for reading and listening? *Canadian Modern Language Review*, 63(1), 59-82.
- Read, J.(2004), Research in teaching vocabulary. *Annual Review of Applied Linguistics*, 24, 146-161.
- Richards, J.(1974), Word lists : Problems and prospects. *RELC Journal*, 5(2), 69-84.
- Shin, D.(2009), *Familizer14 : A program for making word families* [software]. Available at <http://cafe.daum.net/sdhera>
- Ward, J.(1999), How large a vocabulary do EAP engineering students need? *Reading in a Foreign Language*, 12, 309-323.
- Waring, R.(2000), The state rating task : An alternative method of assessing receptive and productive vocabulary. *Immaculata{occasional papers of Notre Dame Seish in University, Okayama}*, 35(1), 125-154.
- West, M.(1953), *A general service list of English words*. London : Longman, Green & Co.
- Zimmerman, C.(1997), Do reading and interactive vocabulary instruction make a difference? : An empirical study. *TESOL Quarterly*, 31(1), 121-140.

<초록>

기본 어휘의 선정 기준

—영어 어휘를 중심으로—

신동광

본 연구는 영어권에선 사용되는 기본 어휘 선정 기준을 소개하고 이를 바탕으로 한국어 기본 어휘 선정에 적용할 수 있는 방안을 모색해 보았다. 영어권에서는 효과적인 어휘교육을 위해 대규모 말뭉치(corpus, 코퍼스)로 부터 빈도수, 사용범위, 사용분포 등의 객관적인 어휘 선정 기준을 적용하여 가장 유용한 어휘를 추출하여 어휘 목록을 제작하고 있을 뿐만 아니라 하나의 단어 기본형이 포괄하는 굴절과 파생형을 모두 명기한 단어군(word family) 목록 은행을 바탕으로 등급화된 어휘 목록을 개발함으로써 어휘수를 통제할 수준별 교재를 체계적으로 제작하고 있다. 또한 본 연구에서는 어휘 선정 기준의 적용 절차가 어휘 목록의 개발결과에 미치는 영향을 분석하기 위해 소규모의 실험연구를 시행하였고 그 결과를 근거로 교육 목적에 따른 효율적인 어휘 선정 기준의 적용절차도 제안하였다. 끝으로 국내에서는 처음으로 이러한 어휘 선정 기준을 적용하여 개발된 국가영어능력평가시험의 어휘 목록 개발 절차를 소개하고 그 과정과 결과를 일반화하여 적용할 수 있는 모델과 시사점을 제시하였다.

【핵심어】 코퍼스, 말뭉치, 빈도수, 사용범위, 사용분포, 단어군 목록, 텍스트 포괄 범위, 국가영어능력평가시험 어휘 목록

<Abstract>

**Criteria for Selecting Basic Words :
Focusing on English Vocabulary**

Shin, Dong-kwang

The present study aims at introducing some valid and reliable criteria for selecting basic English words such as frequency, range and dispersion. By carrying out a small pilot study the effects of the application order of those vocabulary selection criteria on developing an English word family list was examined. In addition, the program Familizer which could be used for making word lists was also introduced. It is the program that was used for developing a basic English word list for the National English Ability Test (NEAT) which would replace the English language section of the College Scholastic Ability Test. The program could contribute to popularizing the skills of making word family lists and developing graded teaching and learning materials. The article concludes with some suggestions and implications for Korean language education.

[Key words] corpus, familizer, frequency, range, dispersion, text coverage, word family list