

LESLLA Symposium Proceedings



Recommended citation of this article

Cucchiarini, C., Dawidowicz, M., Filimban, E., Tammelin-Laine, T., van de Craats, I., & Strik, H. (2015). The digital literacy instructor: developing automatic speech recognition and selecting learning material for opaque and transparent orthographies. *LESLLA Symposium Proceedings*, 10(1), 251–278. <https://doi.org/10.5281/zenodo.10605014>

Citation for LESLLA Symposium Proceedings

This article is part of a collection of articles based on presentations from the 2014 Symposium held at Radboud University, Nijmegen, The Netherlands. Please note that the year of publication is often different than the year the symposium was held. We recommend the following citation when referencing the edited collection.

Van de Craats, I., Kurvers, J., & van Hout, R. (Eds.) (2015). *Adult literacy, second language and cognition*. Centre for Language Studies. Centre for Language Studies.
<https://lesllasp.journals.publicknowledgeproject.org/index.php/lesllasp/issue/view/474>

About the Organization

LESLLA aims to support adults who are learning to read and write for the first time in their lives in a new language. We promote, on a worldwide, multidisciplinary basis, the sharing of research findings, effective pedagogical practices, and information on policy.

LESLLA Symposium Proceedings

<https://lesllasp.journals.publicknowledgeproject.org>

Website

<https://www.leslla.org/>

THE DIGITAL LITERACY INSTRUCTOR: DEVELOPING AUTOMATIC SPEECH RECOGNITION AND SELECTING LEARNING MATERIAL FOR OPAQUE AND TRANSPARENT ORTHOGRAPHIES

Catia Cucchiarini, Radboud University Nijmegen

Marta Dawidowicz, University of Vienna

Enas Filimban, Newcastle University

Taina Tammel-Laine, University of Jyväskylä

Ineke van de Craats and Helmer Strik, Radboud University Nijmegen

Abstract

While learning a new language can now be a lot of fun because attractive interactive games and multimedia materials have become widely available, many of these products generally do not cater for non-literates and low-literates. In addition, their limited reading capabilities make it difficult for these learners to access language learning materials that are nowadays available for free on the web. More advanced course materials that can make learning to read and spell in a second language (L2) more enjoyable would therefore be very welcome.

This article reports on such an initiative, the DigLin project, which aims at developing and testing online basic course material for non-literate L2 adult learners who learn to read and spell either in Finnish, Dutch, German or English, while interacting with the computer, which continuously provides feedback like the most determined instructor. The most innovative feature of DigLin is that in production exercises learners can read aloud and get feedback on their speech production. This is made possible through the use of Automatic Speech Recognition (ASR). In this article we focus on what ASR is and what is needed to employ ASR to develop learning materials for non-literates and low-literates L2 learners. Central issue is the selection of the content for the four languages, which differ in orthographic transparency and present their own specific problems in combination with the mother tongue of the learners.

Keywords: automatic speech recognition (ASR), orthography, CALL, TELL, alphabetic literacy, literacy courseware

Ineke van de Craats, Jeanne Kurvers and Roeland van Hout (eds.)

Adult literacy, second language and cognition

Nijmegen: CLS, 2015, pp. 251-278

1. Introduction

As a result of globalization and internationalisation, learning a second or a third language is becoming increasingly popular. This has led to the development of numerous innovative products and applications that make it easier and more fun to learn new languages. Many of these products, such as Duolingo (<https://www.duolingo.com/>), are even available for free. Although most language learning programs come in different forms and proficiency levels, in general they cater for learners that are able to read. The large groups of LESLLA learners that are insufficiently capable of reading are still overlooked, while it is well known that these learners encounter enormous difficulties in learning new languages, in particular because most of the information and learning materials are available through the written medium and presuppose reading capabilities (Boon, 2012, 2014; Condelli & Spruck-Wrigley 2006; Feldmeier 2008; 2011; Heyn, Rokitzki & Teepker 2010; Kurvers 2002; Kurvers & Ketelaars 2011; Kurvers & Stockmann 2009; Kurvers & Van der Zouw 1990; Onderdelinden, van de Craats, & Kurvers 2009; Pracht 2010; Roll & Schramm 2010; Simpson 2007; Strube 2014; Tammelin-Laine 2011; Tarone 2010; Tarone, Bigelow, & Hansen 2009; Van de Craats & Kurvers 2009; Whiteside 2008; Young-Scholten & Naeb 2010).First, while there is increasing emphasis on language proficiency as a prerequisite for active participation in society, many countries cut down on adult education expenditure (Cooke 2010; Simpson, this volume). Second, in addition to mastering a whole new language system, the (fully) non-literates among the LESLLA learners have to become familiar with new concepts underlying an alphabetic script, such as words, graphemes, phonemes and sounds (see e.g., Kurvers 2007). Given this challenging task, they have to spend considerably long times practicing on their own and performing tedious exercises. Third, LESLLA learners have, in general, limited financial possibilities to buy suitable learning materials. Fourth, they cannot have easy access to the language learning materials that are nowadays available for free on the web, because being able to read is generally a requirement for accessing this information.

Against this background the European project “The Digital Literacy Instructor” (DigLin: <http://diglin.eu>) was started with the aim of developing and testing innovative solutions for LESLLA learners and making them easily accessible to the target groups. DigLin aims to develop and test learning materials that allow LESLLA learners to practice more actively and more frequently. In DigLin Automatic Speech Recognition technology is employed to develop spoken exercises that offer learners the opportunity to practice producing grapheme-to-phoneme correspondences at their own pace, in an anxiety-free setting. The pedagogical approach adopted in DigLin and its

advantages for students and teachers have been discussed in Cucchiari, Van de Craats, Deutekom & Strik (2013) and Van de Craats & Young-Scholten (2015). In the remainder of this paper, we first briefly introduce the DigLin project paying attention to one of its innovative features, i.e. the use of Automatic Speech Recognition to enable practice of L2 speech production and what this requires in terms of technology development and speech material collection. We then go on to discuss recent developments in selecting, designing and developing the content of the learning materials for the four languages involved in the project. Since the orthographies of these languages differ along the opacity-transparency dimension, different choices have to be made and possibly different compromises have to be reached in deciding which words should be practiced in which order. The arguments adduced in favour of the selections made in the DigLin project may be insightful and useful for teachers and researchers who have to deal with similar tasks. Subsequently, we explain how we proceeded to collect information on the reading and pronunciation errors that can be expected in the various L1-L2 combinations and present the information we gathered.

2. The DigLin project and its innovative character

2.1. Background

The DigLin project is funded by the Lifelong Learning Program (LLP) of the European Commission and is aimed at developing and testing L2 literacy learning materials in four different languages. Automatic Speech Recognition (ASR) is employed to analyze the learner's speech output and provide feedback. In line with the LLP requirements, DigLin is also aimed at disseminating the knowledge gathered within the project and at promoting exploitation of the project's results.

DigLin is carried out by a consortium consisting of partners from the Netherlands (Radboud University Nijmegen and Friesland College), Germany (University of Leipzig) and later Austria (University of Vienna), United Kingdom (Newcastle University) and Finland (University of Jyväskylä), and addresses four languages: Dutch, German, English, and Finnish. These languages have been chosen because their orthographies differ along the opacity-transparency dimension: Finnish with its clear correspondence between graphemes and phonemes has a shallow orthography and is therefore transparent, Dutch and German are in between, and English with its deep orthography is opaque.

The DigLin learning materials are not developed from scratch, but from a pedagogically sound basis which is FC Sprint² (Deutekom 2008), a language learning approach for Dutch L2 learners developed at Friesland College in Leeuwarden, The Netherlands. The rationale behind FC Sprint² is that students have to work with their own resources and have to be autonomous. In FC Sprint² students have to find out themselves instead of being told by a teacher. The teacher is the last resort. Further information on the principles underlying FC Sprint² and DigLin can be found in Cucchiari, Van de Craats, Deutekom & Strik (2013) and Van de Craats & Young-Scholten (2015).

The system developed within the FC-Sprint² program for non-literates and low-literates has been adopted for DigLin, and specific content and exercises have been developed for the four languages in the project. Traditional (digital) course material for literacy learning tends to focus on receptive tasks in which learners can listen to audio recordings and perform identification exercises of the drag-and-drop type, while in DigLin we also incorporate production exercises, as will be explained in the following section.

2.2. Automatic Speech Recognition technology

The innovative feature of DigLin is that it employs Automatic Speech Recognition to allow learners to practice L2 speech production through spoken, recoding (blending) exercises to learn grapheme-to-phoneme or graphemes-to-word correspondences in the L2, and to automatize them.

In the past ASR has been employed in reading tutors for children learning to read in their L1 (Mostow 2008). More recently, ASR has also been used in mobile applications for illiterate adults (Al-Barhamtoshy, Abdou & Rashwan 2014) learning to read in their L1. In DigLin, we carry out research on developing dedicated ASR technology for each of the four target languages in question. We study to what extent it is possible to perform speech-to-text conversion for L2 speech of beginner learners and readers and to detect possible reading or pronunciation errors in the learners' L2 utterances with a view to providing feedback on the errors observed.

In order to analyze the learners' responses, ASR technology is first employed to recognize the words and utterances spoken by the learners. When dealing with multiple languages and non-native speakers this can be challenging (Benzeghiba et al. 2007) especially in the case of illiterates (Al-Barhamtoshy, Abdou & Rashwan 2014) and in the case of beginner L2 learners (Van Doremalen, Cucchiari & Strik, 2010). In DigLin, these problems are compounded because we have to deal with (non-literate or low-literate) beginner readers trying to learn an L2, and measures have to be taken to ensure

that the recognition process is as successful as possible. ASR is a stochastic procedure in which speech corpora containing speech signals and their annotations are employed to train the speech recognizer and thus derive information about three 'knowledge sources': the language model, which contains probabilities of words and word sequences, the acoustic models, which model how the speech sounds are realized, and the lexicon, which is the connection between the language model and the acoustic models. During the recognition process (see Figure 1) the incoming speech signal is first analysed to extract the acoustic features, and then a search algorithm converts it into a string of words by using the three information sources. Since language learners may produce errors in pronunciation, vocabulary and grammar, the three knowledge sources (acoustic models, lexicon and language model) have to be adapted in the case of learner speech, for instance by using learner speech material (see below) for training or adaptation. In addition, to limit the difficulties in speech recognition, measures can be taken to constrain the nature of the exercises so that the computer can choose from among a limited number of possible answers.

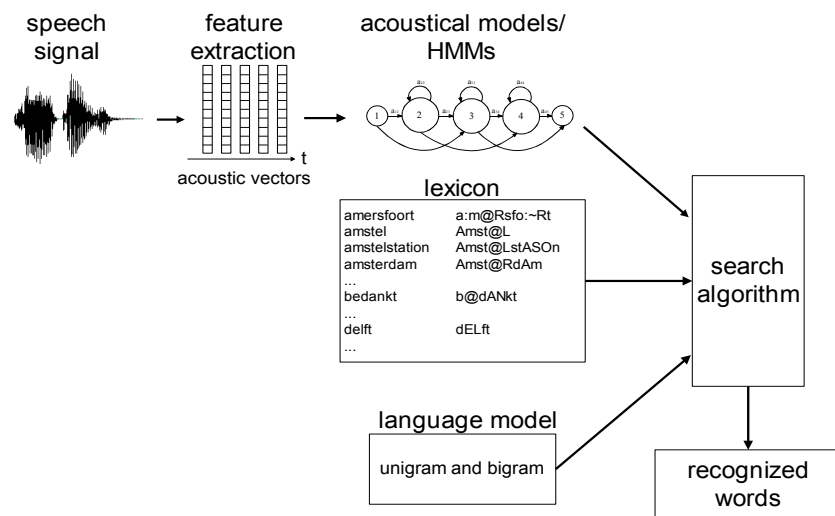


Figure 1: Example of an ASR system for a public transport information system. The lexicon contains station names in two representations: an orthographic and a phone transcription.

In a second stage, after the speech signal has been converted into a representation of the words it contains, ASR-based algorithms using acoustic models trained on native speech are employed to try to detect reading or pronunciation errors in the identified words.

To optimize the performance of the ASR technology, recordings of non-native speech are required. This speech material serves multiple purposes. First, it is used to test the performance of the ASR modules for each of the four languages with speakers from the target group. Second, it can be used to adapt the acoustic models employed by the speech recognizer so as to improve its performance in recognizing which words or utterances the learner is trying to pronounce. Third, to facilitate the process of identifying reading and/or pronunciation errors in the learner's utterances, it is important to know which errors can be expected for each L1-L2 combination. This information can be obtained from the literature, from teachers who work with learners of the specific target groups, from contrastive analyses of the L1 and L2 phonological systems, and directly from data if a sufficient number of speech recordings of the target group is available. If the latter is not the case, limited amounts of L2 speech can be recorded and used to supplement the information from the literature. Fourth, non-native speech recordings can be used to evaluate the accuracy of the reading and pronunciation error detection algorithms and see whether they can detect the errors contained in these recordings.

At the time of writing, the system is being tested (see below) and a detailed account of the performance of ASR in the four languages is not yet possible.

3. Recent developments in the DigLin project

In this section we report on recent developments for the four languages involved in DigLin with respect to the steps identified in Van de Craats & Young-Scholten (in press). These development steps are briefly described below. We then proceed to discussing these steps for each of the four languages which are presented in order of ascending orthographical complexity from Finnish to English.

3.1. Creating a 'sound bar' for each language for use with exercises in each set

The sound bar (see Figure 2) is a supporting tool that gives the learner an overview of the entire alphabet with the single graphemes, digraphs and trigraphs used in the software. Learners can also listen to the sounds corresponding to each grapheme, digraph or trigraph. This is to help them

establish letter-to-sound connections. For some languages the sound bar contains almost all letters of the alphabet, but for English this is not the case. In the following sections we explain which choices were made for the various languages and show how the sound bar looks like for each of them.

3.2. Using the FC-Sprint² Leerbedrijf technology to create fifteen exercise sets for each language

For the DigLin software we implemented five different types of the exercises already contained in FC-Sprint². These address the following sub-skills of the reading process: (1) the meaning and form of a word, (2) establishing phoneme-grapheme correspondences in visual and aural analysis and synthesis (blending), (3) recognizing whole words, (4) recognizing strings of phonemes, and (5) automatizing phoneme-grapheme correspondences and the decoding and recoding of words. The exercises developed for this latter purpose required actions like pushing and hovering over buttons, dragging and dropping letters and words, and typing letters. The exercise of reading with the help of the sound bar and the one of reading without any help (Test yourself) were added at a later stage.

Experience with FC-Sprint² had shown that within one set of exercises –using the same twenty words–, only restricted variation in the phoneme-grapheme repertoire and sufficient repetition would yield success.¹ This entails that only a restricted number of new vowels and consonants per set of exercises could be introduced to meet the first criterion, and that no less and no more than twenty words with those graphemes were required to meet the second criterion of sufficient practice. Moreover, another criterion was that the meaning of the words employed should be depictable as much as possible to avoid that learners keep in mind a wrong meaning. On the other hand, we had to accept that learners do not grasp the exact meaning at once, but first form a basic idea and only later do they store the meaning with more detailed vocabulary knowledge.

Collecting information on possible pronunciation errors and speech recordings for each L1-L2 combination

As explained above, recordings of non-native speakers (potential learners) are required to optimize the speech recognition algorithms. In addition, we need to collect information about the possible errors potential DigLin users are likely to make in the L2 they intend to learn. For the four languages involved in DigLin, this information was gathered from the literature, L2 corpora and recordings of speakers from the target group. In addition, for each language a limited number of recordings of speakers from the target group were collected. These were transferred into PRAAT (Boersma & Weenink 2003) and

phonetically annotated in SAMPA.² The information obtained for the different L1-L2 combinations is presented below for each target language (L2).

4. Finnish

4.1. Creating the Finnish sound bar

In Finnish, letter-sound correspondence is very consistent when compared with many other languages. The Finnish sounds (phonemes) are presented in Figure 2 in their orthographic form (graphemes and digraphs).

a	b	c	d	e	f	g	h	i	j	k	l	m	n	o
p	q	r	s	t	u	v	w	x	y	z	ä	ö	ng	nk

© University of Jyväskylä / © FC-Sprint² Leerbedrijf bronnen

Figure 2: *The Finnish sound bar*

Each phoneme is spelt with the corresponding single letter when short and two letters when long, except for the velar nasal /ŋ/, which only appears before /k/ and is spelt with <n(k)> or <ng>. All the letters of the alphabet are included in the sound bar, but only those used in the software are in black. We decided not to include the graphemes <c>, <q>, <w>, <x>, and <z> in the exercises because they occur only in some loan words and therefore are infrequent in Finnish.

4.2. Using the FC-Sprint² Leerbedrijf technology to create fifteen exercise sets for Finnish

In Finnish there are very few one syllable CV(C) words and short and frequent minimal pairs with concrete meanings. For these reasons exercises like those for Dutch contrasting e.g., e - i - a etc. are not feasible. Because the most simple and very common syllable structure in Finnish words is CVCV we decided to begin with these prototypical words instead and to move on to words of other shapes. Additionally, individual sound segments are less of a literacy problem in Finnish, where the letter-sound correspondence is regular and the total number of phonemes to be acquired is small. Practicing the variety of word shapes (combinations of syllable types), or the rhythm of words, is also thought to work better than practicing with only minimal pair exercises focusing on the long and

short sounds which are very common in Finnish. This is based on the experience of teachers, as there is no research on this topic.

Nearly all the letters and sounds are presented in the first three sets of words. Only /d/, /f/, /g/, and /h/ are presented later because they do not occur in CVCV words. The general selection criteria point us to easy, short, frequent (in the environment of the target group), concrete, depictable words, mainly nouns, with some adjectives and verbs, proper nouns, transparent with other languages where possible. The nouns and adjectives are presented in their nominative forms and the verbs in 3rd person singular forms. We also introduced some compounds of the CVCV+CVCV type towards the end of the set of exercises, to illustrate that a long Finnish word can consist of familiar parts, which can help reduce the anxiety for reading long words. Maximum length of the words was eight letters because of the restrictions in the software.

Table 1: *The syllable structure of the words used in the software*³

Word type	Syllable structure	Example	Translation
Two-syllable words (including minimal pairs where possible), one-syllable words	CVCV	ka-na	'chicken'
	CV ₁ V ₂ CV	tuo-li	'chair'
	CVC ₁ C ₂ V	jal-ka	'foot'
	CV ₁ V ₁ CV	lää-ke	'medicin'
	CVC ₁ C ₁ V	suk-ka	'sock'
	CVC ₁ C ₂ C ₂ V	help-po	'easy'
	CV ₁ V ₁ C ₁ C ₁ V	viik-ko	'week'
	CV ₁ V ₁ C ₁ C ₂ V	juus-to	'cheese'
	CV ₁ V ₂ C ₁ C ₂ V	puis-to	'park'
	CVC ₁ C ₂ C ₂ V	kort-ti	'card'
	CV ₁ V ₂ C ₁ C ₁ V	kaup-pa	'shop'
	V ₁ V ₁ CV(C)	uu-si	'new'
	V ₁ V ₂ CV(C)	au-to	'car'
	VC ₁ C ₁ V(V)	al-la	'under'
	VC ₁ C ₂ V(V)	an-taa	'give'
	one syllable	mies	'man'
	VCV	i-sä	'father'
	ending with C	ken-gät	'shoes'
Longer words (three and four syllables), compound words (four syllables)	CVCVCV	si-pu-li	'onion'
	CVCVCCV	to-maat-ti	'tomato'
	CVCV + CVCV	va-lo-ku-va (va-lo 'light' +ku-va 'picture')	'photo'

4.3. Collecting information on possible pronunciation errors and speech recordings for each L1-L2 combination

The major problem for most learners of Finnish is to learn the distinction between short and long sounds at the phonemic level (e.g., *tu-li* 'fire', *tuu-li* 'wind', *tul-li* 'customs'), as the actual phonetic duration varies notably depending on the speaker, the sound, the length of the word and the utterance, the position of the sound in the word etc. As to the individual segments, the L1 matters, but the following are problematic for most of the learners:

- The large number of vowel sounds (8), particularly /y/, /æ/, /ø/
- Diphthongs (18, e.g., /uo/, /ou/, /yø/, /øy/)
- /r/, /h/, /ŋ/
- Certain combinations of consonants, depending on the L1 (e.g., /ts/, /sk/).

The following information is based on the comparison of phonological inventories and descriptions of Somali and Arabic languages, available on the Internet. There seems to be no research on the pronunciation problems in Finnish by the representatives of these specific languages. The inventory information has been complemented by discussions with some language and literacy teachers, but is nevertheless theoretical and not data-driven.

Arabic

Vowels: Arabic only has three vowels /a, i, u/, while Finnish has eight /a, o, u, æ, ø, y, e, i/. Thus one can predict problems with most vowels. However, in North Africa, where French is commonly spoken, /y/ could be familiar from French. Also /o/, /æ/ and /ø/ appear as locally controlled allophones, so they are not totally unfamiliar per se, but may appear irregularly, depending on the local environment within the word. The most difficult one might be /e/ (> /æ/, /ø/). When reading aloud, the spelling may further confuse, as the spelling of vowels in foreign words varies considerably. However, this is not expected to be a problem for non-literate beginner learners as they learn the grapheme-phoneme correspondence for the first time in Finnish. Additionally, the one-to-one orthography may help them with establishing the distinctions between Finnish vowels.

The diphthong inventory of Finnish is also quite extensive (18 in standard language, with a lot of regional variation). Diphthongs can be analyzed as sequences of two basic vowels and once they are learned, the major problem is to keep apart similar ones. Likely problems: /ou/ vs. /uo/, /øy/ vs. /yø/, /ei/ vs. /ie/.

Consonants: Arabic has a large variety of consonants, while Finnish has relatively few, only /d, h, j, k, l, m, n, p, r, s, t, v and ɲ/ in native words, with /b, f and g/ in common loan words.

Common errors: v>f, p>b, g>k, ɲ>n. Many Arabic consonants are pronounced emphatically, while Finnish ones are usually quite soft (most speakers of English or German tend to hear Finnish /k, p, t/ as voiced, particularly in word initial position, as there is no aspiration) so it is likely that stops will be pronounced as too strong (also /h/ in some positions, particularly in syllable-final position as in words like /kahvi/ 'coffee'. Also /l, r, and s/ are likely to produce qualitative problems as similar sounds exist in Arabic, but there is a lot of variation. Finns tend to interpret correctly any version of these sounds, but obviously using, e.g., [z] or [ʃ] for /s/ marks the speech as accented, as do various versions of /l/ and /r/.

Prosodic errors: Arabic has long vowels and consonants which are qualitatively like the short ones, so this should not be a basic problem. Learning to hear and produce the distinction may be problematic in specific contexts, such as unstressed syllables.

Somali

Vowels: Most short Somali vowels are quite close to the Finnish ones. There is no /y/, so errors of the type y>i or y>u or y>ø are likely. Also the division of the central area is different, there is no /ø/ but several vowels nearby, so substitutions like /ø>/o/ or /e/ are possible. Some Somali vowels (particularly /i/ and /e/) are qualitatively different when long, so errors of the type /ii>/ee/ or /ee>/ææ/ are likely. Diphthongs ending in /i/ or /u/ exist in Somali, but distinctions like /ou/ vs. /uo/, /øy/ vs. /yø/, /ei/ vs. /ie/ are likely to cause errors, as are any diphthongs containing the unfamiliar /y/ or /ø/.

Consonants: Potential errors in consonants are: /p>/b/, /ɲ>/n/, /v>/f/. The quality of /t/ and /d/ is different, but it is hard to say whether this could produce confusion.

Prosody: Somali has long vowels, but their quality is not always the same as that of short vowels. Also geminate consonants exist, but not for all consonants that can be long in Finnish. Predicted errors: kk>k, tt>t, ss>s. Any long and short distinctions may be hard to perceive and pronounce in certain contexts, particularly in unstressed syllables.

We recorded 15 non-native speakers for ASR, all of them had a low to intermediate level of literacy. They were from three adult education centres from Southern Finland with seven Somali speakers (four men, three women) and eight Arabic speakers (five men, three women), and they were divided into two groups according to their age (older and younger speakers). They were asked to read aloud the 300 words we used in DigLin.

5. Dutch

5.1. Creating the Dutch sound bar

The sound bar in Figure 3 is an inventory of the Dutch sounds (phonemes) disguised in their orthographic form (graphemes and digraphs).

a	b	c	d	e	f	g	h	i	j	k	l	m	n	o	p	q	r	s	t
u	v	w	x	ij	z	aa	ee	ie	oe	oo	uu	au	ou	ei	eu	ui	ng	ch	

© 2013 Radboud Universiteit Nijmegen / © FC-Sprint² Leerbedrijf bronnen

Figure 3: *The Dutch sound bar*

All the letters of the alphabet, the diphthongs <ou> and <au> (/au/), <ei> and <ij> (/ɛi/) and <ui> (/œy/), as well as the digraphs <aa>, <ee>, <ie>, <oe>, <oo>, <uu>, <eu>, <ng> and <ch> figure in the sound bar, but only those used in the software are in black. We decided not to include the graphemes <q> and <x> in the primary set of exercises because of their infrequency. Those graphemes are pale and cannot be activated in the sound bar. The <ij> (/ɛi/) is part of the Dutch alphabet and is identically pronounced to the digraph <ei>. The digraphs <aa>, <ee>, <oo> and <uu> represent the vowels /a:/, /e:/, /o:/ and /y:/ respectively, in closed syllables (e.g. *raam* 'window'), while in open syllables simple graphemes are used (e.g. *ramen* 'windows'). The digraph <oe> stands for the vowel (/u'/) and <au> and <ou> are two orthographic representations of the same diphthong /ou/. The schwa (/ə/) is not included in the sound bar separately because the corresponding grapheme is <e>, e.g. in *tafel* 'table'. We could have opted for marking the <e> with a grey button –as was done for English and German– to indicate that we are dealing here with more than one correspondence for that grapheme. We have not done so because (i) it would do harm to the simplicity of the sound bar and (ii) the schwa only occurs in a restricted number of morphemes (-el, -en, -er, -je, -eren) that a learner will recognize rather soon. These morphemes are introduced in Exercise set (or word list 5).

5.2. Using the FC-Sprint2 Leerbedrijf technology to create 15 exercise sets for Dutch

The process of selecting the Dutch words can best be illustrated by presenting the words of the first set in Table 2. Since Dutch has more graphemes (17) for vowels (16) than the native languages of our learners do, we have built up the

wordlists around the central element of the syllable –the vowel- and added the consonants around in a more or less systematic way.

Table 2: *The first set of 20 words (Dutch); basic selection of vowels and consonants*

a (k, p, t, m, n)	oo (k, p, t, m, n, s, b, r, l)	ie (k, p, t, m, n, s, b, r, v)
kam - comb	boom - tree	tien - ten
kat - cat	boon - bean	kies - molar, back tooth
kan - jug	boot - boat	biet - beet
man - man	noot - nut	vies (adjective) - dirty
map - binder	kool - cabbage	vier - four
pan - (sauce-) pan	roos - rose	riem - belt
pak - parcel or suit	rook - smoke	
7 words	7 words	6 words

All the words in Table 2 can be represented by a picture, although *rook* ('vapor') is not really easy to grasp for a learner. In general, nouns are easier to represent than adjectives and verbs. We have postponed their introduction as much as possible.

As for phonology we would like to start with phonemes in CVC words that are known in most languages and are relatively easy to distinguish from each other, in this case at the corners of the vowel triangle. So, typologically frequent phonemes and regular orthography.

The first concession we had to make was to take <oo> (/o:/) instead of <oe> (/u'/) because two times a digraph with the letter <e> might be confusing for our learners. As for the consonants, we would have preferred to start with only plosives and nasals, but this turned out to be impossible. We could not find enough words that also met the other criteria (monosyllabic, with the <a>, <oo> or <ie>). Therefore we added <s>, , <r>, <l>, and <v>, notwithstanding that for Somali and Moroccan learners, who are our target group of learners, /p/ is a new sound that is not distinguished from /b/. We think that early confrontation with new sounds like /p/ and presentation of the two sounds /p/ and /b/ in opposition helps to draw the learner's attention to this specific difficulty. It helps to make them attentive and active learners and stimulates learning.

In the second set of exercises, we could easily make 20 words with three new vowels (<aa>, <oe> and <i>). Successively, all vowels, consonant clusters and diphthongs were introduced and the words became longer. In exercise set 15, disyllabic words like *gebouw* 'building', *koffer* 'suitcase', *koelkast* 'fridge' are found. Words of more than eight (di)graphs cannot be dealt with in this program.

5.3. Collecting information on possible pronunciation errors and speech recordings for each L1-L2 combination

To find information on pronunciation errors made by Moroccan learners of Dutch, we examined existing literature on reading errors (Kurvers & Van der Zouw 1990: 193-199) and the LESLLA corpus (Sanders, Van de Craats & De Lint, 2014) which contains semi-spontaneous speech and a sentence imitation task. For Somali speakers of Dutch Kamphuis & Amer (2013) was consulted and we interviewed speech therapist Coppens, who coached a group of Somali speakers. Together, this resulted in basic list of common errors presented in Table 3.

Table 3. *Inventory of common errors from various sources*

Moroccan	Front vowels	Back vowels	Consonants	Consonant clusters
Kurvers & van der Zouw	<ee> → /ɪ/ <ee> → /ɛ / or /ɪ/ <i> → /i' /	<u> → /y/ <u> → /ɔ/ <oo> → /ɔ/ <au> → <oo>	<g> not pronounced or as /h/.	Substitution, transposition (ts and st),
(read aloud)		<uu>, <u>, <eu>, <o>, and <oo> → /u'/. <ui> → /uu/.	<h> not pronounced. <w> → /v/.	deletion and addition in consonant clusters.
LESLLA corpus (Sanders et al.) (semi-spontaneous and sentence imitation tasks)	/e:/ → /ɛ / or /ɪ/ /ɪ/	/ə/ → /ɔ/, /o:/. /y:/ → /u'/. /œy/→ /au/.		Deletion and addition in consonant clusters.
Somali	Front vowels	Back vowels	Consonants	Consonant clusters
Coppens, p.c. (semi-spontaneous speech)	/y' / → /u' /		/p/ → /b/. /v/ → /w/.	S-cluster in onset and coda. Insertion of /ə/ before the s-cluster, and in between a word-final cluster (-təs , -pət, -lət, -ləs, -kəs etc.). Deletion of /ə/ in word-final position.

One of the problems with the literature was that Moroccan learners were not split up into Berber and Arabic speakers. Therefore we collected data ourselves from Moroccan adults with either an Arabic or a Tarifiyt Berber language background and compared them to what Kurvers and van der Zouw found. It turned out that in general, they make similar common errors.

All non-native speakers recorded for the present DigLin project had a low to intermediate level of literacy. Ten of them were men, and ten were women. There were eight Somali speakers, six Moroccan Arabic speakers and six speakers of Tarifiyt Berber, equally divided over old and young speakers from five different adult education centers or institutions spread over the country. We asked them to read aloud the 300 words we used in the DigLin fifteen word lists.

6. German

6.1. Creating the German sound bar

The sound bar in Figure 4 is a simplified inventory of frequent German sounds (phonemes) disguised in their orthographic form (graphemes, digraphs and trigraphs). Graphemes are followed by digraphs and trigraphs, both groups in alphabetical order, so learners can distinguish what belongs to the standardised alphabet and what is added independently of the standard.

a	ä	ä	ä	b	c	d	e	e	f	g	h	i	j	k	l	m	n	o	o	ö	ö	p
q	r	s	ß	t	u	u	ü	ü	v	w	x	y	z	au	ch	chs	ei	eu	ie	ng	qu	sch

Figure 4: *The German sound bar*

Also for the German sound bar design, the major goal was to offer learners a resource on all correspondences of phonemes and graphemes or multi-graphemes that appear in the DigLin software.⁴ The group of graphemes does not exactly correspond with the standardized alphabet. The *Umlaute* <ä>, <ö>, <ü> and the <ß> do not belong to the standardized alphabet. Since they are commonly used graphemes, we decided to include and organize them by criteria of proximity to letters of the alphabet: the *Umlaute* by graphic similarity and the <ß> by its phonetic closeness to <s> (<ß> always corresponds with /s/, <s> has the same sound in for instance coda positions). The sound file in the sound bar connected to each grapheme was chosen on the basis of the most common realizations of the grapheme. Alternative realizations are marked

within the words by additional grey dots: for example, when <s> precedes <p> or <t> at the beginning of a syllable, it is not realized as /z/, but changes to /ʃ/; therefore it is marked with a grey dot (see Fig. 5). This decision was based on the fact that the different realizations of a grapheme or digraph depend on their position within a word as well as the preceding and following letters. The integration of all alternatives in the sound bar would have made it visually overloaded and unclear and may have led to the assumption that the correspondence is arbitrary. The inclusion of grey dots within the words aims at turning the learner's focus and attention on the specific positions in which certain correspondences appear so that ideally they can deduce the cause and find regularities.

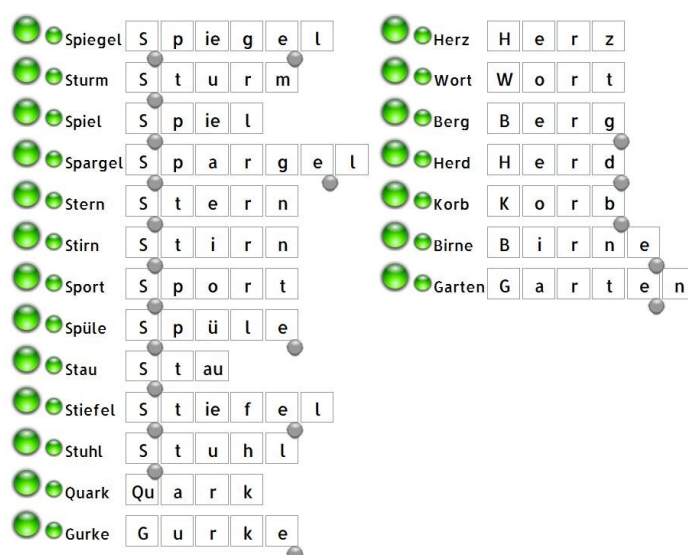


Figure 5: Example for grey dots within the words (*Alle Wörter 9*)

The convention of grey dots, however, was not applied to the vowel system because of the high range of vowels in the words. Marking all variations within the words would have led to too many grey dots and thus a visual overload. Therefore, all vowels and *Umlaute* are represented in their short and long version in the sound bar, whereas <i> and <ie> are kept separately. Solely the schwa (/ə/), as in the unstressed syllable of *Birne* 'pear' (see Fig. 5), is marked with a grey dot, so learners can become aware the relevance of stress patterns for German orthography.

6.2. Using the FC-Sprint² Leerbedrijf technology to create fifteen exercise sets for German

The German word lists only contain words with maximally two syllables. Based on what is known about common errors (see Markov, Scheithauer & Schramm 2015) made by speakers from the target group of the German DigLin version (Arabic and Kurdish speakers), we chose to focus on vowels when compiling the wordlists. The main criterion for distinction is the length of the German vowels: Long vowels are introduced in word sets 1-7; short vowels are implemented in word sets 8-15. As shown in Table 4, within these two groups we separately introduced similar vowels.⁵

The criteria for introducing consonants are based on frequency and similarity. According to Rokitzki, Nestler & Sokolowsky (2013: 99) the consonants <f>, <l>, <m>, <n>, <r>, <s>, <sch> and <w> are particularly easy to hear and pronounce because of their lasting quality. They were therefore integrated into the first word sets. Consonant clusters only appear from word set 6 on and become more complex in the following word sets in terms of their related vowel quality (see Table 4): Word set 6 contains words with consonant clusters in the beginning of a word because this has no impact on the vowel quality. Word set 8 introduces consonant clusters in the coda of a syllable or word that signal the shortening of the preceding vowel. From word set 14 onward, the number of elements in a consonant cluster rises to four.

6.3. Collecting information on possible pronunciation errors and speech recordings for each L1-L2 combination

Due to the high number of Kurdish (22%) and Arabic (14%) speaking literacy learners in Germany (Schuller, Lochner & Rother 2012: 6) these two languages were chosen to be focused on in the ASR system. On the basis of the contrastive overview of phonemes of German and these two languages (Markov, Scheithauer & Schramm 2015: 52ff, see Table 7) a list of possible common errors was deduced. Errors were predicted where the German phonemes have no or a different correspondence in Arabic or Kurdish. There is, for example, no correspondence for the German /ə/, which causes either omitting or replacing it by other phonemes as /e:/ or /i:/, which later on was confirmed on the basis of non-native speech recordings that were made in Berlin and Leipzig in spring 2014 for DigLin. Fourteen learners with L1 Arabic or Kurdish who were of

Table 4: *Overview of increasing orthographic complexity of German word sets 1-15*

Set	Stress first syllable	Vowels and diphthongs	Consonants	Particular grapheme-phoneme correspondences		
				General	Syllable-initial	Syllable-final
1	stressed	Long: a, e, o, au	n, m, l, g, b, t, f, s, r, sch	sch	s, r	o, a, e, en, consonants do not change in coda position
2	stressed	ei	p, k, w, h	ei	h	El, ei
3	stressed	long: u	d v ß	v		Er
4	stressed	long: ie i ö	z	ie, i, ie+r		I
5	stressed	long: double vowel short: a (+r)	j	a+r in the same syllable: vocalic r		d, b, g, s
6	stressed				cons. clusters	
7	stressed	long: ü, vowel+mute h, eu		vocalic r		
8	stressed	short: a, u, I, ö		syllable-final s, double consonants	intervocalic consonant sequences and syllable-final clusters	
9	stressed		qu	vocalic r, st, sp		
10	stressed	short: e o	ck ng			
11	stressed, overlap-	short: ü ä	x chs		Intervocalic ck ng x	
12	ping				double consonants and sch	
13	syllable separation		ch, clusters CCC	a o u + ch other + ch		Ig
14			clusters CCCC			
15	un-stressed					

Table 5: *Contrastive overview of inventory of phonemes in German, Arabic and Kurdish (translated and adapted from Markov, Scheithauer & Schramm 2015: 52ff)*

Examples in German		German	Arabic	Kurdish Kurmancî	Examples in German	German	Arabic	Kurdish Kurmancî
Consonants	Bein	b	b	b	Schule	ʃ	ʃ	ʃ
	Post	p		p	Etage	ʒ		
	Dach	d	d	d	Licht	ç		
	Ton	t			Jacke	j		j
	/t/ (no aspiration)		t	t	lachen	x	x	x
	Gott	g		g	Reise	ʀ ; ʁ; r	r	r; r
	Kind	k	k	k	Haus	h	h	h
	Wanne	v		v	Mann	m	m	m
	Fenster	f	f	f	Nase	n	n	n
	Sonne	z	z	z	singen	ŋ		
	Haus	s	s	s	lesen	l	l	l

Examples in German		German	Arabic	Kurdish Kurmancî	Examples in German	German	Arabic	Kurdish Kurmancî
Affricates	Dschungel	dʒ	dʒ	dʒ	Quark	kʷ		
	Taxi	ks			Zehe	ts		
	Tschüss	tʃ			Psychotherapie	ps		
	Pferd	pʃ						

Examples in German	German	Arabic	Kurdish Kurmancî	Examples in German	German	Arabic	Kurdish Kurmancî
Vowels	Saal	a:	a:	Vowels and Umbaute	Summe	ʊ	ʊ
	Tasse	a	a		---	u	
	--				Söhne	ø:	
	--		æ:		---		
	Beet	e:	e:		Hölle	œ	
	Bett	ɛ			blühen	y:	
	---				Hütte	ʏ	
	Käse	ɛ:			---		
	Liebe	i:	i:		laufe	ə	
	Lippe	i			Winter	e	
	--		i	Diphthongs	Mai	aɪ	
	Lohn	o:	o:		---		ej
	Sonne	ɔ			Eule	ɔʏ	
	---				Träume		
	---				laufen	aʊ	
	Stuhl	u:	u:				

7. English

7.1. Creating the English sound bar

Unlike the other languages involved in DigLin, English has a highly irregular orthography where almost 50% of English words are not regular grapheme-phoneme correspondences (GPC) (see e.g., Carney 1997; Shappek & Welch 2012). The sound bar created for English in Figure 6 loosely follows revised GPC rules for English monosyllabic words proposed by Vainikka (2013) much fewer than previous attempts (Cummins 1988; Bell 2004). It shows English phonemes in their orthographic forms including irregular patterns.

a	b	c	d	e	f	g	h	i	j	k	l	m	n	o	p	q	r	s	t	u	v	w	x	y	z
ae	ai	air	ar	ay	ch	ck	ea	ear	ee	ie	igh	ir	mb	ng	oa	oe	oi	oo	ou	ow	qu	sh	th	ue	wr

Figure 6: *The English sound bar*

Along with 26 single letters in the alphabet are double letters, digraphs and trigraphs (e.g., <oo>; <ch>; <igh>), and the split digraph in words like <gate> where the first vowel's pronunciation is informed by the presence of silent <e>. As with the other three languages, the sound bar shows graphemes that represent more than one sound with a grey dot (e.g., <c> can be /s/ or /k/). Because <q> is infrequent, the letter is grey and it appears on the second line in black as <qu> because it is always followed by a <u> for /kw/.

7.2. Using the FC-Sprint2 Leerbedrijf technology to create 15 exercise sets for English

Learning to read in English takes longer than in most other languages due to its irregular orthography (Goswami 2005). The existence of regular spelling rules alongside irregular spelling has fuelled continuous debate about how children learn to read, i.e. using a phonics, whole word or a whole language method (see Rayner et al.'s 2002). At present in the UK, primary school teachers are directed to combine phonics and whole word methods (Wyse and Styles 2007). Teachers of LESLLA learners vary in method and focus as the recommendations in the ESOL curriculum focus on the next level up and include very little on teaching basic literacy at the sub-CEFR A1 level.

Vainikka starts with monosyllabic words and distils regular and irregular spelling patterns into a set of 43 rules. Each letter is scored for regularity, and the uniform, exceptionless GPCs get the highest score, a 1. This scoring yields the order below, shown for the first ten rules all of which apply in both American and British English. The remaining rules refer to increasingly specific patterns for vowel monographs and vowel and vowel-consonant digraphs and trigraphs such as <oe> and <ew> and <ugh>.

When preparing word lists for the exercise sets, it was impossible to strictly follow the above order because only consonant GPCs are uniform. Vowels were considered in the context of LESLLA learners' developing phonological competence and their assumed ability to more easily distinguishable cardinal vowels, of vocabulary relevant to them and also depictable for the software itself. There were other adjustments including delay of the first GPC that two repeated consonant letters are a single phoneme) until the end of the exercise sets given potential confusion for beginners for whom double letters represent geminates (e.g. Arabic speakers). The second rule was also introduced later as GPCs which involve silence will confuse beginners. 3 was therefore the first rule applied to word sets. Vainikka's rules are based on American English and include final <r>.

In most British English varieties it is not pronounced, but since it influences vowel pronunciation <r> was included, but with vowels.

Exceptions to these are a small set of sight words.

Rule 1. <CC> = C. Two adjacent instances of a consonant are read as one

Rule 2. <b, g, h, k, l, s, w, and gh> can be silent

Rule 3. Uniform single letters: <b, d, f, k, l, m, n, p, r, t, v, z>

Rule 4. Uniform digrapheme <ch, ck, ng, ph, sh>

Rule 5. Uniform clusters/digraphs: bl-, br-, dr-, fl-, fr-, pl- pr-, shr-, tr-, -mp, -nd, -nk, -ft, -nt, -pt

Rule 6: <h, w, y, j, qu> are uniform in onsets, and <x> is uniform in codas

Rule 7: <th> has two uniform pronunciations, voiced and voiceless

Rule 8: <s> has two uniform pronunciations, voiced and voiceless

Rule 9: <c> is [s] and <g> is </dʒ/ before <e, I and y>; <c> is /k/ and <g> is /g/ elsewhere

Rule 10: words ending with vowel + <y> are uniform

Figure 7: *Vainikka's (2013) most uniform rules*

7.3. Collecting information on possible pronunciation errors and speech recordings for each L1-L2 combination

English is difficult not only due to its irregular orthography, but also syllable structure similar in complexity to its Germanic cousins Dutch and German. Pronunciation errors will differ depending on the learner's L1. Two L1s were chosen: Arabic and Bengali. The following shows a description of these languages' phonologies and errors that can be predicted for learners of English.

Arabic

Although a number of varieties of Arabic exist, errors learners produce when learning English are roughly similar and include both a range of vowel errors due to Arabic having only three vowels /a, i, u/ and consonant errors where consonants absent in Arabic cause problems in English and result in confusion between the stops /p/ and /b/, /g/ and /k/ and between nearly all fricative and affricate minimal pairs in English. The many initial and final consonant clusters in English create difficulties for Arabic speakers given that its canonical syllable structure is CV(C). To bring syllables into conformity with Arabic syllable structure, Arabic rules of epenthesis are often applied where a word such as

<price> is realized as /pɪraɪs/, <spring> as /ɪsprɪŋ/ and <next> as /nekɪst/ (Broselow 1976).

Bengali

For speakers of Bengali/Bangla and its varieties there are fewer vowel problems as Bangla has seven. There are also more consonants similar to English (Miller 2008). One difference that may cause difficulty are English labio-dental and dental fricatives which are absent in Bengali (Islam 2004). These tend to be replaced with L1 phonemes as aspiration is important in Bengali, where it is phonemic. According to Swan and Smith (2001), Bengali speakers pronounce the English voiceless consonants /p, t, ʃ, k/ without aspiration in all positions.

According to Sircar and Nag (2014), Bengali allows a large set of consonant clusters in word medial positions, particularly in mono-morphemic words and particularly in onsets where three-member clusters are allowed. The only coda clusters allowed are in loan words. Similar to Arabic, Bengali learners may epenthesize, e.g. in *fr, fl, kr, gr* clusters a vowel is inserted as in /fəlɔr/ 'floor' and *sp, sk, st* where there is vowel insertion before initial consonant cluster as in /ɪskɔl/ 'school'.

To test what the above discussion predicts for Arabic and Bengali speakers of English, recordings were made of 16 participants (four adult male and four adult female native speakers of Arabic and one male and seven female native speakers of Bengali). Their ages ranged between 19–57 and all were living in the UK at the time of the recording. The participants were equally divided into two proficiency groups based on the Common European Framework of Reference: a) eight low-level learners (CEFR A1 or lower) who had been in the UK for less than six months, and b) eight higher proficiency learners (CEFR B1 or higher) who had been in the UK longer. We asked them to read aloud the 300 words we used in the 15 DigLin word lists. The predictions were confirmed.

8. Evaluation of the DigLin system and future perspectives

In the previous sections we have explained the rationale behind the selection of the practice materials for the four languages in the project and have provided an overview of the choices made for each of these languages. Because the four languages in DigLin occupy different positions along the orthographic transparency continuum this information can be helpful not only to teachers and researchers working on these specific four languages, but also to those working with languages that occupy similar positions on this continuum.

For those interested in employing ASR for literacy instruction, the four sections on the pronunciation errors to be expected in each L1-L2 combination indicate how relevant information can be obtained when large amounts of annotated speech data are lacking, which is unfortunately very often the case.

At the time of writing the four versions of DigLin are being tested with LESLLA learners. An important aspect of the DigLin system that has not been discussed so far is its capacity to log learner behaviour during practice. This is an interesting feature for evaluation and research purposes because it allows to gain insight not only in the results of practice, but also in the learning process. In addition, during testing speech recordings are made of all learners and these can in turn be used to gain more information on the errors made and to improve the ASR algorithms. In the near future we will be able to report on this interesting aspects of the project.

Acknowledgements

This project has been funded with support from the European Commission under project: 527536-LLP-1-2012-1-NL-GRUNDTVIG-GMP. This publication reflects the views only of the project consortium members, and the Commission cannot be held responsible for any use which may be made of the information contained therein.

We are indebted to the other members of the DigLin team for their valuable contributions, in alphabetical order: Jan Deutekom, Joost van Doremalen, Vanja de Lint, Maisa Martin, Rola Naeb, Jan-Willem Overal, Karen Schramm and Martha Young-Scholten.

Notes

- 1 In DigLin there are 15 sets of exercises in which a list of 20 words is practised.
- 2 SAMPA (Speech Assessment Methods Phonetic Alphabet) is a computer-readable phonetic script using 7-bit printable ASCII characters, based on the International Phonetic Alphabet (IPA).
- 3 There are two ways of seeing long and short vowels: either we can say that the long vowel (or consonant) is a sequence of two phonemes or that the long vowel (or consonant) is one specific phoneme, distinguished from the short one. Both ways are actually used in Finland. In Table 1 the representations are based on the idea that a long vowel represents two phonemes.
- 4 However, the software does not contain words with <y> and <c> in isolation (only the di- and trigraphs <ch>, <ck> and <chs> where <c> only has the orthographic function of distinguishing between <h> as /h/ or when spelled <ch> as /x/ or /ç/; in the case of <ck> it indicates that the preceding vowel is short while <k> usually follows long vowels). The sound for <c> is /k/ since

- there is no word appearing with an isolated <c>. So <ck> did not get included in the sound bar because learners can deduce how <ck> is pronounced
- 5 Similarity refers to graphic proximity such as in <o> and <ö> and to phonetic closeness in terms of mode and place of articulation according to the vowel diagram.

References

- Al-Barhamtoshy, H., Abdou, S., & Rashwan, M. (2014). Mobile technology for illiterate education. *Life Science Journal*, 11(9), 242-248.
- Ali, E. (2011). *Speech intelligibility problems of Sudanese learners of English*. Doctoral dissertation, Leiden University. Utrecht: LOT.
- Business and Economics*, 3 (2): 191-203.
- Barros, A. (2003). *Pronunciation difficulties in the consonant system experienced by Arabic speakers when learning English after the age of puberty*. Unpublished M.A. thesis, West University, May.
- Bell, M. (2004). *Understanding English spelling*. Cambridge, England: Pegasus Educational.
- Benzeghiba, M., Mori, R. D., Deroo, O., Dupont, S., Erbes, T., Jouvet, D., Fissore, L., Laface, P., Mertins, A., Ris, C., Rose, R., Tyagi, V., & Wellekens, C. (2007). Automatic speech recognition and speech variability: a review. *Speech Communication*, 49, (10-11), 763-786.
- Boersma, P., & Weenink, D. (2003). *PRAAT Phonological Transcription*, version 4.2.34. Available online via <http://www.praat.org> (June 2012).
- Boon, D., & Kurvers, J. (2012). Ways of teaching reading and writing: Instructional practices in adult literacy classes in East Timor. In P. Vinegradov, & M. Bigelow (eds.), *Low-Educated Adult Second Language and Literacy Acquisition, 7th Symposium – Minneapolis 2011*, 67-91.
- Boon, D. (2014). *Adult literacy education in a multilingual context. Teaching, learning and using written language in Timor-Leste*. Doctoral dissertation, Tilburg University.
- Broselow, E. (1976). *The phonology of Egyptian Arabic*. Doctoral dissertation, University of Massachusetts, Amherst, Massachusetts.
- Carney, E. (1997). *English spelling*. Routledge.
- Condelli, L., & Spruck-Wrigley, H. (2006). Instruction, language and literacy. What works study for adult ESL literacy students. In I. van de Craats, J. Kurvers, & M. Young-Scholten (eds.), *Low-Educated Adult Second Language and Literacy Acquisition, Proceedings of Inaugural Symposium – Tilburg 05*, 111-133.

- Cooke M. (2010). ESOL in the United Kingdom. Paper at the inaugural EU-Speak workshop, Newcastle, 6 November.
- Deutekom, J. (2008). *FC-Sprint², Grenzeloos leren* [Learning without borders]. Amsterdam: Boom.
- Doremalen, J. J. H. C. van, Cucchiarini, C., & Strik, H. (2010). Optimizing automatic speech recognition for low-proficient non-native speakers. *EURASIP Journal on Audio, Speech and Music Processing*.
- Feldmeier, A. (2008). The case of Germany: Literacy instruction for adult immigrants. In M. Young-Scholten (ed.), *Low-Educated Adult Second Language and Literacy Acquisition*, 3rd Symposium –Newcastle 2007, 7-16.
- Feldmeier, A. (2011). *Alphabetisierung von Erwachsenen nicht deutscher Muttersprache*. Leseprozesse und Anwendung von Strategien beim Erlesen isoliert dargestellter Wörter unter besonderer Berücksichtigung der farblichen und typographischen Markierung von Buchstabengruppen. Online: http://bieson.ub.uni-bielefeld.de/frontdoor.php?source_opus=1814.
- Goswami, U. (2005). Synthetic phonics and learning to read: A cross-language perspective. *Educational Psychology in Practice*, 21 (4), 273-282.
- Heyn, A., Rokitzki, C., Teepker, F. (2010). Alphabetisierung von Migranten in der Fremdsprache Deutsch – Lernfortschrittsmessung mit dem Marburger Kompetenzrad. *Deutsch als Fremdsprache*, 47 (4), 210-221.
- Islam, A. (2004). L1 influence on the spoken English proficiency of Bengali speakers. *C-Essay in English*, Högskolan Dalarna University.
- Kamphuis, N., and Amer, S. (2013). *Somalisch*. <http://meertaligheidentaalstoornissen.vu.wikipaces.com/Somalisch> (November 2013).
- Kurvers, J. (2002). *Met ongeletterde ogen. Kennis van taal en schrift van analfabeten*. [With illiterate eyes. Knowledge of language and writing of illiterate adults.] Amsterdam: Aksant.
- Kurvers, J. (2007). Development of word recognition skills of adult L2 beginning readers. In N. Faux (ed.) *Low-Educated Second Language and Literacy Acquisition*. Richmond, Virginia: The Literacy Institute at Virginia Commonwealth University, 23-43.
- Kurvers, J., & Ketelaars, E. (2011). Emergent writing of LESLLA learners. In C. Schöneberger, I. van de Craats, & J. Kurvers (eds.), *Low-Educated Adult Second Language and Literacy Acquisition*, 8th Symposium – Cologne 2010. Nijmegen: Centre for Language Studies, 49-66.
- Kurvers, J., & Stockmann, W. (2009). Alfabetisering NT2 in beeld. Leerlast en succesfactoren. [Focus on L2 literacy acquisition. Learning load and determinants of success.] Report, Tilburg University.

- Kurvers, J., & Van der Zouw, K. (1990). *In de ban van het schrift. Over analfabetisme en alfabetisering in een tweede taal*. [In the spell of script: About illiteracy and literacy teaching in a second language]. Amsterdam/Lisse: Swets & Zeitlinger.
- Markov, S., Scheithauer, C., Schramm, K. (2015). *Lernberatung für Teilnehmende in DaZ-Alphabetisierungskursen. Handreichung für Lernberatende und Lehrkräfte*. Münster: Waxmann.
- Miller, D. (2008). *Production of Bangla stops by native English speakers learning Bangla: an acoustic analysis*. Doctoral dissertation, University of North Dakota.
- Mostow J., (2008). Experience from a reading Tutor that listens: Evaluation purposes, excuses, and methods". In C. K. Kinzer & L. Verhoeven (eds.), *Interactive Literacy Education: Facilitating Literacy Environments through Technology* (pp. 117-148). New York: Lawrence Erlbaum Associates, Taylor & Francis Group.
- Onderdelinden, L., I. van de Craats, & J. Kurvers (2009). Word concept of illiterates and low-literates: worlds apart? In I. van de Craats & J. Kurvers (eds.), *Low-Educated Adult Second Language and Literacy Acquisition, 4th Symposium - Antwerp 2008* (pp. 35-48). Utrecht: LOT Occasional Series 15.
- Pracht, H. (2010). Alphabetisierung in der Zweitsprache Deutsch als Schemabildungsprozess. Bedingungsfaktoren der Schemaetablierung und -verwendung auf der Grundlage der „usage-based theory“. Münster: Waxmann.
http://www.waxmann.com/index.php?id=6&no_cache=1&tx_p2waxmann_pi1%5Bautor%5D=PER103239&tx_p2waxmann_pi1%5Bbuchstabe%5D=P
- Rokitzki, C., Nestler, D., & Sokolowsky, C. (2013). Lernabenteuer Schriftsystem. In Feick, D., Pietzuch, A., & Schramm, K. (eds.), *Alphabetisierung für Erwachsene* (pp. 91-109). Berlin: Langenscheidt.
- Roll, H., & Schramm, K. (2010) (eds.). *Alphabetisierung in der Zweitsprache Deutsch*. [Osnabrücker Beiträge zur Sprachtheorie 77]. Duisburg: Gilles & Francke.
- Sanders, E., Van de Craats, I., & De Lint, V. (2014). *The Dutch LESLLA Corpus*. Proceedings of LREC 2014, Reykjavik, 2715-2718.
- Shappeck, M., & Welch, K. (2012). *Linguistics for Pre-Service Educators*, p. 199.
- Schuller, K.; Lochner, S. & Rother, N. (2012). Das Integrationspanel. Entwicklung der Deutschkenntnisse und Fortschritte der Integration bei Teilnehmenden an Alphabetisierungskursen. Working Paper 42. Nürnberg: Bundesamt für Migration und Flüchtlinge.

- Sircar, S., & Nag, S. (2014). Akshara-syllable mappings in Bengali: A language-specific skill for reading. In H. Winskel, & P. Padakannaya (eds.), *South and Southeast Asian psycholinguistics*, 19 (pp. 202-211). Cambridge University Press.
- Strube, S. (2014). *Grappling with the oral skills. The learning and teaching of the low-literate adult second language learner*. Doctoral dissertation, Radboud University Nijmegen. Utrecht: LOT.
- Swan, M., & Smith, B. (2001). *Learner English: A teacher's guide to interference and other problems* (2nd ed.). Cambridge, England: Cambridge University Press
- Tammelin-Laine, T. (2011). Non-literate immigrants – a new group of adults in Finland. In C. Schöneberger, I. van de Craats, & J. Kurvers (eds.), *Low-Educated Adult Second Language and Literacy Acquisition, 8th Symposium – Cologne 2010*. Nijmegen: Centre for Language Studies, 67-78.
- Tarone, E. (2010). Second language acquisition by low-literate learners: An understudied population. *Language Teaching*, 43, 75-83.
- Tarone, E. Bigelow, & Hansen, (2009). *Literacy and second language oracy*. Oxford: Oxford University Press.
- Vainikka, A. (2013). *English reading and spelling rules for short words*. ms Johns Hopkins Virginia University, WV.
- Van de Craats, I., & Young-Scholten, M. (in press). Developing technology-enhanced literacy learning for LESLLA learners. In M. Santos & A. Whiteside (eds.), *Low Educated Second Language and Literacy Acquisition. Proceedings of the 9th symposium – San Francisco 2013*.
- Whiteside, A. (2007). Who is 'You'? ESL literacy, written text and troubles with deixis in imagined spaces. In M. Young-Scholten (ed.), *Low-Educated Adult Second Language and Literacy Acquisition, 3rd Symposium – Newcastle 2007*, 99-108.
- Wyse, D., & Styles, M. (2007). Synthetic phonics and the teaching of reading: the debate surrounding England's 'Rose Report'. *Literacy*, 41(1), 35-42.
- Young-Scholten, M., & Naeb, R. (2010). Non-literate L2 adults' small steps in mastering the constellation of skills required for reading. In T. Wall & M. Leong (eds.), *Low-Educated Second Language and Literacy Acquisition: Proceedings of the 5th Symposium, Banff, 2009* (pp. 80-92). Calgary: Bow Valley College.