

<https://www.doi.org/10.31918/twejer.2031.26>

e-ISSN (2617-0752)

p-ISSN (2617-0744)



Spelling Errors in MA Theses by Non-native English Speakers in Kurdistan Region of Iraq: A Corpus-based Study

Shamal A. Abdullah

Soran University- Iraq

Prof. Dr.Himdad A. Muhammad

Salahaddin University-Erbil, Iraq

shamal.abdullah@soran.edu.iq

himdad.muhammad@su.edu.krd

Abstract:

It is usual for the second language learners to make spelling errors in writing, as many pointed out this issue worldwide. On the other hand, advanced speakers of a second language are supposed to make few or no errors of this type in second language writing productions, especially, in academic writings such as master theses. This is not always true, among fluent speakers, as the committing of linguistic errors surpasses this level of proficiency and continues to happen in academic productions. This paper addresses the issue that academic English speakers in the Kurdistan Region of Iraq still face challenges to produce error-free academic written works, especially spelling errors. The study covers MA theses from major universities in Kurdistan Region written on English linguistics. The frequency of the errors is taken into account, and then the results are qualitatively analysed. The results indicate that MA students make spelling errors in their final version of the MA thesis. Spellings errors usually result from lack of technical skills and typing proficiency, but the case is that the errors are happening after certain revisions.

Keywords: spelling error, corpus-based study, academic writing, non-native English speaker, MA students.

Spelling Errors in MA Theses by Non-native English Speakers in Kurdistan Region of Iraq: A Corpus-based Study

1. Introduction

English has become globally a prominent language for scientific and academic purposes since ESL and EFL writers publish numerous manuscripts. Also, English is defined by Vold (2006) as the “lingua franca of academic discourse”. The demand for a qualified English for academic discourse has been significant over the past years. Thus, the use of academic English, or in other words, advanced English language usage, is essential and prerequisite in different areas. Academic discourse means using language at the academic level linguistically. It is significant, in large part, in many fields and aspects; educational, instructional, scientific. Thus, writers are urged to write in high quality and style, but this academic discourse usually becomes challenging when it is to be written, especially in a foreign and/or second language.

Writings by EFL and ESL learners are prone to having many linguistic deficiencies and errors, and this is usual globally in reference to related studies. Academic writings, written by advanced non-native English speakers, on the other hand, are supposed to be with highly language accurate and free of errors, especially English major fields. This supposition might not be entirely true, and some language learning challenges remain to be challenging in higher and more advanced ESL and EFL skills. To acquire a high and professional written product in English by non-native speakers, experts continuously investigate to improve such acquisition. Corpus linguistics is one of the best recent approaches to be used for language investigation and advancement. It has notably been common, in the latter part of the 20th century,

especially after the rapid technological advancement of computers. Corpus linguistics is another approach for investigation on languages. Linguists and centres for language research depend on corpora to improve the usage of English in the academic field, especially for non-native speakers.

In Kurdistan Region of Iraq, English is the required language when writing a master's degree in linguistics in English language departments. To improve writing in the second language, many manuscripts have been written, but there have not been investigations at advanced levels and covering data from academic productions. This study analyses the linguistic errors and specifically spelling errors in MA theses. The study uses corpus as the main approach and corpus tools for the analysis of the results.

Linguistic errors are the result of imperfect usage of linguistic and grammatical elements. In general, linguistic errors are common worldwide among L2 learners and non-native speakers of a language. Linguistic errors, especially, a type like spelling in academic written productions does not seem common. This study investigates the frequency of spelling errors in academic writings in Kurdistan Region of Iraq by advanced English speakers. For this purpose, the current study works to discover to what degree spelling errors are common in English academic writings, i.e. theses, in public universities in Kurdistan Region of Iraq. The study also looks for highlighting the most common spelling error types. In addition to that, the study looks for the causes of these errors, and suggests relevant solutions.

1.1 The Scope

There are many types of linguistic errors, but the current study deals with the analysis of spelling errors made by the Kurdish scholars in their academic English writings at the three major public universities of the Kurdistan Region. In this study, MA theses are taken as the source of the data. In Kurdistan region, many theses are written in English not merely in English major fields. In order to obtain the problems referred to in the study, only theses written by advanced English speakers are taken into consideration in the major of English Linguistics from 2015 to 2018. The aim of choosing this period is due to the recentness of the data. The texts are collected electronically for the purpose of building the corpus. Thus, 21 individual texts are chosen from three major universities in Kurdistan region, namely Salahaddin, Sulaimani, and Duhok Universities. Seven texts are taken from each university to guarantee the liability of the study. The size of the corpus then consists of 505,000 words.

1.2 Value of the study

In academic and professional writing, slight errors in spelling may cause considerable differences in meaning. Almost in all academic writings, proficiency in language is highly required. For this purpose, highlighting the negative sides of academic writing is urgent so as to be avoided. On the other hand, there are many uses of corpora, based on the type of corpus. Learner or in other words, the non-native corpus is usually used for L2 improvements and some other insights on learner language. Numerous researches can be conducted on this learner corpus to find out how to improve the ESL/EFL classes, academic writing and teaching methodology. It is also significant to compare the writing of local scholars with native

and professional English texts and get great insights for investigation and researching.

In general, the significances of this study are multi-dimensional:

1. For learners: this study provides learners with a list of spelling errors that occur the most and continue to be challenging in higher levels of English proficiency.
2. For instructors: the results are useful for language and linguistic instructors to improve their EFL/ESL classrooms. As this study works systematically on finding spelling errors, and the results are systematic, instructors can change their view about teaching academic English writing.
3. For researchers: the study proves that there is a need for peer-review and spelling checkers, and thus, researchers need to take this case into more consideration.

1.3 The Hypotheses

This study hypothesises that spelling errors continue to be challenging in advanced writing productions, and thus, investigation on advanced texts for linguistic errors is a relevant case for studying. MA students in English linguistics are regarded as advanced EFL/ES. MA students need to use advanced language in their final thesis. Thus, the theses, under scrutiny as the data of the study, have undergone a number of revisions, but this does not indicate an error-free production. The study also hypothesises that errors in compound and hyphenation are the most challenging types as the pilot study highlights it. It is also believed that the

underlying factors of such errors are interlingual effects, intralingual effects, technical skills, and dyslexia.

2. Review of Literature

Non-native speakers of language usually tend to make errors during second language use, although everyone looks for better fluency. As this is normal among the learners, linguists and scholars deeply investigate this and use different ways to highlight the errors as well as the remedies. Literature shows that scholars are interested in linguistic error analysis worldwide. Scholars use different methods and approaches to investigate this phenomenon among non-native speakers. Corpus studies, as a recent approach turns highly prolific in academic studies. Many scholars use corpus-based method to analyse and investigate linguistic errors, especially, committed by learners. Corpus-based studies that depend mostly on data frequency are highly reliable. The current study provides a number of most recent studies on linguistic error analysis based on corpus approach.

A corpus-based study by Hamed (2018) is conducted to identify the common linguistic errors committed by the non-English major Libyan students. The descriptive error analysis method is implemented to analyse the errors, and the researcher reads the texts to identify the errors. The data is gathered from the Language Centre at Omar EL-Mukhtar University by taking 40 sample written texts by the Libyan students at a pre-intermediate level. The researcher employs Hubbard's taxonomy to classify the error types. Later, the occurrence frequency of various sorts of linguistic errors is counted. The results indicate that the substance errors occur with the highest frequency following by grammar, syntax, and lexical errors, respectively. The data analysis also

shows that spelling error is the highest in frequency among the sub classifications of the errors, while the possessive case and relative clause errors obtain the least frequency. The researcher comes to the point that errors can be the result of over-generalisation in the target language, ignorance of rule limitations, and inadequate application of rules which tend to result from the first language that is Arabic.

Hulvova (2017) conducts a study to investigate the linguistic errors committed by EFL learners. This study provides an enlightening knowledge about the language learning process, linguistic errors and corpus studies. It is quantitative and is based on the method of error analysis. The source data is comprised of 36 written texts by 18 Czech high school learners studying the grammar of level B2 in a Czech local grammar school. The learners were asked to answer specific questions in two different tasks. To describe the errors, the researcher uses a modified annotation system developed by the Centre of English Corpus Linguistics at the University Catholique de Louvain in Belgium, in which 32 error tags are used. The results of this empirical study show 710 overall errors belonging to the linguistic domains. The frequency analysis indicates that grammar errors share the highest frequency, much more than any other sort of linguistic error, following by lexis, form, and sentence structure errors, respectively. Each of these domains is analysed deeper into specific sub-types. The more in-depth analysis shows that article errors in grammar domain, dependent preposition errors in lexis domain, spelling errors in form domain, and missing words in sentence structure errors are the highest frequency in the among the sub-types. The researcher concludes that EFL learners need more linguistic awareness and

teachers should give better instructions for the avoidance of such errors, although some errors may seem hard to eradicate.

Divsar and Heydari (2017) conduct a study to identify the linguistic errors made by Iranian EFL learners in IELTS essay writings. A corpus-based error (CEA) analysis approach is implemented along with the coding scheme for error classification. The data consisting of 70 sample essays from “ielts-blog.com” is compiled into a learner corpus. They analyse the errors based on a qualitative method involving a descriptive design. They categorise the errors into 13 different types and identify a code for each type. The results show that word choice comes in the first place that learners tend to make. Following that verb form, spelling, and noun errors are more frequent, respectively. On the other hand, punctuation errors and errors regarding Connective word, definitive statement and linking words errors are the least frequent ones. They point to three main causes for these errors, including; lack of modern grammar teaching methods, learners L1 interference and unfamiliarity with the pragmatic aspects of L2. They find that these corpora-based studies for error analysing are important in three ways: instructors know the language learning progress in the first place. Secondly, researchers gain more language learning strategies, and thirdly, learners get benefit from these errors to enhance their language learning.

Seitova (2016) conducts a study to explore the common English language errors made by Kazakh and Russian First Language students. This study is designed to identify, describe, and evaluate errors. The data includes two corpora: 32 compositions directly written in English and 32 other compositions translated from L1 into English. These two corpora are then analysed for errors and compared. The errors are classified into seven types

under two categories of syntactical and morphological. The data shows that spelling, misuse of like+ v-ing form, and omission or misuse of prepositions is the most frequently occurring errors in the combined data of composition and translation. The researcher finds that the language learning process is a long process, and the occurrence of errors is usually predicted. Furthermore, she finds that the influence of the L1 (Kazakh and Russian) shares the highest portion in the causes of the occurrence of these errors, besides the causes such as language mixing, L2 proficiency levels, literacy of the L1, social factors, and individual variations. She indicates that the error analysis is an excellent support for learners. She also identifies that the implications are worthy for teachers and curriculum designers to enhance the language learning process, especially in determining the learning gaps.

Previous studies show that linguistic errors, in general, are unfortunately occurring with high frequency among second-language speakers of English, especially learners. Scholars seem to work on the learners' productions, and thus, the language used for academic and professional purposes is not taken into much consideration and not investigated deeply. Previous studies also indicate that spelling errors usually occur with high frequency and needs attention. Thus, in contrast to the previous studies, the current study takes MA theses that are supposed to use advanced English as the source data for analysis. The texts are collected in one corpus and analysed through corpus tools including AntConc. This study is thus different from previous studies and sheds light on critical sides of linguistic errors in academic writings.

3. Method

The current study is multidimensional. It is both qualitative and quantitative. A corpus is firstly made out of the selected texts after a certain number of steps, including data selection, text collection, computing, storing, purification and compilation. Then a quantitative account is given to the results. Finally, the results are analysed qualitatively.

3.1 Participants

The participants of this study are not random, but they are anonymously taken. The participants are advanced non-native English speakers studying MA in English linguistics in the three major public universities in Kurdistan Region of Iraq. The participants are then both academic and advanced in the English language, and the data includes 21 theses, seven from each university.

3.2. Procedure

This study primarily depends on a corpus as a source data. Building a corpus is not an easy task, and needs many steps to be handy for investigation. The process of building the current corpus started via identifying the target of the source data. In the second step, the identified texts are collected electronically on a computer as a plain text format. Following that, the compiled texts are purified from not-related contents since the MA theses contain many parts that are irrelevant such as cover page, governmental headings, table of contents, quoted texts, non-English texts, references, and etc. Finally, all the texts are combined to build the study's corpus.

The corpus is then investigated through corpus tools including AntConc. After inquiring the frequency of the errors individually for each type, the results are qualitatively analysed. In addition, the corpus is mainly written and does not contain any data derived from spoken texts.

3.3 Ethical Considerations

Anonymity and confidentiality are two aspects of ethical consideration that researchers need to consider. Walford (2005) states that anonymity indicates referring the population by false names, which means using invented names. On the other hand, Dawson (2007) believes that anonymity indicates the avoidance of the usage of any name, addresses, or indications. For this purpose, formal permission is taken to compile a corpus out of such kinds of texts and investigate. Furthermore, the texts are purified, and the authors' names are omitted as well. The source data is entirely anonymous.

3.4 Methodology and procedure

For the purpose of identifying spelling errors Antconc, a corpus analysing tool, and Grammarly are implemented. The total frequency for each type of error is counted and then normalised. Normalisation is the process of gaining a unified rate for each error, especially in larger corpora.

Normalized result = result X (desired size)/(size of this corpus)

The normalisation factor is derived through the following equation to make the process easier:

Normalization factor = (desired size)/(size of this corpus)

Then the normalised result can easily be found by the following equation:

Normalized result = result X normalization factor

The normalisation number is mostly taken by 10,000 in large corpora. The different corpora are normalised to the frequencies per 10,000 words.

3.5 Error Analysis

Crystal (2008, p173) defines the process of error analysis as a “technique for identifying, classifying and systematically interpreting the unacceptable forms produced by someone learning a foreign language, using any of the principles and procedures provided by linguistics.” Errors are supposed to reflect, in a systematic way, the level of competence obtained by second-language users. For Corder (1975) Error analysis analyses error occurrence through a systematic procedure that includes gathering, identifying, classifying, explaining, assessing and providing solutions.

Currently, error analysis is seen as a method, and it is very widespread globally among the researchers. A common usage of the method is for determining linguistic misuses in second language wirings. Richards & Schmidt (2010, p. 201) state that error analysis “may be carried out in order to:

- A. Identify strategies which learners use in language learning
- B. Identify the causes of learner errors

C. Obtain information on common difficulties in language learning, as an aid to teaching or in the preparation of teaching materials”.

Thus, error analysis can be seen as a methodology that tries to amend the deficiencies challenging in language use. The current study applies this method along with corpus-based methodology towards acquiring this amendment.

3.6 Spelling Error categorisation

The spelling error is described by Laframboise (1996) as a mismatch between the used spelling by the writer and dictionary spelling. Kukich (cited in Liu, 2015, P 1629) specifies two types of errors, typographic and cognitive errors. Typographic errors are related to keyboard adjacencies including insertion, deletion, substitution, and transposition, whereas cognitive errors are results of phonetic similarities, such as *taire includes insertion error by adding <a> inside the word. In addition to this, there is another way of categorizing spelling error. Laframboise (1996) states that spelling errors occur on homophones, consonants, vowel errors, inflected words and reversals. In a general way, Cook (1997) categorises spelling errors into six different types as the following: letter insertion, letter omission, letter substitution, transposition, compounding, hyphenation, and apostrophes.

Based on these categorisations and divisions, in this study, spelling errors are divided into the following types that are individually explained and exemplified:

A. Letter insertion: It includes errors where an additional letter(s) is put inside or one of the ends of a word. This type can

include consonant doubling, specific insertion, and intra-lingual insertions. For example, <ll> in **beautifull*, and <pp> in **appricot*.

- B. Letter omission: It is the leaving out of a letter in a word. The most common sub-types include double consonants, consonant clusters, one consonant in a pair, exclusion of <gh> in certain words, and exclusion of <e> in the final position. For example, <p> in **oposite*, <c> in **exlusion*, <c> in **aquaintance*, <gh> in **drout*, and <e> in **cultur*.
- C. Letter substitution: It is the replacement of one letter by another in a word. It can also include the replacement of multiple letters. According to Cook (1997), the most common substitution error among English non-native speakers occurs with the replacement of <a>, <e>, and <i> by each other. For example, <i> replaced by <e> in **emage*, and <a> is replaced by <e> in **grete*.
- D. Transposition: It is the substitution of two consecutive letters. Two common types include the substitution of <e> with <i>, and <e> with <t>. For example, <e> and <i> are transposed in **recieve*, and <e> is transposed with <t> in *quite*.
- E. Compounding: It occurs when two sperate words are merged to create a new word. The error happens when the produced word is not acceptable. The error also happens when a compound word is split while it should not be; the separate words should come together; however, this composition is sometimes based on the position of the word the context. Examples for this error can be the words *over viewing*, whereas the correct form should be *overviewing*. Another

example can be the word **tabletennis*, where the merged word should be split as *table tennis*.

- F. Apostrophes: It is the omission or addition, and sometimes misplacement of the apostrophe in contractions and possessive manner. For example, <'> is omitted in **theyre* or *lets*, and <'> is misplaced in **mens'*.
- G. Hyphenation: It is the misuse of the hyphen in compound words. There are some compound words that do not require hyphen and vice versa. For example, hyphen is omitted in **long term* and **longterm*. There are three categories in hyphenation errors; lack of hyphen, addition of unneeded hyphen, and adding redundant space after hyphen before the next part of the word.

Although in some cases, the hyphenation and apostrophe are regarded as the subtypes of punctuation errors, the current study includes these two types within spelling as they affect the way the words appear. There are also other types such as addition of spaces and usage of multi-dialect words, but these types are excluded in this study.

4. Results and Discussion

Spelling errors in texts written in advanced English seem to be existent, besides the revisions done on them. Table (1) shows the overall results for the errors, including normalisation and the percentage of each type.

Table (1). The overall result of the spelling errors in MA theses

Error Type	Overall Instances	Normalised Result	Percentage
1. Insertion	23	0.455	5.94
2. Omission	51	1.009	13.18
3. Substitution	33	0.653	8.53
4. Transposition	19	0.376	4.91
5. Compounding	82	1.623	21.19
6. Apostrophes	15	0.297	3.88
7. Hyphenation	164	3.246	42.38
Total	387	7.659	

The above results indicate that spelling errors exist in noticeable frequency. It shows that MA students in writing the thesis still face difficulties in managing English spelling, although computers are great assists nowadays. The results also show that errors in hyphenation are the most frequent, followed by compounding, omission, substitution, insertion, insertion, and transposition, respectively, while errors in apostrophes are the least occurring.

Looking deeper into the results, the errors in hyphenation are divided into two types: lack of hyphen and adding an unneeded hyphen between words. Table (2) shows the detailed results for hyphenation.

Table (2). The detailed results for hyphenation errors

Hyphenation Error Types	Overall Instances	Normalised Result	Percentage
1. Lack of Hyphen	158	3.127	96.332
2. Addition on Unneeded Hyphen	6	0.119	3.658
Total	164	3.246	

The above results show that the writers tend to omit out the hyphen in the places where a hyphen is required much more than adding hyphen to a place where hyphen is not required. This shows that mastering the use of hyphen needs to be taken into more consideration. In most of the cases, where two words together modify another succeeding word, a hyphen is needed, and in the same way, the most errors are made in this regard. Examples for these errors in corpus include the words *role playing* in which a hyphen is needed, on the other hand, hyphen is added to **metaphor* which is not required and it is wrong.

Concerning errors in the compounding words, the errors are categorised into two types, including the separation of words that should be related and the linking of two words that should be separated. Table (3) shows the detailed results for both types.

Table (3) The detailed results for errors in compounding words

Errors in Compounding Words	Overall Instances	Normalised Results	Percentage
Separating Words	57	1.128	69.505
Linking Words	25	0.495	30.485
Total	82	1.623	

The above results show that advanced students of English separate compounding words with higher frequency than linking them in unnecessary cases. This indicates that while writing, writers tend to think of separate words in some situations unaware of changes in meaning, and also in the same sense, they link two words where there is no necessity to do that. Examples for such errors from the corpus can be the words *every day* which should be linked to be used as adjectives, and the word **alot* is misspelt; the real spelling is *a lot*.

Regarding the omission of letters, the detailed results can refer to the letters that are mostly omitted. The results are calculated for every single letter that is omitted in more than one cases, and the letters that are omitted in only one case are calculated together. The aim of this calculation is that one case may not be notable decision on the case. Table (4) shows these detailed results.

Table (4) The detailed results for the most involved letters in omission errors

Omitted Letters	Overall Instances	Normalised Results	Percentage
1.a	4	0.079	7.843
2.c	4	0.079	7.846
3.e	8	0.158	15.691
4.h	2	0.040	3.923
5.i	10	0.198	19.614
6.l	2	0.040	3.923
7.r	4	0.079	7.846
8.s	5	0.099	9.807
9.t	4	0.079	7.846
10.u	4	0.079	7.846
11. All One-time Occurring Letters	4	0.079	7.846
Total	51	1.009	

The results show that all vowel letters excluding <o> are included under the errors in omission. The letter <i> is the most occurring letter followed by <e>. This result indicates that these two letters, which give different pronunciations and they are not pronounced in some cases, undergo the effect of phonology skills in English language. Concerning the consonants, the letter <s> is the most occurring word, and this might be the result of the high usage of the letter in English language. Examples include the omission of <e> in **creat*, omission of <i> in **profle*, and omission of <s> in **posses*.

In insertion, the same strategy as the omission is applied for the detailed results. Table (5) illustrates the distribution of insertion errors among the letters.

The detailed results for the most involved letters in)5(Table
insertion errors

Inserted Letters	Overall Instances	Normalised Results	Percentage
1.c	2	0.040	8.699
2.e	9	0.178	39.146
3.l	5	0.099	21.748
4.r	2	0.040	8.699
11. All One-time Occurring Letters	5	0.099	21.748
Total	23	0.455	

The results indicate that letter <e> is inserted into unrequired places with the highest frequency, and letter <l> comes in the second place. Inserting <e> in most cases may result from phonological concerns since the letter is silent in many cases. There also inter-lingual effects take account, especially in Kurdish, speakers use a kind of short vowel, such as schwa. In the second place, <l> is seen in five cases, but all the cases include doubling the letter where double letters are not required. Example include the insertion of <e> into the words such as **inferenceing* and **startes*, and insertion of <l> into the words **bellow* and **installment*.

In the errors of substitution where two letters are involved, usually, the occurrence of one case is more common. Table (6) shows the detailed results for errors in substitution.

Table (6) The detailed results for the most involved letters in substitution errors

Substituted Letters	Overall Instances	Normalised Results	Percentage
1.a by i	2	0.040	6.061
2.e by a	6	0.119	18.184
3.e by i	3	0.059	9.092
4.d by f	2	0.040	6.061
5.i by e	3	0.059	9.092
6.f by j	2	0.040	6.061
7.n by t	2	0.040	6.061
8. All One-time Occurring Substitutions	13	0.257	39.399
Total	33	0.653	

It can be inferred from the results that vowel letters occur in more cases than consonant letters. There is strange fact about these results; there are no cases where vowel is substituted by a consonant, and vice versa is also true. In all cases, only vowels or consonants are involved. The results also show that the substitutions in vowel letters occur in situations where the pronunciation may not highly change, such as the replacement of <e> by <a> in **weekend*, and the replacement of <i> by <e> in **enformation*.

In the same manner of substitution, errors in transposition involve two letters. In transposition, letters change their order, and the frequency is calculated by the number of cases for each disorder. Table (7) illustrates to what extent the transposed letters occur.

Table (7) The detailed results for the most involved letters in transposition errors

Transposed Letter	Overall Instances	Normalised Results	Percentage
1.t and a	2	0.040	10.527
2.l and e	2	0.040	10.527
3.r and o	3	0.059	15.790
8. All One-time Occurring Transpositions	12	0.237	63.162
Total	19	0.376	

The results again highlight that vowel letters are involved with high frequency in transposition errors. There are different cases in this kind of error; the transposition of <r> and <o> occur with the highest frequency, and in all cases, the error occurs in the word *from*, where the disorder of the letters happens. These cases might result from quick writing, especially the word *form* is found in the dictionary, and thus, computer does not highlight the word as incorrect. In most other cases the transposition happens due to the closeness of letters on computer keyboard. Another example from the corpus for this kind of error can be the disorder of <a> and <i> as in **conatiners*.

The least occurring error type is in the usage of apostrophes. Writers tend to make fewer errors in apostrophes, and this might be the result of lesser usage of apostrophes in academic writings. There are three ways in which the error happens including lack of the apostrophe, addition of the apostrophe, and using a wrong character. There is only one case where the apostrophe is not used, and there are no cases where apostrophe is added to the wrong place. The rest of cases occur in the usage of a wrong character;

instead of <'> writers used <'>. This might be the results of computers which in some cases the apostrophe is doubled like a quotation mark.

In another way, the number of involved letters in all kinds of errors can be calculated. This kind of calculation gives more insight into the letters and causes of the errors. For this purpose, the calculation of the involved letters, in ten and more cases, are taken into account throughout the errors in insertion, omission, substitution, and transposition. The detailed results are shown in table (8).

Table (8) The detailed results for the letters with ten and more error cases

Letters	Overall Instances	Normalised Results	Percentage
1.a	17	0.336	16.664
2.e	36	0.712	35.288
3.i	24	0.475	23.525
4.r	13	0.257	12.743
5.t	12	0.237	11.763
Total	102	2.019	

The above table clearly indicates that vowel letters involve more in errors much more than consonants. The letter <i> takes the lions share followed by <e>, <a>, <r>, and <t> respectively. The letters <i> and <e> both occur with a noticeable frequency, and this tends to be the result of the wide usage of these two letters in English. In another aspect, these two letters are silent, especially

<e> or create different pronunciations, and thus, they are involved a lot in spelling errors.

It seems that correct spelling resists to be challenging for MA students who share advanced skills in English language. The results highlight the issue of typing skills. These results indicate the existence of challenges in typing English among Fluent English speakers in the Kurdistan Region of Iraq.

There are various causes behind committing linguistic errors by non-native English speakers' language as a second language and as a foreign language. Based on the participants of this study, who are advanced non-native English speakers, it causes turn to be different to young learners of English. There are different causes behind spelling errors in English writing, which can be classified under the following main categories:

1. Interlingual: Mother language usually leaves effects on the use of the second language. Since languages vary in structure and orthography, learners and fluent speakers of English language need to master the orthography and structure to avoid errors. Kurdish Language as the mother tongue of the participants of this study is orthographically different to English which is the second language of the participants.
2. Intralingual: English is a language that the pronunciation of the words usually differs from the real sequence of letters. Thus, English learners in the early levels usually write what is heard, which causes spelling errors. However, Advanced English speakers are supposed not be involved in this case; MA students are still influenced by intralingual factors.

3. Technical: Immature in typing, assistance from non-English specialist typist, depending on computer autocorrecting feature, uncomfortable keyboard, un-customised writing software, late submission, quick revisions, mixing English letter buttons on keyboard with local language's, slip of the finger, and lack of revision.
4. Dyslexia: Is a disability which causes the affected people to make disorders in words. This cause may play a role in committing typing errors.

Thus, spelling errors are seen challenging, even in higher levels of English proficiency. This indicates that a deeper revision is advised for producing academic writings. This study suggests EFL and ESL instructors focus more on spelling and implement better modern procedures to decrease this issue. Also, MA students and scholars need to make more revisions before final publication of any product.

Conclusion

English as a globalised language is highly used in writing different academic productions. For the purpose of writing an MA thesis in the field of English linguistics, English language is the required language in public universities of Kurdistan Region of Iraq. This study hypothesised that spelling errors are challenging and are available within advanced writings. The results, finally, indicated that the study's hypothesis is relevant. Spelling errors can be in different forms; for that purpose, the study categorised the errors into different types and individually analysed each one in a detailed way. There are different causes behind errors. Spelling errors can mainly result in interlingual, intralingual, technical

skills, and Dyslexia. Finally, further and deeper investigations, on academic productions, is suggested for other types of linguistic errors.

References

- Anthony, L., 2018. Antconc 3.5.7.0: A free text analysis software. Available on line at <http://www.laurenceanthony.net/software/antconc/>. s.l.:s.n.
- Cook, V. J., 1997. L2 users and English spelling. *Journal of Multilingual and Multicultural*, 18(6), pp. 474-488.
- Corder, S. P., 1975. Error Analysis, Interlanguage and Second Language Acquisition. *Language Teaching and Linguistics: Abstracts*, Volume 8, pp. 201-218.
- Crystal, D., 2008. *A Dictionary of Linguistics and Phonetics*. 6th ed. New Jersey: Blackwell Publishing.
- Dawson, C., 2007. *A Practical Guide to Research Methods, A User Friendly Manual for Mastering Research Techniques and Projects*. 3rd ed. Oxfordshire: How to Books Ltd.
- Divsar, H. & Heydari, R., 2017. A Corpus-based Study of EFL Learners' Errors in IELTS Essay Writing. *International Journal of Applied Linguistics & English Literature*, 6(3), pp. 143-149.
- Hamed, M., 2018. Common Linguistic Errors among Non-English Major Libyan Students Writing. *Arab World English Journal*, 9(3), pp. 219-232.

- Hulvova, R., 2017. Error analysis of texts written by Czech EFL students. Brno: Masaryk University.
- Laframboise, L. K., 1996. Developmental Spelling in Fourth Grade: An Analysis of What Poor Readers Do. *Reading and Horizons*, 36(3), pp. 231-248.
- Liu, Y., 2015. Spelling Errors Analysis in College English Writing. *Theory and Practice in Language Studies*, 5(8), pp. 1628-1634.
- Richards, J. & Schmidt, R., 2010. *Longman Dictionary of Language Teaching and Applied Linguistics*. 4th ed. Great Britain: Pearson Publication.
- Seitova, M., 2016. Error Analysis of Written Production: The Case of 6th Grade Students of Kazakhstani School. *Antalya, Procedia - Social and Behavioral Sciences* .
- Vold, E. T., 2006. Epistemic modality markers in research articles: a cross-linguistic and cross- disciplinary study. *International Journal of Applied Linguistics*, 16(1), pp. 61-87.
- Walford, G., 2005. Research Ethical Guidelines and Anonymity. *International Journal of Research and Method in Education*, 28(1), pp. 83-93.

پوخته:

شتیکی ئاساییه بۆ فیرخوازانى زمانى دووهم كه هه‌له‌ى رینووسى له نوسینه‌كانیان دا ئەنجام بدەن، هه‌روه‌ك زۆر خه‌لكى تریش له سه‌راسه‌رى جیهان باسى ئەم گه‌رفته ده‌كهن. له لایه‌كى تروه‌ه، وا پێشبینى ده‌گریت كه ئاخێوه‌ره پێشكه‌وتووێه‌كانى زمانى دووهم هه‌له‌ى كه‌م یاخود هێچ هه‌له‌یه‌ك له به‌ره‌مه‌هه‌ نوسراوه‌ ئه‌كادیمیه‌كان به‌ زمانى دووهم ئەنجام نه‌ده‌ن بۆ نمونه له ماسته‌رنامه‌كان. ئەمه‌ش هه‌مووچارىك راست نیه له‌نیۆ ئاخێوه‌ره زمانپاراوێه‌كانى زمانى دووهم، چونكه ئه‌نجامدانى هه‌له‌ى زمانه‌وانى ئەم ئاسته‌ى زمانپاراوى تیپه‌رده‌كات و له نوسینه ئه‌كادیمیه‌كانیشدا به‌رده‌وام ده‌رده‌كه‌ون. ئەم توێژینه‌وه‌یه ده‌ستنیسانی ئەم كێشیه ده‌كات له نیۆ ئاخێوه‌ره ئه‌كادیمیه‌كان له هه‌ریمى كوردستان كه به‌رده‌وام پوه‌پرووى ئاسته‌نگه‌كانى ده‌بنه‌وه له پێشكه‌ش‌كردنى به‌ره‌مه‌یه‌كى نوسراوى بێ هه‌له‌ى زمانه‌وانى، له نیۆشیاندا به‌ تایبەت هه‌له‌ى رینووسى. ئەم توێژینه‌وه‌یه ئەم ماسته‌رنامه‌هه‌ له‌خۆ ده‌گریت كه له زانكۆ مه‌زنه‌كانى هه‌ریمى كوردستان و له‌سه‌ر بواری زمانه‌وانى ئینگلیزى نوسراون. تیکرای دووباره‌بوونه‌وه‌ى هه‌له‌كان وه‌رگیراوه‌و ئه‌نجامه‌كان به‌ شیوه‌یه‌كى چه‌ندایه‌تى شیکراونه‌ته‌وه. ئه‌نجامه‌كان ئه‌وه ئاماژه به‌وه ده‌كهن كه قوتابیانی ماسته‌ر هه‌له‌ى رینووسى له دواى وه‌شانى ماسته‌رنامه‌كانیان دا ئەنجام ده‌ده‌ن. به‌شیوه‌یه‌كى زۆر، هه‌له‌ رینووسیه‌كان له ئەنجامى نه‌بوونی به‌ه‌ره‌ى ته‌كنیکى و لیها‌توویى تایپکردن پووده‌ده‌ن، به‌لام كێشه‌كه لی‌ره‌دا ئه‌وه‌یه كه ئەم هه‌لانه دواى چه‌ندین پێداچوونه‌وه هه‌ر پووده‌ده‌ن.

ملخص:

من المعتاد أن يرتكب متعلمو اللغة الثانية أخطاء إملائية في الكتابة ، كما أشار الكثيرون إلى هذه المشكلة في جميع أنحاء العالم. ومن الناحية الأخرى، من المفترض أن يرتكب المتحدثون المتقدمون باللغة الثانية في الإنتاجات الكتابية أخطاء قليلة أو معدومة من هذا النوع باللغة الثانية ، خاصة في الكتابات الأكاديمية مثل أطروحات الماجستير، وهذا ليس صحيحاً دائماً بين المتحدثين بطلاقة حيث أن ارتكاب الأخطاء اللغوية يفوق هذا المستوى من الكفاءة ولا يزال يحدث في الإنتاج الأكاديمي. ويتناول هذا البحث المشكلة التي تواجه المتحدثين باللغة الإنجليزية الأكاديمية في إقليم كردستان العراق، إذ يواجهون دوماً عوائق هذه المشكلة في إنتاج كتاباتهم الأكاديمية الخالية من الأخطاء اللغوية ولا سيما الأخطاء الإملائية. يتضمن البحث رسائل الماجستير التي كُتبت عن اللغويات الإنجليزية من جامعات كبرى في إقليم كردستان العراق، وأخذ الأخطاء المكررة بعين الاعتبار، ثم عرض النتائج وتحليلها نوعياً. والنتائج تشير إلى أن طلاب الماجستير يرتكبون أخطاء إملائية في نسختهم النهائية من رسالة الماجستير. وإن هذه الأخطاء الإملائية تحدث عادة بسبب نقص المهارات الفنية وكفاءة الطباعة لدى الطلبة، ولكن المشكلة هنا تكمن في أن هذه الأخطاء تكرر بعد مراجعات عدة.