

**REPUBLIC OF TURKEY
CUKUROVA UNIVERSITY
INSTITUTE OF SOCIAL SCIENCES
ENGLISH LANGUAGE TEACHING DEPARTMENT**

**A CORPUS-BASED STUDY ON NOUN AND VERB AGREEMENT ERRORS BY
TURKISH LEARNERS OF ENGLISH AS A FOREIGN LANGUAGE**

AHMED ALI MALI SHUWADI

MASTER OF ARTS

ADANA / 2019

**REPUBLIC OF TURKEY
CUKUROVA UNIVERSITY
INSTITUTE OF SOCIAL SCIENCES
ENGLISH LANGUAGE TEACHING DEPARTMENT**

**A CORPUS-BASED STUDY ON NOUN AND VERB AGREEMENT ERRORS BY
TURKISH LEARNERS OF ENGLISH AS A FOREIGN LANGUAGE**

AHMED ALI MALI SHUWADI

Supervisor: Assoc. Prof. Cem CAN

MASTER OF ARTS

ADANA / 2019

Çukurova Üniversitesi Sosyal Bilimler Enstitüsü Müdürlüğüne,
Bu çalışma, jürimiz tarafından İngiliz Dili Eğitimi Anabilim Dalında YÜKSEK LİSANS
TEZİ olarak kabul edilmiştir.

Chair: Assoc. Prof. Cem CAN
(Supervisor)

Member of Examining Committee: Assoc. Prof. Dr. Hasan BEDİR

Member of Examining Committee: Assist. Prof. Dr. M. Pınar BABANOĞLU

ONAY

Yukarıdaki imzaların adı geçen öğretim üyelerine ait olduğunu onaylarım.

I certify that this thesis conforms to the formal standards of the Institute of Social Sciences
.../.../2019

Prof. Dr. Serap ÇABUK
Enstitü Müdürü

PS: The uncited usage of reports, charts, figures, and the photographs in this thesis, whether or original quoted from other sources, is subject to the Laws of Works of Art and Thought NO:5846.

Not: Bu tezde kullanılan özgün ve başka kaynaktan yapılan bildirişlerin, çizelge, şekil ve fotoğrafların kaynak gösterilmeden kullanımı, 5846 sayılı Fikir ve Sanat Eserleri Kanundaki hükümlere tabidir.

ETİK BEYANI

Çukurova Üniversitesi Sosyal Bilimler Enstitüsü Tez Yazım Kurallarına uygun olarak hazırladığım bu tez çalışmada;

- Tez içinde sunduğum verileri, bilgileri ve dokümanları akademik ve etik kurallar çerçevesinde elde ettiğimi,
- Tüm bilgi, belge, değerlendirme ve sonuçları bilimsel etik ve ahlak kurallarına uygun olarak sunduğumu,
- Tez çalışmada yararlandığım eserlerin tümüne uygun atıfta bulunarak kaynak gösterdiğimi,
- Kullanılan verilerde ve ortaya çıkan sonuçlarda herhangi bir değişiklik yapmadığımı,
- Bu tezde sunduğum çalışmanın özgün olduğunu,

bildirir, aksi bir durumda aleyhime doğabilecek tüm hak kayıplarını kabullendiğimi beyan ederim.../.../2019

Ahmed Ali MALI SHUWADI

ÖZET**YABANCI DİL OLARAK İNGİLİZCE ÖĞRENEN TÜRK ÖĞRENCİLERİN
YAPTIKLARI AD VE EYLEM UYUMU HATALARI ÜZERİNE DERLEME
DAYALI BİR ÇALIŞMA****AHMED ALI MALI SHUWADI****Yüksek Lisans Tezi, İngiliz Dili Eğitimi Bölümü****Danışman: Assoc. Prof. Dr. Cem Can****Eylül 2019, 101 sayfa**

Bu çalışma, yabancı dil olarak İngilizce öğrenen Türk öğrencilerin İngilizce ad ve eylem uyumu kullanırken yaptıkları hatalar incelendi. Çalışma Türk öğrencilerin Cambridge Öğrenci Derlemi içerisinde bulunan derlemde, CEFR'ye göre A1 – A2, B1 – B2, ve C1 – C2 dil yeterlilik aralığında yazılı anlatımlarında ilgili hataları derleme dayalı teknikleri kullanarak çözümlendi. Bu çalışma sürecinde betimsel ve deneysel olmak üzere iki farklı araştırma yöntemi kullanıldı. Çalışmanın deneysel bölümünde iki grup üzerinde çalıştık ve deney grubu öğretmeni İngilizce de ad ve eylem uyumunu öğretirken derleme dayalı teknikleri kullanıldı, kontrol grubu öğretmeni. Türk EFL öğrenenlerin kurumu analiz edilmiş, İngilizce ad ve eylem uyumu sıklığı, çalışmaların İngilizce ad ve eylem uyumunun aşırı kullanımı / altında kullanıldığı üzerinde Log-likelihood ve Type Token oranı ile belirlenmiştir, A1 – A2, B1 – B2, ve C1 – C2 yeterlilik seviyelerinin CEFR yeterlilik seviyelerine kıyaslayınca İngilizce ad ve eylem uyumunu aşırı kullandığı bulunmuştur. Ayrıca, derlemede bulunan hata tipolojileri göz önüne alındığında,

Anahtar kelimeler: Derleme dayalı dil Öğretimi, Öğrenci derleme, Hata çözümlemesi, ad ve eylem uyumu, İngilizce yabancı dil olarak öğrenen türk öğrenciler

ABSTRACT**A CORPUS-BASED STUDY ON NOUN AND VERB AGREEMENT ERRORS BY
TURKISH LEARNERS OF ENGLISH AS A FOREIGN LANGUAGE****AHMED ALI MALI SHUWADI****Master of Arts, English Language Teaching Department****Supervisor: Assoc. Prof. Dr. Cem CAN****September 2019, 101 Pages**

In this study, Turkish EFL learner's interlanguage noun and verb agreement errors has been investigated. The study has made use of Error Analysis in analyzing the data for noun and verb agreement errors. The present study has made use of the corpus-based techniques to analyze the use of noun and verb agreement in the Turkish EFL learner's written performance across A1 – A2, B1 – B2, and C1 – C2 CEFR levels of proficiency available in the related corpora of the Cambridge Learner Corpora. Two different research methods were used in this study, both descriptive and experimental method. Two different groups were involved in this study, experimental and the control group. Experimental group were taught noun and verb using the corpus data-driven techniques; and the control group were taught in a traditional way of teaching noun and verb agreement. The result of the experiment was determined by the test on noun and verb agreement taken by the participants after the experiment. In addition, the overuse, underuse and the errors made in using noun and verb agreement by the Turkish EFL learners were analyzed using Log-Likelihood and Log-Dice in investigating and finding out the statistic evidence upon the possible errors committed by the Turkish EFL learners. In this study Turkish EFL learner's use of noun and verb agreement in their academic writing has been checked and comparison of the result have been made between the control and the experimental group, however the result has revealed that Turkish EFL learner overuse noun and verb agreement in their writings. In addition corpus-based implantation on noun and verb agreement errors has been done to see if the participants could master the rules or not. the result have shown that the both group have improved, but the experimental group overused noun and verb agreement more than their counterpart control group. Firstly Overall frequencies have been identified and analyzed in a tabular form both CLC and the Corpora compiled by the

researcher across A1 – C2 proficiency band of CEF. The data were analyzed using the soft wares called Sketch Engine, AntConc, Log-dice and Log-likelihood in order to compare the result. Secondly, the frequencies of all the categories A1 – C2 of the CLC data as well as the corpus compiled by the researcher during his implementation were analyzed respectively, and each proficiency level frequency is shown respectively.

Keywords: Corpus-based language teaching, Learner corpus, Error Analysis, noun and verb agreement, Turkish EFL Learners,



ACKNOWLEDGEMENT

I would like to express my special thanks and gratitude to a number of very important people without whose support and guidance this thesis could not be completed. First of all I would like to use this opportunity to thanks and show my sincere gratitude to my supervisor Assoc. Prof. Dr. CEM CAN for his helps tirelessly and guidance throughout the writing process of this thesis; words are not enough to thank you sir. Apart from thesis he always gives us great support in term of linguistics as well as study materials.

Secondly I would like to use this chance to thank Assoc. Prof. Dr. Hasan BEDİR for being my supervisor in my first year in MA, and also for showing us the quantitative way of doing research, and I would like to show my gratitude and respect to Assoc. Prof. Dr. Neşe CABAROĞLU for her contribution in teaching us the qualitative way of research method. I want to thank a special friend of mine Abdulkadir Abdulrahim for his help and whom we struggled together and always accompanied me during our sleepless nights in our thesis writing process.

Bunch of thanks to Assoc. Prof. Dr. Rana YILDIRIM for guiding us through documentation process in the international student's office. I am also grateful to one of the most amazing person, Research Assistant and an academician Tuğba ŞİMŞEK for always being there for me whenever I needed help.

Special thanks goes to our head of department Professor Dr. Zuhale OKAN whose smiles keep us motivated and who always listen to our complaints and finds a way to help solve the problem.

I am extending my deepest thanks to Professor Dr. Hatice SOFU for her great contribution, guidance and support in the field of language acquisition and linguistics in whole.

My endless gratitude goes to my elder brother Alhassan Ali for always being there to support me Economically, who has never say no to my request for support, he is a person whom without I cannot make it to this position, someone who never gets tired of advising me on what step to take in life. I will not forget to say my Words of thanks to other elder brother of mine Mamman Ali and Abubakar Ali for always calling to ask about my health condition.

I also thank all of the Cukurova University ELT family for making me feel at home away from home; whenever I leave Turkey, I will surely miss this home.

DEDICATION

I dedicated this study to my late father who departed to his better home, the man who was always praying for me and proud of me, but couldn't make it to see the end of my thesis, the one and only strong man who taught us what life is and how to cope with people around us in a peaceful way, the man whom we always learn from. I pray for him, may God shower his infinite mercy upon him, forgive his shortcomings and may He rest in peace. Secondly I am dedicating this study also my Mum, my everything, my world as well as my biggest source of happiness, whose my coming to Turkey is never her wish due to its being far away from her. The one and only strong woman who could not hold her tear whenever I am leaving home for school.

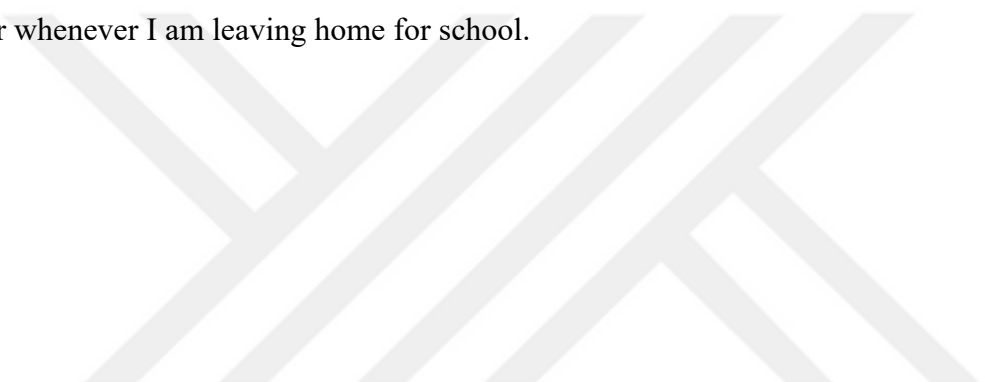


TABLE OF CONTENT

ÖZET	iv
ABSTRACT	v
ACKNOWLEDGEMENT	vii
DEDICATION	viii
TABLE OF CONTENT	ix
LIST OF ABBREVIATION	xii
LIST OF FIGURES.....	xiii
LIST OF TABLES.....	xiv

CHAPTER 1 INTRODUCTION

1. Introduction	1
1.1. Background of the Study	1
1.1.1. Corpus linguistics	1
1.2. Statement of the Problem	2
1.3. Aim of the Study	4
1.4. Importance of the Study	4
1.5. Research Questions	5

CHAPTER 2 LITERATURE REVIEW

2. Introduction:	6
2.1. Corpus Linguistics.....	6
2.2. Corpus-based Approaches and Studies.....	10
2.3. Corpus and Language Teaching	10
2.4. Corpora	11
2.5. Types of Corpora.....	11
2.5.1. Learner Corpora.....	11

2.5.2. General Corpora	13
2.5.3. Parallel Corpora.....	13
2.5.4. Comparable Corpora.....	13
2.5.5. Specialized Corpora.....	14
2.5.6. Historical or Diachronic Corpora	14
2.5.7. Monitor Corpora	14
2.6. Corpus-based Error Analysis.....	14
2.7. Noun (subject) -Verb Agreement	16
2.8. Approaches to Noun and Verb Agreement.....	17
2.9. Data-Driven Learning	18
2.10. Approaches to Data-Driven Learning (DDL).....	20
2.11. Criticism against Data-Driven Learning (DDL).....	21
2.12. Contrastive Analysis.....	21
2.13. Interlanguage	22
2.14. Language Transfer.....	23
2.15. Error Analysis.....	23
2.16. Importance of Error analysis	26
2.17. Computer-aided Error Analysis.....	27
2.18. Traditonal Error Analysis:	28
2.19. Subject-Verb Agreement Error.....	28

CHAPTER III

METHODOLOGY

3.1. Introduction	30
3.2. Research Design	30
3.3. Participants:	31
3.4. Research Context.....	31
3.5. Data Collection Tools.....	32
3.5.1 Log-likelihood	32
3.5.2 Type-token Ratio	34

CHAPTER IV FINDINGS

4. Introduction	36
4.1. Overall frequencies of noun and verb agreement across A1- C2 Proficiency level of ICLE and the Corpus Compiled by the Researcher.....	36
4.2. Finding of the frequency and distribution of Noun and Verb Agreement error in CLC37	
4.3. Finding of the frequency and distribution of Verb Agreement error in CLC.....	43
4.4. The analysis of the data compiled by the researcher from the CONTROL and the EXPERIMENTAL group during the implementation study	51
4.5. Frequencies of Noun and Verb in Pre-essay and Post-essay Type and Token across A1 – C2 proficiency level.....	67
4.5.1. Introduction	67

CHAPTER V DISCUSSION

5.1. Introduction	69
5.2. Discussion.....	69
5.3. Chapter summary.....	71

CHAPTER VI CONCLUSION

6.1. Introduction	72
6.2. Conclusion.....	72
6.3. Limitations of the Study:	73
6.5. Suggestions for Further study:.....	73
6.5. Chapter summary.....	74
REFERENCES	75
APPENDICES.....	80

LIST OF ABBREVIATION

ARCHER	: A Representative Corpus of Historical English Registers
CA	: Contrastive Analysis
CCEC	: Collins Co-build English Collocation
CEA	: Computer-aided Error Analysis
CEFR	: Common European Framework
CIA	: Contrastive Interlanguage Analysis
CLC	: Cambridge Learner Corpus
DDL	: Data-driven Learning
EA	: Error Analysis
EFL	: English as a Foreign Language
ELL	: English Language Learning
ELT	: English Language Teaching
ESL	: English as a Second Language
ESP	: English for Special Purpose
FL	: Foreign Language
IH	: Interface Hypothesis
IL	: Interlanguage
L1	: First Language
L2	: Second Language
L2L	: Second Language Learning
LL	: Log-likelihood
MT	: Mother Tongue
NNS	: Non-native Speaker
NP	: Noun Phrase
NS	: Native Speaker
SLA	: Second Language Acquisition
TL	: Target Language
SOV	: Subject, Object, Verb
PST-1SG1 ...	: First person singular past
PST-2SG	: Second person past
PST-3SG	: Third person past
PST-1PL ...	: First person plural
PST-2PL	: Second person plural past
PST-3PL	: Third person plural past

LIST OF FIGURES

	Pages
Figure 1. Screenshot of Log-likelihood Calculator	33



LIST OF TABLES

	Pages
Table 1. Data Collection Tools, Procedure, and time of Implementation/application period.....	32
Table 2. Contingency Table for Log-likelihood Calculation (adapted from Rayson and Garside, 2000, p.3).....	33
Table 3. Turkish EFL learner's Sub-corpus Breakdown. Demographic information is shown in the table below:.....	35
Table 4. Overall CLC List of noun errors (wrong usage) and the correct form A1 – C2 proficiency level.....	37
Table 5. List of noun errors (wrong usage) frequencies and the correct form in A1 proficiency level of CEFR	38
Table 6. List of noun errors (wrong usage) frequencies and the correct form in A2 proficiency level of CEFR	38
Table 7. List of noun errors (wrong usage) frequencies and the correct form in B1 proficiency level of CEFR	39
Table 8. List of noun errors (wrong usage) frequency and the correct form in B2 proficiency level.....	40
Table 9. List of noun errors (wrong usage) frequency and the correct form in C1 proficiency level.....	41
Table 10. List of noun errors (wrong usage) frequency and the correct form in C2 proficiency level.....	42
Table 11. The CLC overall list of verb errors (wrong usage) frequency and the correct form in A1 – C2 proficiency level of CEFR.	43
Table 12. List of verb errors (wrong usage) frequency and the correct form in A1 proficiency level.....	44
Table 13. List of verb errors (wrong usage) frequency and the correct form in A2 proficiency level.....	45
Table 14. List of verb errors (wrong usage) frequency and the correct form in B1 proficiency level.....	46
Table 15. List of verb errors (wrong usage) frequency and the correct form in B2 proficiency level.....	47
Table 16. List of verb errors (wrong usage) frequency and the correct form in C1 proficiency level.....	48

Table 17. List of verb errors (wrong usage) frequency and the correct form in C2 proficiency level.....	48
Table 18. Frequency and Log – Likelihood ratios of noun agreement errors in ICLE.....	49
Table 19. Frequency and Log – Likelihood ratios of verb agreement errors in ICLE.....	50
Table 20. The overall list of NOUN errors (wrong usage) frequency and the correct form in A1 – C2 proficiency level of CEFR.	51
Table 21. List of noun errors (wrong usage) frequencies and the correct form in PRE and POST CONTROL A1 – C2 proficiency level of CEFR compiled from the implementation by the researcher	52
Table 22. List of noun errors (wrong usage) frequencies and the correct form in PRE and POST EXPERIMENTAL A1 – C2 proficiency level of CEFR compiled from the implementation by the researcher.....	53
Table 23. The overall list of NOUN errors (wrong usage) frequency and the correct form in A1 – C2 proficiency level of CEFR.	53
Table 24. List of Verb errors (wrong usage) frequencies and the correct form in PRE and POST CONTROL A1 – C2 proficiency level of CEFR compiled from the implementation by the researcher	55
Table 25. List of Verb errors (wrong usage) frequencies and the correct form in PRE and POST EXPERIMENTAL A1 – C2 proficiency level of CEFR compiled from the implementation by the researcher.....	56
Table 26. The frequency and the Log-Likelihood ratio values comparison of A1 – A1 of both AGN and AGV errors in the experimental and the control group's corpus.....	61
Table 27. The frequency and the Log-Likelihood ratio values comparison of A2 – A2 AGN and AGV errors in the experimental and the control group's corpus compiled by the researcher	61
Table 28. The frequency and the Log-Likelihood ratio values comparison of B1 – B1 of both AGN and AGV errors in the experimental and the control group's corpus	62
Table 29. The frequency and the Log-Likelihood ratio values comparison of B2 – B2 of both AGN and AGV errors in the experimental and the control group's corpus	62
Table 30. The frequency and the Log-Likelihood ratio values comparison of C1 – C1 of both AGN and AGV errors in the experimental and the control group's corpus	63

Table 31. The frequency and the Log-Likelihood ratio values comparison of C2 – C2 of both AGN and AGV errors in the experimental and the control group’s corpus	63
Table 32. The frequency and Log-Likelihood ratio of A1 – A2 vs B1 – B2 AGN and AGV errors in ICLE.....	64
Table 33. The frequency and Log-Likelihood ratio of A1 – A2 vs C1 – C2 AGN and AGV errors in ICLE.....	64
Table 34. The frequency and Log-Likelihood ratio of A1 – A2 vs C1 – C2 AGN and AGV errors in ICLE.....	65
Table 35. Frequencies and the overall Log-Likelihood ratio statistic of both noun and verb agreement errors across A1 – C2 proficiency level of EXPERIMENTAL group	65
Table. 36. Frequencies and the log-likelihood statistic result of noun and verb errors across A1 – C2 proficiency level of CONTROL group in total.	66
Table 37. Frequencies of Noun and Verb in Pre-essay Type and Token across A1 – C2 proficiency level of PRE- EXPERIMENTAL group.....	67
Table 38. Frequencies of Noun and Verb in Pre-essay Type and Token across A1 – C2 proficiency level of POST- EXPERIMENTAL group.	67
Table 39. Frequencies of Noun and Verb in Pre-essay Type and Token across A1 – C2 proficiency level of PRE – CONTROL group	68
Table 40. Frequencies of Noun and Verb in Pre-essay Type and Token across A1 – C2 proficiency level of POST – CONTROL group	68

CHAPTER 1

INTRODUCTION

1. Introduction

This chapter will give us detail information on the background of the study by taking into consideration, the possible error made by Turkish EFL learners, the definition and the history of corpus linguistics and the recent studies made in the field of corpus, the statement of the problem, the aim of the study and the importance of the study will be discuss as well.

1.1. Background of the Study

1.1.1. Corpus linguistics

Nowadays writing has become a very difficult task not only for students but also for those learning English for some purposes such as: those intended in taking an important English proficiency test, those who are learning English to do business, those learning English for communicating purposes while they travel abroad as well as those learning to improve themselves and get promoted in their various working companies. This can be possible when corpus studies are widely carried out and implemented in various field as well as schools.

Many linguists and researchers define corpus linguistics in many different ways. The definition of corpus linguistics according to Granger (2002 p.4) she states that corpus linguistics can best be defined as a linguistic methodology, which is founded on the use of electronic collections of naturally occurring texts, corpora. Charles (2011) claims that all of the corpora are based on written texts and have used the same design criteria to allow comparisons to be made across the two varieties of English and across time. Granger (2002 p.4) in her study, she states that corpus is not a new branch of linguistics and not a new theory of language; however, the concrete evidence it uses proves it to be a particularly powerful methodology, which is likely to change perspective on language.

Nadja (2005) defines corpus linguistics not as a methodology, but as the analysis of authentic text on the basis of computerized corpora. Najda claims that “mostly, the analysis is done with the aid of the computer, i.e. with a specific software, and measures the frequency of the language material been investigated.”

In addition, this study makes use of corpus data and corpus-based techniques to investigate the errors committed by the Turkish second language learners. Learners mostly make errors due to the intralingual errors, looking into the previous studies; learners commit errors in their placement of verb after the subject. In order to produce meaningful language, noun and verb have to agree in every point. Rodrigue (2015, p. 22), points out that. "The use of singular and, plural verbs and correct noun and verb agreement are more abstract concepts and require wide explanations." So, for every second language learner mistake on agreement is natural. Barkhuizen (2005 p. 5) concludes that the goal of SLA research for many researchers is to describe the language development of L2 learners and to explain how it develops over time. Moreover, nouns and verbs need to agree on number. The number of the noun singular or plural determines the form of verb. In any piece of writing, if the subject appeared to be a singular noun, the verb that follows the subject has to be singular, and if the subject is plural noun, the verb that follows the subject must be plural. Second language learners mostly make mistakes in determining the right position in placement of verb after noun; looking into the previous studies on learner corpora we can easily detect the learner errors.

In English and many other languages, verbs must agree in number with their subjects, as in the contrast between the dog barks and the dogs bark. (Haskell & MacDonald, 2003 p. 760) Gao, et al. (2000) in their study they stated that "In the most straightforward cases the noun-verb agreement rule directs us to use the third person singular inflection if the subject is a singular proper noun, a singular common noun, a mass noun, or a third person singular pronoun."

1.2. Statement of the Problem

Writing is one of the most important skill in language, but it is also one of the field where most of the students have problem and difficulty in. The problem is not limited to only student; it has extended up to some private learners of English as a Foreign Language. As a result noun and verb agreement play a very important role in day to days writings, without giving them a great attention while writing, one can make his reader or audience have difficulty in understanding the message that a writer or speaker is trying to pass. In order to pass a clear message in one's writing, a great attention must be given to noun and verb errors.

This study aims at investigating the noun and verb agreement error made by the Turkish EFL learners in their academic writings compared across their various proficiency levels in reference to CEFR band. In order to find out the frequency of the Turkish EFL learners noun and verb errors a 5 week-corpus-based implementation has been conducted.

According to Granger et al. (2002) they shade light on the importance of input to leaners, she states that for teachers and material designers to be able to provide the most useful input to the leaners, they need large amount of information on what learners may be expected to have learned by what stage and analyzing errors is a valuable source of information.

Can (2017) claims that analyzing interlanguage errors and carrying out error analysis (EA) on texts produced by learner are practicable thanks to readily accessible learner corpora. The emergence of interlanguage corpora has contributed to the SLA field, using electronic learner corpus we could investigate learner's interlanguage errors in detail and analyze them. Corpus gives us the opportunity to obtain learners authentic errors and how frequent they occur along with various proficiency levels, from different pedagogical and linguistics aspects. It is of crucial important to consider these authentic leaner errors, so that a possible solution to tackle these problems could be found. This study aims at analyzing the Turkish EFL learner's noun and verb agreements errors across A1-C2 proficiency levels based on CEFR.

There will be time, of course, where learners will get subject and verb agreement correctly, too, but still illogically. To give a typical example, most of intermediate students will be moving back and forth between correctly use and omission of the third person S, instead of producing a concrete or correct forms. The major problem here is how accuracy also seems to depend on the verb at hand. Have you not seen a situation where students get certain verbs right more frequently than others? "He ates", "She pays" and "Yusuf likes" are much more likely to be used correctly than, the usage of for example "He sees", "She goes" or "Ali watches."

Second language learners most of the times have to think twice before producing the right form. Such problem is always there in the leaners process of second language acquisition, one of the reason might be in the case that, some specific verbs create an easy and suitable phonetics sound that makes them sound more "third person-like" than others. If our students make a noun and verb agreement error, we shouldn't directly think that it is just a performance breakdown at a certain level, thinking that it is just a slip of tongue, for

example, “He like banana”, “She have a toy”, “My friends like”, “Everybody know” and “She always play, and many more examples could be found when research is extended.

1.3. Aim of the Study

This study aims at analyzing the errors related to the use of noun and verb agreement by Turkish EFL learners across six proficiency levels, A1-A2; B1-B2; C1-C2, respectively in reference to the Common European Framework of Reference for language (henceforth CEFR) (council of Europe 2001). According to grammatical rule the verb must agree in number with its subject for example when we take a look at English, present tense rules, by adding an S or ES verbs changes to show agreement in the third person singular form. So how is it possible that student of all levels, from different countries and age groups seem to commit error more frequently than would seem reasonable? It is just adding a letter at the end of the verb; second language learners find it difficult to place verbs after nouns.

The corpus used in this study is the Cambridge Learner Corpus (CLC), the largest defined test performance corpora that lets the investigation of the linguistic and rhetorical features of the learner performances in CEFR proficiency bands. This study also aims at improving Turkish EFL learners noun and verb use with a corpus-driven 5 weeks implementation

1.4. Importance of the Study

This study aims at finding out the noun and verb agreement errors Turkish EFL learners make, and to find out the possible reasons behind those errors. Ellis (1994) emphasizes "the significant of gathering well-defined samples of learner language in order to make a clear statement regarding the types of errors the learners made and under what circumstance" this way, through analyzing collected data of the learners help in figuring out the student's progress. As learner corpora have nowadays become readily and easy to access, it gives us the opportunity to examine and analyze interlanguage errors on learner-generated texts. (Can 2017).

This provides the natural sample of the language, using the learner, corpora helps us in determining student's attainment towards the target language. Having and analyzing the authentic learner produced data gives us the opportunity to see the problems faced by the learners as well as their progress in language learning. To come up with the possible

solution to the problems and errors made by learners, there is a need for more research in the field of learner corpora in terms of pedagogical and academic demands. Therefore, this study has contributed a lot in finding out the problems faced by Turkish EFL learners in connection with noun and verb agreement.

1.5. Research Questions

The Research questions of the study are as follows:

1. What is the distribution of noun and verb agreement error types along with the A1-C2 Proficiency range according to CEFR?
2. What are the most common and persistent noun and verb agreement error types in Turkish EFL learners?
3. Do the corpus-driven techniques implemented in teaching noun and verb agreements in English yield a statistically significant difference between the success levels of experimental and control groups?

CHAPTER 2

LITERATURE REVIEW

2. Introduction:

This chapter provided an overview of the definition of corpus its applications, and to review publications and current study related studies in the area of noun and verb agreement errors as well as that of corpus linguistics and language teaching. Then it has brought to us some background information about learner corpora, corpus analysis, corpus-based studies of noun-verb agreement. In addition to that, it went on by giving more light on interlanguage errors committed by the Turkish EFL learners in their process of second language acquisition/learning in details.

2.1. Corpus Linguistics

Getting into corpus and putting corpus-based materials into practice cannot only help the teacher to benefit with its nature of being authentic, but it can also help the teachers know what actually their students need to know and this can help the teachers bridge the gap in teaching and learning as well as involved both the teachers and the students hand on practically, and this can bring about a great effect in the pedagogic field.

The few years old computer learner corpus (CLC) research as a field of scientific research (it emerged as a field recently in 1980s) rises the disputable assessment of its achievements to some degree premature. There is a need for teachers and language learners to apply corpus data themselves, for teachers in their language teaching and for learners in their learning process, by doing so, it will help both parties providing them with self-reliance. There will be no need for them to look for help from the native speaker of the target language. According to Dazdarevic & Fijuljanin, (2015) they claim that, corpus shows how language is used in real life and puts an end to the so-called necessity of relying on a native speaker's intuition to tell what is commonly or rarely used in English. Granger, (2004) also states, however, many research has been done to show the progress made in the field and the possibility of doing better in future.

According to Granger (2002), corpus-based researches of second language learner varieties have been a recent distinctive: academics and publishers have long started collecting

corpora for non-native since before 1980s and 1990s, which is now referred to as learner corpora.

According to Hunston, (2006) he defines corpus as an electronically stored gathered samples of naturally occurring text” The basic resource for corpus linguistics is a collection of texts, called a corpus. Corpora can be in different sizes, and are gathered for different reasons, and are composed of texts of varying types. Hunston (2006) gives a further important explanation on the sizes of corpora, he claims that most modern corpora are at least 1 million words in size and consist either of full texts or of big collection of extracts from long texts. However, “there are many different ways in which a text corpus can be studied with the use of computer software programs; in other words, the corpus is ‘re-created inside the computer” (Charles, 2011)

There is not one specific definition of corpus, as Charles mention in his study. “A number of researchers have asked and defined corpus from their various perspectives.” Charles, (2011) according to him, one of the most debated questions about corpus linguistics is: is corpus linguistics a branch of linguistics (such as phonology, grammar, semantics, lexical studies) or a methodology of language studies? According to Hunston corpus linguistics is regarded as a sophisticated method of finding answers to the kinds of questions linguists have always asked” (Hunston, 2006)

According to Granger (2002), in her own study states that “corpus linguistics is a linguistic methodology which is based on the use of electronic collections of naturally occurring language, in form of text called corpora” according to her it is not a new field of linguistics and also not a new theory of language, but the reality about it is that it uses makes it a kind of strong methodology, which has the potential to amend views on language. Many researchers believed that Granger’s definition is an informative and one of the widely accepted definitions of corpus linguistics.

Boulton, (2012) highlights that generally corpus is defined as a big collection of reliable texts in electronic form designed represent a language variety. According to Charles increasingly, multi-modal corpora are becoming important in the areas of corpus linguistics. “Corpora also include ‘video, audio and textual records of interaction (and related metadata information) found from recordings of real life dialogues” (Charles, 2011).

Díaz-negrillo, (2013) has also states that corpus linguistics is said to be the collections of naturally occurring text. As these collections are computerized, they can

easily be searched and manipulated, and, because of their size, they provide a reliable basis on which to describe and model learner language use.

Warfield, (2012) states although corpus linguistics is strictly speaking a methodology rather than a theory of language, it has opened up new approaches to the study of language and new fruitful ways of matching theory and data. Charles (2011) added “corpus linguistics is viewed by some as an experienced way of linguistic analysis and description, using natural examples of language data stored in corpora as the beginning.” A famous researcher described a large corpus as “a test bed for hypotheses and can be used to add a quantitative dimension to many linguistic studies” (Hunston, 2006)

In other Word, Corpus linguistics is the computerized study of language authentic text, obtained from native and non-native language of learner corpora. Learner corpora is the natural text produced by language learners, learner corpora had contributed a lot in terms of material development. Therefore, learner corpora had a direct impact on language teaching, language learning and language learners. Taking into consideration the teaching and learning materials being created out of learner corpora in the field of material designing and development, this helps and makes the second language learning easier. In addition to tracking English language changes and production of dictionaries, textbooks and other learning and reference materials, a major application of corpora is to study different aspects of linguistics, including lexis, multi-word units, grammar, discourse, pragmatics, discourse intonation, registers etc. (Charles, 2011). Römer, (2008) emphasizes that “in applied linguistics, most of the researchers and learners treasure what corpus linguistics provides to language pedagogy.”

According to Leech, (2009), the term corpus linguistics can be simply defined “as ‘the study or analysis of language through the use of (computer) Corpora.’” Charles, (2011) also added that, A corpus usually can be defined as a collection of different language text in form of electronic, chosen based on different criteria to serve, as a language or language variety or as a way of finding data for linguistic research.

Some criticism toward corpus approaches. Many researchers note down the short comings of the corpora from their various points of view to start with Römer, (2011) claims that corpora have not been an important resource in the design of language teaching syllabi that emphasize communicative competence. Hunston, (2006) also points out that, it is obvious that corpus methodologies are typically quantitative. Well, corpus linguistics has been criticized for focusing on only the relative quantity and for not widening the explanatory capability of linguistic theory. Moreover Römer, (2011) notes that, “I would,

continually, still be doubtful to say that corpora and corpus tools have been completely applied in pedagogical contexts and would argue that there are still works to be done in filling the gap between research and application.” Here he emphasizes that more research and work need to be done for implementing and application of research result by putting it into practice in the field of language learning and teaching as well as material development. In addition, Römer (2011) claims that the implementation of English language teaching (ELT) till now, is likely to be only marginally affected by the progresses of corpus research, and similarly few teachers and learners know about the presence of useful resources and put into practice the corpus computers or concordances themselves. He is calling on the widely implementation of the corpus knowledge and ideas in all educational levels by both the teachers and the learners, this will enhance the knowledge and practice of corpus in the area of pedagogy.

While most of corpus researchers including Römer, (2011) states that corpus linguistics has a great power to help improve language learning and teaching methodologies, but they do not most of the time make enough efforts to reach the users with the main reason of corpus and to figure out what teachers exactly want and need.

Economic situations can also bring about some drawbacks in the area of corpus development, for example, in this case we can think of some under-developed countries and villages, for them access to the modern technology may be difficult. Boulton, (2012) who notes that, “the Internet has made the corpora and many related tools within reach of researchers teachers and learners” he further claims that however, many criticize that many of them require a lot of capital and effort for the training of learners (and teachers) to understand the justification as well as how their usage are efficiently. (Boulton, 2012)

Hunston, (2006) in his famous study, states that corpora are examine through the use of devoted software, a kind which necessarily reflects assumptions about methodology in corpus research. At the most basic level, corpus software: he listed three methods in corpus investigation as follows:

- . Searches the corpus for a given target item,
- . Counts the number of instances of the target item in the corpus and calculates relative frequencies,
- . Displays instances of the target item so that the corpus user can carry out further investigation.

(Hunston, 2006, p.234)

2.2. Corpus-based Approaches and Studies

Römer, (2011) states that, “for the recent few decades, corpora have not only reformed linguistic research but have also changed perspective in on second language learning and teaching.” The benefit of corpus approach is that learners will learn the form of the foreign language because they are involve in analyzing form of the native language on the basis of natural content. (Dazdarevic & Fijuljanin, 2015)

In applied linguistics, most of the investigators and learners give important to what corpus linguistics can add to the development of language pedagogy? Corpora and corpus tools are yet to be largely applied in pedagogical contexts. (Römer, 2011)

Leech, (2009), states that “corpora have, of course, been extensively used for investigating the history of the language before.” In their study McEnery & Wilson (1996: 12) notes that even though corpora is one of the source of proves among many, complementing rather than changing other data sources such as introspection and elicitation, there is common agreement nowadays that they are the only reliable source of proves for frequency.

2.3. Corpus and Language Teaching

Dazdarevic & Fijuljanin, (2015) note that advancement in electronic and computer sciences have developed the view toward language teaching, One of the most important contributions of the computer sciences to language teaching has been noticed within corpus linguistics in building, implementing, and analyzing language corpora through computer systems. Römer, (2008) points out that in the past two decades, corpora and corpus evidence has not only been used in linguistic studies but also in the pedagogical purposes. Römer, (2008) in his study on corpora and language teaching, further shades more light on the use of corpora, he claims that when we talk of implementation of corpora in language teaching, this involves the use of corpus tools, i. e. the actual text collections and software packages for corpus access, and of corpus methods. However for learners to benefit from the use of corpora, Dazdarevic & Fijuljanin, (2015) emphasize that language teachers must be given a sound knowledge of the corpus-based approach, so the teacher sets the student this task, to explore some patterns in some of the corpora, and gets them to formulate the usage rules for this form. Teachers as the instructors they need to be well trained in order to pass the information and guide their students rightly through the corpus-based methodology in their process of language learning and.

Römer, (2011) claims, “even though the progress that has not been investigated is made in the field of applied corpus linguistics, there is still chance for a better change” Through corpus we can achieve a lot in improving language learning and teaching, but how? Many of tasks can be formulated for the development of pedagogical corpus applications and help corpora makes more powerful contribution to improve language learning and teaching. Among the benefits, we can take vocabulary grammar as some of the great example of corpus positive impact. Dazdarevic & Fijuljanin, (2015) points out that using corpus in language teaching can be an effective tool in teaching vocabulary, grammar and language use to learners of English as a foreign language.

2.4. Corpora

The achievement of developing language corpora that bring about the readily access to quantifiable language use in naturalistic settings have been largely accepted by many scholars in a various set of language fields.(Francom, LaCross, & Ussishkin, 2010)

2.5. Types of Corpora

2.5.1. Learner Corpora

Granger (2004) in her own terminology she defines computer learner corpora as electronically collected spoken or written texts produced by foreign or second language learners. Then the data can be stored in electronic form; a large amount of it can be gathered in a short period of time. Due to its easy collection and non-time consuming, nowadays learner corpora are counted in the millions not even in the hundreds or thousands of words. The basic information that computer software can find in a corpus is the word frequency list, all corpora do is reveal which words are used in their constituent texts and how frequent they are (Charles, 2011).

Granger (2002) defines corpora here in relation to language settings. Computer learner corpora are electronically gathered spoken or written texts uttered by anyone who English serve to be his or her target language or second language learners in a various language settings. When it's compiled in a computer, then the data can be analyzed with linguistic software tools, from the beginner, which are the easier ones to the harder ones, which are the ones from advance language learners.

Corpora can give a lot to teaching and learning, it can help progress our knowledge of how language works as well as how it develops over the time. Which makes the

language describable through various reference materials (Boulton, 2010) Díaz-negrillo, (2013) criticized and overstates the case by claiming that data obtained under experimental conditions do not qualify as learner corpus data on the grounds that such production does not equate to authentic language use.

Learner Corpora are also used to carry out the contrastive interlanguage analysis, by comparing the L1 and L2 students authentic spoken or written text, it can also be used to investigate the errors committed by second language learners. Learner corpora brings a new type of data which brings about thinking both in SLA research, which tries to find out the structure of foreign/second language learning, and foreign language teaching research, the reason is to develop the learning and teaching of foreign languages. (Granger, 2002)

Corpora have provided language learners and teachers with a lot of potentials to do more in the field of language learning and teaching, it can help learning and teaching develop beyond our expectation.

Research on learner corpora helps us in measuring our student's progress in their language learning process. As Barkhuizen (2005), points out that the aim of second language learning research for various researchers is to find out and show the process of the language development of second language learners and to give light on how it develops in time.

A great consideration on language instruction was given in the literature, Corder (1967), and Selinker (1972). Corder focused on the language instruction and brought the idea of an 'internal syllabus' i.e all second language learners would have the same kind of internal system for language acquisition as what children acquiring their L1 do, teaching benefit from the findings of linguistics in many cases, including error analysis and subject-verb agreement

Another thing to do is to give attention to the needs of learners and see which groups of learners may benefit most from which kind of materials. (Römer, 2011) Römer, advises teachers as language instructors on material usage, materials must be carefully selected, considering students level as well as their capabilities. Moreover, a further important thing is problems on student's willingness towards corpora. Römer, (2011) states, "In relation to this problem are questions directed to the learners' eagerness and ability to handle computer corpora, online searches, and the exercises given by their teachers." Granger (2002) points out that the concept of authentic language learner production is problematic as much language learner output is produced in structured learning environments.

2.5.2. General Corpora

In corpus linguistics, a corpus is often described as being either 'general' or 'specialized'. General corpora are usually much bigger than specialized corpora. (Charles, 2011) General corpora' contained general texts, texts that are not belong to a single text type, subject field, or register.

Large general corpora can be define as reference corpora because they usually function as baseline against which judgments about the language classification are held in more specialized corpora.(Evans, 2004) What distinguishes general corpora from specialized corpora is the purpose for which they are compiled. General corpora aim to examine patterns of language use for a language as a whole, and specialized corpora are compiled to describe language use in a specific variety, register or genre. (Charles, 2011) The most interesting type is definitely the general monitor corpora because their size and synchronicity allows them to be regarded as representing the entirety of a given language and are as such interesting to a very wide array of scientists. (Dazdarevic & Fijuljanin, 2015)

2.5.3. Parallel Corpora

Evans, (2004) points out that a parallel corpus involved a collection of texts, which have been translated into one or more other languages. Just like comparable corpora in that they contained two or more collections of texts in various languages. The major difference is in the fact that they have been adjust so that the learner can see all the samples of a particular search term in one language and all the translation equivalents in a second language. (Evans, 2004)

2.5.4. Comparable Corpora

A Comparable Corpus is a collection of "similar" texts in different languages or in different varieties of a language. Two or more corpora are gathered along the same parameters but each involves a different language or a different type of the same language can be said to be comparable corpora. (Evans, 2004, p. 1)

2.5.5. Specialized Corpora

Specialized corpora involve texts from a particular genre or register or a specific time or context. They may involve a sample of this type of text or, if the dataset is finite and of a manageable size (Evans, 2004 p. 1)

2.5.6. Historical or Diachronic Corpora

Evans, (2004, p.2) states “in order to study how language changes over time texts from different time periods can be assembled as a historical corpus.” (Two examples of this type are the Helsinki Diachronic Corpus of English Texts (containing 1.5 million words written between 700 and 1700) and the ARCHER (A Representative Corpus of Historical English Registers) corpus (1.7 million words covering the years 1650 to 1990).

Hunston, (2006) points out that “comparisons between historical periods are difficult because of the lack of truly comparable corpora.” Although substantial numbers of texts from earlier periods are now available electronically, there is nowhere near the same quantity or variation as is available for contemporary texts.

2.5.7. Monitor Corpora

A monitor corpus is one that is ‘topped up’ with new texts on a regular basis. This is done in such a way that “the proportion of text types remains constant” which means that each new version of the corpus is comparable with all previous versions. (Evans, 2004, p. 2), Hunston (2002, p. 16) adds that the best example of this type is the Bank of English held at the University of Birmingham.

2.6. Corpus-based Error Analysis

Learners mostly commit errors due to the interlingual errors; Looking into the literature Sanal, (2008) describes interlanguage errors as the mistakes committed by learners in the TL because of the influence of their mother tongues, the second is intralingual developmental errors, which refer to the TL that is being learned. As Sanal (2008), points out interlanguage errors as the mistakes committed by learners in the TL due to lack of knowledge of that TL’s rules.

Learners commit errors in their placement of verb after the subject. Noun and verb must agree with one another in every point. The basic rule of subject and verb agreement is

that a verb must agree in number and person with its subject. Singular subject should take singular verb and plural subject has to take plural verb.

Teaching benefits from the findings of linguistics in many ways including error analysis, as well as subject and verb agreement. Learners mostly make errors because of the intralingual errors. But it does violate the grammatical rule and just as in the human society, any violation of an important rule becomes a problem to the society. (Length, 2016)

Looking into the literature, learners commit errors in their placement of verb after the subject; noun and the verb have to agree with one another in every point. As Rodrigue (2015) in his famous paper states that, the use of singular and, plural verbs and correct subject and verb agreement are more abstract concepts, which require more detailed elaborations.

Can, (2017, p. 19) states that looking into the learner corpora we can easily detect the learner errors, as learner corpora have presently become readily accessible, it is practicable to examine interlanguage errors and carry out error analysis (EA) on learner generated texts. According to Schachter (1974) also points out that, “EA can only be used to find out and analyze errors that learners naturally produced and ignores potential errors which learners do not make because they avoid difficult structures.” Errors in subject-verb agreement is becoming wide spread and it seems as if many people are either not aware of it or they consider it as less important as it does not affect the message being conveyed. (Length, 2016)

According to Shami, (2013, p. 329) he states that for many years, there have been many studies in the process of the first language acquisition and second language learning. Findings about first language acquisition have been adapted to foreign language learning and it has been concluded that process works in a similar way. In their process of second language acquisition learners commit errors naturally, this might occur due to interlanguage semantic and linguistics incompetent and it has to be detected.

As Can, (2017) stated that, Error detection is the process of determining whether the semantic content and the linguistic form of learners’ communicative performances are erroneous or deviant from the standard average of their background. This process can be executed through classroom observation. And David (2005) adds that errors may be assessed according to the degree of its interference with communication. In some cases, the teachers and the teaching method also have part in the reason behind the student’s error. As Hourani, (2008), when learners make such kind of errors it is believed that either

the learners have a great misunderstandings of the concept or they had been taught by the memorization method instead of practicing. Errors involving subject verb agreement can be classified under local error as it does not affect the message being of communication. (Length, 2016) According to their result, Negro, et al., (2005) they find out that errors significantly increased when unrelated words were added to the sentence recall task.

2.7. Noun (subject) -Verb Agreement

In the field of English as a second language teaching and learning, subject verb agreement problem is becoming more and more obvious across various proficiency levels. Students often generally make mistakes while writing an essay and in their daily communication as well. An example from Nayan & Jusoff, (2009, p. 191) in their study they had an oral interview with one of the participants. Where he was asked a question and he answered this way “It really make me unhappy. Fortunately, my family especially my father need me to help his business. Recently, my father want to expand his business by selling LPG gas. It really tedious to get a license.” this example shows that the learner failed to use the correct rule of third person noun-verb agreement.

Veenstra, et al., (2016, p. 2) in their study they claim that, in English and Dutch, singular subjects require singular verbs and plural subjects require plural verbs (e.g., the dog barks, the dogs bark). Speakers apply this rule fast and automatically. Errors occur occasionally and are often caused by distractor nouns that carry a grammatical number that is different from the head.

Although the subject-verb agreement structure was introduced early to students i.e. when they were in the primary level, they still face problems in acquiring the correct form of it. (Stapa & Izahar, 2010, p. 61) In fact, cases of errors of subject-verb agreement were found not only in student’s essays but even in writings of colleagues in universities. (Length, 2016) Haskell & MacDonald, (2003, p. 760) states that, in recent years a considerable number of studies have explored the mechanisms by which speakers produce such agreement. This interest in agreement has arisen for several reasons. Moreover, in most languages, the rule to make a verb agree with its subject is very simple and can be expressed as: If the subject of the sentence is singular, then the verb has to be singular; if the subject is plural, the verb has to be plural. According to Negro et al., (2005, p. 134) claim that since the human mind is characterized as having limited storage and processing

capacities, the agreement would be more difficult to compute if the subject and the verb were far from each other. (Negro et al., 2005, p. 135)

Al Murshidi, (2014, p. 44) also states that Subject-verb agreement is a basic principle of the English language grammar, it simply denotes that a singular subject needs a singular verb and vice versa. Traditional theories of agreement production assume that verb agreement is an essentially syntactic process. However, recent work shows that agreement is subject to a variety of influences both syntactic and non-syntactic, which raises the question of how these different sources of information are integrated (Haskell & MacDonald, 2003, p. 260)

Subject-verb agreement has provided critical insights into the cue-based memory retrieval system that supports language comprehension by showing that memory interference can cause erroneous agreement with non-subjects: ‘agreement attraction’. (Schlueter, et al., 2018, p.74) More recent work, however, has challenged this view in several domains. First, there is a growing body of evidence that semantic factors may play a more central role in agreement computation. (Haskell & MacDonald, 2003, p.761)

2.8. Approaches to Noun and Verb Agreement

Franck, et al. (2007, p.174) state that linguistic theory developed in the framework of Principles and Parameters/Minimalism has proposed a detailed account of agreement operations taking place during the syntactic derivation of the sentence Mingazova, et al., (2014) points out that “Old English had the category of verb number, but it was lost over time. Old English verbs had the categories of person, number, tense and mood. The category of person had already been weak in Old English:”

Subject-verb agreement is influenced by both notional and grammatical number. yet, the extent to which these two factors are independent remains unclear. (Veenstra et al., 2016, p.1) in some occasion students find it hard to use the correct subject-verb number marked rules. On those occasions when the participant produced a number marked verb such as was or were, the rate of agreement errors was markedly higher in the mismatch condition, indicating that the number of the second noun (conventionally referred to as the local noun) (Haskell & MacDonald, 2003, p.761)

There is an important relationship between the speaker’s memory and the application of noun-verb agreement rules, sometimes they cannot be aware of the mistakes they make, and in some occasion, they cannot keep in mind the rules. Here we need to bring

into account the term proximity. The terms “proximity” or “attraction” denote that speakers or writers are not able to maintain in memory the number of the subject noun in order to specify the number of the verb because the subject and the verb are separated by an interfering noun. (Negro et al., 2005, p. 234) Franck et al., (2004, p. 150) state that linear proximity of the local noun to the verb in the word string is not the critical trigger of agreement failures given that almost no errors were produced with the local noun preceding the verb.

2.9. Data-Driven Learning

According to Guan, (2013, p. 106) DDL (data-driven learning) refers to the discovery and exploration-based learning model based on corpus by using the original data in the corpus or the retrieval results by the corpus retrieval tools. Guan explains that discovery and exploration of knowledge guided or by oneself through the use of corpus data is all what data-drive learning means. The use of linguistic corpora in language learning is often described in terms of data driven learning. (DDL) (Smith, 2010) Most of the researches in DDL support the presence of computer labs; this makes all kinds of activities possible. However, this is not always the reality in every institution: there may be no computer room at all, it may be regularly unavailable at appropriate times (Boulton, 2009, p. 96)

While Jichun, (2015, p. 255) in his study defines data-driven learning as a new research method in language study in the linguistic field which greatly depends on the use of computer and concordance software. Luo, (2016, p. 1) states that data-driven learning (DDL), involving the direct or indirect application of corpus technology, has received considerable attention among researchers and teachers over the past two decades.

In recent years, Western ELT scholarship has emphasized student-centered learning, focusing in particular on Computer-Aided Language Learning (CALL) and the use of linguistic corpora, including Data-Driven Learning (DDL) (Smith, 2010). Boulton, (2009, p. 83) states that DDL is well within the reach of regular teachers and learners in ordinary language teaching contexts, and that a small investment in terms of time and effort can lead to immediate and, more importantly, long-term language learning benefits.

In addition, Boulton emphasizes the importance of independent or learner centered study. Data-driven learning involved the use of ICT in its content, One particular use of ICT which claims to focus on learning rather than teaching is data-driven learning (DDL)

(Boulton, 2009, p. 82), to elaborate data-driven learning is more learner centered in its nature. Boulton notes that many learners may choose to be guided on what to do, summing that it is the teacher's role as expert to show them, and resent having to take any responsibility for their own learning. (Boulton, 2009, p. 86), in contrary to this case, data-driven learning sometimes eased the work of teachers, in which it cuts out the middleman (here the middleman refers to the teacher), whereby learner gets engaged directly to the learning material without the teacher involvement in the process of learning, he or she (learner) gets exposed to the material and carry out the learning independently.

Boulton, (2012, p. 25) in his another study on what data for data-driven learning, notes that it seems likely that many learners around the world are already Googling the Internet in ways not entirely dissimilar to DDL, a practice which may be actively encouraged by their teachers while remaining invisible in the DDL research literature.

DDL increased learners' active involvement in the learning process, which usually requires learners to discover language rules by themselves based on their observation and analysis of the concordance output. (Luo, 2016, p. 1) But in the case of beginners this might be a great challenge. Beginners due to their proficiency level may have problem in their process of independent learning. Guan, (2013, p. 106) claims that learners who have a certain problem will use of retrieval software to discover rules and draw conclusions on the basis of observing and analyzing a large number of real corpus, and master a grammatical structure or word usage through real-time practice. a further important explanation given by Boulton, (2009, p. 81) views and claims on corpora is that, computer corpora have many potential applications in teaching and learning languages, the most direct of which – when the learners explore a corpus themselves – has become known as data-driven learning (DDL).

Hunston (2002, p. 170) also points out that DDL has to do with creating a situations in which learners can answer questions about language themselves by studying corpus data in the form of concordance lines or sentences. Moreover Boulton, (2009, p. 82) adds that DDL generally involves having students get exposed to large quantities of authentic data the electronic corpus, in order to play an active role in exploring the language and detecting patterns in it. Through DDL, the authentic data can help learners in getting used to the target language communication and guide or help them acquire the language use successfully.(Guan, 2013) however Boulton, (2010, p. 14) adds, many activities require learners to guess the meanings of words, or to pay attention on collocations usage patterns

and so on, these are the types of activities commonly associated with DDL can help learner developed their language proficiency.

Compared with traditional foreign language teaching and learning method, data-driven learning is characterized by “autonomic learning”, “authentic language input”, “self-discovery”, and “bottom- up inductive learning”, which will conducive to the formation of the students' personalized learning abilities and the development of their autonomic learning abilities (Guan, 2013, p. 105), This provided the students with greater chance of self-reliance, they will be able to continue their language learning outside the classroom and after their education without the need for teachers or textbooks. (Boulton, 2010) This analysis reflected that, DDL advocates bottom-up, inductive learning. Not in the data-driven learning, students first come into contact with a large amount of authentic language data, but not prescriptive grammatical rules. After their independent observations, they will generalize grammatical rules. (Guan, 2013, p. 107)

Luo, (2016, p. 2) in his study notes that, Data-driven learning (DDL) is classified into direct and indirect. In direct DDL, learners consult corpora directly for solving language problems; whereas indirect DDL refers to learners' use of paper-based concordance lines extracted by teachers or researchers. Since we can clearly see the great effect of DDL in the field of pedagogy, data-driven learning needs to be introduced to learners as early age as possible. For DDL to make any significant impact, it should be introduced in early age, reduce perceived threats to teacher (and learner) roles, circumvent the problems inherent in using computers, and enhance its reputation for direct relevance and ease of use. (Boulton, 2010, p. 4)

2.10. Approaches to Data-Driven Learning (DDL)

It is an approach based on corpus and concordance-based materials to learn language. (Guan, 2013) According to Luo, (2016, p. 1) Data-driven learning has been accepted as an effective approach which aims at helping learners solve vast writing problems such as correcting lexical or grammatical errors, developing the use of collocations and generating ideas in writing DDL. A corpus-assisted language learning method was accepted to be advantageous in finding answer to learners' writing problems. (Luo, 2016, p. 2) Where DDL seems to be most useful is for extending or deepening knowledge of existing language items, distinguishing close synonyms, detecting patterns of usage, collocation, colligation, morphology, and so on. (Boulton, 2009, p. 83)

Boulton, (2012), notes that Data-Driven Learning (DDL) approach is a pedagogic application of corpus linguistics in classroom. Boulton, (2010, p. 7) emphasizes DDL as an approach can seem difficult enough, with its associated elements of discovery learning and induction, not to mention the fuzzy and probabilistic nature of language – quite different from the familiar comfort of rules and “being taught”.

2.11. Criticism against Data-Driven Learning (DDL)

According to Boulton, (2010, p. 2) DDL itself is not as it has been claim it is found to be too difficult, demotivating, irrelevant, and inefficient. DDL can be difficult to leaners of a particular different major. However, DDL tasks are often thought of as part of a student-centered research approach, and some experienced teachers believe that many non-major students lack the motivation or proficiency to engage with such tasks. (Smith, 2010) From the teacher’s point of view, “if DDL has yet to make real inroads to mainstream teaching practices and environments, the problem could lie at any one of three stages:” (Boulton, 2009, p. 91)

- a) Teachers might not know about DDL;
- b) They might know but be unwilling or unable to put it into practice;
- c) They might try it and then reject it.

(Boulton, 2009, p. 91)

According to Boulton, (2009, 2010, p. 89) from time to time it is claimed that learners may have difficulty with the authentic language found in corpora, especially in interpreting the truncated concordance lines of key words in context (KWICs). DDL is assumed to be making things unnecessarily difficult, which may be the case in which learners are introduced all at once to the new approach (DDL), new materials (corpora), and new technology (software).

2.12. Contrastive Analysis

Granger, (2004, p. 127-128) points out that contrastive Interlanguage Analysis, may involve two types of comparison: a comparison of native language and learner language (L1 vs. L2) and a comparison of different varieties of interlanguage (L2 vs. L2). The “compare list” facility in Wordsmith Tools makes it possible to automate these

comparisons: if the two languages and cultures are similar, learning difficulties will not be expected, but if they are different, then learning difficulties are to be expected, and the more they differ, the more the difficulty is expected.” (Lennon, 2008, p.1) & (Kramsch, 2007)

2.13. Interlanguage

Granger, (2004) states that, Interlanguage is a classification in its own right, which can be examine as such without comparing it to any other variety. However, for many reasons, both theoretical and applied, it is very important to compare it to other language varieties to bring out its specificities. The term interlanguage is defined as the language intermediate between the native and the target language. (Lennon, 2008) Zhang, (2013) points out that on the basis of interlanguage analysis, it was believed that second language learners are forming their own self-contained linguistic systems. This has no relation with the system of the native language and not that of the target language, but instead falls between the two. Contrastive analysis assumed that errors have only one cause, namely influence from the mother tongue. Stapa & Izahar, (2010) state that by the late 1960s, second language learning started to be examined in much similar way that first language learning had been studied for some time

Zhang, (2013) notes that we can see that first language interference that occurs in the learners’ writing is simply semantic transfer in which these learners are constantly forming hypotheses of semantic equivalents between English and Chinese words, i.e. Chinese learners tend to make word-for- word translations in their writing, However the implication of this term is that one’s habits in his native language may prevent him from acquiring habits of the second language, not only in writing, it can similarly happen in speaking too. (Lennon, 2008).

For the second language learner, usually the direction of the influence will be from the native language to the target language at the phonological level; this will produce typical foreign pronunciations. Lennon, (2008) states that working with very advanced Dutch learners of English, found that they were more willing to transfer patterns from their native language to English where the meaning was literal, and unwilling to do so in idiomatic or metaphorical environments.

Lennon, (2008) points out that, the contrastive analysis model works best in predicting phonological error. However, “errors of morphology, syntax, lexis and

discourse are imperfectly predicted by contrastive analysis. Lennon further explained that, above the phonological level, language planning is far more under the control of the learner, who may adopt certain strategies to cope with difficulty, more or less consciously.” These include avoidance of difficult forms and simplification of subsystems of the foreign language. (Lennon, 2008)

2.14. Language Transfer

The study of learners’ errors has been developed mainly by two theories. They are Contrastive Analysis (CA) and Error Analysis (EA). When CA lost its initial appeals, EA superseded CA as the dominant approach to studying the learners’ language. (Jichun, 2015 p. 254)

Traditional textbooks had long paid attention to what were felt to be the errors most likely to occur and tried to guard learners from particular pitfalls in phonology, morphology, syntax and lexis.

(Adeeb et al., 2015 p. 439) Shami, (2013 p. 324) points out that, interlingual errors may occur at different levels such as transfer of Phonological, morphological, grammatical lexica- semantic elements of the native language into the target language. Sometime interlingual transfer plays a great role in learner’s language learning process, interlingual transfer errors occur due to the effect of L1. Although, these errors are assumed to exist at the beginning of second language learning, it appeared in among the secondary school students.

2.15. Error Analysis

According to Khansir, (2012 p.1029) Error Analysis is a type of linguistic analysis that focuses on the errors learners make. It consists of a comparison between the errors made in the target language and that target language itself. However several researchers in the field of error analysis have carried out a number of researches of learner’s errors indicated that the influence of the L1 was much less than that said by Contrastive Analysis. Thus, all the mistakes of the language learner are not due to the makeup of his mother tongue. (Khansir, 2012 p.1029)

Can (2017) claims that owing to these limitations of traditional EA, corpus-based computer-aided approach to error analysis has been considered with its pedagogical

significances as a new direction of a useful methodology to guide Foreign Language (FL) teaching in this field.

Error analysis, on the other hand, deals with the learners' performance in terms of the cognitive process they make use of in recognizing or coding the input they receive from the target language. (Shami, 2013 p.325) An error is part of human learning, as can be seen from the English saying “To err is human, to forgive is divine.” (Zhang, 2013) Thus, grammar is an important aspect of language that a learner has to develop mastery among all other aspects of language. Grammar plays the role of skeleton of a language. (Andrade, 2010 p.50) However, “grammatical errors or difficulties usually contribute to lack of competent or effective communication, either verbally or written.” (Adeeb et al., 2015 p.434) There are problems of classification. Classification of errors depends on error being localizable to the domains of phonology/graphology, morphology, syntax, lexis, and discourse. (Lennon, 2008) The quality of a piece of writing is often evaluated by the number of errors so the numerous verb-form errors made by the learners inevitably contributed to the poor quality of writing produced. (Dulay, Burt, & Krashen, 1982) Corder, (1967) classifies errors into two types: performance errors and competence errors. The performance errors are made when learners are tired or hurried, and thus they are not serious because they are not inherent. The competence errors are more serious because they result from improper learning. (Significance & Learner, 1967)

The distinction between “errors” and “mistakes” is highly problematic since in performance correct and incorrect forms of a single target often occur side by side. (Lennon, 2008) In the ESL context, knowledge of grammar becomes an issue of intense community interest, evident in media discussions, especially by those who are concerned with the standard of English language learning and teaching. (Stapa & Izahar, 2010 p.56)

Among the five basic language skills listening, speaking, reading, writing and translation, writing is the most difficult for both native and non-native learners, especially for those non-native non-English majors. (Jichun, 2015 p.254) Adeeb et al., (2015 p.433) note that, “Grammatical errors are one of the major difficulties faced by second language learners.” Among the most common types of grammatical errors are subject-verb agreement errors. Errors are classified according to different levels: modality (i.e., level of proficiency in speaking, writing, reading, listening), linguistic levels (i.e., pronunciation, grammar, vocabulary, style), form (e.g. omission, insertion, substitution), type (systematic errors/errors in competence vs. occasional errors/errors in performance), cause (e.g., interference, interlanguage). (BUSSMANN, 1996) As for communication strategy-based

errors, they result from holistic strategies: e.g. approximation and language switch, and analytic strategies such as circumlocution (expressing the concept indirectly, by allusion rather than by direct reference) (Adeeb et al., 2015 p.436).

Andrade, (2010 p.50) claims that a learner needs to learn certain structure or necessary grammar to produce correct language. Whenever there is deviation from the accepted norms of language, the idea of EA comes in. Andrade further explains. “EA brings a learner to a point at which a learner deviates from the accepted standard of the language (English).” (Andrade, 2010 p.50) however, If the major problems of the language learners were found out and the corresponding remedial teaching methods were put forward, students would surely make less language errors in their writings and their writing proficiency would be greatly improved. (Jichun, 2015 p.254)

In china writing is said to be hard to those students whose major of study is not English language. In Chinese non-English major students’ English writings, some of the sentences are really difficult to understand, “Chinglish” expressions can be found in almost every piece of essay and linguistic errors are also very common. (Jichun, 2015 p.254) In writing, sometimes students make error due to some certain factors. There are many factors, which contribute to boosting such a problem. One of the prominent problems, which aggravate the difficulty in writing, is lack of mastery of English grammatical rules. (Adeeb et al., 2015 p.435)

Errors analysis does not regard them as the persistence of old habits, but rather as signs that the learner is internalizing and investigating the system of the new language. (Shami, 2013 p.324) There is another more fundamental objection to error analysis as a tool for describing transitional competence. However, at best, error analysis can only provide negative evidence concerning certain aspects of the language that have not been acquired by a learner at a particular time. (Lennon, 2008) If learners could not make a sentence of any perfect tense forms, for example, error analysis cannot say whether they have mastered the present perfect tense or not.

Another kind of errors is a result of overgeneralization of TL rules and semantic features. For example, students think that since “s” is added to mark the plural noun, it can also mark the plural verb. Thus, they add “s” to the verb when the subject is plural, whereas they omit such “s” when the subject is singular. (Adeeb et al., 2015 p.439)

Another important thing is that, there are problems in assigning psycholinguistic causes. Assigning psycholinguistic causes to error is by its very nature speculative,

particularly for global errors. In practice, it is often difficult to decide whether an error is caused by language transfer or not. (Lennon, 2008)

Adeeb et al., (2015 p.433) notes on the contribution of the language families to error, when learners' language family is different with the target language, there will be possibility of the learners ending in making error in their process of language learning. So far discrepancies exist between any two languages, in particular, when languages belong to different language families. One example of languages that have such discrepancies is English and Arabic; Arabic is a Semitic language, while English is an Indo-European language.

2.16. Importance of Error analysis

Lennon, (2008) notes that there are problems of error evaluation; studies of errors in English have found great differences in error gravity rankings among individual judges who may employ quite different sets of criteria. Lennon further explained. In particular, native-speaker judges who are not language teachers tend to employ communicative criteria, judging an error as less serious if it does not impede communication and more serious if it does. (Lennon, 2008) Shami, (2013 p.329) also states what a foreign language learner does in operating on the target language is not different from that of a child acquiring his first language.

The subjects' verb-form errors were identified and categorized under four category types: (Dulay et al., 1982)

Classifications of error:

According to Dulay et al., (1982) classified errors into four categories. Below are the categories; as follows:

- (a) Omission
- (b) Addition
- (c) Mis-formation and
- (d) Ordering

Errors of omission are made when compulsory elements are omitted. These occur mainly in tense markers or number markers such as the omission of the grammatical

morphemes, for example, the omissions of the -ed marker in the simple past tense verbs, (Dulay et al., 1982)

Errors of addition are made when unnecessary elements are present with the use of redundant markers, such as, putting the -s marker on verbs after the plural pronouns/nouns in the simple present tense, for example, “They likes...” and “Students wants...”. (Dulay et al., 1982)

Errors of mis-formation occur when the wrong forms of the verbs are chosen in place of the right ones. These commonly occur in cases of subject-verb agreement (SVA) when the wrong verb-forms are selected, for example, “Student are”. (Dulay et al., 1982)

Errors of ordering are made when the correct elements are wrongly sequenced, for example, in the use of phrasal verbs such as, “I pick up her,” instead of “I pick her up.” (Dulay et al., 1982)

Shami, (2013 p.329) notes that children's learning their native mother tongues make plenty of mistakes is a natural part of language acquisition process. They lack great deal of information about subject-verb agreement.

2.17. Computer-aided Error Analysis

The system of CEA developed at Louvain has a number of steps. First, a native speaker of English who also inserts correct forms in the text corrects the learner data manually. Next, the analyst assigns to each error an appropriate error tag (Dagneaux et al., 1998) a further important explanation on analysis of error tagging by Dagneaux et al., (1998) who note that Error-tagged files can be analyzed using standard text retrieval software tools, thereby making it possible to count errors, retrieve lists of specific error types, view errors in context, etc. The error tagging system is hierarchical: error tags consist of one major category code and a series of subcodes. There are seven major category codes: Formal, Grammatical, LeXico-grammatical, Lexical, Register, Word redundant/word missing/word order and Style. (Dagneaux et al., 1998)

Dagneaux et al., (1998) claim that the emergence of a new source of data for SLA research was brought to the life in the early 1990s: the computer learner corpus (CLC). Computer learner corpus (CLC) may be defined as a collection of machine-readable natural language data produced by L2 learners.

2.18. Traditonal Error Analysis:

Dagneaux et al., (1998) claim that EA's exclusive focus on overt errors, in other words, both non-errors such as instances of correct use, and non-use or underuse of words and structures are disregarded. Sridhar (1980) added, "Little attempt was made either to define "error" in a formally rigorous and pedagogically insightful way or to systematically account for the occurrence of errors either in linguistic or psychological terms.

Another important thing to be considered are error typologies. As Dagneaux et al., (1998) lists five limitations pertaining to the nature of data and related error typologies used by traditional EA and:

- Limitation 1: EA is based on heterogeneous learner data;
- Limitation 2: EA categories are fuzzy;
- Limitation 3: EA cannot cater for phenomena such as avoidance;
- Limitation 4: EA is restricted to what the learner cannot do;
- Limitation 5: EA gives a static picture of L2 learning

Cited in (Can, 2018; p. 78)

2.19. Subject-Verb Agreement Error

In the late twentieth century, Error Analysis based on learner corpora, which was initiated by Granger, provides a brand-new perspective into the aspect. (Jichun, 2015) "Error analysis (EA) is a well discussed area in the field of EFL." Analysis of errors in using grammar is a part and parcel of Second Language Acquisition (SLA). Linguists, educationists and researchers talk about knowledge of grammar as significant for developing linguistic competence. (Andrade, 2010) Wang, Zhao, & Shi, (2015) clearly points out that verb errors are one of the most common grammar errors made by non-native writers of English.

Errors in language learning have always been the center of attention, and knowledge of grammar has become one of the most actively discussed questions in language and literacy Pedagogy. (Stapa & Izahar, 2010). Through the errors committed by the learners, the errors can reveal the learner's proficiency level in their language of target over a period of time. They also revealed the gap within the new language. Wang, et al., (2015) state that with the increasing number of people all over the world who study

English as second language (ESL), grammatical errors in writing often occur due to cultural diversity, language habits, and education background.

Nowadays according to some researchers, verb errors are high in percentage. Wang, et al., (2015) According to a rough statistic on essays written by Chinese students, “verb related errors have given a percent as high as 15.6% among all grammatical errors, in which subject-verb agreement errors on the third person singular form cover 21.8%. Different techniques are used to find out errors commit by ESL learners, as for main techniques for the task, most methods can fall into two basic categories, machine learning based and rule-based.” The use of machine learning methods to tackle this problem had shown a promising performance for specific error types. (Wang et al., 2015) Even if the research goal is clear and a precise error definition can be derived from it, it is often unclear how to interpret a learner utterance. Each error is a difference between the utterance and an explicit or implicit ‘correct’ utterance. (Lüdeling, Sauer, Belz, & Mooshammer, n.d.)

CHAPTER III

METHODOLOGY

3.1. Introduction

The aim of this chapter is divided into three: first it provides the research design of the study and introduces the data collection and the analysis procedures, secondly, the participants and the implementation procedures of the study. Finally the instrument used and the procedures are stated and detail explanation of them is given.

3.2. Research Design

The study was conducted through the use of mixed method; both quantitative and qualitative data collection methods were used, purposive sampling has been used, which has provided us with more complete and understanding of the research problem that deals with human beings and their behaviors. This study also comprises descriptive as well as experimental method in its design in order to find a reliable result. It is descriptive because, it provided the researcher with the opportunity to calculate and described to the reader the calculated statistic result of the agreement error frequencies. The study aims at finding out and analyzing the interlanguage form of developmental noun and verb agreement errors of Turkish EFL Learners using Cambridge Learner Corpus (CLC) data, which is most probably the largest learner corpora due to its huge amount. This particular study aims to give a great contribution to the field of learner-corpus-based research on error analysis and foreign language learning with its authentic result of noun and verb agreement errors commit by the Turkish EFL learners.

The study was conducted in one of the biggest private English as a foreign language teaching Centre in Adana, the participants were from different schools, various, professions and departments but we are not so sure about the participant proficiency. The corpus used in the present study is the Cambridge Learner Corpus (CLC). The research questions were formed based on the mixed method and sampling in order to meet the rich answer the research aim or problem.

3.3. Participants:

This study was conducted with 52 students across the various proficiency level, but we are not very clear with their proficiency, the students were from six different classes across A1-C2 levels in one of the biggest English as a Foreign Language (EFL) learning Centre. The classes in total contain 70 students, however, during the first implementation there were 52 students in the classes, as a results the other students joined the classes after the first implementation, but they are excluded from the study on the 52 students who were present at the first implementation session were included.

The participants in this study were chosen according to a case sampling, which is a part of purposive sampling. Two groups were involved in the study, the students were studied in two different major groups as experimental and the control which are split into pre and post group respectively, which are as follow: experimental and the control group.

3.4. Research Context

The study was conducted in 6 different classes across A1-C2 proficiency levels in a private English as a Foreign Language Center the students take all the four skills i.e. writing, reading, listening and speaking

The data obtained from both the Turkish EFL learners and the existing Cambridge corpus data was used in this study. The data was analyzed using Log-Likelihood and Log-Dice to investigate and find out the statistical evidence upon the possible errors committed by the Turkish EFL learners, using the information gathered statistically the study tried to find out whether there is a significant difference across the A1-C2 Turkish EFL learners groups. Through the use of the mixed method, we were provided with the opportunity to obtain and analyzed the data from both quantitative and qualitative traditions.

The participants were Turkish EFL learners across A1-C2 CEFR proficiency levels, currently studying at a private English as a Foreign Language School called Turkish American Association (TAA) June, 2019 academic sessions two hours a day five days a week. The implementation lasted for five weeks, as a part of the writing course. The course instructors carried out the reading, listening and speaking part of the course; the researcher carried out the writing part of each of the levels involved in the study. In the writing part, argumentative essay writing was completed using noun and verb agreements correctly was also highlighted with a corpus-based methodology. In every two weeks after implementation an activity was conducted in the classes.

3.5. Data Collection Tools

Table 1.

Data Collection Tools, Procedure, and time of Implementation/application period

Data Collection Tools	Time of Application during the Research
And Procedure	
1.CLC Cambridge Learner Corpus	13 th – 26 th May 2019
Analysis of noun and verb errors in terms of frequency	two weeks
Pre-essay and post-essay before and after the implementation	13 th May – 16 th June 2019

Table 1 above shows the data collection tools and procedure and the implementation period in the present study, the analysis dates, topics and the activities done in different period during the five weeks implementation.

3.5.1 Log-likelihood

Log-likelihood is a test for statistical significance, similar to the chi-squared measure that is often used in corpus analysis such as for collocation or keyword analysis and it is sometimes called G-square or G score (Baker et al, 2006). In a statistical analysis of texts, to test the frequency distributions, LL test is a reliable alternative to Pearsons' Chi-square (Dunning, 1993). LL test considers word frequencies weighted over two different corpora. It measures higher or lower frequencies than expected. G2 score or LL is Log-likelihood value is as p-value in Pearsons' Chi-square (McEnery, Xiao& Tono, 2006).

Log – likelihood calculator

Log-likelihood and effect size calculator

To use this wizard, type in frequencies for one word and the corpus sizes and press the calculate button.

	Corpus 1	Corpus 2
Frequency of word		
Corpus size		

Notes:

1. Please enter plain numbers without commas (or other non-numeric characters) as they will confuse the calculator!
2. The LL wizard shows a plus or minus symbol before the log-likelihood value to indicate overuse or underuse respectively in corpus 1 relative to corpus 2.
3. The log-likelihood value itself is always a positive number. However, my script compares relative frequencies between the two corpora in order to insert an indicator for '+' overuse and '-' underuse of corpus 1 relative to corpus 2.

How to calculate log likelihood

Log likelihood is calculated by constructing a contingency table as follows:

	Corpus 1	Corpus 2	Total
Frequency of word	a	b	a+b
Frequency of other words	c-a	d-b	c+d-a-b
Total	c	d	c+d

Note that the value 'c' corresponds to the number of words in corpus one, and 'd' corresponds to the number of words in corpus two (N values). The values 'a' and 'b' are called the observed values (O), whereas we need to calculate the expected values (E) according to the following formula:

$$E_i = \frac{N_i \sum_j O_j}{\sum_j N_j}$$

In our case N1 = c, and N2 = d. So, for this word, E1 = c*(a+b) / (c+d) and E2 = d*(a+b) / (c+d). The calculation for the expected values takes account of the size of the two corpora, so we do not need to normalize the figures before applying the formula. We can then calculate the log-likelihood value according to this formula:

$$-2 \ln \lambda = 2 \sum_i O_i \ln \left(\frac{O_i}{E_i} \right)$$

Figure 1. Screenshot of Log-likelihood Calculator

Table 2.

Contingency Table for Log-likelihood Calculation (adapted from Rayson and Garside, 2000, p.3)

	CORPUS ONE	CORPUS TWO	TOTAL
Freq. of words	a	b	a+b
Freq. of other words	c-a	d-b	c+d-a-b
TOTAL			

Source: McEnery, Xiao & Tono, 2006

Note that the value 'c' corresponds to the number of words in corpus one, and 'd' corresponds to the number of words in corpus two (N values). The values 'a' and 'b' are called the observed values (O). We need to calculate the expected values (E) according to the following formula:

$$E_i = \frac{N_i \sum_j O_j}{\sum_j N_j}$$

In our case $N1 = c$, and $N2 = d$. So, for this word, $E1 = c*(a+b) / (c+d)$ and $E2 = d*(a+b) / (c+d)$. The calculation for the expected values takes account of the size of the two corpora, so we do not need to normalise the figures before applying the formula. We can then calculate the log-likelihood value according to this formula:

$$-2 \ln \lambda = 2 \sum_i O_i \ln \left(\frac{O_i}{E_i} \right)$$

This equates to calculating LL as follows: $LL = 2*((a*\log(a/E1)) + (b*\log(b/E2)))$

The word frequency list is then sorted by the resulting LL values. This gives the effect of placing the largest LL value at the top of the list representing the word which has the most significant relative frequency difference between the two corpora (McEnery, Xiao& Tono, 2006).

Source: McEnery, Xiao& Tono, 2006.

3.5.2 Type-token Ratio

Type-token, token is found by counting the number of word in a given text in total, all the word are count include the repeated words. However, type is the total count of the words without the repetition of them. Type/token is used to find out the lexical variety of a corpus. Moreover, “T/t is obtained by dividing the type count, which is always < 1 . A high T/t indicates a high degree of lexical variation while a low T/t indicates a low degree of lexical variation. In addition, T/t can also be expressed as a percentage, multiplying the ratio by 100” The type-token ratio is shown to be a helpful measure of lexical variety within a text. (Thomas, 2015).

In her study, Granger (2002) highlights that it is computed by the means of the formula below and the results obtained from the type/token calculation draw conclusion on lexical richness in learner texts. T/t is calculated as shown below;

$$\text{Type-token} = (\text{number of types} / \text{number of tokens}) \times 100$$

In this present study we used the same formula to make the type/token calculations to find out the ratio.

Table 3.

Turkish EFL learner's Sub-corpus Breakdown. Demographic information is shown in the table below:

L1	CEFR Level	Corpus size	Exam	Age	Register	Years	Exam format
Turkish	A1	31277	BEC1	10	Formal	1993	Advice
	A2		BEC2	-	Informal	-	Articles
			BEC3	50		2012	Application/response
			BECH	+			Argumentative/opinion
	B1	14319	BECP				Business
	B2	1	BECV				Complaint/apology/response
			BULA				Composition/essay
	C1	10150	TS				Descriptive/creative/autobiogr
	C2	1	CAE				aphical
			CPE				Informative/ news
	Total	27596	FCE				Letter/reference
			ICFE				Note/email/memo
			IELTS				story
			AC				Survey/questionnaire/form
			GT				

The table above shows the Turkish EFL learner's sub-corpus breakdown. Including the demographic information. Starting with the illustration of their status as Turkish L1. The Common European Framework of reference (CEFR) followed by the corpus size, type of the exams taken, Age, Register, year of the exams and the exam format.

CHAPTER IV

FINDINGS

4. Introduction

In this chapter, Turkish EFL learners corpus data on noun and verb agreement errors have been analyzed and compared across A1–C2 proficiency level in order to find out some meaningful answers to our research questions. Log-likelihood calculator, Log-dice, as well as Sketch engine have been used in the process. The present study aims at comparing the frequencies of the Turkish EFL learner's noun and verb agreement errors to find out the overuse and the underuse in Cambridge Learner Corpus data, and the result of the implanted study of the experimental and control group data. Both noun and verb error were tagged and counted, their frequencies were written and analyzed.

4.1. Overall frequencies of noun and verb agreement across A1- C2 Proficiency level of ICLE and the Corpus Compiled by the Researcher

In the present study all categories and overall of noun and verb agreement errors have been examined, focus was not given to a specific category. The comparison result of a small noun and verb corpora compiled by the researcher from the implementation with the experimental and the control group have been illustrated.

Firstly, overall frequencies have been identified and analyzed in a tabular form both CLC and the Corpora compiled by the researcher across A1 – C2 proficiency band of CEF. The data were analyzed using the soft wares called Sketch Engine, AntConc, and Log-likelihood in order to compare the result found. Secondly the frequencies of all the categories A1 – C2 of the CLC data as well as the corpus compiled by the researcher during his implementation were analyzed respectively, and each proficiency level frequency is shown individually.

4.2. Finding of the frequency and distribution of Noun and Verb Agreement error in CLC

4.3.1 Finding of the frequency and distribution of Noun Agreement error in CLC

Table 4.

Overall CLC List of noun errors (wrong usage) and the correct form A1 – C2 proficiency level

Category	Error	Correct	Freq.
A1	hour	hours	23
A2	friend	friends	37
B1	kind	kinds	14
B2	kind	kinds	19
C1	kind	kinds	9
C2	sportsmen	sportsman	2

Category=the proficiency level of CEFR

Freq=row frequency of each category

Error=the row of the wrong form of noun used by the student in their writing

Correct=the row of the correct form of noun used by the student in their writing

The table above shows the overall noun agreement errors frequency across **A1 – C2** proficiency level of CEFR commit by the Turkish EFL learners in the Cambridge learner corpora data. As can be seen from the table, we can say that **kind – kinds** is seem to be the highest and the most frequent noun errors committed in **B1, B2** and **C1** across all the proficiency levels respectively, which is the highest and the most common persistent error (wrong usage) **42** in total, followed by **friend – friends** in the **A2** level, which is the second highest frequency of the wrong noun usage in total **37**, then **hour – hours** from the **A1** level wrong usage in total is **23**, and lastly **sportsmen – sportsman** in the **C2** proficiency level of CEFR, wrong usage in total is **2**. The table reveals that, the higher proficient the students get the lower the number of errors made. It is clearly seen that the frequency gets lower and lower as the student's level get higher, as in **B2, C1** and **C2** proficiency levels respectively, except in **A1** and **B1** the frequency is lower than in **A2** and **B2** respectively.

In the table above we decided to bring the highest, common and the most frequent errors from each proficiency band ranging from **A1 – C2** proficiency level in ICLE data. In the description above we decided to combine the total number of words, which are the same even though they are from different proficiency level, Such as **kind – kinds** from **B1, B2** and **C1**.

Table 5.

List of noun errors (wrong usage) frequencies and the correct form in A1 proficiency level of CEFR

Error	Correct	Freq.
hour	hours	23
t-shirt	t-shirt	14
pound	pounds	12
present	present	10
pencil	pencils	8
camera	cameras	7
dollar	dollars	7
game	games	6
day	days	5
megapixel	megapixels	5

Freq=row frequency of each category

Error=the row of the wrong form of noun used by the student in their writing

Correct=the row of the correct form of noun used by the student in their writing

As can be inferred from the table above, the most common is the word **hour-hours**, which appears to be the highest error frequency in the A1 proficiency level of CEFR, which are 23 in total. Students commit the highest errors while they use the word hours, Instead of using hour, they most of the time wrongly used hours.

In the table above we decided to bring the first 10 highest, common and the most frequent noun errors throughout the A1 proficiency level of ICLE.

Table 6.

List of noun errors (wrong usage) frequencies and the correct form in A2 proficiency level of CEFR

Error	Correct	Freq.
friend	friends	37
thing	things	20
friends	friend	6
hour	hours	5
women	woman	5
lots	lot	4
place	places	4
year	years	4
day	days	3
month	months	3

Freq=row frequency of each category

Error=the row of the wrong form of noun used by the student in their writing

Correct=the row of the correct form of noun used by the student in their writing

As we can infer from the table 6 above **friend – friends** appears to be the highest error frequency throughout the A2 proficiency level. In place of friend they tend to use friends, which is the wrong choice due to the number agreement rules. They probably have difficulty on making a decision on the number, i.e. whether singular noun or plural noun should follow the verb (number agreement).

In this table above we decided to bring the first 10 highest, and the most frequent noun errors throughout the **A2** proficiency level of ICLE. Which is **friend – friends** with total of **37** errors.

Table 7.

List of noun errors (wrong usage) frequencies and the correct form in B1 proficiency level of CEFR

Error	Correct	Freq.
kind	kinds	14
student	students	11
friend	friends	9
problem	problems	8
restaurants	restaurants	7
reason	reasons	6
thing	things	6
person	people	5
women	woman	5
week	weeks	5

Freq=row frequency of each category

Error=the row of the wrong form of noun used by the student in their writing

Correct=the row of the correct form of noun used by the student in their writing

As can be inferred from the table seven above, the word **kind – kinds** appears to be the highest error frequency, with the total of 14 wrong usage in the **B1** proficiency level, so this indicates that mostly the students have problem in deciding what noun follows the verb they have chosen in their writings, is it plural or a singular noun? The table reveals that the students mostly use the plural form of noun after a singular verb as a result they tend to end up making errors, making the wrong choices.

So in the table above we decided to bring the first 10 highest, and the most frequent noun errors throughout the **B1** proficiency level of ICLE data. Which is **Kind – Kinds** with total of **14** errors.

Table 8.

List of noun errors (wrong usage) frequency and the correct form in B2 proficiency level

Error	Correct	Freq.
kind	kinds	19
student	students	11
thing	things	9
chart	charts	9
time	times	6
type	types	5
life	lives	5
programme	programmes	5
area	areas	4
friend	friends	4

Freq=row frequency of each category

Error=the row of the wrong form of noun used by the student in their writing

Correct=the row of the correct form of noun used by the student in their writing

As can be seen from the table 8 above, the word **kind – kinds** comes out to be the highest error frequency, 19 in total, this indicates that the student have used it more frequent and wrongly, the table has reveal that most probably the students could not decide on when to use the singular or the plural form of noun after the verb. This shows the difficulty the students have most of the time while using noun and verb agreement in their writing.

So in the table above we decided to show the first 10 highest, and the most frequent noun errors throughout the **B2** proficiency level of ICLE data. Which is **Kind – Kinds** wich is in total **19** errors, students appear to use the wrong form kinds in the moat frequent.

Table 9.

List of noun errors (wrong usage) frequency and the correct form in C1 proficiency level

Error	Correct	Freq.
kind	kinds	9
life	lives	3
meal	meals	3
place	places	3
rates	rate	2
effect	effects	2
category	categories	2
situations	situation	2
clubs	club	2
organisation	organisations	2

Freq=row frequency of each category

Error=the row of the wrong form of noun used by the student in their writing

Correct=the row of the correct form of noun used by the student in their writing

The table 9 above displays the 10 highest frequency in **C1** proficiency level to be **kind – kinds**, witch is 9 in total this word is the most frequent error commit by the **C1** student in their writing task. However this shows that the student wrongly used kinds instead of using its singular form kind, that is why kinds appears to be the highest errors in the C1 proficiency level of CEFR.

So in the table above we decided to show the first 10 highest, and the most frequent noun errors throughout the **C1** proficiency level of ICLE data. Which is **Kind – Kinds** witch is in total **9** errors, students appear to use the wrong form kinds in the moat frequent.

Table 10. *List of noun errors (wrong usage) frequency and the correct form in C2 proficiency level*

Error	Correct	Freq.
sportsmen	sportsman	2
aspects	aspect	1
barmen	barman	1
centre	centres	1
document	documents	1
women	woman	1
guest	guests	1
model	models	1
sport	sports	1
schools	school	1

Freq=row frequency of each category

Error=the row of the wrong form of noun used by the student in their writing

Correct=the row of the correct form of noun used by the student in their writing

The table 10 above shows the description of the highest noun frequency in the **C2** proficiency level of CEFR. This illustrates that the C2 level students errors is low due to their high proficiency in language, this indicates that the students are probably competent in their, writing, use of language as well as noun agreement error in their writing task. In the C2 level as it can be seen **sportsmen – sportsman** appears to be the highest errors which are only 2 in total, instead of **sportsman** they wrongly use the plural form of it **sportsmen**

So in the table above we decided to show the first 10 highest, and the most frequent noun errors throughout the **C2** proficiency level of ICLE. Which are sportsmen – **sportsman** witch is in total **2** errors,

4.3. Finding of the frequency and distribution of Verb Agreement error in CLC

Table 11.

The CLC overall list of verb errors (wrong usage) frequency and the correct form in A1 – C2 proficiency level of CEFR.

Category	Error	Correct	Freq.
A1	play	plays	2
A2	like	likes	10
B1	are	is	14
B2	is	are	52
C1	is	are	23
C2	was	were	11

Category=the proficiency level of CEFR

Freq=row frequency of each category

Error=the row of the wrong form of verb used by the student in their writing

Correct=the row of the correct form of verb used by the student in their writing

As can be seen the table above shows the overall verb agreement errors frequency across **A1 – C2** proficiency level of CEFR, commit by the Turkish EFL learners in the Cambridge learner corpora (CLC). As can be seen from the table, the highest errors is **is – are**. The verb **is – are** seem to be the highest and the most frequent errors commit by **B2**, **C1** and **B1** across all the proficiency levels respectively, which is **89** in total, followed by **was– were** in the **C2** level, which is the second highest error frequency, **11** in total, following that is **like – likes** from the **A2** level **10** in total, and lastly **play – plays** in the **A1** proficiency level of CEFR, which is the lowest wrong usage in the table **2** in total. The table reveals that, from the beginning student's committed errors were few, as their level increases, so does their errors increases too from **A1 2** in total, **A2 10** in total, **B1 14** in total, until level **B2 52** in total, then it started decreasing as they got to **C1** and **C2** level the errors frequency has dropped to **11** in total at the end from the highest **52** in **B2** level.

Table 12.

List of verb errors (wrong usage) frequency and the correct form in A1 proficiency level.

Error	Correct	Freq.
play	plays	1
come	comes	1
go	goes	1
have	has	1
has	have	1
like	likes	1

Freq=row frequency of each category

Error=the row of the wrong form of verb used by the student in their writing

Correct=the row of the correct form of verb used by the student in their writing

Table 12 displays the 6 errors frequencies and we can infer from the table that the highest error in A1 proficiency level is the verb **play – plays** which are **2** in total. Students appear to make errors only twice instead of using play the plural form of the verb plays they wrongly use play which the singular form of the verb, this is called a number agreement error.

So in the table above we decided to show the first 6 highest, and the most frequent verb errors there are actually only 7 throughout the C2 proficiency level of ICLE.

Table 13.

List of verb errors (wrong usage) frequency and the correct form in A2 proficiency level.

Error	Correct	Freq.
like	likes	10
is	are	6
was	were	5
are	is	3
come	comes	2
cost	costs	2
have	has	2
were	was	2
don't	doesn't	1
leave	leaves	1

Freq=row frequency of each category

Error=the row of the wrong form of verb used by the student in their writing

Correct=the row of the correct form of verb used by the student in their writing

As we can infer from the table 13 above **like – likes** appears to be the highest verb error frequency throughout the A2 proficiency level, which is 10 in total. In place of likes the plural form of the verb, the students tend to use like, which is the wrong form of the verb due to the number agreement rules. They probably have difficulty on making a decision on the number, i.e. whether singular form of the verb or plural the plural form of the verb should follow the noun. They found the number agreement difficult.

In this table above we decided to bring the first 10 highest, and the most frequent verb errors throughout the **A2** proficiency level of ICLE.

Table 14.

List of verb errors (wrong usage) frequency and the correct form in B1 proficiency level

Error	Correct	Freq.
are	is	8
is	are	6
was	were	3
come	comes	2
have	has	2
wants	want	2
doesn't	don't	1
has	have	1
help	helps	1
likes	like	1

Freq=row frequency of each category

Error=the row of the wrong form of verb used by the student in their writing

Correct=the row of the correct form of verb used by the student in their writing

As can be inferred from the table 14 above, the word **is** – **are** appears to be the highest verb error frequency, the with the total of 14 wrong usage in the **B1** proficiency level, so this indicates that the students mostly the had problem in deciding what verb follows the noun in their writings, is it plural form of the verb or a singular form? The table reveals that the students mostly use the singular form of the verb after a plural form of noun as a result they tend to end up making errors.

So in the table above we decided to show the first 10 highest, and the most frequent verb errors throughout the **B1** proficiency level CEFR.

Table 15.

List of verb errors (wrong usage) frequency and the correct form in B2 proficiency level

Error	Correct	Freq.
is	are	35
are	is	17
was	were	12
were	was	11
has	have	7
have	has	6
wants	want	5
includes	include	3
costs	cost	2
does	do	2

Freq=row frequency of each category

Error=the row of the wrong form of verb used by the student in their writing

Correct=the row of the correct form of verb used by the student in their writing

As can be seen from the table 15 above, the word **is** – **are** seem to be the highest verb error frequency, which is **35** in total, this indicates that the student have used it most frequently but wrongly, the table has reveal that most probably the students could not decide on when to use the singular or the plural form of the verb is and are after the plural or the singular form of noun. This shows the difficulty the students have most of the time while using noun and verb agreement in their writing.

So in the table above we decided to display the first 10 highest, and the most frequent verb errors throughout the **B2** proficiency level of CEFR. We have analyzed the ICLE data.

Table 16.

List of verb errors (wrong usage) frequency and the correct form in C1 proficiency level

Error	Correct	Freq.
is	are	14
has	have	9
are	is	9
have	has	6
were	was	5
mention	mentions	2

Freq=row frequency of each category

Error=the row of the wrong form of verb used by the student in their writing

Correct=the row of the correct form of verb used by the student in their writing

The table 16 above shows the 10 highest verb frequency in **C1** proficiency level, it can be seen that the word **is – are** is the highest verb errors frequency, witch is 14 in total, the word **is**, is the most frequent error commit by the **C1** student in their writing essay. However this shows that the student wrongly used the verb **is** instead of using its plural form **are**, that is why is appears to be the highest errors in the C1 proficiency level of CEFR.

So in the table above we decided to show the first 10 most frequent verb errors throughout the **C1** proficiency level of CEFR.

Table 17.

List of verb errors (wrong usage) frequencyy and the correct form in C2 proficiency level

Error	Correct	Freq.
was	were	8
are	is	5
is	are	4
has	have	3
were	was	3

Freq=row frequency of each category

Error=the row of the wrong form of verb used by the student in their writing

Correct=the row of the correct form of verb used by the student in their writing

The table 17 above shows the description of the highest verb frequency in the **C2** proficiency level of CEFR. This illustrates that the C2 level students errors is low due to their high proficiency level, this indicates that the students are probably competent in their, use of language as well as verb agreement error in their writing task. In the C2 level as it

can be seen **was – were** appears to be the highest errors which are 11 in total, 8 was errors in as the highest plus 3 were errors which is the one of the two lowest in the table above. Instead of using the plural form **were**, they wrongly use the singular form **was** 8 times in their essays. This reveals that the students probably find it hard to apply the number agreement in their writing.

So in the table above we decided to show the throughout the C2 level we have on five errors categories, since we do not have up to ten we decided to display all of the 5 highest frequency we have, which the most frequent of them is the verb were.

Table 18.

Frequency and Log – Likelihood ratios of noun agreement errors in ICLE

Proficiency Level	Corpus size	Frequency	LL	Significance
A1 – A2	289.888	436	185,17	0,000^{***} -
B1 – B2	360.750	666	114,83	0,000^{***} -
C1 – C2	151.302	141	24,88	0,000^{***} +

- indicates underuse noun agreement errors in A1 proficiency level relative to A2 proficiency level

- indicates underuse noun agreement errors in B1 proficiency level relative to B2 proficiency level

+ indicates overuse noun agreement errors in C1 proficiency level relative to C2 proficiency level

The table above shows the Log – likelihood result across A1 – C2 comparison of noun agreement errors. In this table A1 is compared to A2, B1 to B2, and C1 to C2 respectively. The result obtained from the table above reveals that A1 – A2 LL ratio value is **0,000^{***} -** which indicated that noun agreement is underused, this means that there is not any significant difference in comparison of A1 – A2 Proficiency levels.

However it can be seen that similar result is obtained from B1 – B2 proficiency band too, the LL ratio value in B1 – B2 is **0,000^{***} -** as in A1 – A2 in the same table; it can be seen from the table that, the students in B1 – B2 also underused the noun agreement. This indicates that the level of errors commit is low in both A1 – A2, B1 – B2 in comparison to C1 – C2 proficiency level.

In the C1 – C2 proficiency band in the table above we can infer that, LL ratio value **0,000^{***} +** for this frequency indicated that noun agreement is overused, this shows that the errors committed in C1 – C2 proficiency level is higher in comparison to A1 – A2 and B1 – B2 above.

Table 19.

Frequency and Log – Likelihood ratios of verb agreement errors in ICLE

Proficiency Level	Corpus size	Frequency	LL	Significance
A1 – A2	289.888	441	206,06	0,000^{***} -
B1 – B2	360.750	792	164,04	0,000^{***} -
C1 – C2	151.302	188	3,12	0,077⁺

- indicates underuse verb agreement errors in A1 proficiency level relative to A2 proficiency level

- indicates underuse verb agreement errors in B1 proficiency level relative to B2 proficiency level

+ indicates overuse verb agreement errors in C1 proficiency level relative to C2 proficiency level

It can be seen from the table that the Log – likelihood result across A1 – C2 comparison of verb agreement errors. In this table A1 is compared to A2, B1 to B2, and C1 to C2 respectively. The result obtained from the table above reveals that A1 – A2 LL ratio value is 0,000^{***}- which indicated that noun agreement is underused, this reveals that there is not any significant difference from the comparison of A1 and A2 Proficiency levels.

The similar result is obtained from B1 – B2 proficiency level, the LL ratio value in B1 – B2 is 0,000^{***}- as in A1 – A2 in the same table above; it can be seen from the table that, noun agreement is underused by the B1 – B2 students. This indicates that the level of errors committed is lower in both A1 – A2, B1 – B2 in relevant to C1 – C2 proficiency level.

In the C1 – C2 proficiency band, we can infer that, LL ratio value is 0,077^{***}+ this shows that there is not any significant difference even though it is positive since the ratio is less than one, positive values of frequency indicated that noun agreement is overused, this shows that the errors committed in C1 – C2 proficiency level is higher in comparison to A1 – A2 and B1 – B2 above.

4.4. The analysis of the data compiled by the researcher from the CONTROL and the EXPERIMENTAL group during the implementation study

Table 20.

The overall list of NOUN errors (wrong usage) frequency and the correct form in A1 – C2 proficiency level of CEFR.

Error	Correct	Freq.
people	person	8
engineer	engineers	4
student	students	4
company	companies	3
thing	things	3
women	woman	3
attitude	attitudes	2
behavior	behaviors	2
doctor	doctors	2
information	informations	2

Freq=row frequency of each category

Error=the row of the wrong form of verb used by the student in their writing

Correct=the row of the correct form of verb used by the student in their writing

The table above shows the overall noun agreement errors frequency across **A1 – C2** proficiency level of CEFR from both PRE and POST CONTROL and EXPERIMENTAL GROUP in general, commit by the Turkish EFL learners in the data compiled by the researcher during the implementation. As can be seen from the table, we can say that **people – person** is seem to be the highest and the most frequent noun errors committed, across all the proficiency levels respectively, which is the highest and the most common persistent error (wrong usage **8** in total, followed by **engineer – engineers** and **student – students** is the second highest frequency of the wrong noun usage in total **4** and **4** respectively, then **company – companies**, **thing – things**, and **women – woman** in total is **3**, **3**, **3** respectively and lastly **attitude – attitudes**, **behavior – behaviors**, **doctor – doctors** and **information – informations** in total of **2**, **2**, **2**, **2**. Respectively. The table reveals that, the higher proficient the students get the lower the number of errors they make. It can be seen from the tables that the most of the errors is commit by the experimental groups than the control group/

Table 21.

List of noun errors (wrong usage) frequencies and the correct form in PRE and POST CONTROL A1 – C2 proficiency level of CEFR compiled from the implementation by the researcher

CATEGORIES	PRE – CONTROL				POST - CONTROL		
	Error	Correct	Freq.		Error	Correct	Freq.
A1	students	student	1	A1	women	woman	3
A2	person	people	1	A2	Parent	parents	1
B1	company	companies	1	B1	behavior	behaviors	2
	Thing	things	1		attitude	attitudes	2
	Enemy	enemies	1		persons	people	2
B2	-	-	-	B2	company	companies	1
					Hospital	hospitals	1
					university	universities	1
C1	student	students	1	C1	-	-	-
C2	people	person	1	C2	engineer	engineers	4
	Company	companies	1		doctor	doctors	1
	Hospital	hospitals	1		lesson	lessons	1
	University	universities	1		student	students	1

Freq=raw frequency of each category

Error=the row of the wrong form of noun used by the student in their writing

Correct=the row of the correct form of noun used by the student in their writing

The table above shows the list of the noun agreement error frequencies across A1 – C2 proficiency levels of both PRE and POST CONTROL GROUP in general.

Table 22.

List of noun errors (wrong usage) frequencies and the correct form in PRE and POST EXPERIMENTAL A1 – C2 proficiency level of CEFR compiled from the implementation by the researcher

by the researcher.

CATEGORIES PRE – EXPERIMENTAL				POST - EXPERIMENTAL			
	Error	Correct	Freq.		Error	Correct	Freq.
A1	student	students	1	A1	-	-	-
	Person	people	1				
A2	people	person	1	A2	-	-	1
	Thing	things	1				
B1	man	men	1	B1	thing	things	1
	Person	people	1		information	informations	1
B2	information	informations	1	B2	person	people	1
C1	knowledges	knowledge	1	C1	-	-	-
	Doctor	doctors	1				
C2	-	-	-	C2	-	-	-

Freq=raw frequency of each category

Error=the row of the wrong form of noun used by the student in their writing

Correct=the row of the correct form of noun used by the student in their writing

The table above shows list of the noun agreement error frequencies across A1 – C2 proficiency levels of both PRE and POST EXPERIMENTAL GROUP in general.

Table 23.

The overall list of NOUN errors (wrong usage) frequency and the correct form in A1 – C2 proficiency level of CEFR.

Error	Correct	Freq.
are	is	6
have	has	6
say	says	3
make	makes	3
help	helps	2
ask	asks	2
bring	brings	2
love	loves	2
start	start	2
want	wants	2

Freq=row frequency of each category

Error=the row of the wrong form of verb used by the student in their writing

Correct=the row of the correct form of verb used by the student in their writing

The table above shows the overall VERB agreement errors frequency across **A1 – C2** proficiency level of CEFR from both PRE and POST CONTROL and EXPERIMENTAL GROUP'S errors in general, commit by the Turkish EFL learners in the data compiled by the researcher during the implementation. As can be seen from the table, we can say that **are – is** and **have – has** are seem to be the highest and the most frequent noun errors committed, across all the proficiency levels respectively, which is the highest and the most common persistent error (wrong usage **6 and 6** in total, followed by **say – says** and **make – makes** is the second highest frequency of the wrong noun usage in total **3** and **3** respectively, and lastly **help – helps, ask – asks, bring – brings, love – loves, start – starts and want – wants** in total of **2, 2, 2, 2, 2 and 2**. Respectively. It can be seen from the tables that the most of the errors is commit by the experimental groups than the control group.



Table 24.

List of Verb errors (wrong usage) frequencies and the correct form in PRE and POST CONTROL A1 – C2 proficiency level of CEFR compiled from the implementation by the researcher

CATEGORIES		PRE – CONTROL		POST - CONTROL			
	Error	Correct	Freq.		Error	Correct	Freq.
A1	do	does	1	A1	is	are	2
	Give	gives	1		have	has	2
	Says	say	1		say	says	1
	Make	makes	1		bark	barks	1
	Likes	like	1				
	Loves	love	1				
A2	say	says		A2	loves	love	1
					Blind	blinds	1
					Keep	keeps	1
B1				B1	Think	thinks	1
	ask	asks	1		start	starts	2
	Think	thinks	1				
B2	Spend	spends	1	B2	show	shows	1
	find	finds	1		Improve	improves	1
C1	Bring	brings	1	C1			
	Help	helps	1		saw	see	1
	Has	have	1		Meet	met	1
C2	Wants	wants	1	C2	-	-	-
	is	are	1				
	Want	wants	1				
C2	gives	give	1	C2			
	Kind	kinds	1				

Freq=raw frequency of each category

Error=the row of the wrong form of noun used by the student in their writing

Correct=the row of the correct form of noun used by the student in their writing

The table above shows the list of the VERB agreement errors frequency across A1 – C2 proficiency levels of both PRE and POST CONTROL GROUP in general.

Table 25.

List of Verb errors (wrong usage) frequencies and the correct form in PRE and POST EXPERIMENTAL A1 – C2 proficiency level of CEFR compiled from the implementation by the researcher

by the researcher.

CATEGORIES	PRE – CONTROL			POST - CONTROL			
	Error	Correct	Freq.	Error	Correct	Freq.	
A1	ask	asks	1	A1	take	takes	1
A2	want	wants	2	A2	doesn't	don't	1
	Have	has	1		beat	beats	
B1	brings	thing	1	B1	-	-	-
	Has	have	1				
	Need	needs	1				
	Make	makes	1				
	Want	wants	1				
B2	create	creates	1	B2	-	-	-
	Makes	make	1				
C1	are	is	1	C1	-	-	-
	Have	has	1				
C2	are	is	2	C2	help	helps	1

Freq=raw frequency of each category

Error=the row of the wrong form of noun used by the student in their writing

Correct=the row of the correct form of noun used by the student in their writing

The table above shows the list of the VERB agreement errors frequency across A1 – C2 proficiency levels of both PRE and POST CONTROL GROUP in general.

LESSONS PLANS OF CORPUS–BASED IMPLEMENTATION

FIVE–WEEKS LESSON PLAN AND PROCEDURE FOR THE CORPUS–BASED IMPLEMENTATION ON NOUN AND VERB AGREEMENTS

Week 1 and week 2

Objectives and the flow of the implementation:

- ⇒ Pre essay was given to the students by the researcher.
- ⇒ Students can define Corpus.
- ⇒ Student can explain what argumentative essay is.
- ⇒ Student can show the steps in writing argumentative essay.
- ⇒ Students can identify the number agreement in noun and verb agreement.

Researcher

- Will greet and introduces himself to the student again in every implementation day immediately as he gets into the class.
- Will show the students a sample of argumentative essay and show them the rules of writing an argumentative essay.

Week 3 and week 4

Objectives

- ⇒ Students can remember noun and verb agreement rules.
- ⇒ Students can use noun and verb in the right position.
- ⇒ Students can tell when the number of noun is singular or plural.
- ⇒ Students can tell which verb to follow when the noun is singular.
- ⇒ Students can tell which verb to follow when the noun is plural.
- ⇒ Students can choose appropriate noun or verb to use in which position in their essay.
- ⇒ Students can write argumentative essay.

Researcher

- The researcher will greet the students, introduce himself again and remind the students about the study and the topic to make them not to forget and get familiar with it.
- Will make the students remember all the number agreement rules of noun and verb.
- Will do an exercises with all the noun and verb number agreement rules.
- Will ask student to fill in the blank space in an exercise as a classroom activity.

Week five

- ⇒ Students can remember all the rules and regulations on how to write an argumentative essay.
- ⇒ Student can remember the noun and verb agreement rules on which they are taught.
- ⇒ Students can write an argumentative essay appropriately using noun and verb agreement they have learned.

NOUN AND VERB CLASSROOM ACTIVITIES DURING THE IMPLEMENTATION

PART A

**PLEASE FILL IN THE BLANK SPACE WITH THE WORDS GIVEN IN THE
BRACKET**

- 1) One of the eggs _____ broken (are / is)
- 2) Each of the girls _____ all the regulations (observe / observes)
- 3) The flower _____ good (smells / smell)
- 4) The football players _____ five miles every day (run / runs)
- 5) That red-haired lady in the fur hat _____ across the street (live / lives)
- 6) The weather on the coast _____ to be good this weekend (appear / appears)
- 7) The center on the basketball team _____ the ball too high (bounce / bounces)
- 8) The football player _____ ten miles every day (run / runs)
- 9) Your friend _____ too much (talk / talks)
- 10) The boys _____ to school every day (walk / walks)

PART B

**PLEASE CIRCLE THE CORRECT ANSWER THAT BEST COMPLETE THE
SENTENCES IN THE FOLLOWING QUESTIONS USING (IS ARE WAS
WERE)**

1. Someone in the game (was / were) not hurt.
2. Neither of the men (is / are) not working.
3. Most of the news (is / are) good.
4. Most of the flowers (were / was) yellow.
5. All of the pizza (were / was) gone.
6. All of the children (are / were) late.
7. 6. Some members of the faculty (is-are) present.
8. Everybody (was-were) asked to remain quiet.
9. Neither of the men (is-are) here yet.
10. Several of the sheep (is-are) sick.

PART A KEYS**ANSWERS TO SUBJECT-VERB AGREEMENT EXERCISE**

1. Is
2. Observe
3. Smells
4. Run
5. Lives
6. Appears
7. Bounces
8. Runs
9. Talks
10. Walk

PART B KEYS**ANSWERS TO SUBJECT-VERB AGREEMENT EXERCISE**

1. Was
2. is
3. is
4. Are
5. Was
6. Were
7. Are
8. Was
9. is
10. Are

Table 26.

The frequency and the Log-Likelihood ratio values comparison of A1 – A1 of both AGN and AGV errors in the experimental and the control group's corpus

Word	Freq. in Corpus CONTROL	Freq. in Corpus EXPERIMENTAL	Log-likelihood	Sig.	
	A1	A1			
AGN	4	2	1,35	0,246	+
AGV	12	2	10,95	0,001 ***	+

+ indicates overuse noun agreement errors in A1 proficiency level of **CONTROL GROUP** relative to A1 proficiency level of **EXPERIMENTAL GROUP**

+ indicates overuse verb agreement errors in A1 proficiency level of **CONTROL GROUP** relative to A1 proficiency level of **EXPERIMENTAL GROUP**

AGN indicates noun agreement errors

AGV indicates verb agreement errors

In the table 26 above, it can be seen that in both the AGN comparison of A1 proficiency levels of both CONTROL and EXPERIMENTAL group proficiency levels, the noun agreement errors are been overused, as well as in the AGV comparison of A1 proficiency levels, the verb agreement errors are overused too.

Table 27.

The frequency and the Log-Likelihood ratio values comparison of A2 – A2 AGN and AGV errors in the experimental and the control group's corpus compiled by the researcher

Word	Freq. in Corpus CONTROL	Freq. in Corpus EXPERIMENTAL	Log-likelihood	Sig.	
	A2	A2			
AGN	3	2	0,10	0,751	+
AGV	5	5	0,03	0,852	-

+ indicates underuse verb agreement errors in A2 proficiency level of **CONTROL GROUP** relative to A2 proficiency level of **EXPERIMENTAL GROUP**

AGN indicates noun agreement errors

AGV indicates verb agreement errors

In the table 27 above, it can be seen that in the AGN comparison of A2 proficiency levels of both CONTROL and EXPERIMENTAL group, the noun agreement errors, is been overused, while in the AGV comparison of A2 proficiency levels, verb agreement errors are underused.

Table 28.

The frequency and the Log-Likelihood ratio values comparison of B1 – B1 of both AGN and AGV errors in the experimental and the control group's corpus

Word	Freq. in Corpus CONTROL	Freq. in Corpus EXPERIMENTAL	Log-likelihood	Sig.	
	B1	B1			
AGN	9	4	3,52	0,061	+
AGV	4	5	0,00	0,952	+

+ indicates overuse noun agreement errors in B1 proficiency level of **CONTROL GROUP** relative to B1 proficiency level of **EXPERIMENTAL GROUP**

+ indicates overuse verb agreement errors in B1 proficiency level of **CONTROL GROUP** relative to B1 proficiency level of **EXPERIMENTAL GROUP**

AGN indicates noun agreement errors

AGV indicates verb agreement errors

In the table 28 above, it can be seen that in both the AGN comparison of B1 proficiency levels of CONTROL and EXPERIMENTAL group, the noun agreement errors are been overused, while in the AGV comparison of B1 – B2 proficiency levels, the verb agreement errors are overused too.

Table 29.

The frequency and the Log-Likelihood ratio values comparison of B2 – B2 of both AGN and AGV errors in the experimental and the control group's corpus

Word	Freq. in Corpus CONTROL	Freq. in Corpus EXPERIMENTAL	Log-likelihood	Sig.	
	B2	B2			
AGN	3	2	0,69	0,407	+
AGV	7	2	4,91	0,027	* +

+ indicates overuse verb agreement errors in B2 proficiency level of **CONTROL GROUP** relative to B2 proficiency level of **EXPERIMENTAL GROUP**

AGN indicates noun agreement errors

AGV indicates verb agreement errors

In the table 29 above, it can be seen that in the AGN comparison of B2 proficiency levels of both CONTROL and EXPERIMENTAL group, the noun agreement errors, is been overused, while in the AGV comparison of B2 proficiency levels, the verb agreement errors are overused.

Table 30.

The frequency and the Log-Likelihood ratio values comparison of C1 – C1 of both AGN and AGV errors in the experimental and the control group's corpus

Word	Freq. in Corpus CONTROL	Freq. in Corpus EXPERIMENTAL	Log-likelihood	Sig.	
	C1	C1			
AGN	1	6	2,54	0,111	-
AGV	4	2	1,48	0,224	+

+ indicates overuse verb agreement errors in C1 proficiency level of **CONTROL GROUP** relative to C1 proficiency level of **EXPERIMENTAL GROUP**

AGN indicates noun agreement errors

AGV indicates verb agreement errors

In the table 30 above, it can be seen that in the AGN comparison of C1 proficiency levels of both CONTROL and EXPERIMENTAL group, the noun agreement errors, is been overused, while in the AGV comparison of C1 proficiency levels, the verb agreement errors are overused.

Table 31.

The frequency and the Log-Likelihood ratio values comparison of C2 – C2 of both AGN and AGV errors in the experimental and the control group's corpus

Word	Freq. in Corpus CONTROL	Freq. in Corpus EXPERIMENTAL	Log-likelihood	Sig.	Word
	C2	C2			
AGN	11	0	#NUM!	#NUM!	#NUM! +
AGV	2	3	0,15	0,695	-

+ indicates underuse noun agreement errors in C2 proficiency level of **CONTROL GROUP** relative to C2 proficiency level of **EXPERIMENTAL GROUP**

AGN indicates noun agreement errors

AGV indicates verb agreement errors

In the table 31 above, it can be seen that in the AGN comparison of C2 proficiency levels of both CONTROL and EXPERIMENTAL group, the noun agreement errors, is been overused, while in the AGV comparison of C2 proficiency levels, the verb agreement errors are underused.

Table 30.

The frequency and Log-Likelihood ratio of A1 – A2 vs B1 – B2 AGN and AGV errors in ICLE

Word	Freq. in Corpus	Freq. in Corpus	Log-likelihood	Sig.		
	A1 – A2	B1 – B2				
AGN	436	666	11,21	0,001	***	-
AGV	441	792	39,27	0,000	***	-

- indicates underuse noun agreement errors in A1-A2 proficiency level relative to B1-B2

- indicates underuse verb agreement errors in A1-A2 proficiency level relative to B1-B2

AGN indicates noun agreement errors

AGV indicates verb agreement errors

As we can infer from the table 32 above noun agreement errors frequency in A1 and A2 is compared to the noun agreement errors frequency in B1 and B2. In both AGN and AGV LL ratio values are underused, this indicates that there are not significance differences between noun agreement (AGN) A1 – A2 in comparison to B1 – B2 with the LL ratio value of 0.001 *** - as well as in the verb agreement errors (AGV) A1 – A2 in comparison to B1 – B2 with the LL ratio value of 0.000 *** - of both noun and the verb agreement errors.

Table 31.

The frequency and Log-Likelihood ratio of A1 – A2 vs C1 – C2 AGN and AGV errors in ICLE

Word	Freq. in Corpus	Freq. in Corpus	Log-likelihood	Sig.		
	A1 – A2	C1 – C2				
AGN	436	141	26,32	0,000	***	+
AGV	441	188	5,54	0,019	*	+

+ indicates overuse noun agreement errors in A1-A2 proficiency level relative to C1-C2

+ indicates overuse verb agreement errors in A1-A2 proficiency level relative to C1-C2

AGN indicates noun agreement errors

AGV indicates verb agreement errors

As can be seen from the table 33 above, both AGN and AGV LL ratio values are overused by the learners, this indicates that there are not important significance differences between noun agreement errors frequency in A1 – A2 compared to the noun agreement errors frequency in C1 - C2 with LL ratio 0.000 ***+ and that of verb agreement errors frequency in A1 and A2 compared to C1 and C2 with LL ratio 0.019 * +. However even

though both the LL ratio values in noun comparison of A1 – A2 vs C1 – C2 and that of the verb errors A1 – A2 vs C1 – C2 is positive, it does not mean that there is a significant differences, because the ratio values are less than one.

Table 32.

The frequency and Log-Likelihood ratio of A1 – A2 vs C1 – C2 AGN and AGV errors in ICLE

Word	Freq. in Corpus 1	Freq. in Corpus 2	Log-likelihood	Sig.
	B1 – B2	C1 – C2		
AGN	666	141	62,56	0,000 *** +
AGV	792	188	54,99	0,000 *** +

+ indicates overuse noun agreement errors in B1-B2 proficiency level relative to C1-C2

+ indicates overuse verb agreement errors in B1-B2 proficiency level relative to C1-C2

AGN indicates noun agreement errors

AGV indicates verb agreement errors

In table 34 above, it is seen that, both AGN and AGV LL ratio values are overused by the learners, this shows positive values, but due to the results of the ratio values being less than 1, this indicates that there are not important significance differences, both between B1 – B2 noun agreement (AGN) errors frequency in comparison to C1 – C2 noun agreement errors frequency (AGN) with LL ratio value 0.000 ***+ and that of B1 – B2 verb agreement errors frequency (AGV) compared to C1 and C2 verb agreement errors frequency (AGV) with LL ratio value 0.000 *** +.

Table 33.

Frequencies and the overall Log-Likelihood ratio statistic of both noun and verb agreement errors across A1 – C2 proficiency level of EXPERIMENTAL group

Error category	Noun				Verb			
Level	01	02	LL	Significance	01	02	LL	Significance
A1	2	0	2.81	0.00 +	1	1	0.00	0.18 +
A2	2	0	2.58	0.00 +	3	2	0.11	0.53 +
B1	2	2	0.00	0.37 +	5	0	6.95	0.00 +
B2	1	1	0.01	0.20 -	2	0	2.45	0.00 +
C1	2	4	1.54	0.91 -	2	0	2.15	0.00 +
C2	0	0	0.00	0.00 +	2	1	0.17	0.19 +

01 is observed frequency in corpus 1

02 is observed frequency in corpus 2

+ indicates overuse in 01 relative to 02

- indicates underuse in 01 relative to 02

In table 35 above, frequencies and the log-likelihood ratio of both noun and verb errors frequency of the EXPERIMENTAL group were analyzed in the same table. In the table above we can infer that A1, A2, B1, and C2 noun agreement errors are overused, meanwhile, B2 and C1 are underused, but there is a significant difference only in B1 proficiency level in the verb part, with ratio of 6.95.

In the verb part on the right side it can be seen that, all of the proficiency levels A1 – C2 are overused, no underused at all.

Table. 34.

Frequencies and the log-likelihood statistic result of noun and verb errors across A1 – C2 proficiency level of CONTROL group in total.

Error category	Noun				Verb			
	Level	01	02	LL Significance	01	02	LL Significance	
A1		1	3	0.86 0.70 -	6	6	0.03 1.40 +	
A2		2	1	0.26 0.24 +	1	4	1.93 0.86 -	
B1		3	6	0.46 1.30 -	3	1	1.55 0.22 +	
B2		0	3	4.47 0.83 -	5	2	1.06 0.55 +	
C1		1	0	1.37 0.00 +	2	2	0.00 0.53 -	
C2		4	7	2.82 1.62 -	2	0	1.95 0.00 +	

01 is observed frequency in corpus 1

02 is observed frequency in corpus 2

+ indicates overuse in 01 relative to 02

- indicates underuse in 01 relative to 02

In table 34 above, frequencies and the log-likelihood ratio of both noun and verb errors frequency of the EXPERIMENTAL group were analyzed in the same table. In the table above we can infer that A2, and C1 proficiency level verb agreement errors are overused, while, A1, B1, B2 and C2 are underused, but there is a significant difference only in B2 proficiency level in the noun part, with ratio of 4.47.

In the verb part on the right side it can be seen that A1, B1, B2 and C2 proficiency level verb agreement errors are overused while A2 and C1 proficiency levels are underused.

4.5. Frequencies of Noun and Verb in Pre-essay and Post-essay Type and Token across A1 – C2 proficiency level

4.5.1. Introduction

In this part type/token and ratio (percentage) of both control and the experimental group were listed and calculated separately in a tabular form. The type/token are categorized into pre and post for both control and experimental group respectively, and their entire ratio is calculated in accordance with the categories. Example: pre - control, post – control, pre – experimental, and post – experimental. Below are the tables respectively

Table 35.

Frequencies of Noun and Verb in Pre-essay Type and Token across A1 – C2 proficiency level of PRE- EXPERIMENTAL group.

Proficiency level	Type	Token.	T/t Ratio
A1	224	574	0.39/(39%)
A2	170	447	0.38/(38%)
B1	190	574	0.33/(33%)
B2	231	620	0.37/(37%)
C1	275	634	0.43/(43%)
C2	248	686	0.36/(36%)

n= raw frequency of epistemic bundles/t ratio= Type/token ratio; percentage of number of stance lexical bundles (types) in total of words (tokens) in each corpus

The table above shows the type/token and the ratio percentage of PRE – EXPERIMENTAL group's essays of A1 – C2 proficiency band. The type is divided by the ratio and is been multiplied by 100, this was how we have found the percentage of the type/token ratio

Table 36.

Frequencies of Noun and Verb in Pre-essay Type and Token across A1 – C2 proficiency level of POST- EXPERIMENTAL group.

Proficiency level	Type	Token.	T/t Ratio
A1	226	575	0.39/(39%)
A2	180	406	0.44/(44%)
B1	223	559	0.39/(39%)
B2	215	527	0.40/(40%)
C1	180	472	0.38/(38%)
C2	204	574	0.35/(35%)

n= raw frequency of epistemic bundles/t ratio= Type/token ratio; percentage of number of stance lexical bundles (types) in total of words (tokens) in each corpus

The table above shows the type/token and the ratio percentage of Post – EXPERIMENTAL group's essays of A1 – C2 proficiency band. The type is divided by the ratio and is been multiplied by 100, this was how we have found the percentage of the type/token ratio

Table 37.

Frequencies of Noun and Verb in Pre-essay Type and Token across A1 – C2 proficiency level of PRE – CONTROL group

Proficiency level	Type	Token.	T/t Ratio
A1	194	405	0.47/(47%)
A2	196	486	0.40/(40%)
B1	176	402	0.43/(43%)
B2	181	416	0.43/(43%)
C1	185	405	0.45/(45%)
C2	245	724	0.33/(33%)

n= raw frequency of epistemic bundles/t ratio= Type/token ratio; percentage of number of stance lexical bundles (types) in total of words (tokens) in each corpus

The table above shows the type/token and the ratio percentage of PRE – CONTROL group's essays of A1 – C2 proficiency band. The type is divided by the ratio and is been multiplied by 100, this was how we have found the percentage of the type/token ratio

Table 40.

Frequencies of Noun and Verb in Pre-essay Type and Token across A1 – C2 proficiency level of POST – CONTROL group

Proficiency level	Type	Token.	T/t Ratio
A1	188	461	0.40/(40%)
A2	166	454	0.36/(36%)
B1	167	499	0.33/(33%)
B2	164	394	0.41/(41%)
C1	174	400	0.43/(43%)
C2	200	457	0.43/(43%)

n= raw frequency of epistemic bundles/t ratio= Type/token ratio; percentage of number of stance lexical bundles (types) in total of words (tokens) in each corpus

The table above shows the type/token and the ratio percentage of POST – CONTROL group's essays of A1 – C2 proficiency band. The type is divided by the ratio and is been multiplied by 100, this was how we have found the percentage of the type/token ratio

CHAPTER V

DISCUSSION

5.1. Introduction

In this chapter finding of the study is been discussed, to check if it has similar result with some studies conducted previously or not, the findings are classified in accordance with the research questions. The result have been found to be similar to the previous studies as in, the percentage of the overused has more than double the percentage of the underused noun and verb agreement errors in of the proficiency level of CEFR respectively. The result shows that noun and verb agreement is overused in different proficiency levels 31 times while underused 17 times across all the proficiency levels

5.2. Discussion

Research Question 1: What is the distribution of noun and verb agreement error types along with the A1-C2 Proficiency range according to CEFR?

Noun Agreement Error Distribution and its types

According to the result obtained from the present study, it shows that the overall noun agreement errors types, their distribution and the frequency across **A1 – C2** proficiency level of CEFR in the Cambridge learner corpora data, **kind – kinds** was the highest and the most frequent noun errors committed in **B1, B2** and **C1** across all the proficiency levels respectively, which is the highest and the most common persistent error (most error) **42** in total, followed by **friend – friends** in the **A2** level, which is the second highest frequency of the wrong noun most error in total **37**, then **hour – hours** from the **A1** level second to the least error in total is **23**, and lastly **sportsmen – sportsman** in the **C2** proficiency level of CEFR, least error in total is **2**. The table reveals that, the higher proficient the students get the lower the number of errors they make. It is clearly seen that the frequency gets lower and lower as the student's level get higher, as in **B2, C1** and **C2** proficiency levels respectively, except in **A1** and **B1** the frequency is lower than in **A2** and **B2**.

In the table above we decided to bring the highest, and the most frequent errors from each proficiency band ranging from **A1 – C2** proficiency level in CLC data. In the

description above we decided to combine the total number of words, which are the same even though they are from different proficiency level, Such as **kind – kinds** from **B1, B2** and **C1**.

Verb Agreement Error Distribution and its types

The result obtained from the present study reveals that, in the overall verb agreement errors frequency across **A1 – C2** proficiency level of CEFR, committed by the Turkish EFL learners in the Cambridge learner corpora (CLC). The highest error is **is – are**. The verb **is – are** was the highest and the most frequent errors commit by **B1, B2,** and **C1** across all the proficiency levels of CEFR, with the total of **89**, followed by **was– were** in the **C2** level, which is the second highest error frequency, it is **11** in total, following that is **like – likes** from the **A2** level which is **10** in total, and lastly **play – plays** in the **A1** proficiency level of CEFR, which is the lowest wrong usage in the table 2 in total.

Comparing both the noun and verb agreement errors, it can be seen from the result that, verb agreement errors committed by the Turkish EFL learners are higher than noun agreement errors in total. The verb errors highest frequency is 89 while for the noun agreement errors the highest is 42.

Research Question 2: What are the most common and persistent noun and verb agreement error types in Turkish EFL learners?

It was one among the most important aim of this study to find out the most common noun and verb agreement errors type commit by the Turkish EFL learners in CLC. As in most of the previous studies, it has been seen that the frequency of noun and verb agreement across the proficiency levels is in huge number, so as it is in the present study. The result obtained from the analysis reveals that, the most common and persistent noun errors are **kind – kinds** with the frequency of **42** in total across all the proficiency levels, these most frequent errors were committed in **B1, B2** and **C1** proficiency band of the CLC respectively. The illustration and the detailed analysis can be seen in the table 4 above on page 41

In the other hand, the most common and persistent verb agreement errors are **IS – ARE** with the frequency of **89** in total across all the proficiency levels of CEFR, these most frequent errors were found in **B1, B2** and **C1** proficiency band of the CLC respectively. The illustration and the detailed analysis can be seen in the table 11 above on page 48

Research Question 3: Do the corpus-driven techniques implemented in teaching noun and verb agreements in English yield a statistically significant difference between the success levels of experimental and control groups?

In the present study, it has been predicted that that a corpus-based implementation of teaching noun and verb agreement may have positive effect on Turkish EFL learner's academic writings. The result of this present study reveals that after the implementation of corpus-driven data, the participant's noun and verb agreement errors have reduced, the implementation have helped them find out their weaknesses in the use of noun and verb agreement in their writings.

In table 23 the experimental group's statistics result shows that the students overused noun and verb agreement more than they underused. In noun agreement part they overused noun agreement in A1, A2, B1, and C2 whiled underused only in B2 and C1. However, the overused the verb agreement in all of the proficiency levels A1, A2, B1, B2, C1 and C2. The statistic illustration is given in the table 23 above in chapter 4 pages 65 and 66.

In the control group's statistics, on table 24 page 67 above, the result indicated that the participants underused verb agreement in A1, B1, B2 and C2 proficiency levels. Meanwhile they overused the verb agreement in only A2 and C1 proficiency levels. According to the result they underused the verb agreement more in comparison to overuse of which is only in two proficiency levels.

In the other hand, from the statistics result of the control group, it can be seen that there is a little change when it comes to verb agreement, from the result it can be predicted that the participants had more difficulty in the verb agreement than in the noun agreement, because from the result it can be seen that they overused verb agreement in A1, B1, B2, and C2 proficiency levels, and underused the verb agreement only in A2 and C1 proficiency levels.

5.3. Chapter summary

To sum up, this chapter has summarized and has provided us with the detailed explanation of the answer to the research questions of the study from the result obtained from the study in whole.

CHAPTER VI

CONCLUSION

6.1. Introduction

This study has investigated Turkish EFL learner's use of noun and verb agreement in their academic writing in term of frequency. Corpus-base implementation was carried out with the participants to enhance their use of noun and verb agreement. Lastly in this chapter conclusion of the study is been discussed.

6.2. Conclusion

The aim of this study is to investigating the use of noun and verb agreement in Turkish EFL learner's argumentative essays in CLC in term of frequency and errors. In addition a corpus-based implementation on noun and verb agreement have been conducted with A1 – C2 English as a foreign Language learners. The implementation was carried out in a private English as a Foreign Language teaching center. This implantation aims is to find out if the implementation has improved their use of noun and verb agreement in argumentative essay or not.

The comparison of control and experimental groups as well as the comparison of their pre and post results indicated that, even though according to most of the result there are no significant differences in most of the proficiency levels, but there are little significant differences in two of the proficiency levels indicated in the analysis table. The result has revealed that the experimental group has overused, meaning that they committed more errors in noun and verb agreement more than their counterparts control group, looking at the frequency table it can be clearly seen. However the result also indicated that after the implementation, it can be seen that the controlled group have improved in their use of noun and verb agreement their post argumentative essay writing, but no significant difference is found in the experimental group's post argumentative essays.

Secondly, this study also aims at improving the participant's correct use of noun and verb agreement in their daily academic writings as well as the argumentative essay in term of frequency, different classroom activities on noun and verb agreement and several lectures has been done by the researcher to help the participants improve their skill in writing in the related topic. In some of the activities a number of participant has committed

not even single error while some with little mistakes, this indicates that, the implementation have had a great positive impact on the students.

6.3. Limitations of the Study:

- The present study is limited to only noun and verb errors omitting misspelled words, in the further studies spelling errors and misused noun and verb also should be investigated.
- The study is limited to writing performance
- It is also limited to the small number of the participants
- I am not so sure about the proficiency level of the students
- The present study is limited to only one corpora and the results of the study are limited to the analysis of the Cambridge Learner Corpora (CLC) the study can be enlarge to the comparison of two or more different corpora.

6.4. Implication for English language Teaching

This study has revealed that corpus-based teaching of noun and verb agreement can be very effective and help make EFL learner get better in their academic writings. Using the corpus-based method makes it easier for teacher to bring into class and provide the students with Authentic material and by doing that the student can be provided with opportunity to see and use the correct form of noun and verb agreement in a context. Through the use of corpus, the most and the less used noun and verb in agreement by native can be identified and taught to the students.

6.5. Suggestions for Further study:

- In this study findings show that, a bigger number of the participants can be investigated because, the number of the participant is low in each level in this present study, and that is why there is not very important significant difference.
- A study similar to the present study can be carried out with two or more different group of EFL learners from two or more different language background as well as countries who speak totally different native languages.
- A similar study can be carried out investigating only noun or verb errors with a huge number of participants.

6.6. Chapter summary

To sum up, this chapter has provided us with the general overview of the study in whole, limitations of the study as well as some suggestions for further study has been given.



REFERENCES

- Adeeb, M., Haider, J., Noori, A. L., Hadi, I., Al, K., Subakir, M., & Yasin, M. (2015). Investigating Subject-Verb Agreement Errors among Iraqi Secondary School Students in Malaysia, 3(5), 433–442.
- Al Murshidi, G. (2014). Subject-Verb Agreement Grammatical Errors and Punctuation Errors in Submissions of Male Uae University Students. *European Journal of Business and Innovation Research*, 2(5), 44–47.
- Andrade, R. G. (2010). 3 1 2 3, 3, 1–5.
- Blom, Elma, et al. “Production and Processing of Subject–Verb Agreement in Monolingual Dutch Children with Specific Language Impairment.” *Journal of Speech Language and Hearing Research*, Jan. 2014, p. 952. Doi: 10.1044/2014_jslhr-1-13-0104.
- Boulton, A. (2009). Data-driven Learning: Reasonable Fears and Rational Reassurance. *Indian Journal of Applied Linguistics*, 35(1), 81–106.
- Boulton, A. (2010). Data-driven learning : on paper , in practice . To cite this version : HAL Id : hal-00393809.
- Boulton, A. (2012). What data for data-driven learning? *EuroCALL Review: Proceedings of the EUROCALL 2011 Conference.*, 20(0), 23–27. Retrieved from http://www.eurocall-languages.org/review/20/papers_20/07_boulton.pdf
- BUSSMANN, H. (1996). *Routledge Dictionary of Language and Linguistics*. *Journal of English Linguistics* (Vol. 29). <https://doi.org/10.1177/00754240122005387>
- Can, C. (2017). A Learner Corpus-Based Study on Verb Errors of Turkish EFL Learners. *Journal of Education and Training Studies*, 5(9), 167. <https://doi.org/10.11114/jets.v5i9.2612>
- Can, C. (2018). Agreement Errors in Learner Corpora across CEFR : A Computer-Aided Error Analysis of Greek and Turkish EFL Learners Written Productions, 6(5), 77–84. <https://doi.org/10.11114/jets.v6i5.3064>
- Charles, M. (2011). *The Routledge Handbook of Corpus Linguistics*. *System* (Vol. 39). <https://doi.org/10.1016/j.system.2011.01.004>
- Council of Europe. (2001) Common European Framework of reference for languages: Learning, teaching, assessment. Cambridge, U.K: Press Syndicate of the university of cambridge.

- Corder, S.P. (1967). The significance of learners' errors. Reprinted in J. C. Richards (Ed.) (1974, 1984) *Error analysis: Perspectives on second language acquisition* (pp. 19-27). London: Longman-. Reprinted from *International Review of Applied Linguistics*, 5(4).
- Corder, S. P. (1975). Error Analysis, Interlanguage and Second Language Acquisition. *Language Teaching*, 8(04), 201. <https://doi.org/10.1017/S0261444800002822>
- David ES (2005) classifying language-learning errors. *The Internet TEFL Journal* (volume 58) Retrieved from <https://www.academic.edu/12197473>.
- Adeeb, M., Haider, J., Noori, A. L., Hadi, I., Al, K., Subakir, M., & Yasin, M. (2015). Investigating Subject-Verb Agreement Errors among Iraqi Secondary School Students in Malaysia, 3(5), 433–442.
- Dagneaux, E., Denness, S., & Granger, S. (1998). Computer-aided error analysis. *System*, 26(2), 163–174. [https://doi.org/10.1016/S0346-251X\(98\)00001-3](https://doi.org/10.1016/S0346-251X(98)00001-3)
- Dazdarevic, S., & Fijuljanin, F. (2015). BENEFITS OF CORPUS-BASED APPROACH TO LANGUAGE TEACHING, (June).
- Díaz-negrillo, A. (2013). Automatic Treatment and Analysis of Learner Corpus Data, 59. <https://doi.org/10.1075/scl.59>
- Dulay, H., Burt, M., & Krashen, S. (1982). Language Two.
- Evans, D. (2004). Corpus building and investigation for the Humanities: An on-line information pack about corpus investigation techniques for the Humanities. *Linguistics*, 15–16. <https://doi.org/10.1109/SocialCom.2010.106>
- Franck, J., Cronel-Ohayon, S., Chillier, L., Frauenfelder, U. H., Hamann, C., Rizzi, L., & Zesiger, P. (2004). Normal and pathological development of subject-verb agreement in speech production: A study on French children. *Journal of Neurolinguistics*, 17(2–3), 147–180. [https://doi.org/10.1016/S0911-6044\(03\)00057-5](https://doi.org/10.1016/S0911-6044(03)00057-5)
- Franck, J., Frauenfelder, U. H., & Rizzi, L. (2007). A Syntactic Analysis of Interference in Subject-Verb Agreement. *MIT Working Papers in Linguistics*, (53), 173–190.
- Francom, J., LaCross, A., & Ussishkin, A. (2010). How specialized are specialized corpora? Behavioral evaluation of corpus representativeness for Maltese. *Lrecconforg*, 421–427. Retrieved from http://www.lrec-conf.org/proceedings/lrec2010/pdf/666_Paper.pdf
- Göksel, A., & Kerslake, C. 2005. Turkish: A comprehensive grammar. London, UK: Routledge.

- Granger, S. (2004). Computer Learner Corpus Research : Current Status and Future Prospects. *Applied Corpus Linguistics. A Multidimensional Perspective*, (JANUARY 2004), 123–146.
- Guan, X. (2013). A Study on the Application of Data-driven Learning in Vocabulary Teaching and Learning in China's EFL Class. *Journal of Language Teaching and Research*, 4(1), 105–112. <https://doi.org/10.4304/jltr.4.1.105-112>
- Haskell, T. R., & MacDonald, M. C. (2003). Conflicting cues and competition in subject-verb agreement. *Journal of Memory and Language*, 48(4), 760–778. [https://doi.org/10.1016/S0749-596X\(03\)00010-X](https://doi.org/10.1016/S0749-596X(03)00010-X)
- Hunston, S. (2006). Corpus Linguistics. *Encyclopedia of Language and Linguistics*, (2001), 234–248. <https://doi.org/10.1016/B0-08-044854-2/00944-5>
- Jichun, P. (2015). A Corpus-based Study on Errors in Writing Committed by Chinese Students. *Linguistics and Literature Studies*, 3(5), 254–258. <https://doi.org/10.13189/lis.2015.030509>
- Kornfilt, J. 1997. Turkish. London, UK: Routledge.
- Khansir, A. A. (2012). Error Analysis and Second Language Acquisition. *Theory and Practice in Language Studies*, 2(5), 1027–1032. <https://doi.org/10.4304/tpls.2.5.1027-1032>
- Kramsch, C. (2007). Re-reading Robert Lado, 1957, Linguistics across Cultures. Applied linguistics for language teachers. *International Journal of Applied Linguistics*. <https://doi.org/10.1111/j.1473-4192.2007.00149.x>
- Length, F. (2016). Subject-Verb Agreement Problem among English as Second Language Learners : A Case Study of One Hundred Level Undergraduates of Federal University of Technology , Minna, 2(2), 20–27.
- Lennon, P. (2008). Contrastive Analysis, Error Analysis, Interlanguage 1. *Bielefeld Introduction to Applied Linguistics*, 51–60.
- Lüdeling, A., Sauer, S., Belz, M., & Mooshammer, C. (n.d.). Error Annotation in Spoken Learner Corpora.
- Luo, Q. (2016). The effects of data-driven learning activities on EFL learners' writing development. *SpringerPlus*, 5(1). <https://doi.org/10.1186/s40064-016-2935-5>
- Mingazova, N., Subich, V., & Shangaraeva, L. (2014). The verb-noun agreement in english and arabic. *Journal of Language and Literature*, 5(3), 43–50. <https://doi.org/10.7813/jll.2014/5-3/8>

- McEnery, T., & Hardie, A. (2012). *Corpus linguistics : method, theory and practice. Cambridge Textbooks in Linguistics*, xv, 294 p.
<https://doi.org/10.1017/CBO9781107415324.004>
- Nayan, S., & Jusoff, K. (2009). A Study of Subject-Verb Agreement: From Novice Writers to Expert Writers. *International Education Studies*, 2(3), 190–194.
<https://doi.org/10.5539/ies.v2n3p190>
- Negro, I., Chanquoy, L., Fayol, M., & Louis-Sidney, M. (2005). Subject-verb agreement in children and adults: Serial or hierarchical processing? *Journal of Psycholinguistic Research*, 34(3), 233–258. <https://doi.org/10.1007/s10936-005-3639-0>
- Römer, U. (2008). Corpora and language teaching. *Corpus Linguistics Volume 1*, 112–131.
<https://doi.org/10.1016/j.esp.2003.12.001>
- Römer, U. (2011). Corpus research applications in second language teaching. *Annual Review of Applied Linguistics*, 31, 205–225.
<https://doi.org/10.1017/S0267190511000055>
- Schlueter, Z., Williams, A., & Lau, E. (2018). Exploring the abstractness of number retrieval cues in the computation of subject-verb agreement in comprehension. *Journal of Memory and Language*, 99, 74–89.
<https://doi.org/10.1016/j.jml.2017.10.002>
- Shami, I. A. A. (2013). University Students' Errors in Using Subject Verb Agreement in Writing. Significance, T. H. E., & Learner, O. F. (1967). Full-Text.
- Smith, S. (2010). Corpora in the classroom: Data-driven learning for freshman English. Retrieved from [http://www3.nccu.edu.tw/~smithsgj/Corpora in the classroom.pdf](http://www3.nccu.edu.tw/~smithsgj/Corpora%20in%20the%20classroom.pdf)
- Stapa, S. H., & Izahar, M. M. (2010). Analysis of errors in subject-verb agreement among Malaysian ESL learners. *3L: Language, Linguistics, Literature*, 16(1), 56–73.
<https://doi.org/10.1017/CBO9781107415324.004>
- Underhill, R. 1976. Turkish grammar. Cambridge: MA, MIT Press.
- Veenstra, A., Veenstra, A., & Acheson, D. J. (2016). Semantic integration and subject-verb agreement : Independent effects of notional and grammatical number Independent effects of notional and grammatical number, (January), 1–17.
- Wang, Y., Zhao, H., & Shi, D. (2015). A light Rule-Based approach to English Subject-Verb agreement errors on the Third Person singular forms. *29th Pacific Asia Conference on Language, Information and Computation*, (15), 345–353. Retrieved from <http://www.anthology.aclweb.org/Y/Y15/Y15-2040.pdf>

- Zhang, R. (2013). A Corpus-based Error Analysis of Students' Writing in Graded Teaching Classes. *Journal of Convergence Information Technology*, 8(10), 551–557. <https://doi.org/10.4156/jcit.vol8.issue10.68>
- Warfield, S. "A Taste for Corpora: In Honour of Sylviane Granger." *ELT Journal*, vol. 66, no. 3, 2012, pp. 418–420., doi:10.1093/elt/ccs032.
- Zhang, R. (2013). A Corpus-based Error Analysis of Students' Writing in Graded Teaching Classes. *Journal of Convergence Information Technology*, 8(10), 551–557. <https://doi.org/10.4156/jcit.vol8.issue10.68>
- <https://www.pcc.edu/staff/pdf/645/SubjectVerbAgreement.pdf>



APPENDICES

APPENDIX A

LIST OF THE ERRORS AND THEIR FREQUENCIES

LIST OF THE NOUN ERRORS IN TOTAL AND THEIR FREQUENCIES

Total Corpus size and the frequencies of noun errors across A1 – C2 proficiency level respectively.

Proficiency level	Corpus size	Freq.
A1	92,956	169
A2	196,932	267
B1	165,482	336
B2	195,268	330
C1	106,363	124
C2	44,939	17
TOTAL	801.940	1.243

List of noun errors and the correct form in BREAKTHROUGH A1 proficiency level.

Error	Correct	Freq.
hour	hours	18
t-shirt	t-shirt	14
pound	pounds	12
present	present	10
pencil	pencils	8
camera	cameras	7
dollar	dollars	7
game	games	6
day	days	5
hours	hour	5
megapixel	megapixels	5
month	months	5
euro	euros	4
lots	lot	4
name	names	3

thing	things	3
week	weeks	3
years	year	3
colour	colours	2
film	films	2
friend	friends	2
fridays	friday	2
s-shirts	t-shirt	2
screen	screens	2

A2

List of noun errors and the correct form in WAYSTAGE A2 proficiency level

Error	Correct	Freq.
friend	friends	37
thing	things	20
friends	friend	6
hour	hours	5
women	woman	5
lots	lot	4
place	places	4
year	years	4
day	days	3
month	months	3
people	person	3
weekends	weekend	3
kind	kinds	3
house	houses	2
minute	minutes	2
fridays	friday	2
apple	apples	2
animal	animals	2
child	children	2
girls	girl	2
snowmen	snowman	2
food	foods	2
pizza	pizzas	2

places	place	2
things	thing	2
baby	babies	2
face	faces	2
tree	trees	2
year	years	2
bottle	bottles	2
person	people	2
book	books	2
subject	subjects	2
t-shirt	t-shirts	2
website	websites	2
film	films	2
pencil	pencils	2
monday	Mondays	2
programme	programmes	2
types	type	2
CD	CDs	2
student	students	2
building	buildings	2

List of noun errors and the correct form in THRESHOLD B1 proficiency level

Error	Correct	Freq.
kind	kinds	14
student	students	11
friend	friends	9
problem	problems	8
restaurants	restaurants	7
reason	reasons	6
thing	things	6
person	people	5
women	woman	5
week	weeks	5
year	years	5

book	books	4
lots	lot	4
job	jobs	4
month	months	4
question	questions	3
woman	women	3
chart	charts	3
vehicle	vehicles	3
degree	degree	2
effect	effect	2
company	companies	2
man	men	2
programme	programmes	2
percentage	percentages	2
proportion	proportions	2
language	languages	2
people	person	2
solutions	solution	2
periods	period	2
view	views	2
day	days	2
improvement	improvements	2
scientist	scientists	2
centre	centres	2
region	regions	2
pound	pounds	2
hour	hours	2
delays	delays	2
computer	computers	2
museum	museums	2
human	humans	2
advantages	advantage	2
holidays	holiday	2
tourist	tourists	2

List of noun errors and the correct form in VANTAGE B2 proficiency level

Error	Correct	Freq.
kind	kinds	19
student	students	11
thing	things	9
chart	charts	9
time	times	6
type	types	5
life	lives	5
programme	programmes	5
area	areas	4
friend	friends	4
tourist	tourists	4
place	places	4
member	members	4
people	person	3
advantages	advantage	3
call	calls	3
problem	problems	3
way	ways	3
person	people	3
children	child	3
languages	language	3
lots	lot	3
starts	start	3
night	nights	3
times	time	3
charts	chart	2
differences	difference	2
table	tables	2
programmes	programme	2
types	type	2
punishment	punishments	2
criminal	criminals	2

category	categories	2
risk	risks	2
religion	religions	2
employee	employees	2
room	rooms	2
activities	activity	2
times	time	2
week	weeks	2
restaurants	restaurant	2
mate	mates	2
animal	animals	2
ship	ships	2

List of noun errors and the correct form in EFFECTIVE OPERATIONAL C1 proficiency level

Error	Correct	Freq.
kind	kinds	9
life	lives	3
meal	meals	3
place	places	3
rates	rate	2
effect	effects	2
category	categories	2
situations	situation	2
clubs	club	2
organisation	organisations	2
sports	sport	2
women	woman	2
document	documents	2

List of noun errors and the correct form in MASTER C2 proficiency level

Error	Correct	Freq.
sportsmen	sportsman	2
aspects	aspect	1
barmen	barman	1
centre	centres	1
document	documents	1
women	woman	1
guest	guests	1
model	models	1
sport	sports	1
schools	school	1
subject	subjects	1
trainer	trainers	1
year	years	1
accommodation	accommodations	1
steps	step	1

APPENDICES

APPENDIX B

LIST OF THE VERB ERRORS IN TOTAL AND THEIR FREQUENCIES

Total Corpus size and the frequencies of verb errors across A1 – C2 proficiency level respectively.

Proficiency level	Corpus size	Freq.
A1	92,956	164
A2	196,932	277
B1	165,482	384
B2	195,268	308
C1	106,363	143
C2	44,939	45
TOTAL	801.940	1.421

VERB AGREEMENT ERRORS

List of verb errors and the correct form in BREAKTHROUGH A1 proficiency level.

Error	Correct	Freq.
play	plays	2
come	comes	1
go	goes	1
have	has	1
has	have	1
like	likes	1

List of verb errors and the correct form in WAYSTAGE A2 proficiency level.

Error	Correct	Freq.
like	likes	10
is	are	6
was	were	5
are	is	3
come	comes	2
cost	costs	2
have	has	2

were	was	2
don't	doesn't	1
leave	leaves	1
look	looks	1
take	takes	1
wear	wears	1

List of verb errors and the correct form in THRESHOLD B1 proficiency level

Error	Correct	Freq.
are	is	8
is	are	6
was	were	3
come	comes	2
have	has	2
wants	want	2
doesn't	don't	1
has	have	1
help	helps	1
likes	like	1
look	look	1
include	includes	1
seems	seem	1
show	shows	1
satisfy	satisfies	1
wear	wears	1
watches	watch	1
were	was	1
weren't	wasn't	1
tell	tells	1
think	thinks	1

List of verb errors and the correct form in VANTAGE B2 proficiency level

Error	Correct	Freq.
is	are	35
are	is	17
was	were	12
were	was	11
has	have	7
have	has	6
wants	want	5
includes	include	3
costs	cost	2
does	do	2
get	gets	2
want	wants	2
use	uses	2
uses	use	2

List of verb errors and the correct form in EFFECTIVE OPERATIONAL C1 proficiency level

Error	Correct	Freq.
is	are	14
has	have	9
are	is	9
have	has	6
were	was	5
mention	mentions	2

List of verb errors and the correct form in MASTER C2 proficiency level

Error	Correct	Freq.
was	were	8
are	is	5
is	are	4
has	have	3
were	was	3

APPENDICES

APPENDIX C

BREAKDOWN OF THE CONTROL GROUP'S NOUN AND VERB DATA FREQUENCY

The overall list of NOUN errors (wrong usage) frequency and the correct form in A1 – C2 proficiency level of CEFR.

Error	Correct	Freq.
people	person	8
engineer	engineers	4
student	students	4
company	companies	3
thing	things	3
women	woman	3
attitude	attitudes	2
behavior	behaviors	2
doctor	doctors	2
information	informations	2

Freq=row frequency of each category

Error=the row of the wrong form of verb used by the student in their writing

Correct=the row of the correct form of verb used by the student in their writing

List of noun errors (wrong usage) frequencies and the correct form in PRE and POST CONTROL A1 – C2 proficiency level of CEFR compiled from the implementation by the researcher

CATEGORIES	PRE – CONTROL			POST - CONTROL		
	Error	Correct	Freq.	Error	Correct	Freq.
A1	students	student	1	A1	women	woman 3
A2	person	people	1	A2	Parent	parents 1
B1	company	companies	1	B1	behavior	behaviors 2
	Thing	things	1		attitude	attitudes 2
	Enemy	enemies	1		persons	people 2
B2	-	-	-	B2	company	companies 1
					Hospital	hospitals 1
					university	universities 1
C1	student	students	1	C1	-	-
C2	people	person	1	C2	engineer	engineers 4
	Company	companies	1		doctor	doctors 1
	Hospital	hospitals	1		lesson	lessons 1
	University	universities	1		student	students 1

Freq=raw frequency of each category

Error=the row of the wrong form of noun used by the student in their writing

Correct=the row of the correct form of noun used by the student in their writing

List of noun errors (wrong usage) frequencies and the correct form in PRE and POST EXPERIMENTAL A1 – C2 proficiency level of CEFR compiled from the implementation by the researcher

CATEGORIES	PRE – CONTROL			POST - CONTROL		
	Error	Correct	Freq.	Error	Correct	Freq.
A1	student Person	students people	1 1	A1	- -	-
A2	people Thing	person things	1 1	A2	- -	1
B1	man Person	men people	1 1	B1	thing information	things informations 1
B2	information	informations	1	B2	person	people 1
C1	knowledges Doctor	knowledge doctors	1 1	C1	- -	-
C2	-	-	-	C2	-	-

Freq=raw frequency of each category

Error=the row of the wrong form of noun used by the student in their writing

Correct=the row of the correct form of noun used by the student in their writing

Error	Correct	Freq.
are	is	6
have	has	6
say	says	3
make	makes	3
help	helps	2
ask	asks	2
bring	brings	2
love	loves	2
start	start	2
want	wants	2

Freq=row frequency of each category

Error=the row of the wrong form of verb used by the student in their writing

Correct=the row of the correct form of verb used by the student in their writing

List of Verb errors (wrong usage) frequencies and the correct form in PRE and POST CONTROL A1 – C2 proficiency level of CEFR compiled from the implementation by the researcher

CATEGORIES	PRE – CONTROL			POST - CONTROL			
	Error	Correct	Freq.	Error	Correct	Freq.	
A1	do	does	1	A1	is	are	2
	Give	gives	1		have	has	2
	Says	say	1		say	says	1
	Make	makes	1		bark	barks	1
	Likes	like	1				
	Loves	love	1				
A2	say	says		A2	loves	love	1
					Blind	blinds	1
					Keep	keeps	1
					Think	thinks	1
B1	ask	asks	1	B1	start	starts	2
	Think	thinks	1				
	Spend	spends	1				
B2	find	finds	1	B2	show	shows	1
					Improve	improves	1
	Bring	brings	1				
	Help	helps	1				
	Has	have	1				
	Wants	wants	1				
C1	is	are	1	C1	saw	see	1
	Want	wants	1		Meet	met	1
C2	gives	give	1	C2	-	-	-
	Kind	kinds	1				

Freq=raw frequency of each category

Error=the row of the wrong form of noun used by the student in their writing

Correct=the row of the correct form of noun used by the student in their writing

List of Verb errors (wrong usage) frequencies and the correct form in PRE and POST EXPERIMENTAL A1 – C2 proficiency level of CEFR compiled from the implementation by the researcher

CATEGORIES PRE – CONTROL				POST - CONTROL			
	Error	Correct	Freq.		Error	Correct	Freq.
A1	ask	asks	1	A1	take	takes	1
A2	want	wants	2	A2	doesn't	don't	1
	Have	has	1		beat	beats	
B1	brings	thing	1	B1	-	-	-
	Has	have	1				
	Need	needs	1				
	Make	makes	1				
	Want	wants	1				
B2	create	creates	1	B2	-	-	-
	Makes	make	1				
C1	are	is	1	C1	-	-	-
	Have	has	1				
C2	are	is	2	C2	help	helps	1

Freq=raw frequency of each category

Error=the row of the wrong form of noun used by the student in their writing

Correct=the row of the correct form of noun used by the student in their writing

Corpus size and the frequencies of noun errors across A1 – C2 proficiency level PRE - CONTROL group in total.

Proficiency level	Corpus size	Freq.
A1	389	1
A2	464	2
B1	369	3
B2	400	-
C1	388	1
C2	690	4
TOTAL	2700	11

Corpus size and the frequencies of noun errors across A1 – C2 proficiency level POST - CONTROL group in total.

Proficiency level	Corpus size	Freq.
A1	429	3
A2	424	1
B1	462	6
B2	362	3
C1	380	-
C2	433	7
TOTAL	2492	20

Corpus size and the frequencies of verbs errors across A1 – C2 PRE - CONTROL groups proficiency level in total.

Proficiency level	Corpus size	Freq.
A1	389	6
A2	464	1
B1	369	3
B2	400	5
C1	388	2
C2	690	2
TOTAL	2700	19

Corpus size and the frequencies of verbs errors across A1 – C2 POST - CONTROL group proficiency level in total.

Proficiency level	Corpus size	Freq.
A1	429	6
A2	424	4
B1	462	1
B2	364	2
C1	380	2
C2	433	-
TOTAL	2492	15

APPENDIX D

BREAKDOWN OF THE EXPERIMENTAL GROUP'S NOUN AND VERB DATA FREQUENCY

Corpus size and the frequencies of noun errors across A1 – C2 PRE- EXPERIMENTAL groups proficiency level in total.

Proficiency level	Corpus size	Freq.
A1	535	2
A2	414	2
B1	540	2
B2	581	1
C1	619	2
C2	651	-
TOTAL	3340	9

Corpus size and the frequencies of noun errors across A1 – C2 POST- EXPERIMENTAL groups proficiency level in total.

Proficiency level	Corpus size	Freq.
A1	543	-
A2	375	-
B1	542	2
B2	491	1
C1	440	4
C2	535	-
TOTAL	2926	7

Corpus size and the frequencies of VERB errors across A1 – C2 PRE- EXPERIMENTAL groups proficiency level in total.

Proficiency level	Corpus size	Freq.
A1	535	1
A2	414	3
B1	540	5
B2	581	2
C1	619	2
C2	651	2
TOTAL	3340	15

Corpus size and the frequencies of VERB errors across A1 – C2 POST- EXPERIMENTAL groups proficiency level in total.

Proficiency level	Corpus size	Freq.
A1	543	1
A2	375	2
B1	542	-
B2	491	-
C1	440	-
C2	535	1
TOTAL	2926	4

APPENDICES

APPENDIX E

TYPE AND TOKEN TABLES

CONTROL GROUP

Type and Token across A1 – C2 proficiency level of PRE- CONTROL group.

Proficiency level	Type	Token.
A1	194	405
A2	196	486
B1	176	402
B2	181	416
C1	185	405
C2	245	724
TOTAL	1177	2838

Type and Token across A1 – C2 proficiency level of POST- CONTROL group.

Proficiency level	Type	Token.
A1	188	461
A2	166	454
B1	167	499
B2	164	394
C1	174	400
C2	200	457
TOTAL	1059	2665

EXPERIMENTAL GROUP

Type and Token across A1 – C2 proficiency level of PRE- EXPERIMENTAL group.

Proficiency level	Type	Token.
A1	224	574
A2	170	447
B1	190	574
B2	231	620
C1	275	634
C2	248	686
TOTAL	1338	3535

Type and Token across A1 – C2 proficiency level of POST- EXPERIMENTAL group.

Proficiency level	Type	Token.
A1	226	575
A2	180	406
B1	223	559
B2	215	527
C1	180	472
C2	204	574
TOTAL	1228	3113

C. WORK EXPERIENCE

COMPANY ORGANIZATION	CITY	POSITION	PERIOD
Semiha yücel Akdeğirmen middle school	Adana, Turkey	Intern	23 September 2015 – 13 July 2016

D. HONOURS AND AWARDS

NAME OF HONOURS AND AWARDS	YEAR
CUELT2017 – 3rd Cukurova International ELT Conferences certificate awarded	2016 - 2017
CUELT2018 - 4th Cukurova International ELT Conferences certificate awarded	2017 - 2018