

LINGUISTIC EXPRESSION AND CONCEPTUAL REPRESENTATION OF
MOTION EVENTS IN TURKISH, ENGLISH AND FRENCH: AN
EXPERIMENTAL STUDY

A THESIS SUBMITTED TO
THE GRADUATE SCHOOL OF SOCIAL SCIENCES
OF
MIDDLE EAST TECHNICAL UNIVERSITY

BY

AYŞE BETÜL TOPLU

IN PARTIAL FULFILLMENT OF THE REQUIREMENTS
FOR
THE DEGREE OF DOCTOR OF PHILOSOPHY
IN
THE DEPARTMENT OF FOREIGN LANGUAGE EDUCATION

SEPTEMBER 2011

Approval of the Graduate School of Social Sciences

Prof. Dr. Meliha Altunışık
Director

I certify that this thesis satisfies all the requirements as a thesis for the degree of Doctor of Philosophy.

Prof. Dr. Wolf König
Head of Department

This is to certify that we have read this thesis and that in our opinion it is fully adequate, in scope and quality, as a thesis for the degree of Doctor of Philosophy.

Prof. Dr. Deniz Zeyrek
Supervisor

Examining Committee Members

Assoc. Prof. Dr. Çiler Hatipoğlu (METU, FLE)	_____
Prof. Dr. Deniz Zeyrek (METU, FLE & COGS)	_____
Assist. Prof. Dr. Annette Hohenberger (METU, COGS)	_____
Prof. Dr. Maya Hickmann (PARIS 8 & CNRS)	_____
Assist. Prof. Dr. Fatma Nihan Ketrez-Sözmen (Bilgi U., FLE)	_____

I hereby declare that all information in this document has been obtained and presented in accordance with academic rules and ethical conduct. I also declare that, as required by these rules and conduct, I have fully cited and referenced all material and results that are not original to this work.

Name, Last name: Ayşe Betül, TOPLU

Signature :

ABSTRACT

LINGUISTIC EXPRESSION AND CONCEPTUAL REPRESENTATION OF MOTION EVENTS IN TURKISH, ENGLISH AND FRENCH: AN EXPERIMENTAL STUDY

Toplu, Ayşe Betül

Ph.D., Department of Foreign Language Education

Supervisor: Prof. Dr. Deniz Zeyrek

September 2011, 150 pages

The present dissertation reports the results of a multi-disciplinary experimental study, which combines psycholinguistic and cognitive methodologies in order to achieve two broad objectives. The first objective is providing a comparative psycholinguistic analysis of the expression of motion events in three languages, namely Turkish, English and French, taking Talmy's *verb-framed language* vs. *satellite-framed language* typology (Talmy, 1985) as the framework. The second one is investigating the relationship between linguistic representation and conceptual representation by taking motion events as the testing ground. In order to pursue these two lines of inquiry, five complementary tasks are conducted on three groups of adult subjects. The results of the first two tasks, the language production task and the language comprehension task, verify the Talmyan typology experimentally by showing sharp differences between the data obtained from native speakers of typologically different

languages (English vs. Turkish and French), as well as remarkable similarities between the data obtained from native speakers of typologically similar languages (Turkish and French). On the other hand, the remaining three non-verbal tasks, the categorization task and the two eye-tracking tasks, present valuable insights into the nature of conceptual event representation by revealing a uniform pattern across languages. This latter result is inconsistent with the renowned *linguistic relativity hypothesis* (Whorf, 1956); however in line with the *universalist view* (Jackendoff, 1990, 1996), which suggests that conceptual event representation is language-free and independent of the linguistic encoding preferences of different languages.

Keywords: motion events, verb-framed languages, satellite-framed languages, conceptual representation, linguistic relativity hypothesis.

ÖZ

TÜRKÇE, İNGİLİZCE VE FRANSIZCA'DAKİ DEVİNİM OLAYLARININ SÖZSEL İFADESİ VE KAVRAMSAL TEMSİLİ: DENEYSEL BİR ÇALIŞMA

Toplu, Ayşe Betül

Doktora, Yabancı Diller Eğitimi Bölümü

Danışman: Prof. Dr. Deniz Zeyrek

Eylül 2011, 150 sayfa

Bu tez, iki temel hedefe ulaşmak için psikodilbilimsel yöntemlerle bilişsel yöntemleri birleştiren çok yönlü bir deneysel çalışmanın sonuçlarını sunmaktadır. Bu hedeflerden ilki; Talmy (1985)'nin *eylem-çerçeve*li diller/*uydu-çerçeve*li diller sınıflandırmasını temel alarak, Türkçe, İngilizce ve Fransızca'daki devinim olayları ifadelerinin karşılaştırmalı psikodilbilimsel analizini sunmaktır. İkinci hedef ise; deneysel çalışma alanı olarak devinim olaylarını alarak, dilsel temsille kavramsal temsil arasındaki bağıntıyı incelemektir. Bu iki amaca ulaşabilmek için üç farklı yetişkin denek grubu üzerinde birbirini tamamlayıcı nitelikte beş adet deney uygulanmaktadır. İlk iki deney olan sözsel üretim deneyi ve sözsel anlama deneyi, tipolojik olarak birbirlerinden farklılık ve birbirleriyle benzerlik gösteren dillerin konuşurlarından elde edilen verilerden hareketle, Talmy'nin tipolojisini deneysel olarak doğrulamaktadır. Öte yandan, diğer üç deney, yani sınıflandırma deneyi ve iki

göz-izleme deneyi, ise dillerarası ortak bir kavramsal örüntü ortaya koyarak kavramsal olay temsilinin evrenselliği konusuna ışık tutmaktadır. Bu ikinci sonuç, yalnızca *dilsel görecelik varsayımı*na (Whorf, 1956) ters düşmekle kalmamakta, aynı zamanda kavramsal temsilin dilden bağımsız olduğunu ve farklı dillerin dilsel ifade tercihleri tarafından belirlenmediğini gözler önüne sererek *evrensel yaklaşımı* (Jackendoff, 1990, 1996) da desteklemektedir.

Anahtar Sözcükler: devinim olayları, eylem-çerçeve diller, uydu-çerçeve diller, kavramsal temsil, dilsel görecelik varsayımı.

To my beloved grandfather, İhsan OLCAY

ACKNOWLEDGMENTS

I would like to express my deepest gratitude to my advisor Dr. Deniz Zeyrek for her continuous support, guidance and encouragement throughout my whole PhD. She has always been more than an advisor to me.

I am pleased to thank Dr. Annette Hohenberger for her endless support and patience at every stage of the preparation of this dissertation.

I would like to thank Dr. Wolf König, Dr. Çiler Hatipoğlu, Dr. Maya Hickmann and Dr. Fatma Nihan Ketrez-Sözmen for their valuable suggestions and comments.

I would also like to thank Dr. Maya Hickmann and my dear colleague Efstathia Soroli for their kind support and hospitality.

This study is carried out at the Human-Computer Interaction Research and Application Laboratory of METU Computer Center, and as part of a Research Project (N^o109K281) funded by The Scientific and Technical Research Council of Turkey (TUBITAK). I would like to express my sincere gratitude to both institutions.

I want to thank Kenan Mecit, the secret actor of the videos used in the experiments, and all the people who participated into this study for not only being part of it but also for saying that they enjoyed it.

I would like to thank my dear colleagues Emel Yüksel and Özge Alaçam for their assistance, and my dear friends Gülsüm Özsoy, Derya Kızılay-Kıvanç, Sara Şahin-Açıkgöz and Tarık Açıkgöz for their support.

I want to say ‘thank you very much’ to my dear friend İbrahim Demir for the most valuable technical assistance he has provided, and for the sleepless nights he has spent in this way.

This dissertation would not have been possible without a precious person. She has always been with me when I needed her, she has treated this dissertation as her own, she has tolerated me in every occasion, and above all she has been the best colleague of all. To my best friend and dear sister Gözde Bahadır: Thank you so much for letting me be in your life, I owe you more than you think!

A very special ‘thank you’ goes to two people who have always been there for me no matter what, and who have made me the person I am now: To my beloved mother Saffet Seyyal Baktır and my dear father Ali Baktır.

The last but not the least, I would like to express heartfelt thanks to two people who have changed my life forever by making a beloved wife and a caring mother out of me: To my love, my friend, my husband Murat Toplu for his never-ending patience and guidance during this six years; and to my little daughter who has always cheered me up when I was down with her tiny kicks...

TABLE OF CONTENTS

PLAGIARISM	iii
ABSTRACT	iv
ÖZ	vi
DEDICATION	viii
ACKNOWLEDGMENTS	ix
TABLE OF CONTENTS	xi
LIST OF PICTURES	xvi
LIST OF GRAPHS AND PLOTS	xvii

CHAPTER

1. INTRODUCTION	1
1.1. Aim and Scope	1
1.2. Methodology	2
1.3. Significance of the Study	3
1.4. Outline of the Dissertation	4
2. LITERATURE REVIEW	5
2.1. General Framework	5
2.2. Motion	5
2.2.1. Concept of Motion	5
2.2.2. Motion Events	6
2.2.3. Forming a Typology of Motion Events	8
2.2.3.1. Talmyan Typology	8
2.2.3.1.1. V-Languages	9
2.2.3.1.2. S-Languages	11
2.2.3.2. Slobian Typology	12
2.2.3.3. New Trends of Typology Formation	14
2.2.4. Typological Classification of the Languages in the Study	16
2.2.4.1. Motion Events in Turkish	16
2.2.4.2. Motion Events in English	18

2.2.4.3. <i>Motion Events in French</i>	20
2.3. Language and Cognition	23
2.3.1. The Language and Cognition Debate in General	23
2.3.1.1. <i>Relativity</i>	24
2.3.1.1.1. <i>Strong Version (Lingusitic Determinism View)</i>	24
2.3.1.1.2. <i>Weak Version (Linguistic Relativity View)</i>	25
2.3.1.2. <i>Universality</i>	26
2.3.1.3. <i>Experimental Studies</i>	27
2.3.1.3.1. <i>Color Studies</i>	27
2.3.1.3.2. <i>Spatial Location Studies</i>	28
2.3.2. The Language and Cognition Debate on Motion	29
2.3.2.1. <i>Motion as the Perfect Testing Ground</i>	29
2.3.2.2. <i>Experimental Studies on Motion Events</i>	29
2.3.2.2.1. <i>Categorization Studies</i>	30
2.3.2.2.2. <i>Recognition Memory Studies</i>	33
2.3.2.2.3. <i>Eye-Tracking Studies</i>	35
2.3.2.3. <i>The Gap in the Literature to be Filled with the Present Study</i>	36
.	
3. GENERAL METHODOLOGY	39
3.1. Framework	39
3.2. Procedure	40
4. EXPERIMENTS	43
4.1. Experiment 1: Video Description Task	43
4.1.1. Aim and Scope	43
4.1.2. Research Questions and Hypotheses	44
4.1.3. Design	45
4.1.3.1. <i>Subjects</i>	45
4.1.3.2. <i>Stimuli</i>	46
4.1.3.3. <i>Apparatus</i>	49
4.1.3.4. <i>Procedure</i>	50
4.1.3.5. <i>Encoding of the Raw Data</i>	51
4.1.4. Results	52
4.1.4.1. <i>Main Analysis</i>	52

4.1.4.2. <i>Semantic Density Analysis</i>	54
4.1.5. Discussion	56
4.2. Experiment 2: Acceptability Judgment Task	58
4.2.1. Aim and Scope	58
4.2.2. Research Questions and Hypotheses	58
4.2.3. Design	59
4.2.3.1. <i>Subjects</i>	59
4.2.3.2. <i>Stimuli</i>	59
4.2.3.3. <i>Apparatus</i>	60
4.2.3.4. <i>Procedure</i>	61
4.2.3.5. <i>Encoding of the Raw Data</i>	61
4.2.4. Results	62
4.2.5. Discussion	64
4.3. Experiment 3: Similarity Judgment Task	65
4.3.1. Aim and Scope	65
4.3.2. Research Questions and Hypotheses	66
4.3.3. Design	67
4.3.3.1. <i>Subjects</i>	67
4.3.3.2. <i>Stimuli</i>	67
4.3.3.3. <i>Apparatus</i>	67
4.3.3.4. <i>Procedure</i>	67
4.3.3.5. <i>Encoding of the Raw Data</i>	68
4.3.4. Results	68
4.3.5. Discussion	70
4.4. Experiment 4: Non-Verbal Eye-Tracking Task	70
4.4.1. Aim and Scope	70
4.4.2. Research Questions and Hypotheses	71
4.4.3. Design	72
4.4.3.1. <i>Subjects</i>	72
4.4.3.2. <i>Stimuli</i>	72
4.4.3.3. <i>Apparatus</i>	72
4.4.3.4. <i>Procedure</i>	73
4.4.3.5. <i>Encoding of the Raw Data</i>	74
4.4.4. Results and Discussion	79

4.4.4.1. <i>Main Analysis</i>	80
4.4.4.2. <i>Discussion of the Main Analysis</i>	81
4.4.4.3. <i>Tempoal Linearization Analysis</i>	82
4.4.4.4. <i>Discussion of the Linearization Analysis</i>	84
4.5. Experiment 5: Pre-Production Eye-Tracking Task	86
4.5.1. Aim and Scope	86
4.5.2. Research Questions and Hypotheses	86
4.5.3. Design	88
4.5.3.1. <i>Subjects</i>	88
4.5.3.2. <i>Stimuli</i>	88
4.5.3.3. <i>Apparatus</i>	88
4.5.3.4. <i>Procedure</i>	88
4.5.3.5. <i>Encoding of the Raw Data</i>	89
4.5.4. Results and Discussion	89
4.5.4.1. <i>Main Analysis</i>	89
4.5.4.2. <i>Discussion of the Main Analysis</i>	90
4.5.4.3. <i>Tempoal Linearization Analysis</i>	91
4.5.4.4. <i>Discussion of the Linearization Analysis</i>	93
 5. GENERAL DISCUSSION	 94
5.1. Discussion of the Verbal Data	94
5.1.1. Video Description Data	95
5.1.2. Acceptability Judgment Data	97
5.2. Discussion of the Non-Verbal Data	100
5.2.1. Similarity Judgment Data	101
5.2.2. Non-Verbal Eye-Tracking Data	102
5.2.3. Pre-Production Eye-Tracking Data	104
 6. CONCLUSION	 108
6.1. Verbal Part of the Study	108
6.2. Non-Verbal Part of the Study	109
6.3. Suggestions for Further Research	111
 REFERENCES	 114

APPENDICES

A. INFORMED CONSENT FORM	125
B. BACKGROUND QUESTIONNAIRE	126
C. POST-PARTICIPATION INFORMATION SHEET	127
D. VITA	128
E. TURKISH SUMMARY	131

LIST OF PICTURES

PICTURES

Picture 1 Background used for <i>into</i> & <i>out of</i> scenes.....	47
Picture 2 Background used for <i>up</i> & <i>down</i> scenes	48
Picture 3 Background used for <i>across</i> scenes	48
Picture 4 The testing room.....	49
Picture 5 The control room.....	49
Picture 6 A screenshot from an example video used in the comprehension task	60
Picture 7 The answering window used in the comprehension task	61
Picture 8 Example screen shot from the alternative videos presented consecutively.....	68
Picture 9 Example screen shot from the alternative videos presented consecutively.....	68
Picture 10 Tobii 1750 Eye Tracker	73
Picture 11 Tobii X120 Eye Tracker.....	73
Picture 12 Example screenshot from raw eye-tracking data	75
Picture 13 The path region for <i>into</i> and <i>out of</i> scenes	76
Picture 14 The path region for <i>up</i> and <i>down</i> scenes	76
Picture 15 The path region for <i>across</i> scenes	76
Picture 16 Example manner region for the scene <i>run into</i> (frame 1)	77
Picture 17 Example manner region for the scene <i>run into</i> (frame 60)	78

LIST OF GRAPHS AND PLOTS

GRAPHS

Graph 1 Outline of the experimental sessions	39
Graph 2 Distribution of the verbal descriptions for all groups.....	52
Graph 3 Results of the Video Description Task.....	53
Graph 4 Results of the semantic density analysis across languages.....	56
Graph 5 The comprehension results for all groups.....	62
Graph 6 Similarity Judgment results for all groups.....	69
Graph 7 The eye-tracking results for all groups.....	80
Graph 8 The eye-tracking results for all groups.....	90
Graph 9 A comparative analysis of the two eye-tracking experiments.....	106

PLOTS

Plot 1 The results of the two-factorial ANOVA for the comprehension task.....	63
Plot 2 Means of MMP Ratio across languages.....	83
Plot 3 Variances of MMP Ratio across languages.....	84
Plot 4 Means of MMP Ratio across languages.....	92
Plot 5 Variances of MMP Ratio across languages.....	92

CHAPTER I

INTRODUCTION

1.1. AIM AND SCOPE

The aim of the present dissertation is twofold. The first one is making a comparative psycholinguistic analysis of motion event expressions in three different languages, namely Turkish, English and French, taking the verbalization typology put forward by Leonard Talmy (1985) as the framework. Talmy argues that *path of motion* is the main and the most indispensable semantic component of a motion event, and thus he frames his typology around that component. His typology includes *Verb-framed languages* (V-languages) on one side, and *Satellite-framed languages* (S-languages) on the other side. In this two-way typology, the *path of motion* is either framed by the main verb of a sentence or by a secondary construction called a satellite (*e.g.* a particle, an adposition or an adverbial). Even though this particular typology has been challenged by several researchers during the last two and a half decades, it has certainly been the most commonly used framework in crosslinguistic studies of motion events and it is still the most solid typology that has been proposed in this regard. For this reason, the current study also uses the Talmyan typology to make a detailed experimental inquiry of motion event usages in the three languages mentioned above. The reason for choosing this particular set of languages is again related to our theoretical framework. As Turkish and French are categorized as V-languages, and English as an S-language in the typology; working on this particular set of languages gives us the opportunity to examine the motion event expressions both from an inter-typological (English vs. Turkish and French), and from an intra-typological (Turkish vs. French) perspective.

The second aim of the dissertation is providing a detailed inquiry on the *language and cognition debate* by taking motion events as the testing ground. In other words, the second aim of the current study is questioning the probable link

between the conceptual representation and the linguistic expression of motion events in typologically different or similar languages. This question is closely related to the *linguistic relativity hypothesis* (Whorf, 1956), which has gained new impetus during the last couple of decades under the name of Neo-Whorfian research. According to this hypothesis, the way human beings talk about a certain event or state has a certain influence on the way that event or state is conceptualized in their minds, therefore language has an effect on people's conceptual representation of the world. The present work will challenge this hypothesis by adopting a crosslinguistic perspective, and by making use of various experimental techniques complementing each other.

1.2. METHODOLOGY

In this dissertation, psycholinguistic and cognitive experimentation techniques are used to shed light on the two main lines of investigation, which have been summarized above and which will be detailed in the upcoming chapter. Within the framework of the study, a total of four tasks are conducted:

1. Video Observation Task
2. Similarity Judgment Task
3. Video Description Task (Description Part + Eye-Tracking Part)
4. Acceptability Judgment Task

The participants are native speakers of the three languages in question, *i.e.* Turkish, English and French, who are all chosen on voluntary basis. All the tasks are computer-based, and they have been conducted at the Human-Computer Interaction Research and Application Laboratory at Middle East Technical University, Ankara, Turkey, and at the SFL (Structures Formelles du Langage) Laboratory, Paris, France. Different from most of the studies in the field that use static pictures or animation clips to elicit data, the current study makes use of real-life video sequences exclusively shot for these experiments. As the third task, the Video Description Task, is composed of two parts and as the results of these two parts will be analyzed separately, the current study actually includes five tasks. Two of the five tasks are verbal in nature and make a psycholinguistic

investigation of the expression of motion event patterns in the three languages. One of them is a language production task (Description Part of the Video Description Task), and the other one is a language comprehension task (Acceptability Judgment Task) which complements the first one. The other three tasks, which are non-verbal in nature, address the language and cognition debate. One of these tasks is a non-verbal categorization task (Similarity Judgment Task), and the other two are cognitive tasks that utilize the eye-tracking paradigm (Video Observation Task and Eye-Tracking Part of the Video Description Task).

1.3. SIGNIFICANCE OF THE STUDY

- This is the first study in the literature, which makes a crosslinguistic experimental analysis of this particular group of languages (*i.e.* Turkish, English and French).
- Different from a great number of studies, which base their psycholinguistic analyses solely on language production data, this study also makes use of language comprehension data to obtain a complete picture.
- It is the first study in the literature, which elicits Turkish motion event expressions by using real-life stimuli.¹
- It is a pioneering study that makes a detailed inquiry on Turkish motion events, combining the verbal expression and the non-verbal representation dimensions.
- The present study is a multi-dimensional study; which can be categorized as a *contrastive linguistic* study due to its crosslinguistic approach, as a *philosophy of language* study due to its focus on the language and

¹ There are a few studies in the literature investigating English and French with real-life stimuli (see Gennari *et al.*, 2002; Pourcel and Kopecka, 2006; Soroli and Hickmann, 2010; Soroli, 2011; and Flecken, 2011).

cognition debate, and as a *psycholinguistic* or *cognitive science* study due to the experimental techniques it uses.

1.4. OUTLINE OF THE DISSERTATION

The dissertation is comprised of six chapters. The following chapter, Chapter II, will present a theoretical and experimental review of the literature on the two main research areas related to the study: the literature on *motion events*, and the literature on *the language and cognition debate*. Chapter III will provide the general methodological framework used throughout the study. Then, Chapter IV will present a systematic description of the five tasks used in the study, along with their results. The results of the experiments will be discussed in detail in Chapter V, and Chapter VI will conclude the text by providing some suggestions for future studies.

CHAPTER II

LITERATURE REVIEW

2.1. GENERAL FRAMEWORK

The present dissertation is based on two major veins of research. The first one is the psycholinguistic investigation of the expression of *motion events* across languages, and the second is the inquiry of *the language and cognition interface*, taking motion events as the testing ground. Accordingly, this chapter will be comprised of two sections: The section entitled “Motion” will provide the theoretical framework for motion event studies in the field together with a summary of the prime experimental studies, and the typological analysis of the languages used in the current study (*i.e.* Turkish, English and French). The section entitled “Language and Cognition” will review both the theory and the practice of the Whorfian Debate (Whorf, 1956), which questions the existence of a possible link between conceptual representation and linguistic representation, first in general terms then specifically for motion events.

2.2. MOTION

2.2.1. Concept of Motion

Motion is a basic physics term first introduced by Aristotle in 350 BC (Kosman, 1969) and later challenged by Descartes, Galileo and Newton in many respects. The Oxford Dictionary of Physics (Daintith, 2010, p. 341) defines motion as “a change in the position of a body or system with respect to time, as measured by a particular observer in a particular frame of reference”. Though the term has been defined on several occasions by several researchers in several fields (*e.g.* physics, philosophy, psychology or linguistics), we will contend with a recent definition proposed by Zlatev, Blomberg, and David (2010): “[M]otion ... can be defined as the experience

of continuous change in the relative position of an object (the figure) against a background” (p.5, emphasis original).

2.2.2. Motion Events

The idea of motion plays an important role in human thought, and it is regarded by many as one of the most basic concepts of our thinking system (Goddard, 1998). Renowned psychologists Miller and Johnson-Laird (1976) define motion verbs as “the most characteristically verbal of all verbs” and as the “purest and prototypical of verbs” (p. 527). On the other hand, the modern and systematic treatments of motion from a linguistic perspective start with Leonard Talmy (1985)’s influential work on *basic motion event*. He defines a motion event as “a situation containing motion and the continuation of a stationary location” (2000, p. 25). According to Talmy, a basic motion event consists of four internal elements:

1. FIGURE: Figure is an object / a person moving or located with respect to another object.
2. GROUND: Ground is the spatial reference point, according to which the motion or location of the Figure is determined.
3. PATH: Path is the trajectory followed or the space occupied by the figure during motion. It consists of a source (the starting point), a medium (the intermediate points), and a goal (the endpoint) (Pourcel & Kopecka, 2006).

It has recently been suggested in the literature that there is an asymmetry both in the expression and the conceptualization of the two main components of a Path, Goal being more focused than Source. This so-called *Source-Goal Asymmetry* is said to be dispersed across languages and across linguistic levels (semantics, morphology, syntax and discourse) (Ishibashi & Kopecka, 2011). Aske (1989), on the other hand, suggests that there are two types of path, *telic path* and *atelic path*. A telic path includes a specific end-point, thus it is resultative as in (1). On the other hand, an atelic path does not specify the end-point of the event but just the medium as in (2) (*cf.* Pourcel & Kopecka, 2006).

- (1) The boy ran into the building.
- (2) The boy ran along the park.

4. MOTION: Motion is the presence of motion or locatedness itself.

There are also two external components of a basic motion event:

- I. MANNER: Manner, which can be defined as the way of moving or being located, is an external component of a motion event. Talmy (2000) proposes that manner is external, because linguistic structure demonstrates that it is conceptualized as a separate event (p. 37). Slobin (2004), on the other hand, says that “manner covers an ill-defined set of dimensions that modulate motion, including motor pattern, rate, rhythm, posture, affect, and evaluative factors” (p. 5). Three types of manner have been suggested: a. Default or Unmarked Manner (*e.g.* to go or to come), b. Forced or Marked Manner (*e.g.* to limp or to tiptoe), c. Instrumental Manner (*e.g.* to ski or to drive) (Pourcel, 2004).
- II. CAUSE: Cause expresses the reason of a motion event.

There are mainly three types of motion events: *self-agentive* or *voluntary motion events*, in which the subject moves as in (3); *caused motion events*, in which the object moves as in (4); and *fictive motion events*, in which there is no physical movement included but where the elements of a motion event are present as in (5) (Talmy, 2000).

- (3) The old lady is tiptoeing down the staircase.
- (4) The little girl rolled the ball down the slope.
- (5) The fence goes from the plateau to the valley.²

Pourcel and Kopecka (2006) also make a distinction between a *motion event* and a *motion activity*. In their conceptualization, the sentence expressing a motion activity

² Example 5 is taken from Talmy (2000, p. 99).

provides a locational reading, as opposed to the directional reading provided by a motion event. Moreover, the core component of a motion activity is manner, whereas a motion event centers around path. They also emphasize that a motion activity does not convey path information at all. According to this definition, (6) expresses a motion event, whereas (7) is just an example of a motion activity (*cf.* Aske, 1989).

(6) Jane is running into the garden.

(7) Jane is running (in the garden).

2.2.3. Forming a Typology of Motion Events

2.2.3.1. Talmyan Typology

Leonard Talmy, a well-known contemporary cognitive linguist, says that his work mostly deals with “the systematic relations in language between meaning and surface expression” (1985, p. 57). He defines meaning in terms of semantic elements such as motion, path or manner; and surface expression as the overt linguistic elements like verb, adposition or gerund. He then proposes two ways for exploring this meaning-surface relationship: either holding the semantic entity constant or holding the surface entity constant. In forming his motion event lexicalization typology, Talmy (1991) follows the first track. He takes the *path of motion* as being the core element of a motion event, and he claims that a motion event cannot exist without a trajectory followed by the figure.

Since the figural entity of any particular framing event is generally set by context and since the activating process [the motion] generally has either of only two values, the portion of the framing event that most determines its particular character and distinguishes it from other framing events is the schematic pattern of association with selected ground elements into which the figural entity enters. Accordingly, either the relating function alone or this together with the particular selection of involved ground elements can be considered the schematic core of the framing event... the relating function that associates the figural entity with the ground elements among which the transition takes place constitutes the *path*. The core schema here will then be either the path alone or the path together with its ground locations. (p. 483, emphasis original)

After setting this basic ground, he puts forward a typology of languages based on their dominant motion event lexicalization patterns. In this typology, Talmy splits languages of the world into two main categories according to their ways of framing the core element, namely the *path of motion*, in a clause/sentence. On the one hand, there are *verb-framed languages* (*V-languages*), which frame the path component in the main verb of a clause or sentence, and on the other hand, there are *satellite-framed languages* (*S-languages*) that frame the same component in a peripheral element called a satellite.³

The world's languages generally seem to divide into a two-category typology on the basis of the *characteristic* pattern in which the conceptual structure of the macro-event is mapped onto syntactic structure. To characterize it initially in broad strokes, the typology consists of whether the core schema [framing event] is expressed by the main verb or by the satellite. (Talmy, 2000, p. 221, emphasis added)

The term “characteristic” that is used in the above quotation is actually the key point in Talmyan theory. Talmy never claims that each language has a one-and-only way of expressing motion events but he rather suggests that each language has its most characteristic way of expressing motion events. He also explains what he means by characteristic: it is colloquial in style, frequent in occurrence, and pervasive in use (2000, p. 27).

2.2.3.1.1. *V-Languages*

V-languages prefer embedding the PATH element into the main verb; therefore, in those languages, there is a great number of verbs naturally indicating the path of movement; such as the French verbs *entrer*, ‘to go in’ or *sortir*, ‘to go out’. The MANNER element is expressed by a gerund or an adverbial in those languages, as in the French examples (8a) and (8b.) Romance languages (*e.g.* French, Spanish or Italian), Turkic languages (*e.g.* Turkish), Semitic languages (*e.g.* Hebrew or Moroccan Arabic), Japanese and Basque are the best-known examples of a V-language (Slobin, 2003).

³ Talmy also talks about a third type of language conflating motion with ground; however as this language type is beyond the scope of the current study, it will not be mentioned here.

(8) a. La dame *entre* dans le café en courant.
 ‘The lady is entering into the coffee house (by) running.’

 b. Alex *sort* de la chambre sur la pointe des pieds.
 ‘Alex is exiting from the room on the tip of the toes.’

V-languages also license the use of a manner verb as the main verb but with a certain restriction. It is a dominant characteristic of a V-language to mark a change of state, and while expressing motion events, the change of state includes the crossing of a boundary, as in *enter* or *exit*. Therefore, it is the *boundary crossing constraint* in V-languages that does not allow the use of a construction like “run into/out of”. However, manner verbs may very well be used as main verbs in cases where there is no boundary crossing. That is why (9) is a legitimate use in French, a V-language; whereas (10) is not (Slobin & Hoiting, 1994). Some authors relate the issue of boundary crossing to path telicity (see Aske, 1989; and Montrul, 2001).

(9) Pierre court dans le parc.
 ‘Pierre is running in the park.’

(10) Pierre court dans le parc.
 *‘Pierre is running into the park.’

There is a widely-accepted claim in the literature arguing that manner of motion is an optional component in V-languages, and that those languages express the manner information in very rare occasions (Talmy, 2000; Özçalışkan & Slobin, 2000; Papafragou *et al.*, 2002; Gennari *et al.*, 2002; Özçalışkan & Slobin, 2003; Pourcel, 2005; Slobin, 2004, 2006; Gullberg *et al.*, 2008; Soroli & Hickmann, 2010 among others). In all of these studies, it is concluded that there is nothing that prevents V-language speakers expressing the manner of motion but that they prefer to omit it in most of the cases. Gullberg *et al.* (2008) even conduct a control experiment to prove that V-language speakers can actually use manner verbs in their languages. The main reason for such an asymmetry is suggested to be the *inferability of manner* (Papafragou *et al.*, 2002, 2006, 2007; and Beavers *et al.*, 2010). For example, the manner expressing gerund can be omitted in the French sentence (11), as the manner

of motion can easily be inferred from the context ('flying' as the default manner of motion of birds).

- (11) L'oiseau est sorti de la caverne (en volant).
 'The bird exited from the cave (by flying).'

However, a serious criticism against this common idea comes from Allen *et al.* (2007) who emphasize the methodological shortcomings of the existing studies by saying that:

Many of the existing studies have focused on situations in which the Manner of motion has been optional rather than mandatory and salient in the context (*e.g.*, in Frog Story narratives in Özçalışkan & Slobin, 1999; and in spontaneous speech in Choi & Bowerman, 1991). In cases where both Manner and Path of motion are salient and occur simultaneously, children speaking verb-framed languages have to encode Manner as well as Path in their speech to describe such events. Thus, they have to learn how to syntactically package Manner and Path together. (p. 22) ⁴

2.2.3.1.2. S-Languages

As for S-languages, they prefer integrating the MANNER component into the main verb, and these languages have a great number of verbs conflating manner with motion, such as the English verbs *crawl*, *stagger* or *limp*. In S-languages, the PATH component is lexicalized with a satellite construction, such as a particle or an affix, as in the English (12a) and German (12b) examples. Typical examples of an S-language are Germanic languages (*e.g.* English or German), Slavic languages (*e.g.* Russian or Polish), Finno-Ugric languages (*e.g.* Finnish or Hungarian), and Sino-Tibetan languages (*e.g.* Mandarin Chinese) (Slobin, 2003).

- (12) a. The lady is *running into* the café.
 b. weil da eine Eule plötzlich raus-flattert.⁵
 because there an owl suddenly out-flaps
 'because an owl suddenly flaps out.'

⁴ This point will be questioned and discussed in the following chapters.

⁵ The German example is taken from Slobin (2004, p. 6).

2.2.3.2. *Slobian Typology*

Dan I. Slobin claims that the two-way typology proposed by Talmy does not cover all the languages in the world, and thus proposes a revised typology. His typology introduces a third type of language in between the V-language and S-language categories, which he calls *Equipollently-framed Languages*. In this type of languages, manner and path are expressed by “equivalent grammatical forms” (2004, p. 25). He proposes three construction types under equipollency:

- I. *Serial-Verb Languages* (MannerVerb + PathVerb): It is not clear which one of the two verbs is the main verb of the sentence. Niger-Congo, Hmong-Mien, Sino-Tibetan, Tai-Kadai, Mon-Khmer, Austronesian languages are examples of this type of languages.
- II. *Bipartite-Verb Languages* ([Manner + Path]_{Verb}): The verb includes two morphemes of equal status, one denoting the manner and the other denoting the path of motion. Algonquian, Athabaskan, Hoka, Klamath-Takelman languages are examples of this type.
- III. *Generic-Verb Languages* (MannerCoverb + PathCoverb + GenericVerb⁶): Jaminjung language spoken in Northern Australia is an example of this type of language.

Slobin is not the only one who suggests a three-way typology. Zlatev and Yangklang (2004) who work on Standard Thai, and Ameka and Essegbey (2001) who investigate serial-verb languages in West Africa also propose a third group of language.

When the properties are tallied, we find that serialising languages share more properties with S-languages than with...V-languages...while still possessing a unique property. What this shows is that they cannot be said to belong to either type. Instead, they appear to belong to a class of their own. (Ameka & Essegbey, 2001)

⁶ These are the general verbs, also called neutral verbs or plain verbs, which do not contain any specific manner or path flavour, such as *to come* or *to go* in English.

Even though Slobin's three-way typology idea was welcomed in the literature, as it only applies to a limited number of languages spoken in smaller communities, the two-way version of the typology still preserves its popularity among researchers.

Slobin (2004, 2006) further suggests that languages should not be categorized solely based on lexicalization patterns and by using dichotomies or trichotomies, rather they should be investigated discursively and should be placed on a continuum. He takes the salience of the manner of motion as the starting point for his continuum, and he calls it a *cline of manner salience*. With this idea in mind, he distinguishes between high-manner-salient languages and low-manner-salient languages. Though his new continuum also has two ends, he does not aim for a dichotomic approach. He rather has both a synchronic and a diachronic approach to manner salience. Synchronically a child whose native language is a high-manner-salient language develops a richer lexicon of manner verbs than a child having a low-manner-salient language as her/his native language. Diachronically speaking, a language can move along the cline over time, for example Italian that moves from less-manner-salience to higher-manner-salience most probably due to its contact with German (Hottenroth, 1985, as cited in Slobin, 2004).

The key difference between Talmy's and Slobin's approaches is that Talmy bases his motion event typology on lexicalization patterns and especially those of *path of motion*, whereas Slobin's typology is rather discursive and mostly based on the *manner of motion* (Pourcel & Kopecka, 2006).

Slobin's crucial addition to motion typology has been to restore an empirical focus on usage-based epistemologies aimed at documenting patterns beyond the sentence level in terms of lexical resources, discursive patterns, rhetorical styles and habitual fashions of speaking. His point has been to integrate all levels of language in order to obtain a more holistic understanding of the dynamics of motion expression. In addition, his work illustrates a descriptive tradition concerned with the empirical data collection of language in use. (Pourcel & Kopecka 2006, p. 11)

2.2.3.3. New Trends of Typology Formation

Recent studies, criticizing Talmy's dichotomy and Slobin's continuum to be insufficient in giving a crosslinguistic account of motion event representation, all try to come up with new and broader typologies. These newly proposed typologies have three main characteristics:

- *As opposed to the manner-based classification of Slobin, they are path-based.*
- *Instead of attesting a language into a single category, they emphasize the variation within a language.*
- *They claim to be construction-based (based on particular situations) rather than being language-based.*

Although the present study will take the main *S-language* vs. *V-language* dichotomy as the framework, a brief overview of the main new trends will also be provided here for the sake of chronological and theoretical completeness of the review.

Pourcel (2004, 2005): Pourcel emphasizes the salience of path of motion in response to Slobin's manner salience idea. According to her, Path is the central and cognitively most salient aspect of a motion event, thus languages should be defined based on their path representation (also see Ibarretxe-Antuñano, 2008).

With humans being meaning-seeking creatures, it appears very likely that the purpose-loaded dimension of motion should therefore be the dimension receiving higher levels of cognitive salience across species members and hence across language groups. As a rule, the end justifies the means, and it is possible that the means, or the Manner of motion in this case, is secondary in human actions. (Pourcel, 2004, p. 510)

Pourcel and Kopecka (2006): They claim that it is hard to confine a language to a single category and provide a detailed analysis of French to support their hypothesis. At the end of their article, they propose that at least for French, there is a *mixed-pattern* that includes the use of more than one lexicalization pattern. They also emphasize the importance of taking semantic and pragmatic factors into consideration, and that of observing motion scenarios instead of single motion events.

Zlatev, Blomberg and David (2010): Zlatev *et al.* provide an experiment-based classification of motion events. They take a phenomenological perspective by emphasizing the importance of taking subjective reality as conceived by human beings rather than the objective reality itself as their norm. They start by questioning the basic terms and assumptions taken for granted in the field. For example, they no longer talk about motion events but *motion situations*, which include both motion events and motion activities.

In this chapter we have tried to show that “motion event” typology has suffered for quite some time from conceptual and empirical problems, and despite the undoubtable contributions of scholars such as Talmy and Slobin, it is time that we move beyond, and establish a more coherent framework for describing our experiences of motion. (p. 24)

Beavers, Levin and Tham (2010): Beavers *et al.* support the idea that all languages belong to a mixed-type, where you can observe more than one way of expressing motion events; however they still emphasize the importance of a dominant pattern.

Although a particular language may have multiple options available for encoding manner and path, some may be preferred on independent grounds, for example due to morphosyntactic complexity or to preferences for certain types of lexemes over others within the lexical inventory of a language. (p. 366)

Croft, Barddal, Hollmann, Sotirov and Taoka (2010): Croft *et al.* are among the strong supporters of the construction-based approach, which proposes to typologize particular events instead of typologizing languages. Proponents of this approach concentrate more on variations within a language, and make more fine-grained semantic and syntactic distinctions (also see Fortis & Vittrant, 2011 along the same lines).

Languages make use of multiple strategies to encode complex events, depending of the type of complex event involved. This follows the more general trend in typological research away from typologizing languages as a whole - which usually leads to declaring that all languages are a “mixed” type - to typologizing particular situation types expressed in a language. (p. 231, emphasis original)

2.2.4. Typological Classification of the Languages in the Study

In spite of the rise of the new trends discussed above, it is an undeniable fact that Talmy (1985)'s V-language vs. S-language dichotomy has constituted, and still constitutes the baseline for most of the psycholinguistic typological studies in the field (Choi & Bowerman, 1991; Berman & Slobin, 1994; Slobin, 1996b, 1997, 2004; Naigles *et al.*, 1998; Özçalışkan & Slobin, 1999, 2003; Papafragou *et al.*, 2002, 2007; Selimis & Katis, 2003; Selimis, 2007; Allen *et al.*, 2007 among many others). The reason for this continuing interest of researchers into Talmyan typology is expressed by Slobin (2004), who is one of the main critics of Talmy himself, as follows: "The satellite- versus verb-framed typology has been useful in systematically sorting the world's languages as well as providing a framework for discourse analysis" (p. 24). Therefore, taking Talmy's typology as the baseline does not mean that the world is split into two camps with a sharp knife but just that "the most characteristic ... way of describing motion in the two languages involves manner and path verbs respectively" (Papafragou & Selimis, 2010, p. 227).

Keeping this broad perspective in mind, the present dissertation will also take the Talmyan dichotomy as the framework for the experiments used. Therefore, we will now provide brief information about the place of the three languages that are used in the current study (Turkish, English and French) within this framework.

2.2.4.1. Motion Events in Turkish

Turkish is qualified as a *verb-framed language* or *path language*. The main characteristic of verb-framed languages is that PATH is typically encoded in the main verb, as in the examples (13) and (14). As for the MANNER component, it is generally expressed by alternative means such as the use of a subordinated construction or an adverbial.

- (13) Çocuklar koşarak merdivenlerden *indiler*.
 manner adv. V_{motion+path}
 ‘The children *descended* from the stairs (by) running.’

- (14) Bebek emekleyerek oturma odasına *girdi*.
 manner adv. V motion+path

‘The baby *entered* to the living room (by) crawling.’⁷

Turkish motion events have mainly been investigated by Dan I. Slobin, Şeyda Özçalışkan and Aslı Özyürek; and the following is a concise review of their main studies on motion events in Turkish.

Özçalışkan and Slobin (1999, 2000): In this pioneering psycholinguistic study, Özçalışkan and Slobin tested child and adult speakers of Turkish (V-language), English (S-language) and Spanish (V-language) in a verbal production task. They used the renowned picture book *Frog, where are you?* (Mayer, 1969) as stimulus materials, and asked their participants to describe what was happening in each picture verbally. The results showed that Turkish as well as Spanish conforms to the V-language pattern, using path verbs to a great extent; whereas English speaker descriptions were more in line with the S-language pattern, and they mostly made use of manner verbs together with path satellites. They also found that children’s descriptions reflect a typological effect as early as 3 years of age, which means that children learn the language-specific lexicalization pattern of their own language very early in life.

Özyürek and Özçalışkan (2000): Özyürek and Özçalışkan’s study also involved a description task; however they observed not only the verbal productions but also the accompanying gestures of the participants. They tested child and adult native speakers of Turkish with an animated cartoon, and asked them to narrate what they watched. Based on the results obtained both from the verbal and the gestural data, they concluded that gestural representation of spatial elements become in line with the linguistic encodings in Turkish, as early as 6 years of age. They suggested that further research comparing Turkish with other languages was necessary to arrive at a more precise conclusion, though.

⁷ Turkish is an agglutinative language and it makes use of a great number of suffixes for various purposes. It is quite noticeable in sentence (14) that the dative case marker “-a” contributes to the path information as well as the path meaning embedded in the main verb *girmek* ‘to enter’. As a detailed semantic-syntactic analysis of motion events is beyond the scope of the current dissertation, this interesting point will not be discussed here. However, the readers who are interested in the issue may find Sinha and Kuteva (1995)’s *Distributed Spatial Semantics Theory* valuable in this regard.

Allen, Özyürek, Kita, Brown, Furman, Ishizuka and Fuji (2007): In this crosslinguistic study, Allen *et al.* tested child speakers of Turkish (V-language), English (S-language) and Japanese (V-language) in order to see whether children of a very young age (mean age 3;8 years) could package the semantic elements of manner and path onto syntactic units as adult speakers of the languages do or whether there is a universal pattern of packaging free of any language-specific effects. They made use of the *Tomato Man Movies* prepared by the Language and Cognition Group at Max Planck Institute for Psycholinguistics (Özyürek, Kita, & Allen, 2001), where motion events are performed by a red tomato-like creature. Their results suggested that:

2.2.4.2. Motion Events in English

(15) The man *limped* out of the apartment building.
 ↓ ⊥ path expressed with a satellite
 manner encoded in the verb

18

(17) He is *running* up the staircase.

Motion events in English have extensively been studied, and in the following a brief summary of some recent studies on motion events in English are provided.

Cifuentes-Férez and Gentner (2006): In this study, the researchers investigated motion events in English (S-language) and Spanish (V-language) comparatively. They used an experimental technique called *novel word mapping technique* (Nagy & Gentner, 1990) where the participants are tested for any language-specific effects when inferring the meaning of a new word presented in a text. The following is an example passage given to the subjects in the study, and they were supposed to predict the meaning of the newly-heard word (here, the verb *ransin*) based on the context.

*So she decided to keep walking up the river and look for a bridge. After a while she noticed the river had become shallow and not so dangerous. So she took off her shoes and socks, rolled up her jeans and **ransined** the river. That night she was very happy to be back among friends again.*

Each passage included either a novel verb or a novel noun, and the aim of the novel verb task was to see whether the subjects would make more manner-based or path-based predictions. The results were in line with the V-language vs. S-language typology; English speakers using (predicting) more manner verbs than path verbs, and Spanish speakers using (predicting) more path verbs than manner verbs.

Papafragou and Selimis (2009): Papafragou and Selimis tested child and adult speakers of English (S-language) and Greek (V-language) in a verbal description task. They used short animated motion clips as stimuli. The results showed a clear typological effect both for the child and the adult group. English speakers used a greater number of manner verbs, whereas Greek speakers used a greater number of path verbs. The researchers also used a verb learning task somewhat similar to that of Cifuentes-Férez and Gentner (2006). The results of this second task also showed a language specific effect, English speakers making more manner guesses and Greek speakers more path guesses.

Soroli and Hickmann (2010): Soroli and Hickmann used both real-life video clips and animated cartoons as their stimuli. They tested adult native speakers of English (S-language) and French (V-language) in a verbal description task. They analyzed the sentences produced by their subjects not only from a typological perspective (manner verb use vs. path verb use) but they also looked at the *semantic density* levels of their utterances. The term semantic density (SD) is used to refer to the inclusion of manner-information-only or path-information-only (SD1) in a sentence, compared to the inclusion of both components simultaneously (SD2)⁸. The results indicated that there was indeed a clear typological effect, English speakers giving manner information and French speakers giving path information in the main verb. The data also showed a clear difference in the semantic densities of the two groups' descriptions. 90% of English responses encoded both components (manner and path), whereas only 55% of French responses encoded them both (45% contained only one component, mostly the path of motion).⁹

2.2.4.3. *Motion Events in French*

French is also categorized as a *verb-framed* or a *path language* by most of the researchers (but see Kopecka, 2004; and Pourcel & Kopecka, 2006), along with other Romance languages. Therefore, it demonstrates a lexicalization pattern similar to Turkish in expressing motion events. In French, PATH is typically encoded in the main verb as in (18) and manner is mostly expressed with an adverbial as in (19) (Choi-Jonin & Sarda, 2007).

- | | | | | |
|------|----|---|----------------|--------------------|
| (18) | Il | <i>descend</i> | les escaliers. | |
| | | motion+path | | |
| | | ‘He is <i>descending</i> the staircase.’ | | |
| | | | | |
| (19) | Il | <i>descend</i> | les escaliers | <u>en courant.</u> |
| | | motion+path | | manner adv. |
| | | ‘He is <i>descending</i> the staircase <u>(by) running.</u> ’ | | |

⁸ Here, they refer to the already mentioned claim that V-languages most of the time use path-information-only, whereas S-languages use both manner-information and path-information in their utterances (see section 2.2.3.1.1 for more details).

⁹ See Papafragou *et al.* (2002), Gennari *et al.* (2002), Skordos and Papafragou (2010), and Ibarretxe-Antuñano (in press) for other crosslinguistic studies including English.

Sablairolles (1995), who forms a typology of intransitive motion verbs, distinguishes three categories; change-of-location verbs (*e.g. entrer* ‘to go in’, *sortir* ‘to go out’), change-of-position verbs (*e.g. voyager* ‘to travel’, *courir* ‘to run’), and change-of-posture verbs (*e.g. s’asseoir* ‘to sit down’, *se baisser* ‘to bend down’) (pp. 281-282). According to his classification, a change-of-location (CoL) verb is what is called a *path verb* in Talmyan typology. Thus, it is interesting to note that Sablairolles found 440 CoL verbs in French as a result of a fine-grained analysis. This finding verifies the claim that French, as a V-language, has a large lexicon of path verbs.

Another fine-grained classification is performed by Choi-Jonin and Sarda (in press) who further classify path verbs by giving examples from French. They propose that there are passage verbs (*e.g. traverser* ‘to cross’), orientation verbs (*e.g. monter* ‘to climb’), and distance verbs (*e.g. suivre* ‘to follow’).

The following is a brief review of the recent studies conducted on motion events in French.

Pourcel (2004): Pourcel analyzed motion event uses in French (V-language) and English (S-language). As a result of a verbal elicitation task, she suggested that her results indicated a clear crosslinguistic difference between the two language groups; English speakers using more manner verbs than French speakers and *vice-versa*. She also found “differential foregrounding of manner in French and English” (p. 507), and argued that different from English which expresses both the manner and path components in almost all the sentences, French has a tendency to background the manner element.

Pourcel and Kopecka (2005, 2006): In these two papers, Pourcel and Kopecka challenged the verb-framedness of French by presenting very detailed linguistic analyses. They claimed that the verb-framed pattern is not the only legitimate pattern in French language, as proposed by Talmy (1985). They noticed in their data (composed of elicited written and oral narratives) that there are four more motion event lexicalization patterns other than the classical V-framing pattern (p. 27):

1. The juxtaposed pattern ($V_{\text{Manner}} + V_{\text{Path}}$)
 Il *court* dans une rue puis *rentre* dans une maison.
 ‘He is *running* on a street and *entering* a house.’

2. The hybrid pattern ($V_{\text{Manner+Path}}$)
 Il *dévale* les escaliers.
 ‘He is *rushing down* the stairs.’

3. The satellite-framing pattern ($V_{\text{Manner}} + \text{Satellite}_{\text{Path}}$)
 Il *tire* Charlot *hors de l’eau*.
 ‘He is *pulling* Charlie *out of* the water.’

4. The reverse pattern ($V_{\text{Manner}} + \text{Adjunct}_{\text{Path}}$)
 Il *marche le long de* la route.
 ‘He is *walking along of* the road.’

Gullberg, Hendriks and Hickmann (2008): Gullberg *et al.* investigated how child and adult speakers of French talk and gesture about voluntary motion events. They elicited data *via* a set of short animated cartoons, and the results showed that when path and manner components are equally salient in the scene, both children and adults focus more on the path of motion and less on the manner of motion in their utterances. The data also suggested that speech and gesture are mostly co-expressive at all ages.

Hickmann, Taranne and Bonnet (2009): Hickmann *et al.* tested child and adult speakers of French (V-language) and English (S-language) in a verbal elicitation task. They used a set of short animated cartoons as stimuli. The results showed that English speakers express both manner and path in compact structures, whereas French focuses more on path (see p. 735 for the exceptions). From a developmental perspective, the data also suggested that children are affected from typological factors when formulating their utterances from three years of age on.

Soroli (2011): Soroli administered a video description task on adult native speakers of French (V-language), English (S-language) and Greek (V-language). She used

both animated cartoons and real-life video clips as stimuli. Her results suggested a clear typological pattern between French and English; French subjects using path verbs most of the time and English subjects making dominant use of manner verbs. On the other hand, her data regarding the Greek descriptions revealed an ambiguous pattern depending on how she encoded the prefixes used by the subjects¹⁰.

2.3. LANGUAGE AND COGNITION

2.3.1. The Language and Cognition Debate in General

The relationship between the human cognitive and language system has been a major matter of interest for researchers in several different fields (*e.g.* philosophy, psychology or linguistics). The inquiries questioning a possible effect of one on the other even date back to Greek philosophy (Carroll, von Stutterheim, & Nüse, 2004). The classical approach, which claims that language is just a medium for expressing thought, was challenged by the groundbreaking argument suggesting that language is not just a passive medium but an alive system that can also mediate our ways of thinking. This claim is expressed by Benjamin Lee Whorf (1956), who is regarded as the father of the idea, as follows:

...language is not merely a reproducing instrument for voicing ideas but rather is itself the shaper of ideas, the program and guide for the individual's mental activity, for his analysis of impressions, for his synthesis of his mental stock in trade. Formulation of ideas is not an independent process ... but is part of a particular grammar, and differs, from slightly to greatly, between different grammars. (pp. 212–213)

This research question, which has occupied the minds of a great number of people, is expressed by von Stutterheim and Nüse (2003) in the following terms: “To what extent are the planning processes required for the construction of an informational network expressed in a piece of discourse related to structural properties of the specific language used?” (p. 856). The same authors propose three levels of representation of information organization in the human mind, namely the outside world, conceptual representation and linguistic representation. They suggest that the

¹⁰ See Choi-Jonin and Sarda (2007, in press) for other studies on French motion events.

language and cognition debate should examine the relationship between the second and the third level.

This highly controversial issue has been largely discussed by many researchers from several disciplines (see Gumperz & Levinson, 1996; and Gentner & Goldin-Meadow, 2003 for book-length reviews). These researchers may be categorized under two major camps according to their positions regarding the question of ‘whether language has an effect on thought’: *relativists* and *universalists*. The following two subsections will talk about the main arguments proposed by the two camps, and about their major supporters.

2.3.1.1. Relativity

The relativists who claim that language has a shaping effect on the human thinking system can also be evaluated under two sub-camps, the proponents of the strong view (*linguistic determinism*) and the weak view (*linguistic relativity*).

2.3.1.1.1. The Strong Version (Linguistic Determinism View)

The names Edward Sapir (1884-1939) and Benjamin Lee Whorf (1897-1941) are the ones mostly associated with this line of thought, and that is why this hypothesis is also known as the *Sapir-Whorf Hypothesis*. However, Sapir and Whorf are not actually the first researchers who put forward such a relativistic approach regarding the language and cognition interface. The German philosopher Wilhelm von Humboldt (1767-1835) had voiced this approach about a century before Sapir and Whorf did:

There resides in every language a characteristic world-view. As the individual sound stands between man and the object, so the entire language steps in between him and the nature that operates, both inwardly and outwardly, upon him ... Man lives primarily with objects, [but] ... he actually does so exclusively as language presents them to him. (von Humboldt, [1836] 1999, p. 60)

However, Whorf (1956) was certainly the one who took the subject as a serious matter of investigation, and who looked for connections between semantic-syntactic

structures of a language and the habitual thought patterns revealed by its speakers. He was the first one to adopt an empirical and crosslinguistic perspective of inquiry by comparing the English language and the Hopi language in this respect (Lucy, 1996).

The best known supporters of this deterministic view, which implies that language does not only lead to language-specific conceptual structures but it also binds its speakers' mode of thinking, are Brown and Lenneberg (1954) who tested the Whorfian hypothesis by using experiments on color term uses and who had positive results supporting the idea. However, the studies following those of Brown and Lenneberg yielded controversial results, which led to a period of skepticism vis-à-vis the hypothesis (see Gentner & Goldin-Meadow, 2003 for the details of this period). Chomsky's emphasis on universality (Chomsky, 1957) also contributed to this dark period of relativity research. Then after a long time of neglect, the issue of language and cognition came to the foreground again in the 1970s with great efforts of Talmy (1985) and Langacker (1987) who analyzed the structures of different languages and who demonstrated considerable differences among them. The shift of the domain of investigation from color to space, a much more promising area of investigation, contributed to the revival of the debate, as well (Gentner & Goldin-Meadow, 2003). All these new developments resulted in a great number of fresh inquiries, which came together in the renowned 1991 conference entitled "Rethinking Linguistic Relativity" (Wenner-Gren Foundation Symposium), and the 1996 volume edited by Gumperz and Levinson, a detailed review of the field. The best-known contemporary supporters of the linguistic determinism view are Lucy (1992) and Levinson (1997).

2.3.1.1.2. The Weak Version (Linguistic Relativity View)

The weak version of the Whorfian hypothesis also supports the idea of an effect from the language side to the thinking side; however it claims that it is not a pervasive influence but rather a limited one observed under certain circumstances. Gentner and Loewenstein (2002) express this idea in the following terms:

Languages vary in their semantic partitioning of the world, the structure of one's language influences the manner in which one perceives and understands

the environment, and therefore, the speakers of different languages should have at least *partly incommensurable* world views. (p. 102, emphasis added)

Soroli and Hickmann (2010) also argue that language has a partial effect on our conceptual representations, as language serves as a filter channeling the incoming information (p. 582). But how can we define this partial effect? Or what are those certain circumstances that lead to language-specific representations? The most ground-breaking answer to those questions is given by Slobin (1996a) with his hypothesis of *thinking-for-speaking* (TFS). According to the view, language can only have an effect on non-linguistic conceptual representation if there is a communicative intention involved in the task.

The language and languages that we learn in childhood are not neutral coding systems of an objective reality. Rather, each one is a subjective orientation to the world of human experience, and this orientation **affects the ways in which we think while we are speaking**. (p. 91, emphasis original)

Slobin expresses that he follows the tradition of the anthropological linguist Franz Boas (1938) who also underlined the importance of the linguistic expression aspect. On the other hand, Gennari *et al.* (2002) proposes a new form of the weak version, which they call *Language-as-Strategy View*. This view suggests, also in line with Slobin TFS view, that language may have an influence on non-verbal performance only when language can serve as a mediator.

2.3.1.2. Universality

The universalist approach, mostly represented by the work of Ray Jackendoff (1990, 1996), claims that the human conceptualization system is bound by universal characteristics. According to this classical theory of language and cognition, the differences observed between languages do not reflect the conceptual representation; they are just reflections of syntactic and lexical linguistic representations.

The supporters of universality harshly criticize the view that there may possibly be a language-specific conceptual structure underlying the crosslinguistic differences observed in several psycholinguistic studies. Pinker (1994), who is a universalist himself, expresses this rather widely held skeptical view with the following words:

“Most of the experiments have tested banal ‘weak’ versions of the Whorfian hypothesis, namely that words can have some effect on memory or categorization.” (p. 65, emphasis original).

The same line of thinking is also voiced by Munnich and Landau (2003) who, after a detailed review of the experimental studies that allegedly proved the Whorfian hypothesis, conclude that “[c]areful inspection of these tasks shows that none of them has been carried out in such a way as to rule out the use of verbal encoding” (p. 148).

2.3.1.3. Experimental Studies

It is a highly demanding task to design experiments to test the above summarized points of views regarding the relationship between language and cognition. The main challenge is to be able to create situations where you have no, or the least possible, linguistic intrusion, so that you can measure the cognition aspect properly. This key point has been emphasized by several researchers engaged in the challenging work of proving or refuting the Whorfian Hypothesis, including Whorf (1956) himself:

To compare ways in which different languages differently ‘segment’ the same situation of experience, it is desirable to analyze or ‘segment’ the experience first in a way independent of any language or linguistic stock, a way which will be same for all observers. (p. 162, emphasis original)

The following is a quick review of the empirical studies investigating the language and cognition interface in two main areas of investigation: color and spatial location.

2.3.1.3.1. Color Studies

Brown and Lenneberg (1954) are regarded in the literature as the pioneers of experimentally-controlled relativity studies, as opposed to studies using the natural observation technique. Their studies demonstrated a positive relationship between the codability of color terms in English, and English speaker’s perception of color evaluated by a memory task (Gentner & Goldin-Meadow, 2003). However, another study comparing the color perceptions of English and Dani speakers (Heider &

Olivier, 1972) proved the contrary. They showed that Dani people of New Guinea who have only two basic color terms (dark and light) in their vocabulary, performed in the cognitive tasks similar to English speakers who have 11 terms to describe the same array of colors. Three decades later Roberson, Davies and Davidoff (2000) replicated the study of Heider and Olivier in an extended manner, and they claimed that their results were incompatible with that of Heider and Olivier (see Munnich & Landau, 2003 for a comparative discussion of the two studies). It is obvious that the empirical studies on color terms and color perception revealed contradictory results regarding the *universalist* vs. *relativist* debate.

2.3.1.3.2. *Spatial Location Studies*

After a few decades of intense research on color perception, the need for a more fruitful testing ground became obvious. It was understood that to be able to make a detailed and precise inquiry on the possible effects of linguistic representation on cognitive representation, we should look at higher-level cognitive representations where there is clear linguistic variation (see Papafragou *et al.*, 2002 for a discussion). Thus, the next field of investigation was spatial location and spatial relations. The reason for choosing *space* as the testing ground is explained by Gentner and Boroditsky (2001) in the following terms:

Verbs and other relational terms – including those concerned with spatial relations – provide framing structures for the encoding of events and experience; hence a linguistic effect on these categories could reasonably be expected to have cognitive consequence. (p. 247)

In their pioneering study, Brown and Levinson (1993) had the aim of seeing whether the variation found in the verbal use of spatial frames of reference is also reflected in non-verbal encoding of spatial relations. To this end, they administered both verbal and cognitive tasks on speakers of Dutch and Tzeltal, and showed that the non-linguistic results of both groups were consistent with their verbal responses. The results of a more fine-grained study by Pederson *et al.* (1998) were also similar to those of Brown and Levinson, and were in line with the relativity hypothesis.

On the other hand, Li and Gleitman (2002) replicated the Brown and Levinson study by changing some parameters, and proved them wrong. They found out that the nature of the environment had a strong effect on people's choice of the frame of reference in non-linguistic tasks; therefore they manipulated the spatial environment and saw that the richness of the environment had a direct effect on people's performances. They tested speakers of English and showed that their results are compatible with that of Tzeltal speakers, which led them to the conclusion that the variation found by Brown and Levinson were actually not due to a language-specific effect but rather a task effect (Munnich & Landau, 2003)¹¹.

2.3.2. The Language and Cognition Debate on Motion

Before moving on to a review of the experimental motion studies conducted to shed light on the intricacies of *the language and cognition interface*, we will discuss the rationale behind choosing motion as the perfect testing ground for neo-Whorfian inquiries.

2.3.2.1. Motion as the Perfect Testing Ground

Motion is an ideal domain for such a challenging line of investigation for several reasons. First of all, the representation of motion is a fundamental but also a high-level human cognitive ability, which may lead researchers deep into cognition (Landau & Jackendoff, 1993). Second, motion terms are acquired rather early in childhood (Choi & Bowerman, 1991). Third, every language has a way of speaking about motion; however there is variability in the linguistic encoding of motion events. Fourth, its expression extends beyond the lexical level, it reaches to sentence and even discourse level (Slobin, 2004).

2.3.2.2. Experimental Studies on Motion Events

The *verbal expression and non-verbal representation of motion events* has been a popular and rich area of psycholinguistic and cognitive experimentation during the

¹¹ Munnich and Landau (2003) point out another flaw in the Brown and Levinson tasks by suggesting that their experimental designs do not “rule out the use of verbal encoding” (p. 148).

last decade. Lucy (1996), who underlines the flaws of the relativity studies administered until 1990s, calls for a new comparative approach which is; (1) dealing with two or more languages, (2) examining a significant language variable, (3) measuring cognitive performance free from linguistic encoding situations, (4) dealing with categories from real-life relations. The following is a review of the main types of recently-conducted experimental studies based on the above principles.

2.3.2.2.1. *Categorization Studies*

One of the major experimental methodologies used to investigate the language and cognition debate is administering a task called *Similarity Judgment* or *Categorization*. The task is composed of triads of pictures or video clips depicting motion events. Even though there may be nuances in the application of the experimental procedure, the general pattern of the task is as follows: The first item of each triad is presented as the main / prime / target item, and the subjects start the task by seeing that main item. Then, they are exposed to two candidate / alternative items, between which to perform their similarity judgment. One of the alternatives shares the same manner with the main item, and the other alternative shares the same path with it. At the end of each triad, the subjects are asked to choose (orally, in written or with the help of a mouse-click) the alternative item, which, in their opinion, is more similar to the main item. The number of triads may change from one study to the other. This task, as well as the recognition memory task that will be mentioned in the following subsection, is a good way of assessing the non-linguistic spatial representations of people (Gennari *et al.*, 2002). This technique has been used by a great number of researchers during the last decade, and some of them will be detailed here (and some of them will just be named due to space limitations) so as to show the methodological variation included in the application of the task.

Gennari, Sloman, Malt and Fitch (2002): Gennari *et al.*'s study was one of the milestone studies that used the similarity judgment technique. In their study, they made use of short real-life videotapes, and to maximize the homogeneity of the stimulus materials, all motion events were performed by the same male actor with the same clothing and in a small number of settings. They tested native speakers of English (S-language) and Spanish (V-language) under three different conditions. The

first one is the “Naming First Condition” in which the subjects first verbalize the motion events that they see and then perform the similarity task. The second one is the “Shadow Condition” where the subjects repeat a series of nonsense syllables while watching the videos; and the third one is the “Free Encoding Condition” in which the similarity task is performed without any additional components¹². As a result, they found that the nature of the encoding had a major effect on similarity results in both languages. They found a language-specific effect in the Naming First Condition, with Spanish speakers making more same-path choices than English speakers, but not in other conditions. Gennari *et al.* believe that these results are consistent with the *Language-as Strategy View*, which claims that the language effect only occurs after verbal encoding of the material.

Papafragou, Massey and Gleitman (2002): Papafragou *et al.* also performed a similar study, where they tested child and adult speakers of English (S-language) and Greek (V-language) in a similarity judgment task. They made use of a picture book with sets of motion events. The pictures were not drawings but digital color photographs. In order to preserve the dynamic nature of the events, each motion scene was represented by a series of three pictures (one showing the agent at the beginning of the action, one in the middle, and the other at the end) instead of a single picture. The target scene was always presented on the left-hand page of the book, and the two alternatives on the right-hand page. The results showed an almost identical pattern for the two language groups, both English and Greek speakers choosing the same-manner alternates in almost half of the sets and the same-path alternates in the other half. Therefore, no language-specific effect was found in that non-verbal task.

Bohnenmeyer, Eisenbeiss and Narasimhan (2006): The group tested the speakers of a large set of languages (*i.e.* 12 V-languages, four S-languages and one Serial-language) in a categorization task. They made use of the Tomato Man Movies (Özyürek, Kita, & Allen, 2001), where motion events are depicted by a tomato-like creature, as stimuli. Their results do not suggest a simple categorical V-framed vs. S-framed distinction but rather a continuum where most of the S-languages stand in the middle and the V-languages are distributed across the whole scale. Bohnemeyer *et al.*

¹² In their design, all subjects, regardless of the condition they are exposed to, take a recognition memory task before taking the similarity judgment task.

performed additional post-hoc tests to see the effects of different aspects of the stimuli on the results, as well. They found a significant effect of direction of movement (the overall same-manner choices for the ramp scenes were lower than the horizontal movement scenes), manner of movement (the overall same-manner choices for the bouncing-scenes were higher than the rolling-scenes and sliding-scenes), and the ground objects (the tree-rocks scenes yielding lower same-manner percentages than the hut-cave scenes).

Papafragou and Selimis (2010): They tested child and adult speakers of English (S-language) and Greek (V-language) crosslinguistically within the Whorfian paradigm. In their paper, Papafragou and Selimis report the results of three experiments, all of which used the same stimuli and the same categorization technique with slight methodological changes. In the first two versions of the experiment, the stimuli (silent animation clips) were presented on two identical laptop computer screens. In the first version, after watching the target scene (with an audio sentence such as “Look! The turtle is doing something!”), the subjects were instructed with a sentence such as “Do you see the turtle doing the same thing now?”, emphasizing the manner of motion and thus biasing the participants towards that component in their choices. Therefore, in the second version of the experiment, the instruction was given with a more neutral sentence like “Do you see the same now?” without the use of a biasing verb. Even though there seemed to be a language-specific effect in the first version of the experiment, this effect was erased in the second version, and both groups behaved identically, having balanced responses between same-manner and same-path choices. In the third version of the experiment, instead of presenting the stimulus videos consecutively, the researchers presented them simultaneously. They used three laptop computers for that, and the videos were run at the same time in a loop (which made the task cognitively more demanding than the original design). The results of this final design also demonstrated a common and balanced pattern for the two language groups, and thus supported the universalist view (*cf.* Papafragou *et al.*, 2002).

Soroli and Hickmann (2010): In a recent study on motion event conceptualization, Soroli and Hickmann used a similarity judgment task, as well. They tested native speakers of French (V-language) and English (S-language), and their stimulus set

was composed of real-life video sequences. As a result of the experiment, they could not find a significant difference between the two language groups, both having a tendency to choose more same-path alternates than same-manner alternates (also see Naigles & Terrazas, 1998; Papafragou *et al.*, 2001; Finkbeiner *et al.*, 2002; Zlatev & David, 2004; Pourcel, 2005; Zlatev *et al.*, 2010; and Cardini, 2010 for similar studies).

When the results of the similarity judgment (categorization) studies presented above are evaluated as a whole, it can be seen that most of them are supportive of the universalist view, suggesting similar results across different languages (but see Bohnemeyer *et al.*, 2006 for mixed results). However, it is also evident that the design of the experiment may have a certain effect on the results. For example, the three conditions used in the Gennari *et al.* (2002) study clearly shows that the encoding type has a significant effect on people's categorization performances. On the other hand, we also believe the nature of the stimuli used to assess the non-verbal representation of the subjects (real-life videos or animation clips) is influential in this regard, and may yield different ratios between manner-based and path-based choices of the participants (see the path-biased pattern in Soroli & Hickmann, 2010 in comparison to the balanced manner-path patterns in Papafragou *et al.*, 2002 and in Papafragou & Selimis, 2010).

2.3.2.2.2. *Recognition Memory Studies*

Another popular experimentation technique in the field is using a *Recognition Memory* or *Memorization Task*. The task basically consists of watching a first set of motion event depictions, then watching a second set after a while, and choosing the ones seen in the first set. Most of the time, there is an extra task between the two sessions to make the recognition harder (*e.g.* an unrelated problem solving task). It is hypothesized that the motion event component (*manner* or *path*) to which more attention is paid will be remembered more accurately, and that this component will change from one language group to another based on the canonical motion event verbalization pattern in that language. In the following, a few studies using the recognition memory task are reported.

Gennari, Sloman, Malt and Fitch (2002): Gennari *et al.*'s memorization task was performed within the same experimental design as the categorization task mentioned above. The recognition memory task was performed, in all three conditions (*i.e.* free encoding, shadow and naming-first conditions), after the encoding of the stimuli. In order to make the task harder, the encoding was followed by a distractor task unrelated to the study, and then came the memorization task. In the present task also, the same triads used in the similarity task were used in random order. The results did not show any language-based effect, neither strong nor weak, between the speakers of English (S-language) and Spanish (V-language), but a common pattern for the two language groups. Another finding of the study was that memory performance varied according to the encoding condition used in the task (see pp. 68-69 for the details).

Papafragou, Massey and Gleitman (2002): Papafragou *et al.* administered their memory task along with the categorization task mentioned above. They tested the same group of participants (child and adult speakers of English and Greek) but they used a different set of stimuli. Their materials for the present task consisted of a set of black-and-white drawings adapted from the picture book *Frog, where are you?* (Mayer, 1969). The subjects were tested in two sessions two days apart, and the results showed no language-specific effect between English (S-language) and Greek (V-language) speakers.

Pourcel (2005): Pourcel also used a memory task to evaluate the non-linguistic representation of motion by English (S-language) and French (V-language) speakers. However, different from other studies in the literature, she made use of a *motion scenario* instead of individual motion events as the stimuli. She claimed that the contextual effects are ignored when we use single motion scenes, and thus presented her subjects a 4.5-minute extract from a Charlie Chaplin movie named "City Lights". Her results suggested clear relativistic recall effects. English speakers had significantly higher error rates in path recall, and French speakers in manner recall. This language-specific effect was also persistent in a late recognition task performed 24 hours later (also see Papafragou *et al.*, 2001; and Oh, 2003 for similar studies).

An overall look at the recognition memory tasks presented above presents us with mixed results, two of them supporting the universal approach and one the relativist

approach. This variation may again be due to experimental design or stimulus effects, because all the three studies mentioned use different task designs and each makes use of a different type of stimulus set.

2.3.2.2.3. *Eye-Tracking Studies*

Most of the crosslinguistic studies on the Whorfian hypothesis has focused on off-line methodologies (categorization or recognition memory), thus one can object that the online nature of the real-life event encoding process has been neglected (Papafragou *et al.*, 2008). Event apprehension is a very quick process, and it is an extremely challenging work to be able to monitor it in real-time. The eye-tracking methodology is a rather recent (Griffin & Bock, 2000) and promising one in this respect. Papafragou (2007) also refers to this technology while talking about language-specific effects in cognitive tasks by saying that “the eye-tracking technology seems a particularly apt tool for studying such effects, since it allows direct insight into the process of how attention is distributed onto elements of a scene during both linguistic and non-linguistic tasks” (p. 23).

The following is a brief review of two experimental studies using the eye-tracking paradigm in order to question both the *linguistic relativity hypothesis* (Whorf, 1956), and the *thinking-for-speaking hypothesis* (Slobin, 1996a).

Papafragou, Hulbert and Trueswell (2008): In their study, Papafragou *et al.* compared eye-movements of English (S-language) and Greek (V-language) speakers both in a *verbal* production and a *non-verbal* memorization task. The stimuli were short animated clips depicting instrumental motion events. They were all simple motion scenes where there was an animate agent performing the act with an instrument (*e.g.* skates). Half of the scenes showed bounded events (with a physical end-point, such as a tent) and the other half unbounded events. The eye-movements were recorded with a remote table-top eye-tracker. They defined the instrument region as the *manner region*, and the end-point as the *path region*. The findings suggested that speakers’ eye-movements followed a language-specific pattern only in the description task, and that no such effect was observed in the non-verbal task. They conclude, with these results, that “when inspecting the world freely, people are

alike in how they perceive events, regardless of the language they speak” (p. 155). However, their eye movements may be guided by their language when the task is to verbally describe the scene.

Soroli and Hickmann (2010): Soroli and Hickmann tested native speakers of English (S-language) and French (V-language) in a verbal description task coupled with an eye-tracking paradigm. They used a portable eye-tracker to record their subjects’ eye movements. Their stimuli consisted of a set of real-life video sequences and a second set of animation video clips. They were looking for any language effects that may be observed in the eye movement patterns of the two language groups. The results suggested a difference in eye-tracking performance depending on the stimulus set being exposed to. No significant difference was found in the participants’ attention allocation patterns with the real-life videos; however there was a difference with the animated cartoon set. The data gathered *via* the animation clips revealed that English speakers paid equal attention to the manner and path regions, whereas French speakers started focusing more on the path region towards the middle of the task. This is only a partial language effect, thus it is hard to relate this effect to any of the language and cognition hypotheses already discussed.

As the number of eye-tracking studies investigating the language and cognition interface in motion situations is still rather low, it is hard to make to an overall evaluation of the currently available results. However, tracking people’s eye-movements is a highly promising online tool for such an inquiry, and upcoming similar studies, including the present study, will certainly shed more light on the debate.

2.3.2.3. The Gap in the Literature to be Filled with the Present Study

As already summarized in Chapter 1, the present study is an experimental study which has two main objectives: making a comparative psycholinguistic investigation of motion event expressions in three languages, namely Turkish, English and French, by adopting the Talmyan motion event verbalization dichotomy (Talmy, 1985) as the framework; and inquiring the relationship between conceptual representation and

linguistic expression (Whorf, 1956) by taking motion events as the testing ground and using cognitive experimental tasks.

The first goal will be achieved *via* two psycholinguistic tasks, a language production task and a language comprehension task. This is a pioneering study in this regard, because as seen in the review presented above, almost all of the similar studies in the literature depend solely on language production data but do not use language comprehension tasks. Actually, production and comprehension are two faces of the same coin, and both are strongly needed in a detailed crosslinguistic inquiry on motion event expressions. Moreover, contrary to the psycholinguistic motion event studies reviewed above, the present study does not only look at the issue from an inter-typological perspective by just testing two languages (one from each end of the dichotomy) but also from an intra-typological perspective by also testing two languages from the same typological class. Therefore, the single-faceted approach adopted by most of the previous studies in the literature will be complemented with a double-faceted approach. Another point of interest in the present study is that it does not rely on others' work at any stage; rather it forms its own original research questions, it creates its own solid experimental design from scratch, and it meticulously collects its own data in response to its own real-life stimulus materials.

Our second goal of questioning the interdependence of linguistic representation and conceptual representation will be realized with three non-linguistic (cognitive) tasks. Most of the studies in the relevant literature use offline experimental techniques (categorization task or recognition memory task) to investigate the language and cognition interface. Even though the value of those techniques cannot be denied, they should be complemented with online techniques, as well. Conceptual event representation is an abstract phenomenon which is almost impossible to fully investigate even with recent high-technology methodologies. However, as the eye-tracking methodology gives us the opportunity to track human visual perception in real-time, it is the methodology that is most appropriate for such an investigation. The present study combines a classical offline methodology (categorization) with an online methodology (eye-tracking) in order to overcome the possible drawbacks of relying on the results of a single task. Moreover, an additional technique called *Articulatory Suppression* (Murray, 1967; Baddeley & Hitch, 1974) is also used in the

above mentioned eye-tracking task to ensure the purely cognitive nature of the task, and thus the main criticism against neo-Whorfian experimental studies is eliminated in the present study. Another major contribution of the present study is that it uses two complementary eye-tracking tasks, one observing the eye-movement patterns of the subjects in a purely non-verbal task (mentioned above) and the other in a task where there is communicative intention. Therefore, it is not just the Whorfian *linguistic relativity hypothesis* (Whorf, 1956) that is investigated in this study but also the Slobin *thinking-for-speaking hypothesis* (Slobin, 1996a).

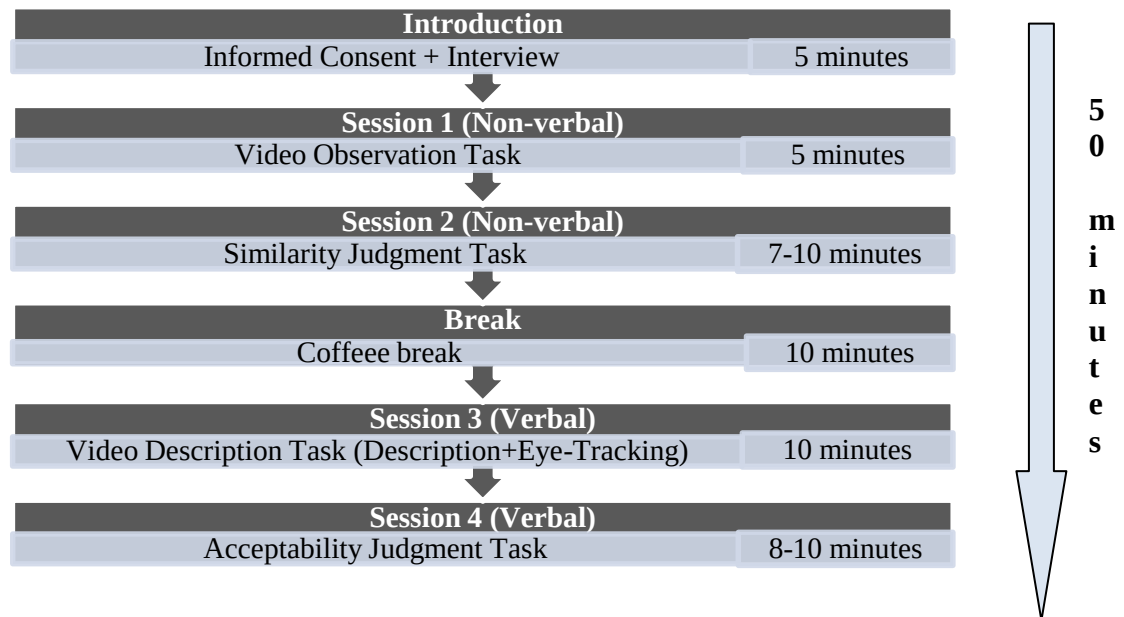
In summary, the main contribution of the present study both to the ‘psycholinguistic investigation of motion events’ literature and to the ‘language and cognition debate’ literature is in its multi-faceted perspective coupled with solid experimental methodologies.

CHAPTER III

GENERAL METHODOLOGY

3.1. FRAMEWORK

The present dissertation aims at making a psycholinguistic inquiry on *motion event* expressions in three languages (Turkish, English and French), and at investigating the *language and thought* debate by using motion events as the testing ground and cognitive experimentation techniques as the methodology. The study is comprised of four experimental sessions, taking place consecutively on the same day. Graph 1 shows a brief outline of the sessions. The whole study is computerized, and takes place at the Human-Computer Interaction Research and Application Laboratory at Middle East Technical University, Ankara, Turkey, and at the SFL (Structures Formelles du Langage) Laboratory, Paris, France. In the tasks, 73 short video clips (60 main items and 13 trial items), each depicting a different voluntary motion event and exclusively shot for the purpose of the current study, are used as stimuli.



Graph 1: Outline of the experimental sessions

All the motion events are performed by the same male actor and with the same clothing in order to preserve the homogeneity of items. A total of 79 subjects from three linguistic backgrounds, i.e. native speakers of English, French and Turkish, participate in the study.

3.2. PROCEDURE

The whole procedure, including the break, approximately takes 50 minutes per subject. Each subject is tested alone in an isolated and sound-proof testing room. Subjects need to use nothing other than the computer mouse during the sessions. When each subject is taken to the testing room, s/he is first informed about the experimental procedure and presented the Informed Consent Form (see Appendix A). This form gives brief information about the study and it is signed by the participant before starting the experiment to testify that s/he takes part in the study voluntarily. The subject is not informed about the actual goal of the study to preserve the naturalness of her/his responses. Before moving on to the experiment itself, each subject has a five-minute long, casual interview with the experimenter. In this interview, the subject is questioned about her/his educational and language background in order to determine a common subject profile (see Appendix B). Then, as the first task necessitates the use of the eye-tracker, the experimenter explains the functioning of the eye-tracking system to the subject, and together they perform the eye calibration procedure, which is a prerequisite step for any eye-tracking experiment. When the calibration is over, the experimenter explains the procedure of the first experiment to the participant and she leaves her/him alone in the testing room. The experimenter stays in the control room, which is located right next to the testing room and separated from the testing room by a one-way mirror, during the sessions. She comes to the testing room at the beginning of each session to clarify the instructions that the subjects already have on the computer screen.

The first experimental session includes a *Video Observation Task*, which has a non-verbal nature and which has the aim of testing the visual attention allocation and conceptual linearization patterns of the subjects. In other words, in this task, it is questioned whether subjects speaking different languages attend to the same components of a motion event (*manner of motion*, *path of motion* or *other*) while

watching the same set of motion event depictions, and whether they attend to them in the same order. The task consists of watching 25 real-life motion event videos, while repeating a series of numbers during the whole session. This is a technique called *Articulatory Suppression*, which is used to prevent sub-vocal verbalizations of the subjects (Murray, 1967; Baddeley & Hitch, 1974). The crucial part of this session is that, while the subjects are watching the scenes presented on the screen, their eye-movement patterns are tracked and recorded for later analysis with a screen-internal eye-tracking system.

The second session includes a *Similarity Judgment Task*, which takes place right after the *Video Observation Task*. This task has the aim of testing the visual perception patterns of the subjects by presenting them a categorization task. The videos in this task are presented in triads. The subjects start by watching the first video, which is the *main* or *target video* of the triad. Then they watch the other two, which are called *alternative* or *candidate videos*, and they are supposed to say which one of the alternative items is more similar to the main item. One of the alternative items depicts a motion event that shares the same manner with the main item, and the other depicts a motion event that shares the same path with it. Thus, it will be observed which component (*manner of motion* or *path of motion*) is taken as the dominant criterion of similarity by the subjects in each language group. The subjects give their answers by using the computer mouse, and their answers are recorded automatically by the system. The goal here is to see which one of the two components of a motion event, *manner* or *path*, is attended more by the speakers of typologically different languages. After this second task, there is a short break, during which the participants are served tea or coffee by the experimenter.

The third task is the *Video Description Task*, which is a language production task aiming at comparing the motion event expressions of speakers of English, French and Turkish. The descriptions are made just after watching each single video, and the subjects have a 10-second-long description period for each video. Their descriptions are video-taped for later analysis. This task also has an eye-tracking part, the results of which will be evaluated separately from the description results. Different from the one in the *Video Observation Task*, this eye-tracking experiment observes the eye movement patterns of the subjects while they are watching motion event videos with

the aim of describing them after watching. In other words, while watching the motion event depictions presented on the screen, the subjects are aware of the fact that they will soon put what they see into words, thus they investigate the scene accordingly. The results of this *pre-description eye-tracking experiment* are compared with the results of the *non-verbal eye-tracking experiment* in the Video Observation Task, in order to shed light on Slobin (1996a)'s *thinking-for-speaking hypothesis*.

The last session includes an *Acceptability Judgment Task*, which is a language comprehension task. In this task, each video is presented with a single sentence describing what the man on the screen is doing, and the subjects are supposed to evaluate the acceptability of those sentences on a 5-item rating scale. There are three types of sentences presented along with the videos: *manner verb + path satellite* sentences (the canonical verbalization pattern in S-languages), *path verb + manner satellite* sentences (the canonical verbalization pattern in V-languages), and *other* sentences used as distractors. This task has the aim of complementing the language production task in the third session by looking at the issue from the comprehension side. The subjects give their answers by using the computer mouse, and their answers are recorded automatically by the system. At the end of this last session, each subject is provided with a Post-Participation Information Sheet (see Appendix C) that gives brief information about the aim of the study and the contact information of the experimenters for further questions, and a little souvenir.

The reasons for that particular ordering of the tasks can be explained as follows: the non-verbal tasks (Video Observation and Similarity Judgment) should be performed before the verbal tasks (Video Description and Acceptability Judgment) so as to avoid any linguistic verbalization effects from the verbal tasks to the non-verbal ones. The Video Observation Task is presented as the very first task, because it is supposed to be the purest-cognitive task thanks to the use of the articulatory suppression technique. The rationale behind putting the Video Description Task before the Acceptability Judgment Task is to avoid any biasing effect from the already presented sentences in the comprehension task to the free descriptions in the production task. The following chapter will present the details of each experiment in a systematic manner.

CHAPTER IV

EXPERIMENTS

This chapter will elaborate on the five experiments used in the current study, and each experiment will be explained in a separate section. Each section will give information about the aim of the experiment, the research questions it deals with and the hypotheses regarding those questions, the design and the procedure of the experiment, the encoding of the data, the results and a brief discussion of the results. The ordering in which the experiments are presented is not the same as the testing order that has been mentioned in the previous chapter. The reason for using a new ordering is to make it more practical for the readers to keep track of the liaisons among the research questions. The verbal experiments are presented here before the non-verbal ones to preserve the consistency of the results. The new ordering of the experiments is as follows:

Experiment 1: Video Description Task

Experiment 2: Acceptability Judgment Task

Experiment 3: Similarity Judgment Task

Experiment 4: Non-Verbal Eye-Tracking Task (in the Video Observation Task)

Experiment 5: Pre-Production Eye-Tracking Task (in the Video Description Task)

4.1. EXPERIMENT 1: VIDEO DESCRIPTION TASK

4.1.1. Aim and Scope

The present task is a language production task during which the subjects describe the motion events depicted in the real-life video sequences presented one-after-another on a computer screen. It aims at investigating the motion event descriptions of the speakers of the three languages, namely English, French and Turkish, crosslinguistically. It tests the S-language vs. V-language dichotomy experimentally (Talmy, 1985) by analyzing the elicited productions in two V-languages (Turkish

and French), and one S-language (English); therefore, this task presents us both an inter-typological and intra-typological perspective.

4.1.2. Research Questions and Hypotheses

The research questions to be answered with the help of the data obtained in this task, and the hypotheses put forward regarding those questions are as follows:

Question 1: Do native speakers of Turkish and French behave similarly and as predicted by the theoretical motion event typology proposed by Talmy (1985) in their verbal descriptions of motion events? In other words, will the V-language pattern be experimentally verified in our language production data with Turkish and French sentences indicating the path of motion in the main verb, and the manner of motion in a separate component (a satellite)?

Hypothesis 1: It is predicted that the speakers of the two languages will behave in a similar way and as theoretically predicted by the typology, and that they will predominantly use sentences with the *PathVerb* + *MannerSatellite* pattern while describing the motion events that they watch.

Question 2: Do the language production data obtained from native speakers of English conform to the S-language pattern of Talmy (predominant use of manner verbs accompanied with path satellites) in contrast to the V-language pattern predicted above for Turkish and French?

Hypothesis 2: It is hypothesized that the sentences produced by native speakers of English will verify the typology by reflecting the canonical S-language verbalization pattern, *i.e.* *MannerVerb* + *PathSatellite*, much more frequently than other patterns.

Question 3: Do native speakers of S-languages express the manner information more often than native speakers of V-languages, as claimed in the literature (see section 2.2.3.1.1 for a review of the idea)? In other words, do speakers of Turkish and French give only-path information in their descriptions, while the speakers of English express both the manner and the path components in their utterances?

Hypothesis 3: It is hypothesized that there will be no such semantic packaging differences between the speakers of S-language and those of V-languages, as both of the motion event components (*manner* and *path*) are salient enough in the stimuli used in the experiment. Therefore, both groups will use both components in their productions.

4.1.3. Design

4.1.3.1. Subjects

A total of 79 subjects (32 native speakers of Turkish, 23 native speakers of English, and 24 native speakers of French) were tested in this task, however data for 17 of them had to be eliminated for technical reasons. Therefore, the data for 20 native speakers of Turkish (9 females and 11 males), 20 native speakers of English (10 females and 10 males) and 22 native speakers of French (14 females and 8 males) were included in the analyses. Each language group is comprised of participants from comparable educational backgrounds, aged between 18 and 35, and they are all chosen among the monolingual speakers of the languages to prevent any possible crosslinguistic influences.

Native speakers of Turkish were either university students or new graduates. Their age range was 18-27 years ($m=19.7$). Native speakers of English were all college students from the USA, who were in Turkey for a short time period for academic purposes. Their age range was 20-30 years ($m=21.9$). Two of the native speakers of French were young teachers who work at an international school in Ankara. In order to increase the number of French participants, it was decided to contact a lab which has comparable equipment in France, the “Structures Formelles du Langage (SFL-*Eng.* Formal Structures of Language)” Lab in Paris¹³, which is part of the University of Paris 8 and the National Center for Scientific Research. Therefore, the other 20 French subjects were university students or new graduates, and they were tested at

¹³ The lab is directed by Dr. Sophie Wauquier and Dr. Maya Hickmann of University of Paris 8 & CNRS, and our host there was the PhD candidate and research assistant Efstathia Soroli.

the SFL Lab. The age range of the French subjects was 23-35 years ($m=28.9$). All the subjects from all language groups were chosen on voluntary basis.

4.1.3.2. Stimuli

Up until the last decade, most of the studies in the field had been using static images as stimuli to analyze motion event verbalizations across languages. The most common material used was the wordless picture book *Frog, where are you?* (Mayer, 1969), which tells the story of a young boy who tries to find his frog. Even though it is common-sensical to use dynamic scenes to test a dynamic element like *motion event*, static materials were and are still in use due to their testing practicality (see Özçalışkan & Slobin, 1999, 2000; Özyürek & Özçalışkan, 2000; Papafragou *et al.* 2001, 2007; Zlatev & Yangklang, 2004; Ibarretxe-Antuñano, 2004a, 2004b, 2008; Ameka & Essegbey, 2001 among many others). On the other hand, there has recently been a number of studies which changed this paradigm by using animated clips or real-life video sequences to be able to better instantiate motion (see von Stutterheim & Nüse, 2003; Allen *et al.*, 2007; Papafragou & Selimis, 2009; Skordos & Papafragou, 2010; Bunger *et al.*, 2010 for the use of animation clips, and see Gennari *et al.*, 2002; Pourcel & Kopecka, 2006; Soroli & Hickmann, 2010; Soroli, 2011; Flecken, 2011 for the use of real-life videos).

In the present study, short real-life video sequences that were exclusively shot for the purpose of these experiments were used. The reason for using videos, instead of pictures or animations was twofold. First of all, as mentioned above, motion is a dynamic event in nature and cannot be illustrated with static images. Secondly, it is thought that having a real person performing real actions in the real world will make the subjects perceive the depicted motion events more clearly. As pointed out by Pourcel (2005), “human motion is the main type of motion conceptualized and expressed in language by speakers” (p. 6). Having real-life materials would also have an effect on the results, especially on the eye-tracking results, by presenting the participants with more salient manner of motion and path of motion components. Slobin (2004), a researcher who has made extensive use of the *Frog Story* materials, also underlines the need for using more dynamic materials: “The frog story research

makes it clear that we need audio and video data—along with grammars and dictionaries, texts and corpora—in order to carry the work forward” (p. 29).

A total of 50 videos were used in this task. There were 10 actions (crawl, dance, hop, limp, march, run, stagger, tiptoe, whirl and zigzag), each depicted with 5 different directions (into, out of, up, down and across). The length of the videos varied from 3 seconds to 22 seconds, depending on the nature of the motion event depicted. All the motion events were performed by the same male actor, with the same clothing and with three fixed backgrounds: 1- an apartment building door for *into* and *out of* scenes (see Picture 1), 2- an indoor staircase for *up* and *down* scenes (see Picture 2), and 3- a narrow street for *across* scenes (see Picture 3). The raw videos were edited on the software Ulead Video Studio 11 to be able to obtain a format suitable for our testing procedure.¹⁴



Picture 1: Background used for *into* & *out of* scenes

¹⁴ See www.ulead.com for more information on the software.



Picture 2: Background used for *up* & *down* scenes



Picture 3: Background used for *across* scenes

In order to avoid an item effect, two versions of the same task were prepared. Version 1 included the actions *run*, *hop*, *crawl*, *whirl* and *limp*, and Version 2 included *march*, *stagger*, *dance*, *tiptoe* and *zigzag*, where each was again depicted with five different directions. Therefore, each version was comprised of 25 videos. Half of the subjects in each language group (*i.e.* 10 speakers of Turkish, 11 speakers of English, and 11 speakers of French) saw the first version of the task, and the remaining half of the subjects saw the second version.

4.1.3.3. Apparatus

Most of our subjects were tested at the Human-Computer Interaction (HCI) Research and Application Laboratory at Middle East Technical University, Ankara, Turkey, with the exception of 20 people tested at the Formal Structures of Language Lab, Paris, France. In this section, first the equipment used in the main lab in Ankara will be detailed, and then the one at the lab in Paris.

In the main lab, the videos were displayed on a 17" computer screen with a screen resolution of 1024x768 (see Picture 4). The whole procedure took place in a sound-proof lab room equipped with two video cameras and a ceiling microphone as well as the main computer. Apart from the testing room, there was also a control room located right next to the testing room. It was separated from the testing room by a one-way mirror, and there the experimenter could observe the testing sessions without distracting the subjects (see Picture 5).



Picture 4: The testing room



Picture 5: The control room

In Paris, the experiment took place in a silent lab room, and a Dell Vostro 1320 laptop computer (15" screen size and 1280x800 screen resolution¹⁵) was used to test the subjects. As there was no video camera or voice recorder readily available in the

¹⁵ The size and resolution differences between the two computer screens were adjusted before the eye-tracking analyses.

lab, the sound recording facility of the software Tobii Studio (provided with the eye-tracking machine) was used to record the verbal descriptions of the participants in Paris.

4.1.3.4. Procedure

The aim of the present task was, as already indicated, to compare motion event descriptions of English, French and Turkish speakers based on the V-language vs. S-language dichotomy. In the task, the subjects were asked to watch a series of short motion event videos presented on a computer screen one after another, and at the end of each video they were asked to describe what the man in the video was doing, preferably in a single sentence. Each subject watched 25 videos. The experiment ran as follows: The video clips were displayed automatically, thus the participants did not need to press anything to move on to the next video. At the end of each video, the written question “What was the man doing?” was seen on the screen for two seconds, and it was followed by a brief beep sound indicating the beginning of the answering period. The participants had 10 seconds to give their answers, and at the end of this time period, there was a second beep sound indicating the end of the answering time. The following video started right after the closing beep. There were three irrelevant videos at the beginning of the task used for a warming-up session (*e.g.* a man listening to an iPod, a man picking up his backpack or a man putting a letter into a mailbox), whose descriptions were not included in the analyses. Each participant described the scene s/he watched in her/his respective native language. The answers of the participants were video-taped for later analyses¹⁶.

This task took about ten minutes for each participant, and during this time period the experimenter was in the control room following the subject’s performance. The experimenter re-entered the room when the task was over.

¹⁶ The answers of the 20 French speakers tested in Paris were audio-recorded.

4.1.3.5. Encoding of the Raw Data

First of all, the recorded descriptions were transcribed for each subject. Almost all of the participants followed the instructions and used a single sentence to describe the whole scene. Then, the sentences were encoded in two different ways:

A) *The Detailed Encoding*: In this detailed encoding process, both the main verbs and the satellites were taken into consideration. The sentences were coded as follows:

1: Manner Verb + Path Satellite

(e.g. The man staggered out of the apartment building.)

2: Path Verb + Manner Satellite

(e.g. L'homme est sorti du batiment en titubant - 'The man exited from the building by staggering'.)

3: Manner-Verb-Only

4: Path-Verb-Only

5: Other

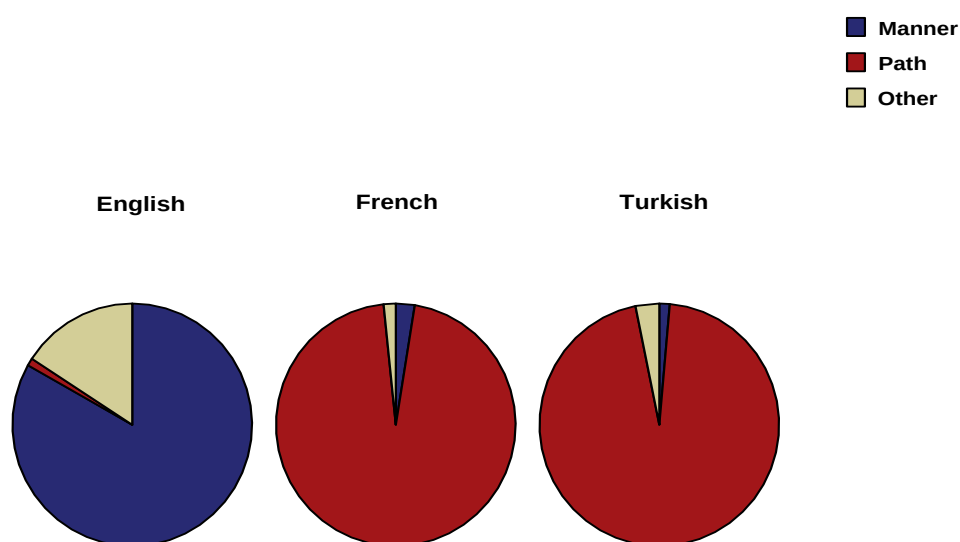
B) *The Concise Encoding*: The encoding type mentioned above was a comprehensive one; however it had too many categories to be included in a neat statistical analysis. Therefore, it was simplified by putting them all into three categories. The sentences encoded above as 1 and 3 were coded as *manner* (1), the ones encoded as 2 and 4 as *path* (2), and the ones encoded as 5 were coded as *other* (3). This encoding was much more practical for our crosslinguistic analysis and it was the one used for the main analysis. However, the detailed encoding items will also be used to examine the semantic densities of speaker utterances in section 4.1.4.2.

4.1.4. Results

4.1.4.1. Main Analysis

In this section, we will look for answers to our Research Question 1 and Research Question 2 (see section 4.1.2) regarding the applicability of the Talmyan motion event typology on our experimental data. The distribution of the verbal descriptions of the three language groups', obtained *via* the *concise encoding*, is as shown in Graph 2. It is quite clear from the graph that all the groups behaved as predicted by the S-language vs. V-language dichotomy (Talmy, 1985).

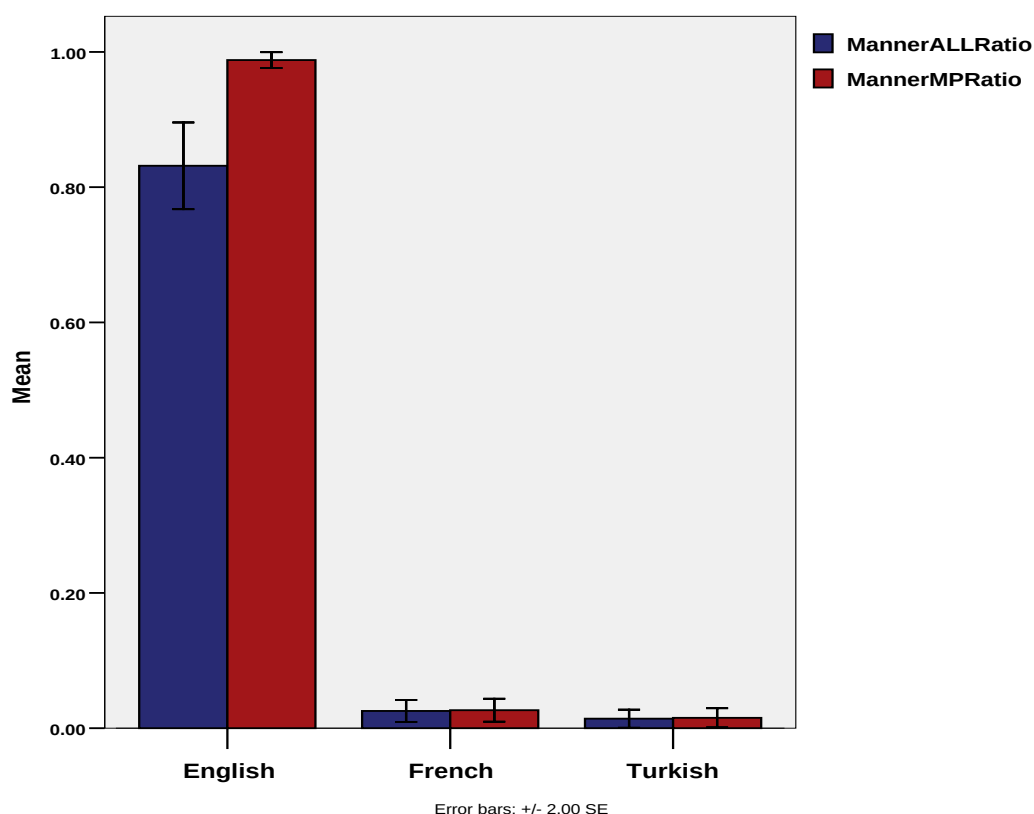
On the one hand, there are English speakers who have a considerably high number of manner answers; 83.2% manner sentences, 1% path sentences, and 15.8% other sentences. On the other hand, there are Turkish speakers with 95.4% path sentences, 1.4% manner sentences, and 3.2% other sentences; and there are French speakers with their 95.8% path sentences, 2.56% manner sentences, 1.64% other sentences.



Graph 2: Distribution of the verbal descriptions for all groups

For the statistical analysis of the data, first of all the encodings were transferred to an SPSS sheet, and the numbers of *manner*, *path* and *other* sentences were counted.

Then, two ratios were formed from this counting: the *Manner-to-ALL Ratio*, the ratio of a participant's manner answers to all of her/his answers; and the *Manner-to-MP Ratio*, which is the ratio of a subject's manner answers to the total of her/his manner and path answers. Choosing the manner component while calculating the ratio was a random choice, the results would not change if path ratio was taken instead of manner ratio. Two Uni-variate ANOVAs were run on the data; one taking *Manner-to-MP Ratio* as the dependent variable and *language group* (English, French or Turkish) as the independent variable; and a second which takes *Manner-to-ALL Ratio* as the dependent variable and again *language group* as the independent. The analyses revealed a significant effect of *language group* on both ratios: $F(2, 59)=5700.421$, $p=.00$, $\eta_p^2=.995$ ¹⁷ for *Manner-to-MP Ratio*; and $F(2,59)=602.306$, $p=.00$, $\eta_p^2=.953$ for *Manner-to-ALL Ratio* (see Graph 3).



Graph 3: Results of the Video Description Task

¹⁷ **Partial eta squared (η_p^2):** A version of eta squared that is the proportion of variance that a variable explains when excluding other variables in the analysis. Eta squared is the proportion of total variance explained by a variable, whereas partial eta squared is the proportion of variance that a variable explains that is not explained by other variables (Field, 2009, p. 791).

These results show that the manner ratios of native speakers of the three languages were significantly different from each other; English speakers having ratios very close to 1.00 ($m=.83$ for *M-to-ALL Ratio*, and $m=.99$ for *M-to-MP Ratio*), whereas Turkish and French speakers having ratios highly close to .00 ($m=.01$, $m=.02$ for *M-to-ALL Ratio*; $m=.01$, $m=.03$ for *M-to-MP Ratio*).

We also used a Helmert contrast to analyze the three language groups in pairs. We first had an inter-typological contrast (English vs. Turkish and French), and it revealed a significant difference ($p=.00$) both for *Manner-to-MP Ratio* and *Manner-to-ALL Ratio* which means that the difference between the typologically different languages are statistically highly significant. Then, we had an intra-typological contrast, which did not reveal any significant results between Turkish and French language groups for any of the ratios. Both the main analysis and the contrast results are perfectly in line with the Talmyan typology (Talmy, 1985), which classifies English as an S-language, and both Turkish and French as V-languages.

4.1.4.2. Semantic Density Analysis

In the main analysis section, we have looked for answers to our first two research questions, inquiring the compatibility of our subjects' verbal descriptions with the theoretical motion event typology of Talmy. In this section, we will try to find an answer to our third research question regarding the S-language and V-language speakers' divergent preferences for the semantic packaging of manner and path information. As already mentioned in section 2.2.3.1.1, a great number of researchers claim that V-language speakers have a significant tendency to omit manner of motion in their motion event verbalizations (see Talmy, 2000; Özçalışkan & Slobin, 2000; Papafragou *et al.*, 2002; Gennari *et al.*, 2002; Özçalışkan & Slobin, 2003; Pourcel, 2005; Slobin, 2004, 2006; Gullberg *et al.*, 2008; Soroli & Hickmann, 2010 among others). Soroli and Hickmann (2010) call this phenomenon *semantic density*. They claim that S-language speakers and V-language speakers diverge in the semantic density (SD) of the information that they provide in their motion event descriptions. S-language speakers provide both the manner of motion and path of motion in their utterances (SD2), whereas V-language speakers give path of motion information but predominantly omit the manner of motion information (SD1). On the

other hand, some researchers like Allen *et al.* (2007) oppose this claim by saying that this is a fake tendency observed as a result of the methodologically-deficient experimental studies, and he adds that when both the path and the manner components are salient enough in the stimuli which are used to elicit data, then V-language speakers make use of the manner information as frequently as S-language speakers do. This is also what we have hypothesized for our own data, and in this section we will evaluate this hypothesis.

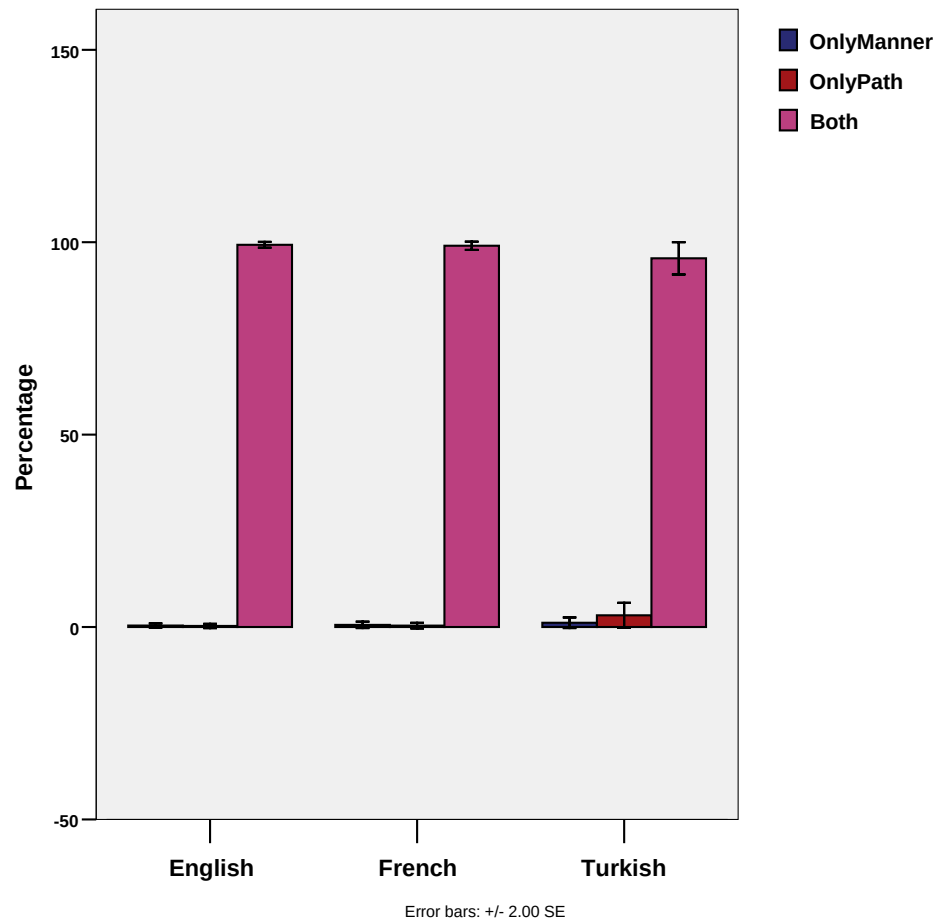
In order to give an answer to that specific question, we will use the *detailed encoding* explained in section 4.1.3.5. We will form three categories out of the original seven categories. The category 3 will be labeled as *only manner* sentences, where there is only the mention of the manner of motion and no mention of path of motion. Likewise, the category 4 will be labeled as *only path* sentences, as only the path of motion is given in those sentences, not the manner of motion. Finally, the categories 1 and 2 will be labeled as *both*, because these sentences give both types of information (*i.e.* manner of motion and path of motion) in a conflated manner¹⁸.

Graph 4 on the next page illustrates the distribution of *manner-only*, *path-only* and *both* responses across languages. As can be clearly noticed, V-language speakers have a tendency to include both types of information in their sentences as much as S-language speakers do. 92.8% of Turkish sentences, 97.4% of French sentences, and 83.4% of English sentences include conflated manner and path information. This result clearly contradicts the idea that V-language speakers have a great tendency to omit manner information.

In order to analyze the data statistically, a Repeated-Measures ANOVA, taking *answer type* (percentage of manner-only answers, percentage of path-only answers or percentage of both answers) as the within-subject variable and *language group* as the between-subject variable, was run. The results of the analysis did reveal neither an *answer type*language group* interaction nor a main effect of *language group*, which shows that the answer types do not differ significantly from one language to the other. There was only the significant main effect of *answer type* ($F(2,59)=7100.564$, $p=.00$,

¹⁸ Other sentences (category 5) are eliminated from the semantic density analysis.

$\eta_p^2=.992$), which supports the clear pattern observed in Graph 4: the percentage of conflated manner and path answers are much higher than the percentages of manner-only and path-only answers for all language groups.



Graph 4: Results of the semantic density analysis across languages

4.1.5. Discussion

The results presented above are perfectly in line with the three hypotheses put forward in section 4.1.2. Speakers of English language, which is qualified as an S-language, predominantly use manner sentences (83.2%) while describing motion events. On the other hand, speakers of Turkish and French, which are both categorized as V-languages, use path sentences in almost all of their motion event descriptions (95.4% and 95.8% respectively). This typological pattern which is

clearly observed in our data verifies the motion event expression dichotomy proposed by Talmy (1985), as well. The data also verifies our third hypothesis, which suggested that the semantic densities of the sentences produced by the speakers of V-language and S-language speakers would not diverge, as both the manner of motion component and the path of motion component are salient enough in our stimuli. It is quite clear from the data that V-language speakers conflate manner information and path information in a single sentence, as frequently as S-language speakers do.

Another point which may be discussed regarding the results of the production data is the higher number of *other sentences* in English. This fact can be observed from the pie chart in Graph 2, and from the discrepancy between the *Manner-to-MP Ratio* and the *Manner-to-ALL Ratio* in Graph 3. In order to find an explanation to this pattern, we went deep into the descriptions made by native speakers of English, which provided us with an interesting picture. Most of the sentences qualified as *other* in English descriptions were composed of a hybrid pattern like *plain verb*¹⁹ + *path satellite* + *manner satellite* as in (20) or like *plain verb* + *manner satellite* + *path satellite* as in (21). As we could not categorize those sentences as being either *manner sentences* or *path sentences*, they were qualified as *other sentences*.

(20) He was walking down the stairs moving right or left every few steps (Subject10).

(21) He was walking zigzagly up the stairs (Subject8).

It is evident from the whole analysis that typologically different languages express motion events in different ways in language production. However, this difference most probably boils down to the differences found in the lexicons of these languages, *i.e.* what kind of verbs (manner or path verbs) speakers find in their lexicon. It is not the personal choice of the speakers, it is just what the lexicon of their respective languages provide them with²⁰. Accordingly, in our first analysis, we found a

¹⁹ Here the term *plain verb* is used to indicate verbs that have the basic motion content without a salient manner or path meaning, like *go*, *walk* or *come*. These verbs are called *neutral verbs* by Özçalışkan and Slobin (2000).

²⁰ We would like to thank Dr. Annette Hohenberger for drawing our attention to this interesting point.

typological effect between manner and path languages (English vs. Turkish and French). However, in our second analysis we saw that the speakers of all three languages are equally apt at providing both components in their verbalizations, distributed over verbs and satellites. In summary, they gave the same information (same semantic density), only they packaged this information differently.

The following task will question the same typology but this time from the reverse aspect, the language comprehension aspect.

4.2. EXPERIMENT 2: ACCEPTABILITY JUDGMENT TASK

4.2.1. Aim and Scope

The present task is an acceptability judgment task, which has the aim of complementing the *language production* aspect presented above by approaching the issue from a *language comprehension* perspective. Like the production task, it aims at comparing the verbal performances of the subjects, however different from the production task, in this task the motion event descriptions are already given and it is the participants' task to evaluate them in terms of their acceptability. One third of the sentences presented in the task use the canonical motion event expression pattern in the language of the task (*i.e.* V-language pattern for Turkish and French subjects, and S-language pattern for English subjects), one third use the contrastive pattern (*i.e.* S-language pattern for the Turkish and French group, and V-language pattern for the English group), and one third use a pattern other than the main two (a distractor pattern). The task aims at investigating whether the speakers of typologically different languages diverge and typologically similar languages converge in their language comprehension patterns as they do in the language production task.

4.2.2. Research Questions and Hypotheses

The specific research questions to be answered with this task and the hypotheses put forward are as follows:

Question 1: Do the experimental data obtained from native speakers of Turkish and French conform to the V-language pattern in the language comprehension task, as well? More specifically, do they really prefer (by giving higher acceptability scores) the sentences with *path verb + manner satellite* conflation pattern to the ones with *manner verb + path satellite* pattern?

Hypothesis 1: It is predicted that native speakers of both languages will give significantly higher scores to the sentences conforming to the V-language pattern than the other ones, thus they will favor the *path verb + manner satellite* sentences.

Question 2: Do native speakers of English prefer the *manner verb + path satellite* pattern to others in a language comprehension task, as well?

Hypothesis 2: It is hypothesized that the judgments of the native speakers of English will perfectly conform to the S-language pattern, and that they will give much higher scores to the *manner verb + path satellite* sentences than the other sentences.

4.2.3. Design

4.2.3.1. Subjects

In the present task, the same 79 subjects were tested and the data for all of them were used in the analyses: 23 native speakers of English (13 females and 10 males), 24 native speakers of French (15 females and 9 males), and 32 native speakers of Turkish (16 females and 16 males).

4.2.3.2. Stimuli

The visual materials used in the task are chosen from the same pool of real-life motion event videos exclusively shot for the purposes of the current study. In this task, the videos are not presented in the plain format but each with a written sentence describing the motion event in the video. The sentences are written at the bottom of the screen, like subtitles (see Picture 6). 10 different video clips are presented in this task, each with three different types of sentences: an S-language pattern sentence (*e.g.*

He is staggering into the building), a V-language pattern sentence (*e.g.* He is entering the building by staggering), and a distractor pattern sentence (*e.g.* He is staggering and entering the building)²¹. Therefore, each subject sees and evaluates a total of 30 videos and sentences.



Picture 6: A screenshot from an example video used in the comprehension task

4.2.3.3. Apparatus

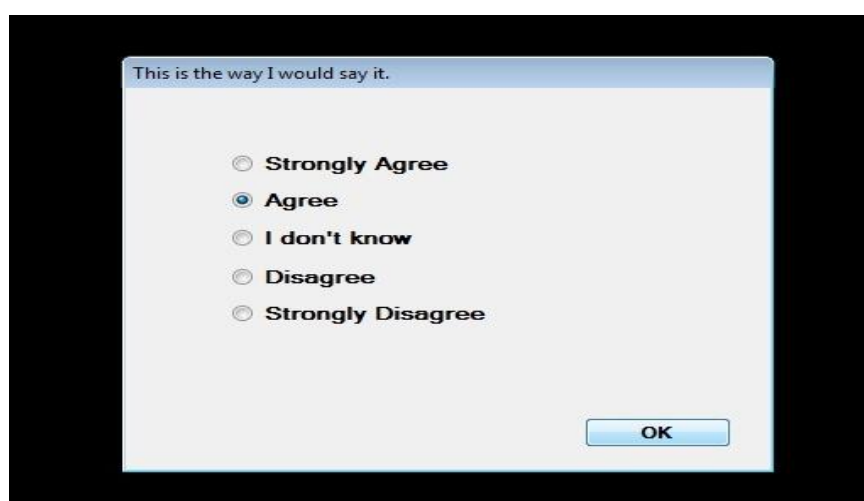
In the present task, the stimuli are again displayed on the same computer screen; however this time there was need for an interface where the subjects may enter their answers, and also for a system that records these answers. Therefore, a specific kind of software just to be used in this task was prepared by an expert²². The video clips with sentences and the answering window where the subjects enter their acceptability ratings were easily displayed by this software. The same system also kept the answers of each subject in a file.

²¹ This is a pattern legitimately used in Modern Greek (Papafraou & Selimis, 2010, p. 227).

²² This software was developed by İbrahim Demir.

4.2.3.4. Procedure

This task was the last one administered, and it approximately took 8-10 minutes. The participants watched 30 videos, each with a written sentence describing the motion event depicted in the video in their respective languages. Their task was to evaluate the acceptability of each sentence on a 5-item scale. After each video, there appeared a window at the top of which it says “This is the way I would say it”, and there were 5 choices underneath: *Strongly agree*, *agree*, *I don’t know*, *disagree*, *strongly disagree* (see Picture 7). The subject was supposed to click on her/his choice with the computer mouse, and the answers were automatically recorded by the system. After each mouse click, there came the next video to evaluate.



Picture 7: The answering window used in the comprehension task

There were again three warm-up items at the beginning of the task. The experimenter explained the procedure before the task, and left the room not to disturb the subject during the task.

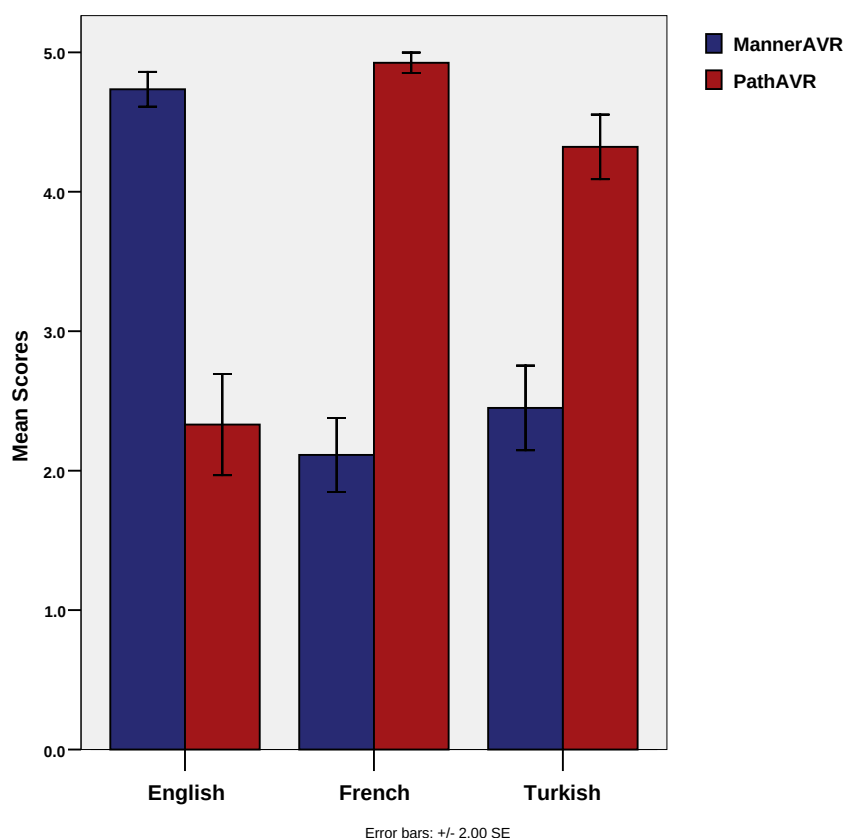
4.2.3.5. Encoding of the Raw Data

The encoding of the data obtained from the present task was clear-cut. The sentences evaluated with the expression *strongly agree*, thus rated as perfectly acceptable, were coded as 5. The *agree* answers were coded as 4, the *I don’t know* answers were

coded as 3, the *disagree* choices were coded as 2, and the *strongly disagree* answers were coded as 1. As each sentence belonged to one of the three sentence types, S-language pattern (*e.g.* He is running out of the building), V-language pattern (*e.g.* He is exiting the building by running) or other pattern (*e.g.* He is running and exiting the building), the answers were entered into the data sheet accordingly. Then, the average values for each sentence group, which showed the average score that each subject gave to each sentence type, were calculated. Therefore, three main values were obtained for each subject: *MannerAverage*, *PathAverage* and *OtherAverage*.

4.2.4. Results

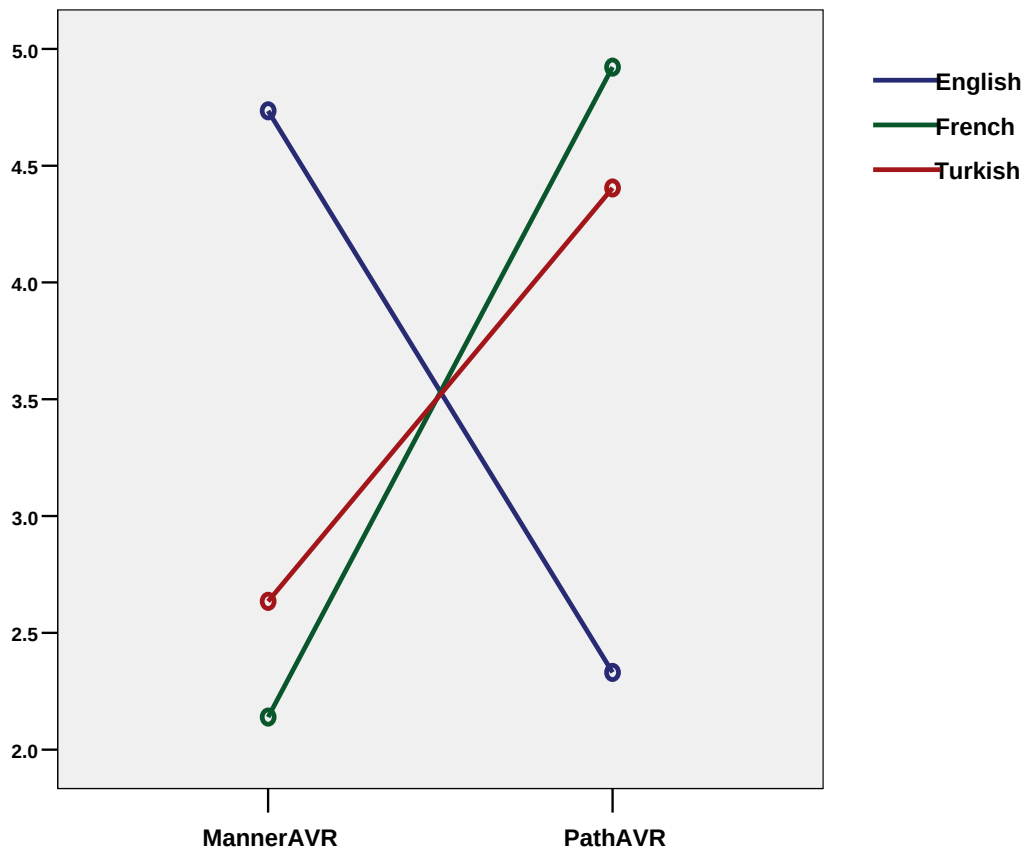
Of the three values calculated, we decided to use the main two in the analyses and to leave out the *OtherAverage*, because *other sentences* were only used as distractors and their ratings were not important for our purposes. The overall results demonstrated a clear typological pattern in the comprehension task (see Graph 5).



Graph 5: The comprehension results for all groups

On the one hand, there were English speaking subjects who preferred S-language pattern sentences (manner sentences) to V-language pattern sentences (path sentences): $m=4.74$, $sd=.2994$. On the other hand, there were Turkish and French speaking subjects ranking V-language pattern sentences (path sentences) as much more appropriate than S-language pattern sentences (manner sentences): $m=4.32$, $sd=.6554$ for Turkish; and $m=4.93$, $sd=.1800$ for French.

For the statistical analysis of the data, a Two-Factorial ANOVA was run with *sentence type* as a two-level within-subject variable (manner average and path average) and *language group* as a three-level variable (English, French and Turkish). The results revealed a significant main effect of *sentence type* ($F(1,22)=39.660$, $p=.00$, $\eta_p^2=.643$), and a significant interaction between *sentence type* and *language group* ($F(2,44)=268.286$, $p=.00$, $\eta_p^2=.924$) (see Plot 1).



Plot 1: The results of the two-factorial ANOVA for the comprehension task

The main effect shows that the overall manner average differs significantly from overall path average, most probably because the language set consists of two languages (Turkish and French) that give higher scores to path sentences but only one language (English) that gives higher scores to manner sentences. On the other hand, the significant interaction points to different manner and path ratings in different languages; English speakers rating manner higher, and Turkish and French speakers rating path higher, as can be seen in the plot. This interaction is the most important effect in this analysis. It means that manner and path ratings are dependent on the language type.

In order to compare the language groups in pairs, a Helmert contrast was also run. There were two comparisons, an inter-typological one (English vs. Turkish and French) and an intra-typological one (French vs. Turkish). The inter-typological contrast revealed a perfectly significant interaction between *sentence type* and *language group* ($F(1,22)=607.756$, $p=.00$, $\eta_p^2=.965$), which means that the ratings given by native speakers of English (S-language) are significantly different from the ratings given by native speakers of French and Turkish (both V-languages). This result is in line both with Talmyan dichotomy and with our hypotheses. On the other hand, the intra-typological contrast that we have run on Turkish and French data revealed a surprising pattern. The *sentence type*language group* interaction was also highly significant for these two languages ($F(1,22)=15.877$, $p=.001$, $\eta_p^2=.419$), which means that native speakers of French and Turkish rated the motion event sentences presented to them significantly different from each other, French speakers giving higher scores to path sentences and lower scores to manner sentences than Turkish speakers. This result is not compatible with our hypothesis, which expected to have similar results for the two languages.

4.2.5. Discussion

The results of the main analysis and the inter-typological contrast are totally in line with our expectations, and with the motion event typology proposed by Talmy (1985). There is an obvious inter-typological pattern, speakers of Turkish and French with high path averages and speakers of English with a high manner average.

However, the intra-typological contrast presented in the previous section demonstrates a different pattern than assumed. Even though the path averages of Turkish and French groups are considerably higher than their manner averages, the difference between the two groups is also significant. A closer look at the data and the plot reveals that this significant difference is mainly due to the difference in their path averages. In other words, the two languages seem to be in different places on a cline of path (Ibarretxe-Antuñano, 2008), French subjects giving higher scores to the canonical motion event expression pattern in their language (*path verb* + *manner satellite*) than Turkish subjects. This is a debatable issue, because their performances were very close in the production task. Here two possible explanations may be suggested and both of these explanations are based on the unexpected performance of Turkish speakers, as French speakers behaved as exactly predicted. One reason for Turkish native speakers to give lower scores to *path verb* + *manner satellite* confluents may be the morpho-syntactic flexibility of the Turkish language which presents its speakers a wide range of choices, contrary to the more strict usage array provided by the French language. A second reason for the discrepancy between the two language groups may be the morphological productivity of the –arak suffix (the manner adverbial making suffix) in Turkish, which may also led to differential choices by different speakers of the same language. These two hypotheses will be discussed in detail in Chapter 5.

4.3. EXPERIMENT 3: SIMILARITY JUDGMENT TASK

4.3.1. Aim and Scope

The first two tasks were both verbal tasks, which made crosslinguistic inquiries on the motion event expressions used in typologically similar and different languages. However, the present task will make a different type of inquiry. This is a non-linguistic categorization task, during which the subjects evaluate the motion events presented to them on the screen based on their perceived similarity. As mentioned in section 2.3.2.2.1, this is a widely-used experimental technique to shed light on the *language and cognition interface* debate, which has regained popularity during the last few decades (see Gentner & Goldin-Meadow, 2003 for a detailed review).

In the task, the video clips depicting motion events are presented in triads, the first one being the main/target video and the other two the alternative/candidate videos. One of the alternatives shares the same manner of motion and the other the same path of motion with the main video. The aim of the task is to see how the speakers of each language group, *i.e.* English, French and Turkish, categorize the visually presented motion events by observing which semantic component (*manner* or *path*) is taken as the dominant criterion of similarity and thus attended more; and to understand whether the canonical motion event expressions used in their language (experimentally verified by the first two tasks) have any effects on their conceptual event representations or not.

4.3.2. Research Questions and Hypotheses

The specific research question and the suggested prediction about its answer are as follows:

Question: Is the conceptual event representation universal (*universalist view*) or is it bound by linguistic encoding preferences of certain languages (*relativist view*)? More specifically, will the non-linguistic categorization performances of the native speakers of the three languages reflect any language-specific effects or not?

Hypothesis: As the present task has a non-verbal nature, in other words as it does not include any linguistic components or linguistic aims; we do not expect to find any language-specific effects in native English, French and Turkish speaker performances. Thus, it is hypothesized that their similarity judgments will not be influenced at all by the canonical motion event verbalization patterns in their respective languages. We do not have a hypothesis regarding the relative choice of manner vs. path, we just predict that they will present a uniform pattern across languages.

4.3.3. Design

4.3.3.1. Subjects

The same 79 subjects (32 native speakers of Turkish, 23 native speakers of English, 24 native speakers of French) were tested in the present task.

4.3.3.2. Stimuli

The videos used in this task are also taken from the same video pool. The subjects watch 10 triads, and in each triad there is a main video and two candidate videos. One of the candidates shares the same manner with the main video, and the other the same path. Thus, each subject is exposed to a total of 30 video sequences.

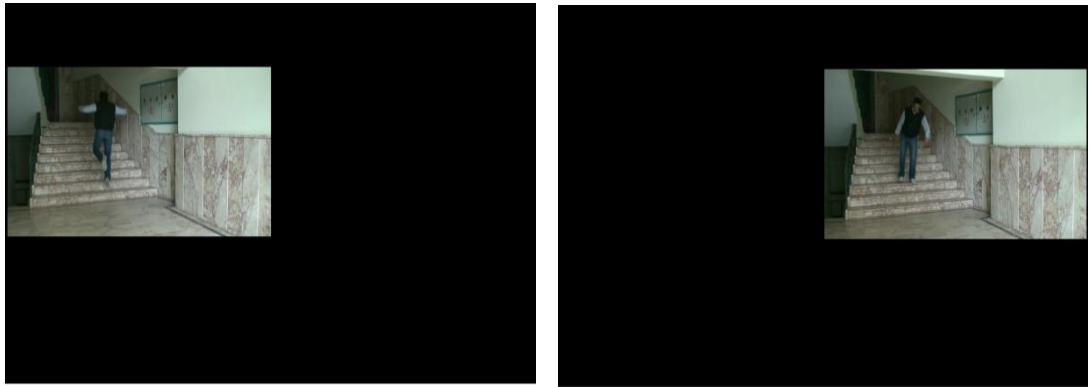
4.3.3.3. Apparatus

The materials were displayed on the same computers as in the other tasks. In the present task, we used again the before-mentioned software, prepared by an expert, which helped us display the stimuli on the screen and record the answers of each subject.

4.3.3.4. Procedure

The aim here was to investigate the conceptual event representations of our subjects with the help of a similarity judgment task. The subject first saw the main video of the triad, then the main video disappeared and the two alternatives appeared side-by-side consecutively (see Pictures 8 and 9). One of the alternatives had the same path as the main video but a different manner; and the other alternative had the same manner as the main video but a different path. For example, if the target item was “hopping down”, then the two alternatives would be “hopping up” and “limping down”. The task was to watch the three videos carefully, and to choose the alternative which is “more similar” to the main video by clicking on the buttons (1) and (2), which appeared at the end of each triad. The mouse-clicks were recorded by the system for later analysis. The display order of same-manner and same-path

alternates were scrambled for all subjects. The whole procedure took about 7-10 minutes. There were three practice triads at the beginning of the task, and the subject was again alone in the testing room during the task.



Picture 8 and 9: Example screen shots from the alternative videos presented consecutively

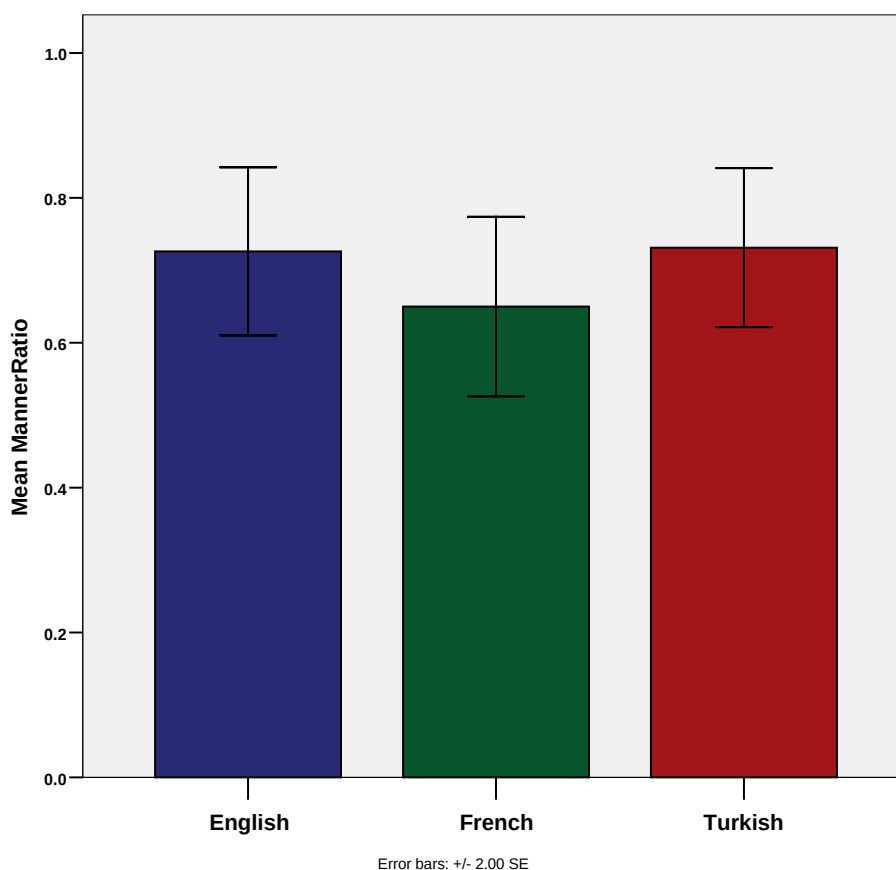
4.3.3.5. Encoding of the Raw Data

First of all, the mouse-clicks of each subject for the 10 triads were coded as either M (the same-manner alternate) or P (the same-path alternate). The next step was to count the number of M and P answers to see the tendency for each language group. For the sake of statistical analysis, a *MannerRatio* (the ratio of the manner choices of a subject to all of her/his choices) was calculated for each subject. Taking the number of manner answers to calculate the ratio was a random choice. As the same-manner alternate numbers and same-path alternate numbers are complements of each other, we would certainly have the same results if we chose to use the *PathRatio* instead of *MannerRatio*.

4.3.4. Results

The overall analysis showed a uniform pattern across the three language groups (see Graph 6). All the groups chose more same-manner alternates than same-path alternates as being more similar to the main item, which means that regardless of their native languages, subjects in the three language groups took *manner of motion*

as the dominant criterion of similarity in this categorization task ($m=7.3/10$, $sd=.2783$ for English; $m=6.5/10$, $sd=.3036$ for French; and $m=7.3/10$, $sd=.3105$ for Turkish).



Graph 6: Similarity Judgment results for all groups

For the statistical analysis, a Uni-variate ANOVA was run with the value *MannerRatio* as the dependent variable and the *language group* as the independent variable; and neither the overall relation between the variables nor the pairwise contrasts did reveal any significant results. These results are in line with our hypotheses.

4.3.5. Discussion

The results of the present task are non-significant in all combinations; however they are totally in line with our hypotheses. As stated in section 4.3.2, we were expecting to find uniform results for the three language groups in line with the *universalist approach* (Jackendoff, 1990, 1996). As this was a non-verbal task, independent of any communicative goals, the parallel choices revealed by our subjects regardless of their native languages were not a surprise.

As can be observed in the graph, speakers of the three languages all went for manner as the criterion for similarity, which can be explained by the dominant salience of the manner component. Actually, both the manner and the path components were visually salient in our stimuli, which is a fact verified by the language production performances of our subjects. Almost all of the subjects in all groups used both components in their descriptions; 83.4% of English sentences, 97.4% of French sentences, and 92.8% of Turkish sentences included conflated manner-path patterns. On the other hand, as the manner of motion is represented by the act of the agent, and as it is a moving and animate agent (as opposed to the static nature of the path of motion), the manner-bias of the subjects may be explained by the animacy and dynamicity effects. This issue will be discussed further in Chapter 5.

The following task is an eye-tracking task, which has the aim of complementing the offline nature of the present task with an online methodology, and of consolidating the data obtained from this task.

4.4. EXPERIMENT 4: NON-VERBAL EYE-TRACKING TASK

4.4.1. Aim and Scope

This task has the aim of consolidating the results of the previous task by shedding more light on the debate on *language and cognition interface* by using a purely cognitive task and an online methodology. To be more specific, in the present task, the eye-movement patterns of the speakers of the three languages in question, namely Turkish, English and French, are recorded while they are watching real-life motion

event sequences. It is investigated whether the speakers of typologically similar and different languages focus on the same semantic components - manner of motion and path of motion - of the motion events presented to them (attention allocation patterns), and whether they do that in the same order (temporal linearization patterns). Eye-movement research does not only provide information on human perceptual system but also valuable insights into human cognitive system via investigating the underlying attentional mechanisms (Richardson and Spivey, 2004a; Murata and Furukawa, 2005; Rothkopf *et al.*, 2007; Rayner and Castelano, 2007). It is argued that people's eye gaze is synchronized with what they are attending to during online information processing. In order to ensure the purely non-linguistic nature of the task, any possible sub-vocal verbalization is inhibited via an experimental technique called *Articulatory Suppression* (Murray, 1967; Baddeley & Hitch, 1974).

4.4.2. Research Questions and Hypotheses

The specific research questions that are asked in the present task and the hypotheses that are put forward for each of them are as follows:

Question 1: Is the conceptual event representation, insofar as this representation is revealed by subjects' eye movements, uniform across languages or do the linguistic motion event expression patterns of a language (*e.g.* V-language pattern or S-language pattern) have any effects on conceptual representation? In other words, do native speakers of Turkish and French attend more to the presumably central component in their linguistic expression pattern (*i.e.* the path of motion), and native speakers of English to their own (*i.e.* the manner of motion) or do they all reveal a common attention allocation pattern in line with the *universalist view*?

Hypothesis 1: Our own expectations are in line with the universalist view, thus we hypothesize that there will be no differences in the attention allocation patterns of the three language groups as revealed in subjects' gaze patterns.

Question 2: Are the temporal linearization patterns of the native speakers of the three languages also uniform across languages? In other words, do native speakers of the three languages in question attend to the manner and path components in the same

order or are there language-specific linearization patterns for different languages, for example English speakers attending to the manner of motion first as it is presumably the central component in the language, whereas Turkish and French speakers prioritizing path of motion for the same reason?

Hypothesis 2: As this is a strictly non-verbal task, no language effect is predicted. Therefore, we hypothesize that the temporal linearization preferences of the three groups will not differ at all, and we also predict that manner component will be prioritized, in other words attended first, by all groups; because it is visually more salient due to its dynamic and animate nature.

4.4.3. Design

4.4.3.1. Subjects

Here again, the same 79 subjects were tested, however due to technical reasons, the data for 18 of them had to be eliminated from the analyses. Therefore, the data for 24 native speakers of Turkish (12 females and 12 males), 22 native speakers of English (12 females and 10 males), and 15 native speakers of French (13 females and 2 males) were used.

4.4.3.2. Stimuli

The present task also made use of the 50 motion event videos used in the production task. Here again, half of the subjects in each language group were tested with the first 25 of the videos and the other half with the remaining 25. The reason for having two versions of the same task is again to avoid an item effect.

4.4.3.3. Apparatus

As explained before, at the main testing lab in Ankara, the videos were all displayed on a 17" computer screen. The eye-tracking system, Tobii 1750²³, is integrated in this

²³ Tobii Technology AB, Sweden.

screen, thus the subjects did not need to use an extra apparatus (see Picture 10). The whole procedure took place in the same testing room. The eye-tracking system at the lab in Paris was not a static but a portable one, a Tobii X120²⁴ (see Picture 11). Thus, there was need for a computer to connect the eye tracker and to use it together. We used the Dell Vostro 1320 laptop computer that we used in other tasks and placed the portable eye-tracker in front of the computer screen, so that they could function together.



Picture 10: Tobii 1750 Eye Tracker



Picture 11: Tobii X120 Eye Tracker

4.4.3.4. Procedure

The aim of the present task was to investigate the attention allocation and temporal linearization patterns of the three language groups crosslinguistically, and to see whether the conceptual event representations - as revealed by eye-movement patterns - are free from any language-specific effects or whether the typological motion event expression pattern of each language has an effect on its speakers' non-linguistic performance. In this task, the subjects were asked to watch a series of short motion event videos presented on the computer screen one after another. There was a two-second-long black screen between each of the individual items, and the subjects watched a total of 25 videos. Their task was to watch those videos, and to repeat the numbers 1-2-3 all along the task. The speed and the rhythm of the repetition were up to the subjects. This technique, which is called the *Articulatory Suppression Technique* (Murray, 1967; Baddeley & Hitch, 1974), had the aim of preventing the

²⁴ See www.tobii.com for more information.

subjects from silently verbalizing the events they are watching, and thus preserving the non-linguistic nature of the task. This same technique was also used by Gennari *et al.* (2002), and they argued that this technique helped them “to minimize linguistic processing of the events and to decrease memory performance by loading verbal working memory during encoding” (p. 56). During the whole task, eye-movements of the subjects were recorded by the system for later analysis. Before starting the task, each subject was informed about the functioning of the eye-tracker, so that they would be more conscious users. Then, the position and the posture of the subject was arranged and fixed, so that the system could detect the eye movements of the participant without difficulty, and thus there would be no eye-track loss in the data²⁵. The calibration of the system and the subject’s eyes was performed by the experimenter together with the subject. The calibration consisted of following a red ball moving along the screen with the eyes in order for the system to determine where exactly the subject is looking. The experimenter again left the room after giving the necessary instructions to the subject. The task took 5 minutes.

4.4.3.5. Encoding of the Raw Data

The data obtained from this experiment were the eye-movement patterns of each subject recorded by the Tobii 1750 (or Tobii X120) eye-tracking system. The raw data included the video observed by the subject plus the eye traces that s/he left on the video, marked by a moving and enlarging red dot (see Picture 12). The larger the dot is, the longer the focus of the subject on that point.

The software supplied by the eye-tracking system, Tobii Studio, has its own statistical analysis tools, however it only supports static data. As our data were dynamic, we could not make use of Tobii Studio for the purpose of our analyses. Therefore, we needed some complicated software for our own specific purposes and for this aim we got help from a software company. The software was prepared by a software engineer, and its complex usage was explained to the experimenter by the same person. The experimenter needed to go through three stages to make use of the software:

²⁵ Despite all the efforts of the experimenter at this stage, and although each subject was told to try not to change the fixated position and their posture much during the experiment, there was a considerable number of subjects whose eye data had to be eliminated due to eye-trace loss.



Picture 12: Example screenshot from the raw eye-tracking data
(The two red dots indicate two fixation points of the subject)

Stage 1: Determining the manner and path regions on each video

As already mentioned, the subjects watched a total of 25 short motion event videos in this task. Both the manner and the path components were salient in all videos. The aim was to find out whether the subjects looked more to the manner component or to the path component throughout the task. To calculate those ratios, it was first necessary to determine the regions that represent the manner of motion and the path of motion for each video. It was more practical to determine the path regions, because they were all fixed static regions. As path is defined as the region between the source and the goal, the following regions were determined as the *path of motion regions* for the three backgrounds used in the videos: 1- for *into* and *out of* scenes, the region between the point where the agent starts/ends the motion event and the apartment building door (see Picture 13), 2- for *up* and *down* scenes, the region between the point where the agent starts/ends the motion event and the elevator door at the top of the staircase (see Picture 14), and 3- for *across* scenes, the region between the two sides of the pavement (see Picture 15). The software has a window where you may visualize an example frame of each video (which is chosen by the system itself), and then you may manually mark your path of motion region to be used in the analyses (with a green lining).



Picture 13: The path region for *into* and *out of* scenes

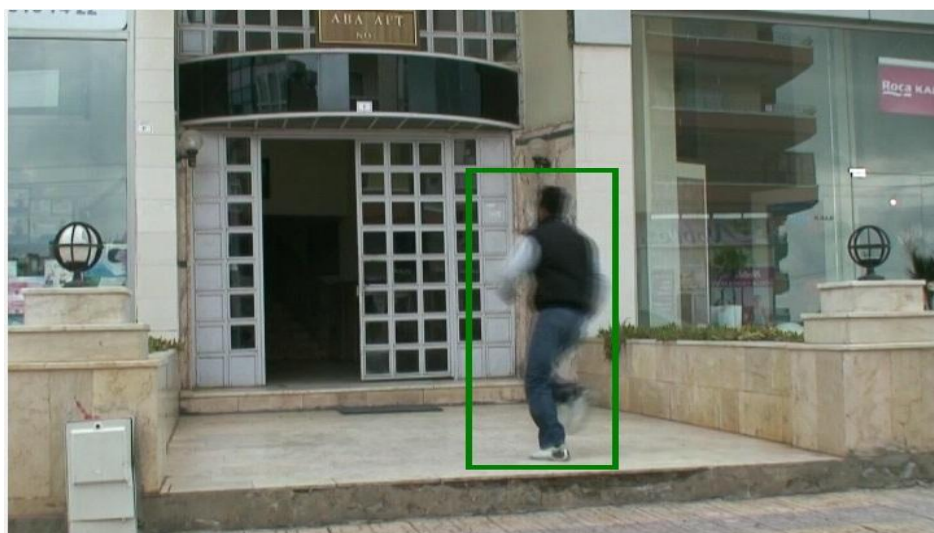


Picture 14: The path region for *up* and *down* scenes



Picture 15: The path region for *across* scenes

The manner region marking was a much more complicated process. We decided that determining the whole body of the agent as the *manner of motion region* was the most appropriate choice; so that every muscle involved in the motion will be included in the region, and therefore more reliable results will be obtained (*cf.* Soroli & Hickmann, 2010). As the manner region is not a static region like our path region, but a dynamic one, each video had to be cut into frames, and the manner regions (*i.e.* the whole body of the agent) had to be manually marked for each and every single frame with a green lining (see Pictures 16 and 17). Here, it is necessary to note that there is an overlap between the manner regions and the path regions determined for the eye-tracking data analyses. We would like to emphasize and ask the reader to keep in mind throughout the whole study that the region that is called the “manner region” for practical purposes may not purely represent the manner of motion, but is rather a confounded representation consisting of the manner inherent to the body of the agent and those sections of the path in which that body is located at certain points in time. This is an inevitable confound that we have to tolerate and acknowledge for the sake of taking the whole body of the agent to represent manner, as justified elsewhere. It is inevitable, because in the 2D representation of the motion events in our videos, the mass of the body of the agent occupies a certain location in space that partially or completely overlaps with the trajectory that we define as the path region.



Picture 16: Example manner region for the scene *run into* (frame 1)



Picture 17: Example manner region for the scene *run into* (frame 60)

Stage 2: Preparing the configuration sheet

The raw eye-tracking data was composed of a complete single video, including the 25 short clips, for each participant, thus it was necessary to determine the starting and the end points for each single clip. For example, the first video of the series is *run into*, and it starts at second 20.048 and ends at second 22.860. As we used two versions of the same task to avoid an item effect, the time points changed from one version of the task to the other. To analyze the raw data, the system needed the starting and ending points for each and every video, thus that information for each version was entered into an .xml file, called the *configuration sheet*, manually.

Stage 3: Running the technical analysis

The technical analysis of the raw data was performed based on the information entered into the system, *i.e.* the manner and path region markings, and the starting and ending points for each single video. It was then necessary to choose the subjects to be included in the analysis and the appropriate configuration for them, and to give the command “run” to the system.

As a result of this three-staged process, we obtained two ratios for each subject and for each video; a *manner look ratio* (the ratio of the eye gazes that fall upon the manner region²⁶ to all eye gazes), and a *path look ratio* (the ratio of the eye gazes that fall upon the path region to all eye gazes). These ratios were used for the statistical analysis of the data, which will be presented in the upcoming section.

4.4.4. Results and Discussion

The technical analysis, which is performed with the help of the software, actually produced a four-column excel sheet for each subject. The first column shows the subject's *manner look ratio* (the ratio of manner looks to all looks), the second one her/his *path look ratio* (the ratio of path looks to all looks), and the third one her/his *other look ratio* (the looks that fall outside the manner and path regions, *e.g.* the shop window in the apartment building entrance scene) for each of the 25 videos watched. The last column is the "no trace" column, which gives the ratio of the missed looks; in other words, the time points where the system could not detect an eye-trace, most probably because of an abrupt movement of the subject. The subjects who had less than 50% eye-traces were eliminated from the analyses.

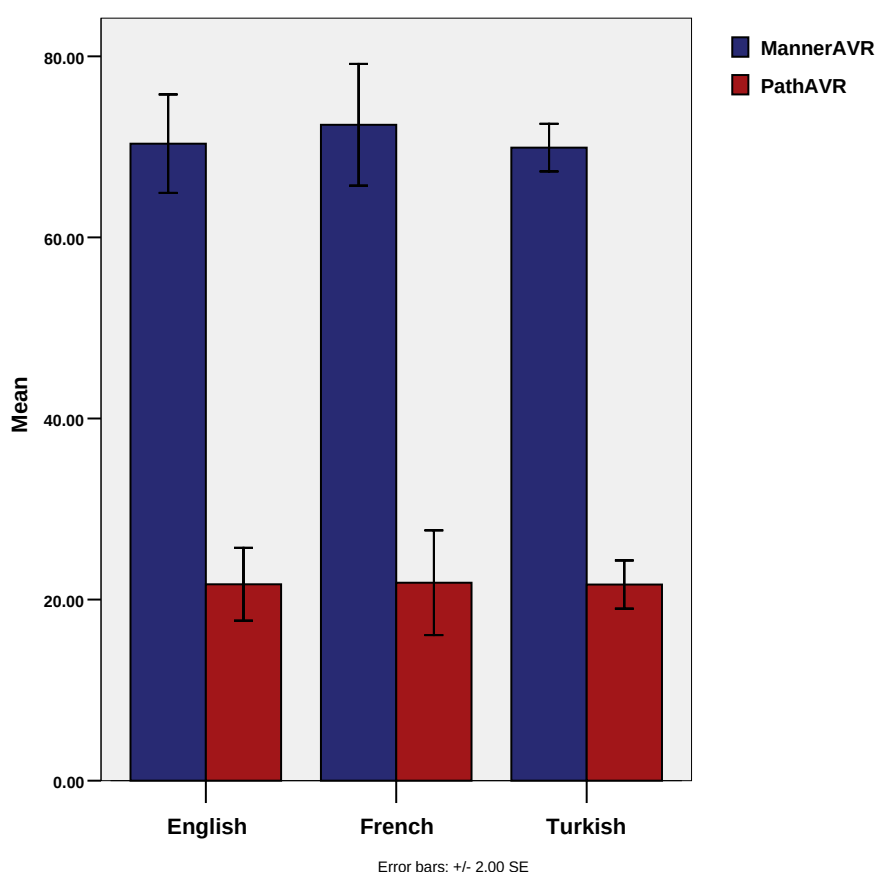
We started by taking the mean values of each column, so that we would obtain an overall *manner ratio*, *path ratio* and *other ratio* for each subject. Then the other ratio was eliminated from the analysis for two reasons; first of all the other looks were not part of the data that we were interested in. Secondly, as using only manner and path region looks would give us more precise ratios, we thought that making the analysis with those ratios would provide more reliable results. We calculated two ratios out of manner and path looks; namely *Manner-to-Manner+Path (M-to-MP) Ratio* and *Path-to-Manner+Path (P-to-MP) Ratio*. As the M-to-MP and P-to-MP ratios were complements of each other, we only used one of them in our analyses: the *M-to-MP ratio*²⁷.

²⁶For methodological reasons, the whole body of the agent will be taken as the *manner region* and the looks that fall upon that region will be labeled as *manner looks* in our analyses of the eye-data. However, this is not a pure representation of manner as there is an overlap between the manner regions and the path regions, as acknowledged.

²⁷ The exact same results would be obtained, if we used the P-to-MP ratio instead.

4.4.4.1. Main Analysis

Graph 7 shows the overall eye-movement patterns for the three language groups. From this graph, it is highly clear that there is a uniform, almost identical, pattern across languages. The speakers of English, French and Turkish, they all have very high manner ratios very close to each other ($m=76\%$, $sd=12.501$ for English; $m=76.6\%$, $sd=12.305$ for French; and $m=76.4\%$, $sd=6.498$ for Turkish), which indicates that they all looked significantly more to the manner of motion than the path of motion while observing the motion event videos presented to them.



Graph 7: The eye-tracking results for all groups

Statistically speaking, a Uni-variate ANOVA which was run taking *M-to-MP Ratio* as the dependent and *language group* as the independent variable did not reveal any significant relations. It means that the ratio of manner looks did not change to a

significant degree from native speakers of one language to those of the other. The Helmert contrast making both an inter-typological (English vs. Turkish and French) and an intra-typological comparison (Turkish vs. French) did not yield any significant results, either, which can be interpreted as pointing to a uniform conceptual event representation pattern across languages.

4.4.4.2. Discussion of the Main Analysis

The proponents of the *linguistic relativity hypothesis* (Whorf, 1956) claim that the native language that we speak shapes -or at least has an effect on- our way of seeing the world. On the other hand, the supporters of the *universalist view* (Jackendoff, 1990, 1996) do not believe in the interdependence of linguistic and cognitive representation. As can be seen in the graph above, there are not any significant differences between the three groups, they even have almost identical patterns. It is a result which is in line with the hypothesis that we put forward at the beginning, and which also seems in line with the universalist approach. As this task does not include a communicative context or any components related to the linguistic representation of motion events, we were not expecting to find any effects of the verbal motion event expressions used in the languages in question. Bunker *et al.* (2010) also points to the very same point by suggesting that when language is not involved in the task or when accessing to language is somehow impeded, the language-specific effect disappears. Our case has both conditions described by Bunker *et al.*; language is not involved in our task and the unconscious access to language is blocked by the use of the articulatory suppression technique.

Obviously, the speakers of the three languages all paid considerable attention to the manner component, while watching real-life videos depicting motion events. This may be either due to the animacy and dynamicity effects, which made our agent (the manner region) much more salient than the non-agent and static background (the path region); or due to the fact that there was partial spatial overlap between the manner region and the path region in our stimuli, and that subjects were also obtaining information regarding the path of motion while they were looking at the manner region. These two propositions will be detailed in Chapter 5.

4.4.4.3. Temporal Linearization Analysis

Linguistically expressing Manner and Path is a classic linearization problem in Levelt's (1989) sense. There is no natural order between Manner and Path; in fact, they are simultaneous aspects of an event. However, languages typically encode Manner and Path in separate lexical items, and they need to be ordered and, more importantly, to be in a particular syntactic relationship with each other. (Allen *et al.*, 2007, p. 22)

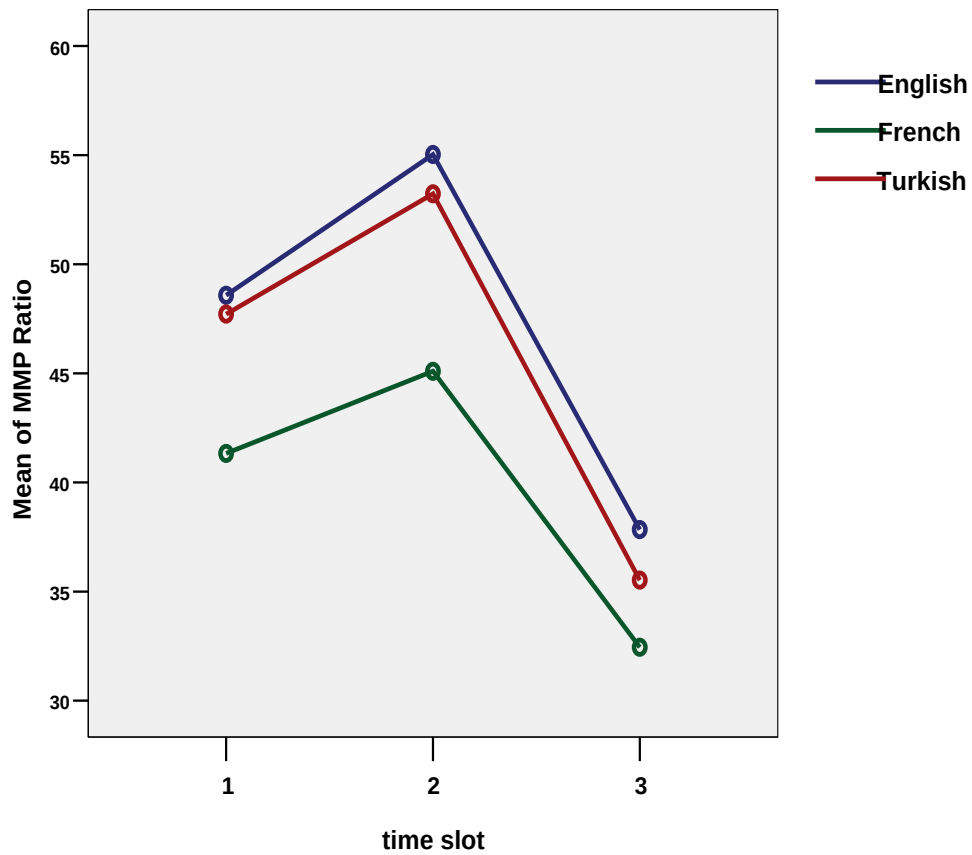
As expressed by Allen *et al.* (2007), even though there are canonical syntactic organizations of the manner and path components in a sentence expressing a motion event in a given language, the perceptual ordering of those items is a matter of debate. In the second part of our analyses, we will deal with this *linearization problem*, the ordering of manner and path components in perception by the speakers of different languages. Eye-tracking methodology is a fruitful technique on the way to solve this problem with the moment-by-moment information it provides.

The *temporal linearization analysis* presented here is a secondary analysis which has the aim of complementing our main eye-tracking analysis by providing clues about the timeline of attention allocation during a motion event scene. For this separate analysis, the eye-movement data were re-encoded as follows: The whole time-frame (the total duration of a single video clip) was cut into three parts; *the beginning*, *the middle* and *the end*²⁸. Then, a *Manner-Minus-Path (MMP) Ratio* was formed to be used in the analyses. This ratio was calculated by subtracting the path look duration from the manner look duration in a certain frame. The overall logic is that if the ratio is above zero, it indicates a manner-dominant look; however if it is below zero, then it is a path-dominant look. The mean and the variance values of the *MMP Ratio* for each single video, for each language group and for each of the three time slots were calculated. Different from the subject-wise design of our main analysis, this timeline analysis was an item-wise one.

A two-factorial ANOVA was run with *time slot* (1, 2 and 3) and *language group* (English, French and Turkish) as three-level within-subject variables. A first analysis

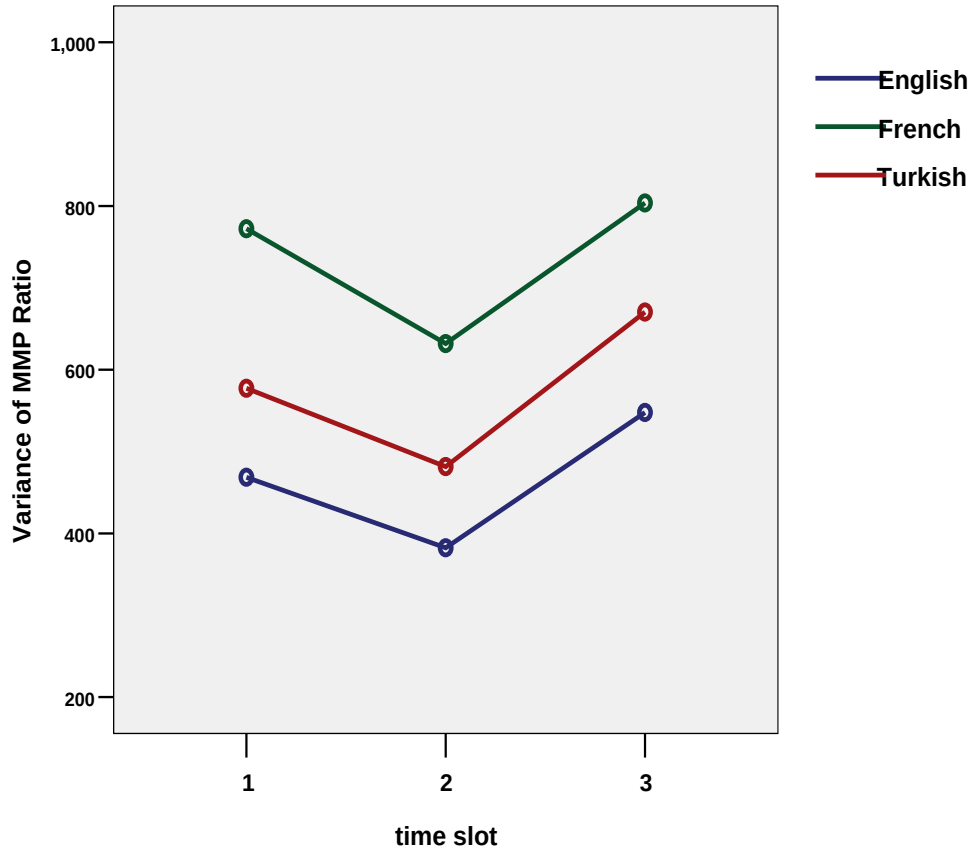
²⁸ The first 10% of each video clip was eliminated from the timeline analysis with the idea that the first looks of the subjects will be random and will not reflect their overall tendency. Therefore, the *beginning* part will start from 10% line and go on till the 40% line, the *middle* part will be between 40% and 70% lines, and the *end* part will be between 70% and 100% lines.

was performed with *mean values* and a second with *variance values*²⁹. As a result, neither the mean analysis nor the variance analysis revealed a significant interaction between *time slot* and *language group*. However, there were significant main effects of *time slot* ($F(2,98)=5.215$, $p<.05$, $\eta_p^2=.096$) and *language group* ($F(2,98)=6.309$, $p<.01$, $\eta_p^2=.114$) in the mean analysis, and a significant main effect of *language group* in the variance analysis ($F(2,98)=14.874$, $p=.00$, $\eta_p^2=.233$) (see Plots 2 and 3).



Plot 2: Means of MMP Ratio across languages

²⁹The variance analysis was performed to see whether that temporally extended process varies across time (e.g. if subjects vary more in the beginning and less in the end).



Plot 3: Variances of MMP Ratio across languages

The main effects in the mean analysis show that the MMP Ratios were significantly different from each other both across languages and across time slots. On the other hand, the main effect in the variance analysis means that the variances across subjects differ significantly from one language group to the other.

4.4.4.4. Discussion of the Linearization Analysis

It can be observed from the plots that all the three languages reveal a similar pattern both in their mean analysis and variance analysis. If we start with the means plot (Plot 2), we can say that there is an overall and dominant manner-bias throughout the whole timeline, which can be inferred from the above-zero values that we have in each time-slot. This pattern shows us that the manner-dominant looks observed in our main analysis for the speakers of the three languages (see Graph 7) are persistent from the beginning to the end. The mean values of the MMP Ratio are obviously the

highest in the second time-slot, which means that manner-looks reach their highest point at around the middle point of the videos. Then in the third time-slot, the value is at its lowest point, which points at the lowering of manner-looks towards the end of the video. The whole pattern in Plot 2 can be interpreted as follows: The subjects, regardless of their native languages, focus more on the manner of motion than the path of motion throughout the whole video. However, their manner-dominant looks make a peak in the middle of the scene, which may suggest a shift from a more balanced manner-path observation that they have at the beginning of the video to a more deliberate and detailed analysis of the manner information. This is probably the point where the subjects form their whole package of manner of motion information. Then, towards the end of the scene, we see a significant decrease, which can point out to a focus-shift from manner of motion to path of motion. It may be suggested by looking at the data that path information is obtained rather late, towards the end of the video by all subjects. It should be noted that even at the end of the timeline, where the subjects turn their attention from manner of motion to path of motion, manner-look ratios are still higher than path-look ratios, which can again be explained by the dominant saliency of the moving and animate agent, or the simultaneous observation of manner and path by a single look.

The variance analysis provides us a rather similar picture. As can be observed from Plot 3, the variance between subjects is in its lowest value in the second time slot, which represents the middle part of the scene. The lowering of variance at that stage may again be interpreted as a common focus on manner of motion by the speakers of the three languages; and again the rise of the variance towards the end may be interpreted as a focus-shift from manner of motion to path of motion.

Another point that should be discussed here is the difference between the performance of native speakers of French, and those of native speakers of English and Turkish, especially in the mean analysis. It can be seen in the plots that French subjects' overall means are significantly lower than the other subjects and that their overall variances are higher than the others. We do not think that this discrepancy has anything to do with any linguistic factors, because they also share the same rising-and-then-falling pattern (and just the opposite in the variance analysis) with other

subjects. Therefore, a plausible explanation may be the low number of French subjects tested in the task.

4.5. EXPERIMENT 5: PRE-PRODUCTION EYE-TRACKING TASK

4.5.1. Aim and Scope

The aim of the present task is to investigate the relation between the language production system and the cognitive processes underlying this system by observing the eye-movements of subjects from three different language backgrounds. The present inquiry is closely related to the *thinking-for-speaking* proposal of Slobin (1996a), which he defines as “a special kind of thinking that is intimately tied to language- namely, the thinking carried out on-line, in the process of speaking” (p. 75). In this task, the eye-movements of the speakers of the three languages, *i.e.* Turkish, English and French, are recorded while they are watching a series of motion event depictions with the aim of describing them just after watching. Here, we question whether the speakers of typologically different languages attend to the motion event components expressed in the main verb in their respective languages more than and also before the other components, in a task which has a verbal aim. In other words, it is observed whether their motion event verbalization patterns influence their attention allocation and temporal linearization patterns while watching motion events with the explicit aim of describing them after watching. Flecken (2011) suggests that observation of eye-movements before or while speaking can tell us a lot about the speech planning process (*cf.* Griffin, 2004). The previous two non-linguistic tasks showed us that conceptual event representation is uniform across languages and free from any language-specific effects. What is questioned further in the present task is whether this representation is still intact – free from any linguistic effects – when there is a linguistic goal involved in the task (here, describing the motion events after watching).

4.5.2. Research Questions and Hypotheses

Question 1: Is the conceptual event representation bound by any language-specific effects when the event representation is formed with the aim of utterance formulation

or is it still uniform across languages? In other words, as proposed by Slobin (1996a), are speakers' conceptual representations affected by the linguistic encoding preferences of their native languages when they are thinking with the aim of speaking?

Hypothesis 1: As this experiment is part of the production task introduced above, and as the subjects are presumably in the process of preparing their verbal utterances while watching the videos, it is predicted that Turkish and French speakers' eye gazes will focus more on the path of motion (the central component in their languages) and that English speakers will attend more to the manner of motion (the central component in their language), in line with Slobin's *thinking-for-speaking* proposal. As put forward by Bungler *et al.* (2010), "when viewing an event while preparing to describe what they see, adults very quickly *direct their attention to components of the event that they plan to talk about...*" (p. 58, emphasis added).

Question 2: Is there a language-specific effect in the temporal linearization patterns of the subjects in this pre-production eye-tracking task?

Hypothesis 2: We expect to observe a language effect in terms of the ordering of motion event components, as well; however we do not have clear-cut hypotheses in this regard. It is possible to assume that English speakers will attend first to the manner of motion as it is the semantic component expressed in the main verb of the sentence and thus presumably the central one for them, and that Turkish and French speakers will attend first to the path of motion as it is the component embedded in the main verb in those languages. On the other hand, if we take the language-based effect for granted, we can also assume that the language-specific word-orders in the linguistic expression of motion events may also play a role in its speakers' perceptual ordering of those components. In other words, it is possible that people attend to the semantic components of motion events "in the order they plan to mention them" (Bunger *et al.*, 2010, p. 58). In this case, English speakers would again attend first to the manner of motion as it is embedded in the main verb of the sentence which is in a post-subject position, and French speakers would again attend first to the path of motion expressed by the main verb that comes just after the subject. Turkish speakers, however, would attend first to the manner of motion as the manner expressing

adverbial usually comes before the path expressing verb in a sentence, because Turkish is a head-final language.

4.5.3. Design

4.5.3.1. Subjects

As this task is part of the language production task presented above, the same 79 subjects mentioned in section 4.1.3.1 were tested. However, data for 18 of them had to be eliminated for technical reasons. Therefore, the data for 24 native speakers of Turkish (12 females and 12 males), 22 native speakers of English (12 females and 10 males), 15 native speakers of French (13 females and 2 males) were included in the analyses.

4.5.3.2. Stimuli

The stimulus materials used in the present task have already been elaborated in section 4.1.3.2.

4.5.3.3. Apparatus

The subjects were tested in the same testing rooms with the same eye-tracking systems mentioned in the previous task.

4.5.3.4. Procedure

As this eye-tracking experiment is part of the Video Description Task, the general procedure was as explained in section 4.1.3.4. As for the eye-tracking procedure, the subjects were again informed that there will be an eye-tracking component in the task, they were again asked not to change their posture during the task, and the calibration was again performed by the experimenter together with the subject. In the present task, as opposed to Flecken (2011) who instructed her subjects to start speaking as soon as they recognize what was happening in the video clip, our subjects were instructed not to speak before the video was over for two reasons. First of all, the

eye-gaze registration systems are so delicate that even the movement included while talking may have an effect on the results. Secondly, we wanted to observe the whole process of pre-verbal message conceptualization without any verbal interruption.

4.5.3.5. Encoding of the Raw Data

The data encoding procedure of the present task is exactly the same as that of the previous eye-tracking experiment.

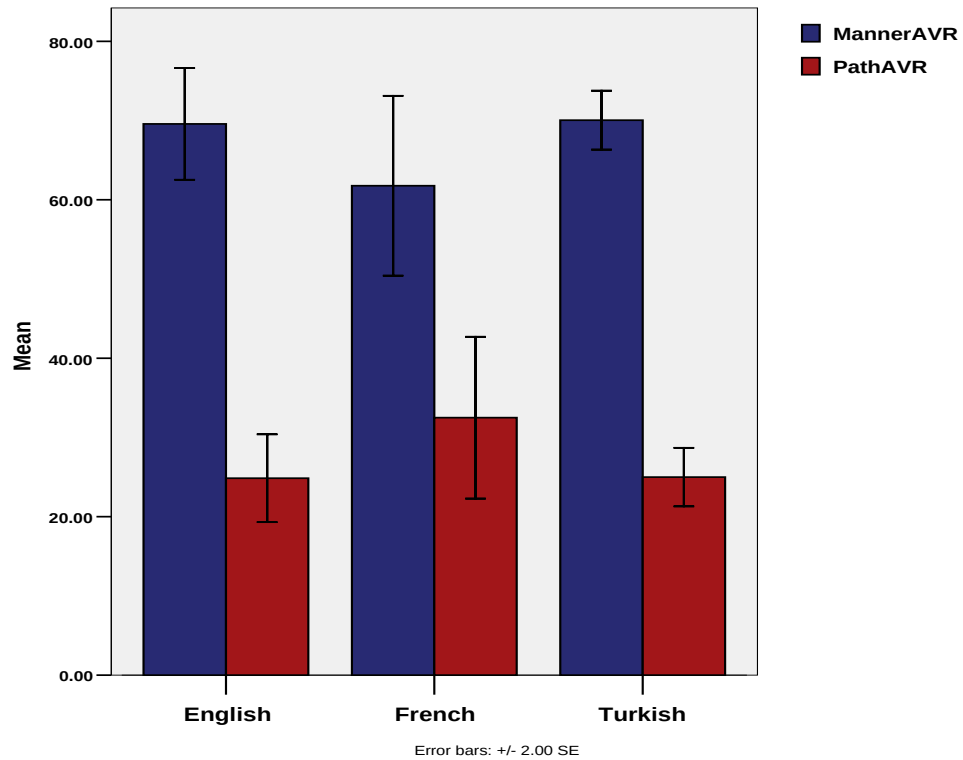
4.5.4. Results and Discussion

Here again, the *Manner-to-MP Ratio* detailed in section 4.4.4 was used in the analyses.

4.5.4.1. Main Analysis

Graph 8 shows the overall eye-movement patterns separately for the three language groups. It is quite evident that there is a uniform pattern across languages; members of the three groups, regardless of their native languages, dominantly attend to the *manner of motion* while they are watching motion events with the aim of describing them (m=69.6%, sd=16.529 for English; m=61.8%, sd=21.971 for French; and m=70%, sd=9.108 for Turkish).

As for the statistical analysis, a Uni-variate ANOVA was run with *Manner-to-MP Ratio* as the dependent, and the *language group* as the independent variable. Neither the main analysis, nor the pairwise comparisons revealed any statistically significant effects. This result suggests a common pattern of conceptual event representation even in a pre-production task.



Graph 8: The eye-tracking results for all groups

4.5.4.2. Discussion of the Main Analysis

The results indicate that there was no significant difference between the attention allocation patterns of the speakers of the three language groups, they all attended more to the manner region than the path region when their eye gaze patterns were observed while they were watching motion event videos with the aim of describing them. When these results are evaluated in the light of the *thinking-for-speaking* hypothesis (Slobin, 1996a), they are not at all supportive. Because, according to the hypothesis, if one observes an event with a communicative intention, then the conceptualization pattern (*i.e.* the attention allocation pattern) of the speaker should conform to the canonical verbalization pattern of her/his language. Our own hypothesis was also in this line. As the participants had the aim of putting what they had watched into words, we were expecting to find language-specific attention allocation patterns, S-language speakers focusing more on the manner of motion and V-language speakers more on the path of motion. However, the results suggested that there was a uniform pattern in favor of the manner component across the three

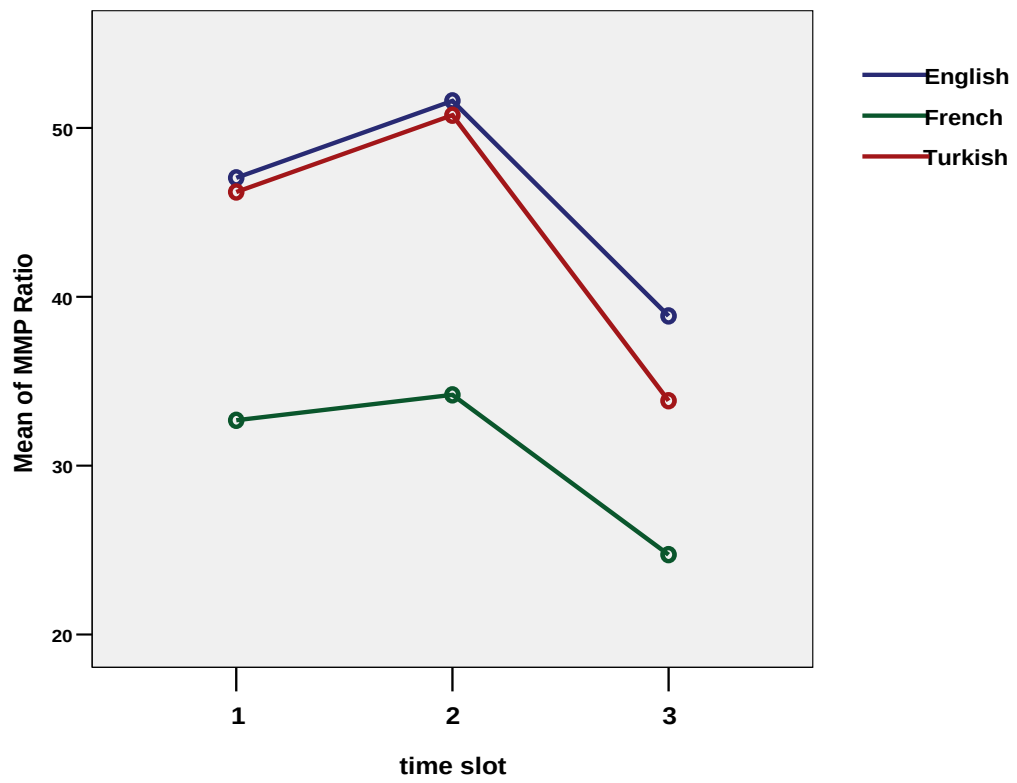
language groups. Therefore, it is obvious that there are no language-specific effects, as hypothesized by Slobin, in our data. These results, which are more or less identical with the results of the previous task, are again in line with the *universalist view*.

Our pre-production eye-tracking results also indicate that the speakers of the three languages all pay more attention to the manner region than the path region, while watching motion event videos. This is again most probably because the manner region is represented by an animate and dynamic agent, whereas the path region by a static background; or because the two regions have partial spatial overlap in the stimuli presented to the subjects. These two hypotheses will be discussed in detail in Chapter 5.

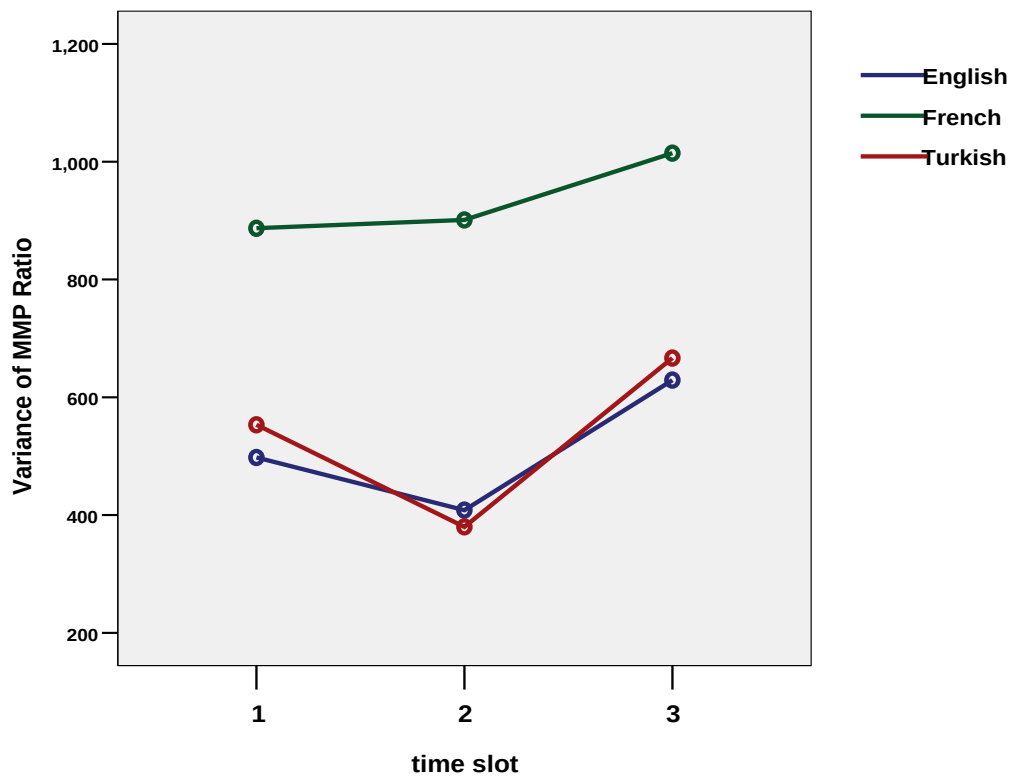
4.5.4.3. Temporal Linearization Analysis

The main analysis presented in the previous section was looking for an answer to our first research question, whereas the analysis in the present section has the aim of investigating our second research question: Do native speakers of different languages attend to manner and path components of a motion event in a common order or in a language-specific order? For this inquiry, we again made both a mean analysis and a variance analysis by using the *Manner-Minus-Path (MMP) Ratio*.

As a result, neither the mean analysis nor the variance analysis revealed a significant interaction between *time slot* and *language group*. However, there was a main effect of *language group* in the mean analysis ($F(2,98)=16.567$, $p=.00$, $\eta_p^2=.253$), and main effects of both *language group* ($F(2,98)=29.830$, $p=.00$, $\eta_p^2=.378$) and *time slot* ($F(2,98)=3.412$, $p<.05$, $\eta_p^2=.065$) in the variance analysis (see Plots 4 and 5). The main effect of language group in the mean analysis suggests that the mean of the MMP Ratios of the three language groups are significantly different from each other, whereas the main effects observed in the variance data show that the variance of the MMP Ratios significantly differ both across languages and across time slots.



Plot 4: Means of MMP Ratio across languages



Plot 5: Variances of MMP Ratio across languages

4.5.4.4. Discussion of the Linearization Analysis

If the results of the linearization analysis of the present task are compared to the linearization results of the previous task (the non-verbal eye-tracking task), it is seen that they present us almost identical patterns. Here again, in the mean analysis, the MMP Ratio makes a peak in the second time slot and a fall in the third. The peak can again be interpreted as a dominant focus on manner information, and the decrease as a focus-shift from manner to path. The results of the variance analysis in this task also revealed a very similar pattern to those of the previous task, decreasing in the second time slot and rising again in the third slot. This result is also in line with the mean analysis result, and it may suggest a focus-shift from manner to path between time slot two and three. The whole picture demonstrates us that manner information is gathered rather early in the timeline, whereas path information is gathered towards the end. As already argued in section 4.4.4.4 for the previous experiment, the persistent focus on the manner component may either be explained by the dominant saliency of the manner component or by the spatial overlap between the manner and path regions, and thus by the simultaneous gathering of manner and path information. This is a point that will be brought up again in Chapter 5.

CHAPTER V

GENERAL DISCUSSION

The present dissertation reports five experiments conducted with the aim of shedding light on the broad question of *whether speakers of typologically different languages verbalize and conceptualize motion events in different ways*. The study is comprised of two lines of investigation; the *verbal* experimentation line which investigates the language production and language comprehension of motion events across languages, and the *non-verbal* experimentation line which inquiries the conceptual representation of motion events by speakers of different languages. Therefore, this chapter will be organized based on these two lines of inquiry. First, the verbal part of the study and its results will be discussed in the light of the theories already introduced. Secondly, the non-verbal part will be discussed in association with the hypotheses already put forward.

5.1. DISCUSSION OF THE VERBAL DATA

Two of the five tasks elaborated in the previous chapter, *i.e.* the Video Description Task and the Acceptability Judgment Task, aimed at making a psycholinguistic analysis of motion event expressions in Turkish, English and French based on the two-way typology introduced by Talmy (1985). The two tasks were complements of each other, one analyzing the issue from the *production side* and the other from the *comprehension side*. The main hypothesis questioned was regarding the canonical motion event verbalization patterns in those languages; whether Turkish and French which are categorized as verb-framed languages and English which is qualified as a satellite-framed language experimentally conform to the verbalization particularities of their own typologies. There have been quite a number of studies analyzing the verbal productions of people speaking those three languages, comparatively with those of the speakers of other languages (see Özçalışkan & Slobin, 1999, 2000, 2003; Papafragou *et al.*, 2002; Gennari *et al.*, 2002; Pourcel & Kopecka, 2005, 2006;

Gullberg *et al.*, 2008; Soroli & Hickmann, 2010 among others); however, the present study has a particularity which makes it unique in the field. Different from those studies, which base their analyses solely on language production data, this study also makes use of language comprehension data to have a complete view.

5.1.1. Video Description Data

The results of the Video Description Task are totally in line with the Talmyan typology and with our hypotheses. We were expecting that native speakers of Turkish and French would dominantly express the path information in the main verb of their sentences, as a common particularity of V-languages. Likewise, we were assuming that native speakers of English would use the main verbs of their sentences to give the manner information, just like the speakers of other S-languages. These hypotheses are clearly verified by our language production data: 95.4% of Turkish descriptions and 95.8% of French descriptions were expressed with path sentences, and 83.2% of English descriptions were formulated with manner sentences (see section 4.1.4.1). These results are also in line with the results of a large number of previous studies in the field (Choi & Bowerman, 1991; Berman & Slobin, 1994; Slobin, 1996b, 1997, 2004; Naigles *et al.*, 1998; Özçalışkan & Slobin, 1999, 2000, 2003; Papafragou *et al.*, 2002, 2005; Selimis & Katis, 2003; Selimis, 2007; Allen *et al.*, 2007 among many others).

On the other hand, our results are not in line with the *semantic density (SD)* idea, which claims that S-language speakers express both the manner and the path information in their utterances (SD2); whereas V-language speakers are contented with the expression of path-information-only (SD1), omitting the manner information most of the time, in case it is highly required in a given context (see Talmy, 2000; Özçalışkan & Slobin, 2000; Papafragou *et al.*, 2002; Gennari *et al.*, 2002; Özçalışkan & Slobin, 2003; Pourcel, 2005; Slobin, 2004, 2006; Gullberg *et al.*, 2008; Soroli & Hickmann, 2010 among others). We hypothesized at the beginning of the study that our data would not reveal such a differential density pattern, and we were assuming that the speakers of the three languages, regardless of their typology, would express both semantic components (*manner* and *path*) in most of their utterances. We suggested such a hypothesis on the grounds that both the manner and

path components were highly salient in the video stimuli by which we elicited our language production data. The results of the semantic density analysis are in line with our expectations, and reveal that V-language speakers express both components of a motion event as frequently as the S-language speakers do (see section 4.1.4.2). We now have two propositions regarding the possible reasons for other researchers to assume a semantic density difference between V-languages and S-languages:

A) We believe that the *differential semantic density hypothesis*, as we call it, theoretically came out based on misleading examples provided. For example, let's look at the example sentences used in Özçalışkan and Slobin (2000) to argue that V-language speakers usually omit the manner information in their sentences: “The owl *flew off* the hole.” vs. “Ağaç kavuğunun içinden bir baykuş *çıkıyor* (*Eng.* An owl *exits* from inside of the tree hole).” They claim that S-languages use the first type of sentences to provide both the manner information (in the main verb) and the path information (in the satellite). On the other hand, V-languages mostly omit the manner information and only use the path verb as in the second sentence. However, it is clear from the context that the owl is getting out of the hole by flying, as it is the canonical manner of motion for a bird. Thus, it is not that V-language speakers omit manner information; it is that manner information is already embedded in the context and needs not to be expressed again. Other examples provided by Pourcel and Kopecka (2005) also have the same deficiency:

(22) Titi est *sorti* de sa cage (en volant).

‘Tweety exited from its cage (by flying).’

(23) Les abeilles sont sorties de la ruche (en volant).

‘The bees exited from the hive (by flying).’

(24) Le bateau est arrivé au port (en navigant).

‘The boat arrived at the harbour (by sailing).’

It is again highly obvious in those sentences that the bird and the bees get out of the cage flying, and the boat arrive at the harbor sailing, thus there is absolutely no need for the V-language speakers to mark it with a manner adverbial (*en volant* ‘by flying’

or *en navigant* ‘by sailing’). If both semantic components were marked enough, then V-language speakers as well as S-language speakers would express both types of information in their sentences.

B) Our second concern is methodological. When we have a look at the studies claiming such a difference (*e.g.* Özçalışkan & Slobin, 2000 or Gullberg *et al.*, 2008), we notice that they made use of either static pictures or animated cartoons to elicit verbal data. As already discussed elsewhere in the present study, eliciting motion event data with a stimulus set which does not properly reflect the dynamic and realistic nature of the events will not provide naturalistic utterances. It is thus necessary to have real-life scenes or scenes that are closest to real-life in order to be able to elicit reliable data in this regard. Therefore, the reason for those researchers having different results between the speakers of V-languages and S-languages regarding the semantic density of the sentences they used may very well be this methodological-deficiency. In other words, if both manner of motion and path of motion are not properly reflected in the stimulus materials used (*e.g.* an animation figure cannot demonstrate the difference between the manners of *staggering* and *limping* as clearly and realistically as a real human figure), then the verbal descriptions may not include both semantic components (*cf.* Allen *et al.*, 2007).³⁰

5.1.2. Acceptability Judgment Data

The results of the Acceptability Judgment Task are mostly in line with the Talmyan motion event typology and with our hypotheses, as well. Turkish and French speakers give much higher scores to the sentences reflecting the V-language pattern ($m=4.32$ for Turkish, and $m=4.93$ for French), and English speakers give higher scores to the sentences with the S-language pattern ($m=4.74$). However, there is

³⁰ On the other hand, Soroli and Hickmann (2010) who used real-life video shootings to elicit production data also found a differential semantic density effect between English and French speakers, which cannot be explained with our *type of stimulus* hypothesis. It may be possible that the path and manner components were not salient enough in their stimulus materials, so that they did not obtain manner expressions from V-language speakers. This result may be due to the fact that expressing manner in V-languages necessitates a higher cognitive-load than expressing it in S-languages, as it is expressed with an external element in V-languages (Slobin, 2004).

another result obtained from the same task which is against our expectations: there is a statistically significant difference between the ratings of native speakers of Turkish and those of native speakers of French ($F(1,22)=15.877$, $p=.001$, $\eta_p^2=.419$). In order to better understand this intra-typological variation, we went deep into the data and saw that the difference mainly emanates from the discrepancy between the ratings that the two groups gave to the V-language pattern sentences (path sentences). In other words, they did not quite agree on the acceptability of the *path verb + manner satellite* sentences; French speakers gave a mean score of 4.93 over 5.00 to those sentences, whereas Turkish speakers only gave a mean score of 4.32. However, when we have a look at the production data, we can easily notice that the *path verb + manner satellite* pattern is the most commonly used one by native speakers of Turkish, as well. Therefore, it was not the typological pattern of the sentences that created that difference between the two groups. Here are two propositions that we put forward regarding the plausible linguistic reasons for that surprising result:

1. One reason for Turkish native speakers to give lower scores to *path verb + manner satellite* confluents than their French peers may be the morpho-syntactic flexibility of the Turkish language. Let's take the example sentence "The man is running into the building". The French equivalent of this sentence would be "L'homme entre dans le bâtiment en courant (*Eng.* The man is entering into the building (by) running)." Native speakers of French do not have many choices in this regard, thus this is more or less this sentence that was used by each of our French subjects in the production task. On the other hand, native speakers of Turkish have a larger range of choices in this respect. Although all the possible choices still include the canonical motion event expression pattern, there may be slight morphological and syntactic differences. For example, the Turkish sentence used as the counterpart of the above sentence in the task was "Adam koşarak binaya giriyor (*Eng.* The man is entering to the building (by) running)." This was also the fixed pattern that was used in all of the *path verb + manner satellite* sentences presented in Turkish in the Acceptability Judgment Task. However, it is also perfectly possible to express the same event with the sentences (25) and (26), which keep the same typological pattern but use different case markings or with the sentence (27), which again keeps the same pattern but uses another form of manner adverbial (-a...-a repetition marker instead of the -arak suffix).

(25) Adam koş-arak bina-**nın** **içine** giriyor.
 Man-Nom. run-MAdv. building-**Gen.** **inside-Dat.** enter-Pr.Prog.3sg.
 ‘The man is entering to the inside of the building (by) running.’

(26) Adam koş-arak bina-**dan** **içeri** giriyor.
 Man-Nom. run-MAdv. building-**Abl.** **inside** enter-Pr.Prog.3sg.
 ‘The man is entering the inside from the building (by) running.’

(27) Adam koş-**a** koş-**a** binaya giriyor.
 Man-Nom. run-**Rep.MAdv.** run-**Rep.MAdv.** building-Dat. enter-Pr.Prog.3sg.
 ‘The man is entering to the building (by) running.’

The fact is that none of the Turkish subjects gave a score below 3 to any of the canonical pattern sentences (*e.g.* Adam sendeleyerek binaya giriyor - ‘The man is entering to the building by staggering’), they never told that they were unacceptable sentences. However, their ratings were not as high as the French speakers’. Therefore, it may be hypothesized that the morphological nuances detailed above may have influenced their judgments. For example, it is highly possible that they were looking for one of the structures listed above and as they could not find it but a close one, they gave a 4 instead of a 5. The descriptions in the language production task verify this hypothesis, as well, by presenting morphologically-varied responses.

The same may also be true on a syntactic basis. In all of the sentences presented in the task, the manner adverbial was used at the post-subject position as in the above examples. However, it is also possible to use it at the pre-verbal position as in (28)³¹. Therefore, this alternative may again be the canonical preference for the Turkish speakers who gave lower scores than expected.

(28) Adam binaya/binanın içine/binadan içeri **koşarak** giriyor.
 Man-Nom. **run-MAdv.** enter-Pr.Prog.3sg

³¹ The canonical position for adverbs in Turkish is just before the verb (Wilson & Saygın, 2001).

2. The second reason for the discrepancy between the two language groups may be the morphological productivity of the manner adverbial making suffix “-arak” in Turkish. When we have a look at the production data of the French speakers, we see that the manner adverbials used while describing the motion events do not change much from one subject to another. For example, almost all of the subjects used the adverbial “en titubant (*Eng.* (by) staggering)” for the scenes where the agent staggers in/ out of/ up/ down/ across, or the adverbial “à cloche-pied (*Eng.* (by) hopping)” when or wherever the agent hops. On the other hand, Turkish has a larger collection of adverbials that can be used in similar circumstances. For example, one can either use “sendeleyerek” or “sallanarak” to say that someone is staggering; or the adverbials “hoplayarak”, “zıplayarak”, “sekerek”, “tek ayak üzerinde zıplayarak” can all be used to tell that one is hopping. As a result, it is also possible that the Turkish speakers who gave lower scores to our target sentences preferred one of the other alternative adverbials to the one used in the experiment, and this prevented them to give the highest score to our sentences.

5.2. DISCUSSION OF THE NON-VERBAL DATA

The other three experiments, *i.e.* the Similarity Judgment Task, the Non-Verbal Eye-Tracking Task and the Pre-Production Eye-Tracking Task, have the aim of investigating the possible effects of the canonical motion event verbalization pattern in a language on its speakers’ conceptualization of motion events. There are actually two main hypotheses examined in these tasks; the *linguistic relativity hypothesis* (Whorf, 1956), and the *thinking-for-speaking hypothesis* (Slobin, 1996a). Two of the tasks, namely the Similarity Judgment Task and the Non-Verbal Eye-Tracking Task (Experiments 3 & 4), aim at observing the non-verbal conceptual representation of motion events by the speakers of the three languages in question, thus at shedding light on the Whorfian Debate. The Pre-Production Eye-Tracking Task, Experiment 5, on the other hand, has the aim of testing whether the eye-tracking performances of our subjects reflect any language-specific effects when they have a verbal task (description task) to complete, which is closely related to the Slobian Hypothesis.

5.2.1. Similarity Judgment Data

The goal of the Similarity Judgment Task is to see which of the alternative motion scenes, the same-manner alternatives or the same-path alternatives, are dominantly categorized as more similar to the main scenes. The results would show us whether it is the manner of motion or the path of motion which is taken as the primary criterion of similarity by our subjects, and thus which is attended more in a non-verbal task. The main motive of the task is to observe whether there is a common or language-specific pattern of attention allocation, and the results reveal a uniform pattern across languages in line with our hypotheses. All of the subjects, regardless of their native languages, choose more same-manner alternates than same-path alternates, therefore it can be inferred that they all attended more to *manner of motion* ($m=7.3/10$ for English, $m=6.5/10$ for French, and $m=7.3/10$ for Turkish). As it was a non-verbal task, free from any communicative intentions, we were expecting to find such a common structure. However, the reason why they focus more on manner of motion but not path of motion should be taken as a matter of discussion.

Our main explanation regarding the manner-dominant results is related to the nature of the two semantic components (manner and path). Manner of motion is mostly represented by the movement of the agent in the scenes (*e.g.* running, staggering or dancing), and thus has two particularities not present in path of motion: *animacy* and *dynamicity*. In the experiment, all the motion events are depicted by an animate (human) agent and they are all dynamic events. Therefore, the animate and dynamic nature of the manner component may have attracted more attention than the path component, which is represented by a static and inanimate background (*e.g.* the staircase or the street). Pourcel (2005) also demonstrated with a set of experiments that *type of manner* (default, forced/marked or instrumental) has a significant impact on the ratio of same-manner choices. It is the forced/marked and instrumental manner items (compared to default manner items) that resulted in higher manner scores in her study. As all the manners included in our stimuli can be classified as forced/marked manners, our results may be considered to be compatible with Pourcel's in this regard, as well.

The Similarity Judgment Task presented us new insights about the relationship between linguistic representation and conceptual representation by demonstrating that the speakers of the three languages behave alike in a non-verbal task. Using triad materials and a categorization task is a very common way of investigating the language and cognition debate. However, it also has its own methodological deficiencies. First of all, the binary nature of the task limits the results as well as our expectations. Secondly, even though it is a non-verbal task free from any linguistic encoding components, there is nothing during the task that prevents the subjects to make sub-vocal verbalization of the motion event scenes they watch. Thirdly, the relationship between language and cognition is a complicated and delicate process, which should best be inquired online, and the categorization technique lacks this particularity. In order to overcome these methodological drawbacks, we also conducted a *non-verbal eye-tracking experiment*, the results of which shed more light on our research questions regarding the possible effects of linguistic encoding preferences on conceptual representation.

5.2.2. Non-Verbal Eye-Tracking Data

The Video Observation Task, coupled with an Articulatory Suppression Paradigm (Murray, 1967; Baddeley & Hitch, 1974), was the first task presented to our subjects. It was a pure non-verbal task tapping into the conceptual representation of motion events across languages. In this task, the eye-movements of native speakers of English, French and Turkish were analyzed to see whether there were common or language-specific *attention allocation* and *temporal linearization* patterns in the data. The Articulatory Suppression Technique helped us to ensure that there was no sub-vocal verbalization of motion events during the task.

The results suggested a uniform pattern across languages, and thus were in line with our hypotheses just like in the categorization task. All of the subjects, regardless of their native languages, focused more on *manner of motion*³² than path of motion (m=76% for English, m=76.6% for French, and m=76.4% for Turkish). The ratios

³² The overlap between the manner region and the path region is acknowledged.

were very close, which indicates a uniform attention allocation pattern across languages.

On the other hand, the reason for focusing more on manner of motion (methodologically represented by the whole body of the agent.) instead of path of motion may again be explained by the *animacy* and *dynamicity* effects already discussed in the previous section. However, in this experiment, we have an additional hypothesis regarding the possible reason for attending more to the region labeled as manner. As already detailed in Chapter 4, the eye-tracking analysis is based on the manner and path regions determined by the experimenter. Path regions are marked as the space between the source and the goal, and manner regions as the whole body of the agent. As these are all dynamic scenes replicating real-life situations, there are constant overlaps between the two regions throughout the videos. Therefore, it is possible that the subjects who concentrated on our manner regions in the task were simultaneously gathering information about the path of motion, as well. In other words, while they were following the agent with their eyes, they were not just focusing on manner of motion but also indirectly on path of motion.

Another point to be discussed is the *temporal linearization patterns* observed in our data. The results suggested that all of the subjects, regardless of their language, focused more on the area defined as manner at the beginning and at the middle of the scenes; however the level of focus lowered towards the end of the scenes, which was interpreted as a relative shift of focus from manner of motion to path of motion (see section 4.4.4.4). Thus, why do people first focus on manner and then path? Does it mean that they gather the manner information before the path information? It can be argued that the reason for manner to be attended first, as well as more, is again the dominant visual salience of that component. It may be possible that subjects start following the agent from the beginning of the video, and it is just towards the end of the scene that they turn their attention to the path component to have the whole picture of the motion event. However, it is also obvious from the results of the linearization analysis that the focus shift that we have mentioned is not an absolute one. In other words, subjects never cease to look at the manner region, it is just the ratio that drops and that gives us the idea that there is a shift. On the other hand, it is possible that people do not need to focus solely on the path region to gather the path

information. They can always have that information indirectly while they are focusing on the manner region, due to the partial spatial overlap between the two regions. It is also possible that people do not need much time to gather the path information and thus an indirect look may suffice to have the necessary information.

To sum up, both the categorization task and the non-verbal eye-tracking experiment reveal uniform attention allocation and temporal linearization patterns across the three languages examined, thus there are no language-specific effects observed in the data.

5.2.3. Pre-Production Eye-Tracking Data

In this task, the eye-movements of the subjects were tracked and recorded while they were watching a series of video clips depicting motion events with the aim of describing them just after watching. Therefore, they were presumably at the stage of utterance preparation while their eye-movements were observed. Slobin (1996a)'s *thinking-for-speaking hypothesis* (TFS) is closely related to this task, because he claims that people's conceptual planning stages reflect a language-specific effect if and only if there is a verbal objective (a communicative intention) involved in the task. Papafragou *et al.* (2006) also argue that motion event components are "in different stages of conceptual readiness in the minds of speakers of Manner vs. Path languages *immediately prior to verbalization*" (p. B77, emphasis added). That is why we observe our subjects' eye-gaze patterns before they start verbalizing what they see (*cf.* Flecken, 2011).

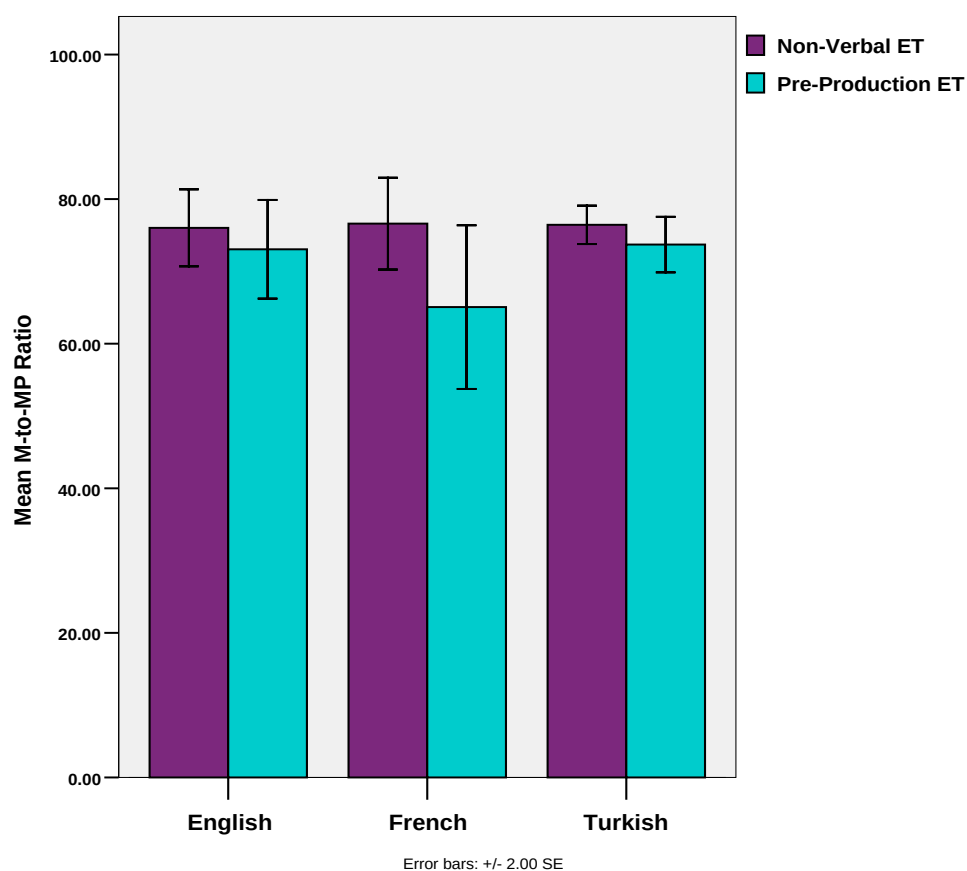
In the research questions and hypothesis section, we talked about our predictions regarding the results of this experiment (see section 4.5.2). It was based on Slobin's argument suggesting that the central semantic component (*manner* or *path*) of a motion event, which is expressed in the main verb in a given language, is attended to more by the speakers of that language during a task where there is a communicative intention. Thus, if we took that claim to be true, we would expect our English subjects to attend more to *manner of motion*, whereas our Turkish and French would focus more on *path of motion* in the current task. However, our results do not support this language-based argument by suggesting a uniform attention allocation pattern

across the three languages. All our subjects, regardless of their native languages, attend more to the region encoded as *manner of motion* (m=69.6% for English, m=61.8% for French, and m=70% for Turkish). This result points at an identical utterance preparation behavior, free from any language-specific effects, and leads us to re-consider Slobin's TFS hypothesis with a critical eye. It may be possible to claim that there is absolutely no effect of verbal typology on conceptual event representation, even when there is language-mediation involved. Or it can also be argued that even though there may be a certain link between linguistic verbalization and conceptual representation, the motion event verbalization pattern of a specific language may not need to be confirmed repetitively in the online attention allocation patterns of its speakers.

The reasons why the subjects focus more on *manner of motion* (represented by the agent) than path of motion can again be explained by the *animacy*, *dynamicity* and *spatial overlap* arguments. The linearization results were uniform for the three language groups, as well. All the subjects first focused on the region labeled as manner, and then relatively more on the region labeled as path (see section 4.5.4.3). To sum up, our pre-production eye-tracking data suggest uniform attention allocation and linearization patterns, free from any language-specific effects. These results are in line with the *universalist view* (Jackendoff, 1990, 1996) reviewed in section 2.3.1.2 of the present study.

If we compare these results with those of the previous eye-tracking experiment, we can see that they point to the same direction. Graph 9 presents a comparative analysis of the average manner and path-looks across the two experiments: the pre-production eye-tracking experiment, which observes our subjects' eye-gaze patterns in a verbal description task, before description; and the non-verbal eye-tracking experiment, which analyzes the eye-movements of the same subjects in a purely non-linguistic task with articulatory suppression. The results do not demonstrate a language-specific pattern but rather a uniform pattern across languages, in which the speakers of the three languages attend significantly more to *manner of motion* (represented by the whole body of the agent for methodological purposes) than path of motion throughout the two experiments.

However, a Mixed ANOVA where *task* is taken as a two-level within-subject variable and *language group* as a between-subject variable, shows that there is a significant main effect of *task*: $F(1,58)=6.471$, $p<.05$, $\eta_p^2=.100$). This result suggests that the M-to-MP Ratios across languages significantly differ from one eye-tracking task to the other, which means that the manner ratios of all subjects are higher in the non-verbal eye-tracking experiment than the pre-production one.



Graph 9: A comparative analysis of the two eye-tracking experiments

Even though manner-ratios are still much higher than path-ratios in the pre-production task, a lowering in the ratio may indicate a relatively increasing focus on the path region. It is probable that it is the descriptive nature of the task that triggered this relative decrease of manner looks and relative increase of path looks. We have already argued that manner of motion is more salient than path of motion due to animacy and dynamicity effects. However, in the task where subjects are to verbalize

what they observe, there is also an increasing ratio of attention on the path component, which can be explained by the fact that subjects want to gather information on the path of motion as they will also need it in their descriptions. As discussed in section 5.1.1, almost all of the descriptions made by our subjects (regardless of their languages) include both the manner information and the path information. Therefore, the need to gather information about path of motion as well as the manner of motion may be a good explanation for the significant task effect observed.

To conclude, there is clearly a common pattern in the three non-verbal experiments, which does not reveal any language-specific effects that may either be related to Whorfian *linguistic relativity hypothesis* (Whorf, 1956) or to Slobin *thinking-for-speaking hypothesis* (Slobin, 1996a). The results provide support to the *universalist view* (Jackendoff, 1990, 1996), which suggests that conceptual event representation is not bound by the linguistic encoding of these events, even when there is a verbal medium included in the task. A final question to be raised here is: How are motion events, mentally represented in a uniform manner by speakers of different languages, expressed by different verbal expressions in those languages? In other words, where is the locus of the shift from language-universal to language-specific? As already discussed in section 4.1.5, this shift most probably lies in the lexicon. It is a matter of the lexical choices that the lexicons of speakers of different languages make available to their speakers. How different languages end up with such varied lexicons is a most intriguing question that immediately presents itself in this line of thinking, albeit well beyond the scope of the present dissertation.

CHAPTER VI

CONCLUSION

The present dissertation was composed of five main chapters. In the first introductory chapter, the aim and scope of the study, the methodology used, and the significance of the study were summarized. The second chapter presented a detailed review of the literature on the two main lines of investigation of the dissertation, namely the literature on motion events and the one on the language and cognition debate. The third chapter was where the general methodological procedure of the study was explained, and it was the fourth chapter which elaborated on the experiments used to find answers to the research questions introduced and on their results. It was in the fifth chapter where there was a thorough discussion of the results obtained from the whole study. The present chapter, the sixth and the last one, will present a recapitulation of the whole study and will conclude the dissertation with some suggestions for future research.

6.1. VERBAL PART OF THE STUDY

The overall goal of the study was to find answers to two broad and interrelated research questions. The first one was *whether native speakers of typologically different and similar languages produce and comprehend motion events in the same way*, and it was answered by the two verbal experiments, namely the Video Description Task and the Acceptability Judgment Task. The second broad question was *whether the language-specific motion event verbalization patterns have an influence on conceptual event representation of the speakers of those languages*, which was examined with the help of the three non-verbal experiments, *i.e.* the Similarity Judgment Task, the Non-Verbal Eye-Tracking Task and the Pre-Production Eye-Tracking Task.

The Video Description Task was a language production task, and it was the one where Talmy (1985)'s renowned two-way typology was questioned with motion event descriptions of Turkish, English and French native speakers elicited by real-life video clips. The results were clear: native speakers of Turkish and French used path sentences in almost all of their descriptions, which was a distinctive feature of V-languages; whereas native speakers of English used a remarkably high ratio of manner sentences in their descriptions, which was a particularity attributed to S-languages. The Acceptability Judgment Task, which is a language comprehension task, also had the goal of inquiring Talmyan verbalization typology, however this time from the reverse perspective by making the subjects rate the acceptabilities of a set of motion event sentences presented along with a set of motion event videos. The results of this comprehension task were in line with those of the production task, as well: native speakers of the V-languages (Turkish and French) gave significantly higher mean scores to path sentences, and native speakers of the S-language (English) gave considerably higher mean scores to manner sentences. The sole unexpected point in the verbal comprehension data was that the mean scores given by native speakers of Turkish to path sentences were significantly lower than those given by native speakers of French, which was explained by the morpho-syntactic flexibility of the Turkish language and the morphological productivity of the *-arak* suffix (see section 5.1.2), and which did not have any considerable effects on our overall results. Therefore, it can be concluded that the data obtained from both verbal experiments conducted within the framework of the present study experimentally verified the Talmyan motion event typology, as well as our hypotheses, with clear results.

6.2. NON-VERBAL PART OF THE STUDY

The non-verbal part of the study was actually dependent on affirmative answers from the experiments in the verbal part. In other words, what was actually questioned in the non-verbal section of the dissertation was whether the typological motion event verbalization patterns observed in the verbal experiments would also be reflected in the non-verbal experiments or not. There were two separate hypotheses investigated in the three tasks conducted within the framework of the non-verbal part of the present dissertation, namely the Whorfian *linguistic relativity hypothesis* (Whorf,

1956) and the Slobian *thinking-for-speaking hypothesis* (Slobin, 1996a). The goal of the categorization task and the non-verbal eye-tracking task coupled with articulatory suppression technique was to see whether the linguistic encoding preferences of a language have any effects on conceptual event representation of the native speakers of the language in question. The categorization task was based on the similarity judgments made by the subjects according to the sameness of manner or path information in video triads depicting motion events. The data obtained from the native speakers of the three languages, *i.e.* Turkish, English and French, did not reveal any language-specific results, all the subjects chose significantly more *same-manner alternates* than same-path alternates regardless of their languages. The manner-dominant choices were argued to be related to the animate and dynamic nature of the manner component compared to the static and inanimate nature of the path component. Although this experiment presented us valuable results regarding our research questions, another experiment was also performed to overcome the methodological shortcomings of the categorization task and to consolidate the results in hand. This second experiment made use of an online methodology, namely the eye-tracking methodology, to better investigate the hypothesis claiming the interdependence of linguistic representation and conceptual representation. For this aim, the eye-movement patterns of the three groups were recorded and analyzed while they were watching a set of real-life motion event depictions with no verbal intention. The *articulatory suppression technique* (Murray, 1967; Baddeley & Hitch, 1974), which has the goal of suppressing possible sub-vocal verbalizations, was also used to ensure the non-verbal nature of the task. The data obtained from the experiment suggested a uniform pattern of attention allocation and temporal linearization across languages, and all the subjects in the task attended more and first to that region that we defined as *manner of motion*³³. These results consolidated the results of the categorization experiment. Thus, it is possible to conclude that there were no relativistic effects observed in the data obtained from both tasks.

On the other hand, the pre-production eye-tracking task analyzed the eye-movement patterns of the speakers of the three languages comparatively during their utterance planning processes. The subjects' eye-movements were examined while they were

³³ We would like to once again note that there is an inevitable spatial overlap between the manner regions and the path regions determined for the analysis of the data.

presumably planning to describe the motion event scenes that they were watching. Therefore, the task included *communicative intention*, which was indicated as the pre-requisite of a possible *thinking-for-speaking* effect in the Slobin sense (Slobin, 1996a). However, the results did not reveal such an effect but rather pointed to a common direction with the previous two non-verbal tasks by demonstrating a uniform pattern of attention allocation and temporal linearization across languages. The data again presented manner-dominant results (insofar as manner is methodologically represented by the whole body of the agent) in all groups, and demonstrated that the canonical motion event expressions used in the languages that we observed were not reflected in the conceptual event representation patterns of the speakers of those languages, even when there was a communicative task to be completed. Therefore, we can conclude that the results of the three non-verbal tasks conducted within the framework of the present study are all in line with the *universalist view* (Jackendoff, 1990, 1996), which argues that conceptual event representation is not bound by the linguistic encodings of these events.

6.3. SUGGESTIONS FOR FURTHER RESEARCH

The present dissertation has undertaken the challenging work of inquiring both the linguistic expression and the conceptual representation of a certain domain, namely motion events, with a crosslinguistic perspective and with various experimental techniques. Apart from being the first study in the literature that investigates this particular group of languages, *i.e.* Turkish, English and French; it is also a pioneering study which elaborates on Turkish motion events by taking both the verbal and the non-verbal dimension into consideration, and by using real-life stimuli to elicit data. It is quite natural that it also has its own limitations. Here are some practical suggestions for the future work, which deals with the same or similar research questions:

- As it was beyond the scope of the current study, the elicited productions obtained from the Video Description Task have not been semantically and syntactically analyzed. A follow-up study may further explore the data, and make a thorough linguistic analysis out of them. Especially the analysis of the

Turkish data from a Distributed Spatial Semantic perspective (Sinha and Kuteva, 1995) would present valuable results.

- The results of the Acceptability Judgment Task were interpreted based on the language production data in hand, which gave us useful insights. Another strategy, next time, may be to have a follow-up interview with each subject after the task to question the reasons for the low and high ratings given to certain sentences.
- The Similarity Judgment Task was a non-verbal task, which has the aim of examining the conceptual event representation of the subjects. Next time, it may be a good idea to have an articulatory suppression component in that particular task, as well, to make sure that there is no sub-vocal verbal encoding involved.
- There are no clear-cut criteria in the literature regarding the determination of manner and path regions in an eye-tracking experiment. Thus, each researcher or research group determines their own regions. We have taken the space between the source and the goal as the path region, and the whole body of the agent as the manner region. Next time, it may be interesting to have a comparative analysis of different manner and path regions. For example, one can just take the lower part of the body of the agent (where most of the manner distinctions take place) as the manner region, and compare the results of that analysis with the one where the whole body is taken.
- There were no fixation points at the beginning of each video clip in our eye-tracking experiments. In the upcoming studies, it will be a practical idea to have them in order to eliminate random looks mostly occurring at the beginning.
- The last but not the least, the same experimental design may very well be used to test language learners in order to investigate possible linguistic and cognitive cross-language influences. In fact, the present author has already collected a large amount of data from second and third language learners

(from native speakers of Turkish learning English as a second language, and native speakers of Turkish learning English as a second language and French as a third language) by using the same five experiments; however as these were beyond the scope of the current study, they were not presented in the dissertation.

REFERENCES

- Allen, S., Özyürek, A., Kita, S., Brown, A., Furman, R., Ishizuka, T., et al. (2007). Language-specific and universal influences in children's syntactic packaging of manner and path: A comparison of English, Japanese, and Turkish. *Cognition*, 102(1), 16-48.
- Ameka, F. K., & Essegbey, J. (2001, March 22-25). *Serialising languages: satellite-framed, verb-framed or neither*. Paper presented at the 32nd Annual Conference on African Linguistics, University of California, Berkeley.
- Aske, J. (1989). Path predicates in English and Spanish: a closer look. In *Proceedings of the Fifteenth Annual Meeting of the Berkeley Linguistics Society* (pp. 1-14). Berkeley, CA: Berkeley Linguistics Society.
- Baddeley, A. D., & Hitch, G. J. (1974). Working memory. In G. H. Bower (Ed.), *The Psychology of Learning and Motivation Vol. 8* (pp. 47-89). New York: Academic Press.
- Beavers, J., Levin, B., & Tham, S. W. (2010). The typology of motion expressions revisited. *Journal of Linguistics*, 46, 331-377.
- Berman, R. A., & Slobin, D. I. (1994). *Relating Events in Narrative: A Crosslinguistic Developmental Study*. Hillsdale, NJ: Lawrence Erlbaum Associates.
- Boas, F. (1938). Language. In F. Boas (Ed.), *General Anthropology* (pp. 124-145). New York: Heath.
- Bohnemeyer, J., Eisenbeiss, S., & Narasimhan, B. (2006). Ways to go: Methodological considerations in Whorfian studies on motion events. *Essex Research Reports in Linguistics*, 50, 1-19.
- Brown, P., & Levinson, S. C. (1993). "Uphill" and "downhill" in Tzeltal. *Journal of Linguistic Anthropology*, 3, 46-74.
- Brown, R., & Lenneberg, E. H. (1954). A study in language and cognition. *Journal of Abnormal and Social Psychology*, 49, 454-462.

Bunger, A., Trueswell, J., & Papafragou, A. (2010). Seeing and saying: The relation between event apprehension and utterance formulation in children. In K. Franich, K. M. Iserman & L. L. Keil (Eds.), *Proceedings of the 34th Annual Boston University Conference on Language Development* (pp. 58-69). Somerville, MA: Cascadilla Press.

Cardini, F. E. (2010). Evidence against Whorfian effects in motion conceptualisation. *Journal of Pragmatics*, 42(5), 1442-1459.

Carroll, M., von Stutterheim, C., & Nüse, R. (2004). The language and thought debate: a psycholinguistic approach. In C. Habel & T. Pechmann (Eds.), *Approaches to Language Production* (pp. 183-218). Berlin: Mouton de Gruyter.

Choi, S., & Bowerman, M. (1991). Learning to express motion events in English and Korean: The influence of language-specific lexicalization patterns. *Cognition*, 41, 83-121.

Choi-Jonin, I., & Sarda, L. (2007). The expression of semantic components and the nature of ground entity in orientation motion verbs: A cross-linguistic account based on French and Korean. In M. Aurnague, M. Hickmann & L. Vieu (Eds.), *The Categorization of Spatial Entities in Language and Cognition* (pp. 123-149). Amsterdam: John Benjamins.

Choi-Jonin, I., & Sarda, L. (in press). Transitive motion verbs in French and in Korean. In H. Cuyckens, W. de Mulder, M. Goyens & T. Mortelmans (Eds.), *Variation and Change in Adpositions of Movement (Studies in Language Companion Series)*. Amsterdam: John Benjamins.

Chomsky, N. (1957). *Syntactic Structures*. The Hague: Mouton.

Cifuentes-Férez, P., & Gentner, D. (2006). Naming motion events in Spanish and English. *Cognitive Linguistics*, 17(4), 443-462.

Croft, W., Barddal, J., Hollmann, W., Sotirova, V., & Taoka, C. (2010). Revising Talmy's typological classification of complex events. In H. Boas (Ed.), *Contrastive Studies in Construction Grammar* (pp. 201-236). Amsterdam: John Benjamins.

Daintith, J. (Ed.). (2010). *The Oxford Dictionary of Physics*. New York: Oxford University Press.

Field, A. (2009). *Discovering Statistics Using SPSS*. London: Sage.

Finkbeiner, M., Nicol, J., Greth, D., & Nakamura, K. (2002). The role of language in memory for actions. *Journal of Psycholinguistic Research*, 31(5), 447-457.

Flecken, M. (2011). Event conceptualization by early Dutch-German bilinguals: Insights from linguistic and eye-tracking data. *Bilingualism: Language and Cognition*, 14(1), 61-77.

Fortis, J. M., & Vittrant, A. (2011, May 23-24). *Toward a typology of motion event constructions*. Paper presented at the Pre-AFLiCo “Trajectoire” Workshop, Lyon, France.

Gennari, S., Sloman, S., Malt, B., & Fitch, T. (2002). Motion events in language and cognition. *Cognition*, 83, 49-79.

Gentner, D., & Boroditsky, L. (2001). Individuation, relational relativity and early word learning. In M. Bowerman & S. C. Levinson (Eds.), *Language Acquisition and Conceptual Development* (pp. 215-256). Cambridge: Cambridge University Press.

Gentner, D., & Goldin-Meadow, S. (Eds.). (2003). *Language in Mind: Advances in the Study of Language and Thought*. Cambridge, MA: MIT Press.

Gentner, D., & Loewenstein, J. (2002). Relational language and relational thought. In E. Amsel & J. P. Byrnes (Eds.), *Language, Literacy and Cognitive Development: The Development and Consequences of Symbolic Communication* (pp. 87-120). Mahwah, NJ: Lawrence Erlbaum Associates.

Goddard, C. (1998). *Semantic Analysis: A Practical Introduction*. New York: Oxford University Press.

Griffin, Z. (2004). Why look? Reasons for eye movements related to language production. In J. Henderson & F. Ferreira (Eds.), *The Integration of Language, Vision, and Action: Eye Movements and The Visual World* (pp. 213-247). New York: Taylor & Francis.

Griffin, Z., & Bock, K. (2000). What the eyes say about speaking. *Psychological Science*, 11, 274-279.

Gullberg, M., Hendriks, H., & Hickmann, M. (2008). Learning to talk and gesture about motion in French. *First Language*, 28(2), 200-236.

Gumperz, J. J., & Levinson, S. C. (Eds.). (1996). *Rethinking Linguistic Relativity*. Cambridge: Cambridge University Press.

Hayward, G., & Tarr, M. (1995). Spatial language and spatial representation. *Cognition*, 55, 39-84.

Heider, E.R., & Olivier, D. (1972). The structure of the color space in naming and memory for two languages. *Cognitive Psychology*, 3, 337-354.

Hickmann, M., Taranne, P., & Bonnet, P. (2009). Motion in first language acquisition: Manner and path in French and in English. *Journal of Child Language*, 36 (4), 705-741.

Ibarretxe-Antuñano, I. (2004a). Language typologies in our language use: The case of Basque motion events in adult oral narratives. *Cognitive Linguistics*, 15(3), 317-349.

Ibarretxe-Antuñano, I. (2004b). Motion events in Basque narratives. In S. Strömquist & L. Verhoeven (Eds.), *Relating Events in Narrative: Typological and Contextual Perspectives* (pp. 89-111). Hillsdale, NJ: Lawrence Erlbaum Associates.

Ibarretxe- Antuñano, I. (2009). Path salience in motion events. In E. Lieven, S. Ervin-Tripp, J. Guo, N. Budwig, K. Nakamura & S. Özçalışkan (Eds.), *Crosslinguistic Approaches to the Psychology of Language: Research in the Tradition of Dan Isaac Slobin* (pp. 403-414). NY: Psychology Press.

Ibarretxe-Antuñano, I. (in press). Linguistic typology in motion events: Path and manner. *Anuario del seminario de Filología Vasca 'Julio de Urquijo' - International Journal of Basque linguistics and Philology*.

Ishibashi, M., & Kopecka, A. (2011, May 23-24). *The (a)symmetry of source and goal*. Paper presented at the Pre-AFLiCo "Trajectoire" Workshop, Lyon, France.

Jackendoff, R. (1990). *Semantic Structures*. Cambridge, MA: MIT Press.

Jackendoff, R. (1996). The architecture of the linguistic-spatial interface. In P. Bloom, M. A. Peterson, L. Nadel & M. F. Garrett (Eds.), *Language and Space* (pp. 1-30). Cambridge, MA: MIT Press.

Kopecka, A. (2004). *Étude typologique de l'expression de l'espace: localisation et déplacement en français et en polonais*. Unpublished doctoral dissertation, University Lumière Lyon 2, Lyon, France.

Kosman, L. A. (1969). Aristotle's definition of motion. *Phronesis*, 14 (1), 40-62.

Landau, B., & Jackendoff, R. (1993). "What" and "where" in spatial language and spatial cognition. *Behavioral and Brain Sciences*, 16, 217-238.

Langacker, R. W. (1987). *Foundations of Cognitive Grammar Volume I: Theoretical Prerequisites*. Stanford: Stanford University Press.

Levelt, W. (1989). *Speaking*. Cambridge, MA: MIT Press.

Levinson, S. C. (1997). From outer to inner space: linguistic categories and nonlinguistic thinking. In J. Nuyts & E. Pederson (Eds.), *Language and Conceptualization* (pp. 13-45). Cambridge: Cambridge University Press.

Li, P., & Gleitman, L. (2002). Turning the tables: Language and spatial reasoning. *Cognition*, 83, 265-94.

Lucy, J. (1992). *Grammatical Categories and Cognition: A Case Study of The Linguistic Relativity Hypothesis*. Cambridge: Cambridge Univ. Press.

Lucy, J. A. (1996). The scope of the linguistic relativity: An analysis and review of empirical research. In J. J. Gumperz & S. C. Levinson (Eds.), *Rethinking Linguistic Relativity* (pp. 37-69). Cambridge: Cambridge University Press.

Mayer, M. (1969). *Frog, Where Are You?* New York: Dial Press.

Miller, G. A., & Johnson-Laird, P. N. (1976). *Language and Perception*. Cambridge, MA: Belknap Press of Harvard University Press.

Montrul, S. (2001). Agentive verbs of manner of motion in Spanish and English as second languages. *SSLA*, 23, 171-206.

Munnich, E., & Landau, B. (2003). The effects of spatial language on spatial representation: Setting some boundaries. In D. Gentner & S. Goldin-Meadow (Eds.), *Language in Mind: Advances in the Study of Language and Thought* (pp. 113-155). Cambridge, MA: MIT Press.

Munnich, E., Landau, B., & Doshier, B. A. (2001). Spatial language and spatial representation: A cross-linguistic comparison. *Cognition*, 81, 171-207.

Murata, A., & Furukawa, N. (2005). Relationships among display features, eye movement characteristics, and reaction time in visual search. *Human Factors* 47, 598-612.

Murray, D. J. (1967). The role of speech responses on short-term memory. *Canadian Journal of Psychology*, 21, 263-276.

Nagy, W., & Gentner, D. (1990). Semantic constraints on lexical categories. *Language and Cognitive Processes*, 5(3), 169-201.

Naigles, L., & Terrazas, P. (1998). Motion-verb generalizations in English and Spanish: Influences of language and syntax. *Psychological Science*, 9, 363-369.

Naigles, L., Eisenberg, A., Kako, E., Hightner, M., & McGraw, N. (1998). Speaking of motion: Verb use in English and Spanish. *Language and Cognitive Processes*, 13, 521-549.

Oh, K. J. (2003). *Language, cognition, and development: Motion events in English and Korean*. Unpublished doctoral dissertation, Department of Psychology, University of California, Berkeley.

Özçalışkan, Ş., & Slobin, D. I. (1999). Learning how to search for the frog: Expression of manner of motion in English, Spanish, and Turkish. In A. Greenhill, H. Littlefield & C. Tano (Eds.), *Proceedings of the 23rd Annual Boston University Conference on Language Development* (pp. 541-552). Somerville, MA: Cascadilla Press.

Özçalışkan, Ş., & Slobin, D. I. (2000). Climb up vs. ascend climbing: Lexicalisation choices in expressing motion events with manner and path components. In S. C. Howell, S. A. Fish & T. Keith-Lucas (Eds.), *Proceedings of the 24th Annual Boston University Conference on Language Development: Vol. II* (pp.558-520). Somerville, MA: Cascadilla Press.

Özçalışkan, Ş., & Slobin, D. I. (2003). Codability effects on the expression of manner of motion in Turkish and English. In A. S. Özsoy, D. Akar, M. Nakipoğlu-Demiralp, E. Erguvanlı-Taylan & A. Aksu-Koç (Eds.), *Studies in Turkish Linguistics* (pp. 259-270). İstanbul: Boğaziçi University Press.

Özyürek, A., Kita, S., & Allen, S. (2001). *Tomato Man movies: Stimulus kit designed to elicit manner, path and causal constructions in motion events with regard to speech and gestures*. Nijmegen, The Netherlands: Max Planck Institute for Psycholinguistics, Language and Cognition Group.

Özyürek, A., & Özçalışkan, Ş. (2000). How do children learn to conflate manner and path in their speech and gestures? Differences in English and Turkish. In E. V. Clark (Ed.), *The Proceedings of the Thirtieth Annual Child Language Research Forum* (pp. 77-85). Stanford, CA: CSLI Publications.

Papafragou, A. (2007). Space and the language-cognition interface. In P. Carruthers, S. Laurence & S. Stich (Eds.), *The Innate Mind: Foundations and The Future* (pp. 272-290). Oxford: Oxford University Press.

Papafragou, A., Hulbert, J., & Trueswell, J. (2008). Does language guide event perception? Evidence from eye movements. *Cognition*, 108, 155-184.

Papafragou, A., Massey, C., & Gleitman, L. (2001). Motion events in language and cognition. In A. H.-J. Do, L. Domínguez & A. Johansen (Eds.), *Proceedings of the 25th Annual Boston University Conference on Language Development* (pp. 566-574). Somerville, MA: Cascadilla Press.

Papafragou, A., Massey, C., & Gleitman, L. (2002). Shake, rattle, 'n' roll: The representation of motion in language and cognition. *Cognition*, 84, 189-219.

Papafragou, A., Massey, C., & Gleitman, L. (2006). When English proposes what Greek presupposes: The cross-linguistic encoding of motion events. *Cognition*, 98, B75-B87.

Papafragou, A., Massey, C., & Gleitman, L. (2007). Motion event conflation and clause structure. In J. Cihlar (Ed.), *Proceedings of the 39th Annual Meeting of the Chicago Linguistics Society*. Chicago, IL: University of Chicago Press.

Papafragou, A., & Selimis, S. (2009). On the acquisition of motion verbs cross-linguistically. In M. Baltazani, G. K. Giannakis, T. Tsangalidis & G. J. Xydopoulos (Eds.), *Proceedings of the 8th International Conference on Greek Linguistics* (pp. 351-365). Ioannina: University of Ioannina.

Papafragou, A., & Selimis, S. (2010). Event categorisation and language: A cross-linguistic study of motion. *Language and Cognitive Processes*, 25(2), 224-260.

Pederson, E., Danziger, E., Wilkins, D., Levinson, S. C., Kita, S., & Senft, G. (1998). Semantic typology and spatial conceptualization. *Language*, 74(3), 557-589.

Pinker, S. (1994). *The Language Instinct*. New York: Morrow.

Pourcel, S. (2004). *What makes path of motion salient?* *Proceedings of the Berkeley Linguistics Society*, 30, 505-516.

Pourcel, S. (2005, October 23-25). *Linguistic Relativity in Cognitive Processes*. Paper presented at the 1st UK Cognitive Linguistics Conference: New Directions in Cognitive Linguistics, University of Sussex, England, UK.

Pourcel, S., & Kopecka, A. (2005). *Motion expression in French: typological diversity*. *Durham & Newcastle Working Papers in Linguistics*, 11, 139-153.

Pourcel, S., & Kopecka, A. (2006). *Motion events in French: Typological intricacies*. Unpublished ms., University of Sussex & Max Planck Institute for Psycholinguistics, Brighton, UK & Nijmegen, The Netherlands.

Rayner, K., & Castelhana, M. S. (2007). Eye movements during reading, scene perception, visual search, and while looking at print advertisements. In R. Pieters & M. Wedel (Eds.), *Visual Advertising*. Hillsdale, NJ: Erlbaum.

Richardson, D. C., & Spivey, M. J. (2004a). Eye tracking: Characteristics and methods. In G. Wnek & G. Bowlin (Eds.), *Encyclopedia of Biomaterials and Biomedical Engineering* (pp. 1028-1032). Marcel Dekker, Inc.

Roberson, D., Davies, I., & Davidoff, D. (2000). Color categories are not universal: Replication and new evidence from a stone-age culture. *Journal of Experimental Psychology: General*, 129, 369–398.

Rothkopf, C. A., Ballard, D. H., & Hayhoe, M. M. (2007). Task and context determine where you look. *Journal of Vision* 7 (14):16, 1–20.

Sablayrolles, P. (1995) The Semantics of Motion. In *Proceedings of the 7th Conference of the European Chapter of the Association for Computational Linguistics* (pp. 281-283). Dublin: EACL.

Selimis, E. (2007). *Linguistic coding of the concept of motion: Literal and metaphorical expressions in adult and child Greek*. Unpublished doctoral dissertation, University of Athens, Greece.

Selimis, S., & Katis, D. (2003). Reference to physical and abstract motion in child language: A cross-linguistic comparison between English and Greek. In G. Catsimali, E. Anagnostopoulou, A. Kalokerinos & I. Kappa (Eds.), *Proceedings of the 6th International Conference of Greek Linguistics*. Rethymno: Linguistics Lab.

Sinha, C., & Kuteva, T. (1995). Distributed spatial semantics. *Nordic Journal of Linguistics*, 18, 167–199.

Skordos, D. & Papafragou, A. (2010). Extracting paths and manners: Linguistic and conceptual biases in the acquisition of spatial language. In K. Franich, K. M. Iserman & L. L. Keil (Eds.), *Proceedings of the 34th Annual Boston University Conference on Language Development* (pp. 397-408). Somerville, MA: Cascadilla Press.

Slobin, D. I. (1996a). From 'thought and language' to 'thinking for speaking'. In J. Gumperz & S. C. Levinson (Eds.), *Rethinking Linguistic Relativity* (pp. 70-96). Cambridge: Cambridge University Press.

Slobin, D. I. (1996b). Two ways to travel: Verbs of motion in English and Spanish. In M. Shibatani & S. Thompson (Eds.), *Grammatical Constructions: Their Form and Meaning* (pp. 195–219). Oxford : Oxford University Press.

Slobin, D. I. (1997). Mind, code, and text. In J. Bybee, J. Haiman & S. A. Thompson (Eds.), *Essays on Language Function and Language Type: Dedicated to T. Givón* (pp. 437-467). Amsterdam/Philadelphia: John Benjamins.

Slobin (2003). Language and thought online: Cognitive consequences of linguistic relativity. In D. Gentner & S. Goldin-Meadow (Eds.), *Language in Mind: Advances in the Study of Language and Thought* (pp. 157-191). Cambridge, MA: MIT Press.

Slobin, D. I. (2004). The many ways to search for a frog: Linguistic typology and the expression of motion events. In S. Strömquist & L. Verhoeven (Eds.), *Relating Events in Narrative Vol. II: Typological and Contextual Perspectives* (pp. 219-257). Mahwah, NJ: Lawrence Erlbaum Associates.

Slobin, D. I. (2006). What makes manner of motion salient? Explorations in linguistic typology, discourse and cognition. In M. Hickmann & S. Robert (Eds.), *Space in Languages: Linguistic Systems and Cognitive Categories* (pp. 59-81). Amsterdam: John Benjamins.

Slobin, D. I., & Hoiting, N. (1994). Reference to movement in spoken and signed languages: Typological considerations. *Proceedings of the Berkeley Linguistics Society*, 20, 487-505.

Soroli, E. (2011, May 24-27). *Encoding and allocating attention to motion events in English, French, and Greek: Typological perspectives*. Paper presented at the 4th International Conference of the French Cognitive Linguistics Association, Lyon, France.

Soroli, E., & Hickmann, M. (2010). Spatial language and cognition in French and English: Some evidence from eye-movements. In G. Marotta, A. Lenci, L. Meini & F. Rovai (Eds.), *Proceedings of the Pisa International Conference "Space in Language"* (pp. 581-597). Florence: Edizioni ETS.

Talmy, L. (1985). Lexical Typologies. In T. Shopen (Ed.), *Language Typology and Syntactic Description Vol. III: Grammatical Categories and the Lexicon* (2nd ed., pp. 66-168). New York: Cambridge University Press.

Talmy, L. (1991). Path to realization: A typology of event conflation. *Proceedings of the Berkeley Linguistics Society*, 17, 480-519.

Talmy, L. (2000). *Toward a Cognitive Semantics Volume II: Typology and Process in Concept Structuring*. Cambridge, MA: MIT Press.

von Humboldt, W. ([1836] 1999). *On Language: On the Diversity of Human Language Construction and Its Influence on the Mental Development of the Human Species*. Losonsky, M. (Ed.), Heath, P. (Translation from German). Cambridge Texts in the History of Philosophy. Cambridge: Cambridge University Press.

von Stutterheim, C., & Nüse, R. (2003). Processes of conceptualisation in language production: Language-specific perspectives and event construal. *Linguistics (special issue: Perspectives in language production)*, 41, 851-881.

Whorf, B. L. (1956). *Language, Thought and Reality: Selected Writings of Benjamin Lee Whorf*. Carroll, J. B. (Ed.). Cambridge, MA: MIT Press.

Wilson, S., & Saygin, A. P. (2001). Adverbs and functional heads in Turkish: Linear order and scope. In L. Carmichael, C.-H. Huang & V. Samiian (Eds.), *Proceedings of WECOL 2001: Western Conference on Linguistics*. Fresno: California State University.

Zlatev, J., Blomberg, J., & David, C. (2010). Translocation, language and the categorization of experience. In V. Evans & P. Chilton (Eds.), *Language, Cognition, and Space: The State of the Art and New Directions* (pp. 389-418). London: Equinox Publishing.

Zlatev, J., & David, C. (2004, July 18-20). *Do Swedes and Frenchmen view motion differently?* Paper presented at the 1st Conference on Language, Culture and Mind, Portsmouth, UK.

Zlatev, J. & Yangklang, P. (2004). A third way to travel: The place of Thai in motion event typology. In S. Stromqvist & L. Verhoeven (Eds.), *Relating Events in Narrative: Cross-Linguistic and Cross-Contextual Perspectives* (pp. 159-190). Mahwah, NJ: Lawrence Erlbaum Associates.

APPENDIX A

INFORMED CONSENT FORM

This is a research conducted by Prof. Dr. Deniz Zeyrek and the PhD student, Ayşe Betül Toplu. The research aims to experimentally investigate Turkish, English and French language users' expression and conceptualization of certain events shown on video clips. The experiments we would like you to participate in are summarized below. All the experiments involve observing video clips and some are enriched by the eye-tracking device. By means of this technique, we will be able to understand where your eye-gaze focuses while you are watching the video clips and compare this data with your linguistic responses. You will be video-taped during the experiments, which will allow us to analyze your linguistic responses better.

EXPERIMENTS			APPROXIMATE TIME SPAN
Video Observation Task	nonverbal	will take place simultaneously	5 min.
Eye-tracking Experiment			
Similarity Judgment Task	nonverbal		10 min.
INTERMISSION			10 min.
Video Description Task	verbal	will take place simultaneously	10 min.
Eye-tracking Experiment			
Acceptability Judgment Task	verbal		10 min.
TOTAL TIME			45 min.

Participation in the experiments is totally voluntary and your responses will remain confidential. Your responses will not be classified as right or wrong; we will investigate them in an objective manner and solely for academic purposes. All your questions and queries will be answered when the experiments are completed.

Thank you in advance for participating in this experiment. Should you need more information about the experiments, do not hesitate to contact Deniz Zeyrek at Department of Foreign Language Education, METU at: dezeyrek@metu.edu.tr or Ayşe Betül Toplu at: aysebet@yahoo.com.

I realize that I am participating in this experiment voluntarily and I know that I can stop the experiments any time and leave. I accept that the information I will provide will be used in academic research. (Please sign below and return the form to the experimenter.)

Name:

Date: ---/---/----

Signature:

APPENDIX B

Background Questionnaire³⁴

Participant number:

Language group:

Name and Surname:

Age:

Education details:

Languages spoken:

Have you ever lived (more than 6 months) in another country other than your home country? Yes / No

If yes: Where, when and for how long?

³⁴ The experimenter asks the questions orally, and writes down subject responses herself. The subject never sees the questionnaire.

APPENDIX C

POST-PARTICIPATION INFORMATION SHEET

This is a research undertaken by Prof. Deniz Zeyrek and Ayşe Betül Toplu. The research is concerned with how motion events are verbalised and conceptualised in three languages, i.e. Turkish and French - typologically Verb-Framed languages, and English –typologically a Satellite-Framed language. Two verbal experiments (one having to do with language production, the other with language perception) were conducted to investigate verbalisation patterns. To investigate conceptualisation patterns, two nonverbal experiments were conducted. By analysing the data obtained in the experiments, the researchers will try to understand what the speakers of the respective languages pay attention to in motion events, namely, path or manner, and in what ways they resemble or differ from each other.

The data and the results of the experiments will be solely used for academic purposes. Should you want to learn the results of the research, please contact the researchers at their e-mail addresses below.

Thank you again for participating in the experiments.

Prof. Dr. Deniz Zeyrek (e-mail: dezeyrek@metu.edu.tr)

Ayşe Betül Toplu (e-mail: aysebet@yahoo.com)

APPENDIX D

VITA

PERSONAL INFORMATION

Surname, Name: Toplu, Ayşe Betül

Nationality: Turkish (TC)

Date and Place of Birth: 8 February 1981, Kayseri/TURKEY

Marital Status: Married

Phone: +90 505 8181100

E-mail: aysebet@yahoo.com

EDUCATION

Degree	Institution/Department	Year of Graduation
M.A.	University of Sorbonne Nouvelle, Paris/ Foreign Language Teaching	2005
B.A.	Bilkent University, Ankara/ Translation & Interpretation (Eng.-Fr.-Tur.)	2004

WORK EXPERIENCE

Year	Place	Enrollment
2007-2009	Department. of Western Lang. & Lit., Harran University, Şanlıurfa	Instructor
2006-2007	English Club Language School, Ankara	Teacher of English

LANGUAGES

Turkish: native

English: advanced

French: advanced

German: intermediate

Turkish Sign Language: beginner

ACADEMIC WORK

1. **Toplu, A. B., & Zeyrek, D. (2011c).** Linguistic and Non-Linguistic Investigation of Motion Events: An Experimental Study on French and English. In A. Botinis (Ed.), *Proceedings of the 4th ISCA Tutorial and Research Workshop on Experimental Linguistics* (pp. 27-30). University of Athens Press.
2. **Toplu, A. B., & Zeyrek, D. (2011b, May 24-27).** *Motion Events in Turkish & French: An Intra-Typological Investigation*. Paper presented at the 4th International Conference of the French Cognitive Linguistics Association (AFLiCo), University of Lyon II, Lyon, France.
3. **Toplu, A. B., & Zeyrek, D. (2011a, May 5-7).** *Devinim Olaylarının İfadesinde Eylem Çerçevesi Bir Dil Olarak Türkçe*. Paper presented at the 25th National Linguistics Conference, Çukurova University, Adana, Turkey.
4. **Toplu, A. B., & Zeyrek, D. (2010b, September 21-22).** *Inter- and Intra-Typological Investigation of Motion Events in Language Production and Comprehesnion: An Experimental Study on Turkish, English and French*. Paper presented at Cognitive Science Seminars: Psycholinguistics and Cognitive Science, Middle East Technical University, Ankara, Turkey.
5. **Toplu, A. B., & Zeyrek, D. (2010a, September 4-6).** *Linguistic Expression and Non-Linguistic Representation of Motion Events: An Experimental Study on Turkish and English*. Poster presented at the 16th Annual Conference on Architectures and Mechanisms for Language Processing, York, England, UK.
6. **Toplu, A. B. (2009, March 12-13).** *A Psycholinguistic Study on Motion Event Expressions in a Third Language Acquisition Situation*. Paper presented at the 2nd Mediterranean Graduate Students Meeting in Linguistics, Mersin University, Mersin, Turkey.
7. **(Baktır) Toplu, A. B. (2008).** Is the Wholist Analytic Style Ratio a Good Predictor of Students' Study Habits and Instructional Preferences? A case study on Turkish students learning English as a foreign language. *The Psychology of Education Review*, 32 (1), 38-44.
8. **Toplu, A. B. (2007, October 25-26).** *A Comparative Study on Two Discourse Markers in Turkish: Veya & Ya da*. Paper presented at the 1st

Mediterranean Graduate Students Meeting in Linguistics, Mersin University, Mersin, Turkey.

9. **Baktır, A. B. (2007, June 12-14).** *Is the Wholist- Analytic Style Ratio a Good Predictor of Students' Study Habits and Instructional Preferences? A case study on Turkish students learning English as a foreign language.* Paper presented at the 12th Annual Conference of European Learning Styles Information Network (ELSIN), Trinity College, Dublin, Ireland.
10. **Baktır, A. B. (2005).** *Le Style Cognitif Dans Les Recherches Sur L'Appropriation Des Langues Etrangères: Une étude empirique sur les étudiants turcs qui apprennent l'anglais comme langue étrangère.* Unpublished master's thesis, Department of Foreign Language Education, University of Sorbonne Nouvelle, Paris, France.

HONORS AND GRANTS

December 2009 - December 2010 : The Scientific and Technological Research Council of Turkey (TÜBİTAK), Project Assistant Fellowship.

October 2004 - October 2005 : Government of France, Full Scholarship for the Master's Degree.

June 2004 : Bilkent University, Ankara, Turkey , Graduated with high honors

September 1999 - June 2004 : Bilkent University, Ankara, Turkey, Full Scholarship for the Undergraduate Education.

RESEARCH PROJECTS

December 2009 - December 2010 : TÜBİTAK Short Term Project entitled « Conceptualization and Verbalization of Motion Events in Turkish, English and French : A Comparative Psycholinguistic Study », Project Assistant.

CURRENT RESEARCH INTERESTS

Crosslinguistic Investigation of Space and Motion

Language and Cognition

Psycholinguistic Experimentation

Eye-Tracking Methodology

APPENDIX E

TURKISH SUMMARY (TÜRKÇE ÖZET)

TÜRKÇE, İNGİLİZCE VE FRANSIZCA'DAKİ DEVİNİM OLAYLARININ SÖZSEL İFADESİ VE KAVRAMSAL TEMSİLİ: DENEYSEL BİR ÇALIŞMA

1. GİRİŞ

Devinim olaylarının dilbilimsel olarak incelenmesi 1980'lerde hız kazanan bir çalışma alanıdır. Bu konunun öncüleri arasında yer alan Leonard Talmy, 1985 yılında yazdığı “Söze Dökme Örüntüleri (*Lexicalization Patterns*)” adlı eserinde, bir devinim olayının temel öğelerini ortaya koyar. Talmy çalışmalarında bu öğelerden ikisini temel alır, bunlar *devinimin yolu* ve *devinimin tarzı* öğeleridir. Dünya dillerini de bu öğeleri ifade ediş biçimlerine göre iki gruba ayırır: Eylem-Çerçevesel Diller ve Uydu-Çerçevesel Diller.

Eylem-çerçevesel dillerde, devinimin yolu ögesi esas eylemin içinde verilirken, devinimin tarzı genellikle bir belirteç tümceciğiyle ifade edilir. Örneğin, “Bebek emekleyerek odadan çıktı” cümlesindeki *çıkma* fiili bu hareketin yolunu, yani yönünü, kendi içinde barındırmaktadır. Çıkma eyleminin tarzı, yani yapılaş şekli, ise bir belirteç olan *emekleyerek* sözcüğüyle ifade edilmektedir. Romans Dilleri (Fransızca, İspanyolca, İtalyanca gibi), Türk Dilleri (Türkçe gibi), Sami Dilleri (Arapça, İbranice gibi), Japonca ve Bask Dili eylem-çerçevesel diller grubunda yer almaktadır.

Uydu-çerçevesel dillerde ise, esas eylem devinimin tarzını içerirken, devinimin yolu çoğunlukla bir ilgeçle verilmektedir. Örneğin, “The baby crawled out of the room (*Birebir çevirisi: Bebek odanın dışına emekledi*)” cümlesindeki *crawl* eylemi devinimin ne şekilde yapıldığını, yani tarzını da içermektedir. Devinimin yolu ise,

dışarıya anlamına gelen *out of* ilgeciyle ifade edilmektedir. Cermen Dilleri (İngilizce, Almanca gibi), Slav Dilleri (Rusça, Polonyaca gibi), Ural Dilleri'nin Fin-Ugur koluna mensup diller (Macarca, Fince gibi) ve Çin-Tibet Dilleri (Mandarin gibi) uydu-çerçeve dillere örnektir (Slobin, 2003).

Bu çalışma kapsamında, eylem-çerçeve diller olarak kabul edilen Türkçe ve Fransızca ile uydu-çerçeve bir dil olarak sınıflandırılan İngilizce incelenmektedir. Çalışmaya anadili Türkçe olan 32 kişi, anadili İngilizce olan 23 kişi ve anadili Fransızca olan 24 kişi gönüllü olarak katılmıştır. Katılımcıların yaş aralığı 18-35'dir ve her bir katılımcı toplam 45-50 dakika süren beş adet ardıl deneye katılmışlardır. Çalışmanın iki temel amacı bulunmaktadır. Bunlardan ilki, devinim olaylarının bu üç dilin konuşurları tarafından farklı şekilde ifade edilip edilmediğini hem dil üretimi, hem de dili anlama boyutlarıyla incelemektir. Bunu test etmek için birbirini tamamlayan iki adet sözlü deney uygulanmıştır ve sonucunda aynı tipolojik gruba mensup Türkçe ve Fransızca'da devinim olaylarının aynı şekilde ifade edileceği, karşıt tipolojik gruba ait bir dil olan İngilizce'de durumun farklı olacağı öngörülmüştür. İkincisiyse, bu farklılığın yalnızca ifade boyutunda geçerli olup olmadığını çeşitli deneysel yöntemlerle sorgulayarak; bu olayların zihinde oluşturulması ve düzenlenmesi, yani kavramsallaştırması sürecinde de bu dillerin konuşurları arasında benzerlik ve farklılıklar gözlenip gözlenmediğini, diğer bir deyişle düşünsel boyutun evrensel olup olmadığını anlamaya çalışmaktır. Bu konu, dil felsefesi literatüründe önemli yer tutan *Dilsel Görecelik Varsayımı* (Whorf, 1956) ile birebir ilintilidir. Bu amaçla da, biri sınıflandırma, diğer ikisiyse göz-izleme tekniğini kullanan üç adet sözsüz bilişsel deney uygulanmıştır. Herhangi bir dil etkileşimini önlemek amacıyla, katılımcılar önce sözsüz deneylere ardından sözlü deneylere katılmışlardır.

Elde edilen sonuçlar göstermiştir ki, dilsel bir amaç içeren deneylerde katılımcılar, aynen öngördüğümüz gibi, dillerinin mensup olduğu tipolojik gruba uygun hareket etmişlerdir. Yani Türkçe ve Fransızca dillerinin konuşurları gerek dil üretimi deneyinde gerekse dili anlama deneyinde eylem-çerçeve dillere uygun sonuçlar verirken, anadili İngilizce olan katılımcılar uydu-çerçeve dil özelliklerine sıkı sıkıya bağlı kalmışlardır. Öte yandan, sözsüz bilişsel deneylerde üç grubun katılımcılarının performansları arasında anlamlı herhangi bir farklılık tespit edilememiştir. Diğer bir

deyişle, katılımcılarımız devinim olaylarını kavramsallaştırırken dillerinden etkilenmemişlerdir. Bu sonuç, kavramsal temsille dilsel temsilin birbirinden bağımsız olduğu görüşüne dayanan *Evrensel Yaklaşım* (Jackendoff, 1990, 1996)’ı desteklemektedir.

2. LİTERATÜR ÖZETİ

Bu çalışmanın arka planını oluşturan iki temel alan bulunmaktadır. Bunlardan ilki devinim olaylarının incelemesi; ikincisi ise, dil ve düşünce bağıntısını inceleyen dilsel görecelik/evrensellik varsayımlarıdır. Dolayısıyla, bu kısımda sırayla her iki alanın yazınına da kısaca değinilecektir.

2.1. Devinim Olayları Literatürü

Devinim eylemlerinin dilbilimsel açıdan sistematik değerlendirmesini, ilk olarak 1985 yılında Leonard Talmy yapmıştır. Talmy bu çalışmasında, devinim olayının en temel iki ögesinin *devinimin tarzı* ve *devinimin yolu* olduğunu söylemekte ve dünya dillerini bu ögeleri işleme biçimlerine göre iki gruba ayırmaktadır. Talmy tipolojisinde *devinimin yolu* ögesini temel aldığı için, sınıflandırma temelde dillerin bu ögeyi nasıl dile getirdiklerine dayanmaktadır. Dünya dilleri, devinimin yolunu ifade şekillerine göre iki gruba ayrılmaktadır: Eylem-Çerçevesel Diller ve Uydu-Çerçevesel Diller.

2.1.1. Eylem-Çerçevesel Diller: Romans Dilleri, Sami Dilleri, Türk Dilleri, Japonca ve Bask Dili’nin de aralarında bulunduğu E-Çerçevesel Diller’de devinimin yolu esas eylemin içinde gizlidir; eylem bu ögeyi bünyesinde barındırır. Diğer bir deyişle, devinim olayı eylemin etrafında çerçevelenmiştir. E-Çerçevesel Diller’de devinimin tarzı ifade edilecek olursa esas eylemle değil, eylem dışı bir ögeyle ifade edilir. (1) ve (2) numaralı Türkçe örneklerdeki *çıkmaq* ve *inmek* eylemleri bünyelerinde sırasıyla “dışarı” ve “aşağı” ögelerini, yani devinimin yolunu barındırmaktadırlar. Devinimin tarzıysa çoğunlukla bir belirteç tümcecisiyle ifade edilmektedir (örn: *sendeleyerek* ya da *koşarak*). Bu diller aynı zamanda “devinim+yol” dilleri olarak da adlandırılmaktadırlar.

(1) Çocuk *sendeleyerek* binadan çıktı.

(2) Adam *koşarak* merdivenlerden indi.

2.1.2. **Uydu-Çerçevesel Diller:** Hint-Avrupa Dilleri'nin birçoğunun (Romans Dilleri hariç) ve Çince gibi bazı dillerin aralarında bulunduğu U-Çerçevesel Diller'in en önemli özelliği devinimin yolunu esas eylemle değil de, uydu adı verilen ikincil bir ögeyle vermeleridir. Bu dillerde devinimin tarzı ögesi genellikle eylemle beraber verilir. Örneğin, U-Çerçevesel bir dil olarak nitelendirilen İngilizce'de çok sayıda devinim tarzı ifade eden eylem bulunmaktadır. (3) ve (4) numaralı örneklerdeki *stagger* ve *run* bunlardan yalnızca ikisidir. Örneklerde de görüldüğü gibi, devinimin yolu birer ilgeçle belirtilmiştir. Bu dilleri “devinim+tarz” dilleri olarak adlandırmak da mümkündür.

(3) The child *staggered out of* the building.

‘Çocuk sendeleyerek binadan çıktı (*Birebir çevirisi: Çocuk binanın dışına sendeledi*)’.

(4) The man *ran down* the stairs.

‘Adam koşarak merdivenlerden indi (*Birebir çevirisi: Adam merdivenlerin aşağısına koştu*)’.

İnceleyeceğimiz son dil olan Fransızca; Türkçe ve diğer Romans Dilleri gibi E-çerçevesel diller grubunda yer almaktadır. Dolayısıyla, bu dilde de devinimin yolu esas eylemin içinde verilmektedir (bkz. (5) ve (6) numaralı cümleler). Aşağıdaki cümlelerin Türkçe çevirileri birebir çeviriler olup, Türkçe’de de aynı şekilde kullanılan ifade biçimleridir.

(5) L’enfant *est sorti* du bâtiment *en titubant*.

‘Çocuk sendeleyerek binadan çıktı’.

(6) L’homme *est descendu* les escaliers *en courant*.

‘Adam koşarak merdivenlerden indi’.

2.2. Dil ve Düşünce Literatürü

Bu çalışmada, devinim olayları yalnızca dilsel ifade boyutunda değil, aynı zamanda düşünsel boyutta da ele alınmakta ve düşüncenin evrenselliği sorusuna değinilmektedir. “Gereç ve Yöntem” kısmında ayrıntılandırılacak olan sözsüz deneylerimizin amacı da, işte bu düşünsel boyutu, yani kavramsallaştırma boyutunu inceleyebilmektir. Bu nedenle, bu konuda yapılmış olan çalışmalardan da kısaca bahsetmek yerinde olacaktır.

2.1.2.1. Dil ve Düşünce Bağntısı: Dil ve düşünce arasındaki ilişki, yani düşüncenin mi dili oluşturduğu, yoksa dilin mi düşüncüyü şekillendirdiği sorusu, yüzyıllardır araştırmacıların ilgi odağı olmuştur. Son yıllarda yeniden popüler hale gelen bu soruya yanıt arayan iki temel yaklaşım bulunmaktadır:

- A. Evrensel Yaklaşım: Temelde Jackendoff’un (1990, 1996) çalışmalarıyla şekillenmiş olan bu görüşe göre, kavramsal temsil (*conceptual representation*) evrenseldir, yani dillerarası farklılık göstermez. Dolayısıyla, karşılaştırmalı dil çalışmalarıyla ortaya konan farklılıklar düşünsel boyuttaki farklılığı değil, tamamen dilsel boyuttaki farklılığı yansıtmaktadır. Diğer bir deyişle, dilin düşünce üzerinde şekillendirici bir etkisi bulunmamaktadır.
- B. Dil-Temelli Yaklaşım: Bu yaklaşımın temelinde, kavramsal temsilin evrensel olmadığı savı bulunmaktadır. Bu yaklaşımı temsil eden araştırmacılar, dille düşünce arasındaki ilişkinin şeklini ve yoğunluğunu değerlendirmeleri bakımından iki alt grupta ele alınabilirler. *Güçlü Dil-Temelli Yaklaşım* olarak adlandırabileceğimiz ilk görüş, dilsel görecelik savını benimsemekte ve dilsel örüntülerin insanın düşünme kalıpları üzerinde belirgin bir etkisi olduğunu savunmaktadır (bkz. Sapir-Whorf Varsayımı, Whorf 1956). *Zayıf Dil-Temelli Yaklaşım* ise, dilin düşünce üzerindeki etkilerini özel durumlarla sınırlandırmaktadır. Örneğin, Slobin (1996a) ‘konuşmak-için-düşünmek’ varsayımında, dilsel amaçla (konuşmak ya da resim betimlemek gibi) yapılan düşünme eylemlerindeki kavramsallaştırmaların dillerarası farklılıklar gösterebileceğini ifade etmektedir. İnsanların yalnızca dil kullanımının yapmakta oldukları işi kolaylaştıracağı düşünsel eylemlerde kullandıkları

dilden etkilenebileceklerini savunan ‘strateji-olarak-dil’ yaklaşımı da bu ikinci grupta sayılabilir (Gennari ve diğ., 2002).

2.1.2.2. Devinim Olaylarının İncelenmesinde Dil ve Düşünce İlişkisi: Devinim olaylarının bu genel tartışmayla bağlantısı üzerine çalışma yapan belli başlı araştırmacılar ve yaptıkları deneysel çalışmaların sonuçları şu şekildedir:

Gennari, Sloman, Malt ve Fitch (2002): Bu çalışmada araştırmacılar, devinim olaylarının İngilizce ve Fransızca’daki farklı ifadelerinin o dilin konuşurlarının dilsel-olmayan deneylerdeki performanslarını etkileyip etkilemediğini incelerler. Bunun için biri hafıza, diğeri de benzerlik değerlendirmesi olmak üzere iki bilişsel test kullanırlar. Deneylerin sonunda, dilsel performansla dilsel-olmayan performansın birçok durumda birbirinden bağımsız olduğunu, fakat bazı özel durumlarda etkileşebildiklerini bulurlar.

Papafragou, Massey ve Gleitman (2002, 2006): Bu iki çalışmada Papafragou ve ekibi, İngilizce ve Yunanca gibi tipolojik olarak farklı iki dilde devinim eylemi kavramsallaştırmalarını incelerler. Ekip, bu iki dilde devinim olaylarının hem dilsel ifadelerini hem de dilsel-olmayan temsillerini deneysel olarak ele alırlar. Yapılan sözlü ve sözsüz deneyler sonucunda, İngilizce anadil konuşurlarıyla Yunanca anadil konuşurlarının devinim ifadelerinin kesin olarak farklılık göstermesine karşın, dilsel-olmayan deneylerdeki performanslarının farklı olmadığını tespit ederler. Ekip bulgularıyla, dil-düşünce ilişkisinde evrensel yaklaşımı desteklemiş olurlar.

von Stutterheim ve Nüse (2003) ; Carroll, von Stutterheim ve Nüse (2004): Bu çalışmalarında yazarlar, Levelt (1989)’in dil üretim modelinin ihmal edilen ögesi olan “kavramsallaştırıcı”yı (*conceptualizer*) ve onun ne kadar dil-temelli olduğunu değerlendirirler. Bu amaçla, anadili Almanca ve İngilizce olan iki gruba *Quest* (Stellmach, 1997) adlı kısa bir animasyon film izletirler ve bu kişilerin devinim olgusunu zihinlerinde ne şekilde düzenlediklerini incelerler. Çalışmanın sonucunda, iki dil konuşurlarının aynı eylemleri zihinlerinde farklı şekillerde düzenlediklerini ve dolayısıyla kavramsallaştırma aşamasının tamamen de evrensel niteliklere sahip olmadığını ortaya koyarlar.

Papafragou, Hulbert ve Trueswell (2008): Bu çalışmada araştırmacılar, farklı tipolojik gruplara ait olan İngilizce ve Yunanca'daki devinim olaylarının kavramsallaştırılmalarını göz izleme tekniği kullanarak incelerler. Bu amaçla iki deney gerçekleştirirler. Birinde katılımcılar sözlü tasvirler yapmak amacıyla, diğerinde ise bir hafıza testine katılmak amacıyla devinim olayı görüntülerini izlerler. Çalışmanın sonucunda; dil amaçlı deneyde Amerikalı ve Yunan katılımcıların performansları arasında fark gözlemediği, bu farkın dil amacı gütmeyen deneyde gözlenmediği ortaya konur.

2.3. Temel Araştırma Soruları ve Hipotezler

- Temel Araştırma Sorusu I: Anadilleri Türkçe, İngilizce ve Fransızca olan yetişkinlerin aynı devinim olayını **ifade etme şekilleri** birbirinden farklı mıdır?

Hipotez I: U-çerçeveveli bir dil olarak nitelendirilen İngilizce ile E-çerçeveveli diller olarak nitelendirilen Türkçe ve Fransızca'nın devinim olayı ifade biçimleri arasında kesin farklılıklar beklenmektedir. İngilizce'de devinimin tarzının esas fiille, devinim yolunsa ayrı bir ögeyle (uydu) ifade edildiği örüntü daha baskınken; Türkçe ve Fransızca'da devinim yolunun esas fiille verildiği ve devinim tarzının bir çeşit uyduyla ifade edildiği örüntünün çok daha baskın olacağı öngörülmektedir. Fakat aynı tipolojik grupta olmalarına rağmen, Türkçe ve Fransızca ifadelerin de bazı farklılıklar göstereceği düşünülmektedir.

- Temel Araştırma Sorusu II: Anadilleri Türkçe, İngilizce ve Fransızca olan yetişkinlerin aynı devinim olayını **kavramsallaştırma örüntüleri** birbirinden farklı mıdır?

Hipotez II: Düşüncenin evrenselliği ilkesinden hareketle, bu dillerin anadil konuşurlarının kavramsallaştırma örüntüleri arasında belirgin bir farklılık gözlenmeyeceği düşünülmektedir.

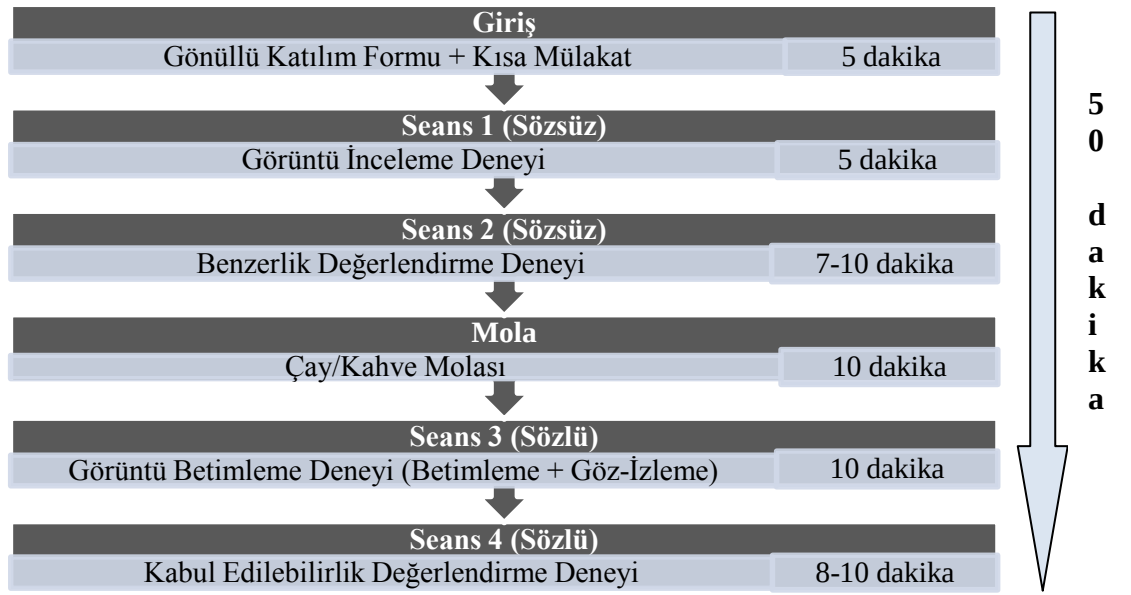
3. GEREÇ VE YÖNTEM

Alanyazında benzer parametreler kullanılarak yapılan çalışmalarda çoğunlukla, devinim olaylarını gösteren çeşitli resimler (örn. Özçalışkan ve Slobin, 2000; Papafragou ve diğ. 2002, 2006) ya da animasyonlar/kısa çizgi filmler (örn. Allen ve diğ. 2007; Gullberg ve diğ. 2008) kullanılmaktadır. Bu çalışmada ise, hem devinim unsurunu hem de psikolojik gerçeklik unsurunu daha iyi yansıttığı gerekçesiyle gerçek videolar kullanılmasına karar verilmiştir. Bu amaçla, toplam 12 fiilin (*koşmak, sendelemek, emeklemek, topallamak, dans etmek, tek ayak üzerinde zıplamak, iki ayak üzerinde zıplamak, zigzag çizerek yürümek, parmaklarının ucuna basarak yürümek, asker gibi rap rap yürümek, dans etmek ve ayaklarını sürüyerek yürümek*) beş farklı devinim yönüyle (*içeri, dışarı, aşağı, yukarı ve karşıdan karşıya*) harekete dökülmesiyle, 60 farklı devinim olayının görselleştirildiği bir video bankası oluşturulmuştur. Videoların tektip olmaları amacıyla da, bütün çekimlerde aynı oyuncu aynı kıyafetlerle ve üç sabit arka planla görüntülenmiştir. Çalışmanın büyük bir kısmı ODTÜ İnsan-Bilgisayar Etkileşimi Araştırma ve Uygulama Laboratuvarı'nda, kalan kısmıysa Paris 8 Üniversitesi'ne bağlı 'Structures Formelles du Langage (SFL)' adlı laboratuvarında gerçekleştirilmiştir. Çalışma her bir katılımcı için 45-50 dakika sürmüştür.

Çalışmaya, 18-35 yaş arası toplam 79 kişi gönüllü olarak katılmıştır. Katılımcıların 32'sinin anadili Türkçe, 23'ünün anadili İngilizce ve 24'ünün anadili Fransızca'dır. Deney takvimi, her bir katılımcı için, birer saatlik dilimler halinde düzenlenmiştir. Her deneyin başında, o bölümün içeriğini ve katılımcıdan neler beklendiğini anlatan bir açıklama sayfası yer almaktadır. Fakat yine de her bir katılımcıya, her deney öncesi (anadillerinde) ayrıntılı sözlü açıklamalar da sunulmuştur. Deneyin yapıldığı ana laboratuvar (ODTÜ İBE Laboratuvarı), katılımcıların deney materyallerini izledikleri bir ana bilgisayar bulunmaktadır. Yapılan beş deneyden ikisinde kullanılmakta olan göz-izleme sistemi de bu ana bilgisayar ekranı içinde yer almaktadır. Laboratuvarda aynı zamanda 24 saat kayıt yapan iki kamera ve bir ses kayıt sistemi bulunmaktadır. Deney odasının bitişiğinde yer alan kontrol odasında da, deneycinin deneyin gidişatını takip edebilmesi ve herhangi bir sorun anında müdahale edebilmesi için bir kontrol

bilgisayarı yer almaktadır³⁵. Katılımcılar çalışmaya teker teker alınmışlardır. Her bir katılımcı deney odasına alındıktan sonra, kendisine ilk olarak Gönüllü Katılım Formu sunulmuş ve okuyup imzalaması rica edilmiştir. Ardından birinci deneyin işleyişi hakkında kısaca bilgi verilip, katılımcının dikkatini dağıtmamak için kendisi odada tek başına bırakılmıştır.

Çalışma dört seanstan ve beş farklı deneyden oluşmaktadır. Bu deneyler, uygulama sıraları ve yaklaşık süreleri aşağıdaki tabloda gösterilmektedir.



3.1. Görüntü İnceleme Deneyi (Sözsüz Göz-İzleme Deneyi)

Bu deney, 5 dakika süren sözsüz bir deneydir. Bu kısımda katılımcılardan ekranda otomatik olarak birbiri ardına belirecek olan 25 adet kısa videoyu izlemeleri istenmiştir. Bu videoları izlerken, bir yandan da sesli ve devamlı olarak 1-2-3-1-2-3-1-2-3 ... şeklinde saymaları gerekmektedir. Söyleyiş Baskılaması Tekniği (*Articulatory Suppression Technique*) olarak adlandırılan bu yöntemin kullanım amacı, katılımcıların videoları izlerken burada geçen

³⁵ Çalışmada yer alan 20 Fransız katılımcının test edildiği SFL Laboratuvarı'nda, ana laboratuvardan farklı olarak, taşınabilir bir göz-izleme cihazı kullanılmış ve kayıtlar video kaydı olarak değil ses kaydı olarak alınmıştır.

eylemleri içlerinden dahi olsa sözcüklere dökmelerini engellemek ve dolayısıyla çalışmanın “sözsüz/bilişsel” doğasını korumaktır (Murray, 1967; Baddeley ve Hitch, 1974). Bu deneyin amacı, katılımcıların bir devinim olayını gözlemlerken ögeleri hangi sırayla izlediklerini ve hangi ögelere (devinimin yolu ya da devinimin tarzı) daha fazla dikkat ettiklerini tespit etmektir. Böylece, farklı anadil gruplarına mensup insanların aynı devinim olaylarını farklı ifade etmelerinin ötesinde, farklı bir biçimde kavramsallaştırıp kavramsallaştırmadıkları incelenecektir.

3.2. Benzerlik Değerlendirme Deneyi

İkinci deney de ilki gibi sözsüz bir deneydir ve 8-10 dakika sürmektedir. Katılımcılar 10’u ana video, 20’si aday video olmak üzere toplam 30 adet kısa video görüntüsü izlemişlerdir. Bu kısımdaki videolar üçlü setler halinde ekrana gelmiştir. Her setin ilk videosu tam ekran olarak gösterilmiş ve bu o setin *ana videosu* olarak adlandırılmıştır. Ardından ikinci ve üçüncü videolar (*aday videolar*) yan yana, küçük ekran olarak ve ardıl bir biçimde sunulmuştur. Soldaki aday video karardıktan sonra, sağdaki aday video başlatılmıştır. Aday görüntülerden biri ana görüntüyle aynı devinim yolunu, diğeri ise aynı devinim tarzını paylaşmaktadır. Katılımcıya aday videolardan hangisinin (1 ya da 2) ana videoya daha benzer olduğu sorulmuş ve seçimini adayların ikisini de izledikten sonra elindeki fareyle uygun videoya tıklayarak belirtmesi istenmiştir. Süreçte anlaşılmayan bir nokta kalmadığına emin olmak için, bu bölümün başına ufak bir deneme seansı konmuştur. Deneme seansında katılımcılara sunulan videolar çalışmaya dahil değildir. Katılımcıların yanıtları sistem tarafından kayıt altına alınmaktadır. İkinci deneyden sonra kısa bir ara verilmiş ve arada katılımcılara çay ya da kahve ikram edilmiştir.

Katılımcının bu deneydeki seçimlerinin, devinim yoluna mı yoksa devinim tarzına mı daha çok dikkat ettiğini/önem verdiğini göstermesi beklenmektedir. Bilindiği gibi, U-çerçevesel dillerde ve E-çerçevesel dillerde bu iki öge farklı şekillerde ifade edilmektedir ve bu farkın ifade öncesi kavramsallaştırma aşamasında da görülüp görülmediği cevap bekleyen bir sorudur.

3.3. Görüntü Betimleme Deneyi & Betimleme Öncesi Göz-İzleme Deneyi

Üçüncü deney sözlü bir deney olup, 10 dakika sürmektedir. Katılımcılardan yine otomatik olarak birbiri ardına ekrana getirilen 25 adet kısa videoyu izlemeleri ve her videodan sonra videodaki kişinin ne yaptığını sözlü olarak söylemeleri istenmiştir. Katılımcılara her video bitiminde, başlangıç ve bitişi birer sinyal sesiyle belirtilen, 10 saniyelik bir tanımlama süresi verilmiştir. Süre bitiminde otomatik olarak bir sonraki video ekrana gelmektedir. Katılımcıların süreçte alışmalarını sağlamak için bu deneyin başına da bir deneme seansı konmuştur. Bu deney esnasında da göz izleme sistemi devrededir. Ayrıca katılımcıların yanıtları deney odasındaki kameralar ve ses kayıt sistemi tarafından kayıt altına alınmaktadır.

Bu deneyde, farklı anadillere sahip kişilerin aynı devinim olayını farklı şekillerde dillendirip dillendirmedikleri incelenmektedir. Öte yandan, bu deneyde de göz-izleme tekniği kullanılarak, sözlü ve sözsüz performanstaki göz hareketi benzerlik ve farklılıklarını inceleme şansı yakalanabilecektir. Bu seanstaki iki deneyin (betimleme ve göz-izleme) sonuçları ayrı ayrı değerlendirilecek ve yorumlanacaktır.

3.4. Kabul Edilebilirlik Değerlendirme Deneyi

Son deney de sözlü bir deneydir ve 8-10 dakika sürmektedir. Üçüncü deneyi tamamlama amacıyla olan bu dördüncü deney, bir anlama deneyidir. Bu deneydeki 30 adet videonun her biri birer cümleyle birlikte katılımcıya sunulmuş ve katılımcıdan bu cümleleri değerlendirmesi istenmiştir. Süreç şöyle işlemiştir: her bir videonun altında, o videodaki kişinin ne yapmakta olduğunu tanımladığı iddia edilen bir cümle belirlemek ve bu cümle video süresince ekranda kalmaktadır. Videonun bitiminde ekranda bir pencere belirlemek ve pencerenin üst kısmında “Böyle bir durumda ben de bu cümleyi kullanırdım” şeklinde bir yargı verilmektedir. Katılımcılardan bu yargıya ne ölçüde katılıp/katılmadıklarını kendilerine sunulan beş seçenekten birine³⁶ tıklayarak belirtmeleri istenmektedir.

³⁶ Kesinlikle katılıyorum, katılıyorum, emin değilim, katılmıyorum ya da kesinlikle katılmıyorum.

Bu bölümde sunulan 30 cümleden 10 tanesi o dildeki doğru kullanımı (*örn.* Adam topallayarak binadan çıktı), 10 tanesi karşılaştırılan diğer dildeki kullanımı (*örn.* Adam binanın dışına topalladı – İngilizce’den birebir çeviri-) ve 10 tanesi de bambaşka bir yapıyı (*örn.* Adam topalladı ve binadan çıktı) ifade etmektedir. Deney yine bir deneme seansıya başlamaktadır.

Bu deneyin amacı, farklı dillerde farklı şekillerde ifade edilen devinim olaylarının anlama boyutunda da bir farklılık gösterip göstermediğini tespit etmektir.

Çalışmanın en sonunda her bir katılımcıya çalışmanın amacını açıklayan bir Katılım Sonrası Bilgi Formu sunulmuştur. Bu formun en altında, katılımcıların daha fazla bilgi almak istemeleri ya da çalışmanın sonucu hakkında bilgi sahibi olmak istemeleri halinde araştırmacılara ulaşabilecekleri e-posta adresleri yer almaktadır. Ayrıca, çalışma bitiminde, her bir katılımcıya ufak bir teşekkür armağanı sunulmuştur.

4. BULGULAR

Bu bölümde öncelikle, yukarıda ayrıntılandırılan beş deneyden elde edilen verilerin ne şekilde düzenlendikleri ve kodlandıkları açıklanacaktır. Ardından, elde edilen sonuçlar özetlenerek, bir sonraki “Tartışma ve Sonuç” bölümü için uygun altyapı oluşturulacaktır. Deneylerin bu bölümdeki sıralaması uygulama sıralamasından farklı olacaktır. Bunun nedeni, araştırma soruları ve elde edilen sonuçlar arasındaki bağıntıları okuyucu için daha anlaşılır kılmaktır.

4.1. Görüntü Betimleme Deneyi

Bu deneyden iki farklı tip bulgu elde edilmiştir. Bunlardan ilki, katılımcıların izledikleri videolardaki devinim olaylarını sözlü olarak tanımlamalarını içeren dil üretimi verisidir. Üç dilin konuşurlarına aynı devinim olayı görüntüleri izletilerek, onlardan bu görüntüdeki olayları kendi dillerinde tanımlamaları istenmişti. Böylece, Talmy (1985)’nin tipolojik sınıflandırması da test edilmiş olacaktı. İkinci veri tipiye,

göz hareketi verileridir ve bu veriler ‘4.5. Betimleme Öncesi Göz-İzleme Deneyi’ başlığı altında aşağıda değerlendirilecektir.

Bu deneyden elde edilen dil üretimi verilerini değerlendirebilmek için, ilk olarak katılımcıların ses kayıtları dinlenerek, her bir devinim olayını ne şekilde tanımladıkları listelenmiştir. Ardından bu cümleler “devinim tarzı”, “devinim yolu” ya da “diğer” olarak sınıflandırılmıştır. Son olarak da, analizlerde kullanmak üzere tek bir oran elde etmek amacıyla, her bir katılımcı için devinim tarzı yanıtlarının sayısının toplam yanıtlara oranı olan *devinim tarzı oranı* hesaplanmıştır.

Elde edilen veriler, Tek-Değişkenli Varyans Analizi (*Univariate ANOVA*) yöntemi kullanılarak analiz edilmiştir. Bağımsız değişken olarak *dil grubu*, bağımlı değişken olarak da *devinim tarzı oranı* alınmıştır. Sonuçlar, dil grupları arasında belirgin bir tipolojik farklılık olduğunu ortaya koymuştur. Anadili Türkçe ve Fransızca olan katılımcıların devinim tarzı oranı ile anadili İngilizce olan katılımcıların devinim tarzı oranı arasında istatistiksel olarak anlamlı bir farklılık tespit edilmiştir ($F(2,59)=602.306$, $p=.00$, $\eta_p^2=.953$). Anadili İngilizce olan grubun devinim tarzı oranı çok yüksekken (hemen hemen her video devinim tarzı fiili içeren bir cümleyle tanımlanmış), anadili Türkçe ve Fransızca olan grubun devinim tarzı fiili kullanım oranı çok düşüktür (videoların çoğu devinim yolu içeren bir fiille tanımlanmıştır).

4.2. Kabul Edilebilirlik Değerlendirme Deneyi

Bir dili anlama deneyi olan bu deneyimizin amacı, dil üretiminde tespit edilmiş olan tipolojik farklılıkların dili anlamada da mevcut olup olmadığını sorgulamaktır. Psikodilbilim alanında yapılan çalışmaların çoğu işledikleri konuya ya dil üretimi ya da dili anlama açısından bakarken, biz bu çalışmada birbirini tamamladığına inandığımız bu iki açıyla da konuyu incelemeyi uygun bulduk. Bu bölümde katılımcılar, videoların her birini birer cümleyle birlikte görmekteydiler. Örneğin, adamın sendeleyerek karşıdan karşıya geçtiği videonun altında “Adam sendeleyerek karşıdan karşıya geçti” şeklinde yazılı bir cümle verilmekteydi. Toplam 30 adet video bulunmaktaydı ve bunlardan 10 tanesi o dilde doğru olarak kabul edilecek birer cümleyle (Türkçe için bir devinim yolu cümlesi), 10 tanesi karşıt tipolojideki cümle yapısına uygun fakat hedef dilde uygunsuz olan birer cümleyle (Türkçe için

devinim tarzı cümlesi), kalan 10 tanesi ise farklı birer tür cümleyle beraber verilmekteydi. Katılımcılardan bu cümleleri 1'den 5'e kadar bir puan vererek değerlendirmeleri istendi. Hedef dile uygun kullanım içeren cümlelerin yüksek puanlar (5'e yakın), karşıt tipolojiye uygun olan cümlelerin ise düşük puanlar (1'e yakın) almaları beklenmekteydi.

İlk olarak, katılımcıların cümlelere vermiş oldukları puanlar hesaplandı. Ardından her bir katılımcının 10 adet devinim tarzı cümlesine verdiği puanların ortalaması "tarz ortalaması" ve 10 adet devinim yolu cümlesine verdiği puanların ortalaması da "yol ortalaması" olarak hesaplandı.

Bu deneyden elde edilen veriler, Tekrarlı-Ölçüm Varyans Analizi (*Repeated Measures ANOVA*) yöntemi kullanılarak analiz edilmiştir. Denekler-arası değişken olarak *dil grubu*, denekler-içi değişken olarak da *cümle tipi* (devinim tarzı ortalaması ve devinim yolu ortalaması) kullanılmıştır. Sonuçlar, dil grupları arasında anlamlı bir tipolojik farklılık olduğunu ortaya koymuştur ($F(2,44)=268.286$, $p=.00$, $\eta_p^2=.924$). Fakat eldeki veriler daha ayrıntılı bir biçimde incelendiğinde görülmüştür ki, bu deneyin sonuçları dil üretimi deneyinin sonuçlarıyla birebir örtüşmemektedir. Bu deneyde de anadili İngilizce olan katılımcılar devinimin tarzına, anadili Türkçe ve Fransızca olan katılımcılar da devinimin yoluna daha yüksek puan vermişlerdir ve bu Talmy (1985)'nin teorisini tipoloji-arası bir değerlendirmeye doğrulamaktadır. Öte yandan, konuya tipoloji-içi bir değerlendirmeye bakacak olursak, aynı tipolojik sınıfa ait olan Türkçe ve Fransızca'daki sonuçlar da anlamlı bir biçimde farklılık göstermektedir³⁷. Bu beklenmedik sonucun muhtemel nedenlerinden "Tartışma ve Sonuç" bölümünde bahsedilecektir.

4.3. Benzerlik Değerlendirme Deneyi

Bu deney sözsüz bir deneydi ve amacı üç dilin konuşurlarını bilişsel bir deneye tabi tutarak, devinim olayı algılarını karşılaştırmalı olarak incelemektir. Bu deneyde herhangi bir dil etkisi görülmeyeceği öngörülmekteydi.

³⁷ $F(1,22)=15.877$, $p=.001$, $\eta_p^2=.419$

“Gereç ve Yöntem” bölümünde de açıklandığı üzere, bu deneyde katılımcılara üçlü setler halinde videolar gösterilmiş ve kendilerinden bir benzerlik değerlendirmesi yapmaları istenmişti. Bilgisayar faresi yardımıyla yapılan seçimler sistem tarafından kaydedilmişti. Devinimim tarzı benzerliğine dayanan seçimler “1”, devinimin yolu benzerliğine dayananlar ise “2” olarak kodlandı ve ardından her bir katılımcının devinim tarzı yanıtlarının toplam yanıtlara oranı olarak açıklanabilecek olan *devinim tarzı oranı* hesaplandı. Bu deneyin analizinde de Tek-Değişkenli Varyans Analizi (*Univariate ANOVA*) yöntemi kullanıldı. *Dil grubu* yine bağımsız değişken olarak alınırken, bağımlı değişken olarak da *devinim tarzı oranı* alındı.

Analiz sonucunda, her üç grubun da baskın olarak *devinimin tarzını* benzerlik ölçütü olarak almış oldukları gözlemlenmiştir ve herhangi bir dillerarası farklılık bulunmamıştır. Türkçe, İngilizce ve Fransızca konuşan grupların yanıtları birbirleriyle kıyaslandığında da, hiçbir dil çifti arasında anlamlı bir farklılık tespit edilmemiştir. Bu sonuçlar, alanda daha önce benzer hipotezlerle yapılan çalışmaların sonuçlarıyla (örn. Gennari et al., 2002; Papafragou et al., 2002; Papafragou and Selimis, 2010) ve *evrensel yaklaşımın* (Jackendoff, 1990, 1996) varsayımlarıyla örtüşmektedir.

4.4. Görüntü İnceleme Deneyi (Sözsüz Göz-İzleme Deneyi)

Bu deneyin temel amacı, çeşitli devinim olaylarını inceleyen farklı anadillere sahip katılımcıların performanslarını karşılaştırmaktı. Deneyde dil kullanımını gerektiren herhangi bir unsur bulunmadığı, hatta gizli dil kullanımını önlemek amacıyla Söyleyiş Baskılaması Tekniği kullanıldığı için; anadili Türkçe, İngilizce ve Fransızca olan üç grubun göz izleme sonuçları arasında istatistiksel olarak anlamlı bir fark bulunmayacağı öngörülmekteydi.

Tobii Göz İzleme Cihazı’ndan elde edilen ham veri, her bir katılımcının göz hareketlerini video üzerinde hareket eden kırmızı yuvarlaklar olarak ifade etmekteydi. Kırmızı yuvarlakların çapının artması, kişinin o noktaya odaklanma süresinin arttığını göstermekteydi. Ham göz hareketi verilerin işlenmesi için kullanılan Tobii Studio yazılımı yalnızca sabit imajların analizini desteklediği için, hareketli videolardan oluşan verilerimizi analiz edebilmek için farklı bir çözüme ihtiyaç

duyulmaktaydı. Bu amaçla bir yazılım firmasıyla bağlantıya geçildi ve burada çalışan uzmanlara ihtiyacımıza yönelik bir yazılım hazırlatıldı. Bütün göz izleme verileri için kullanılmış olan bu yazılımın işleyişi şu şekildedir: amacımız gereği katılımcıların bir devinim olayının hangi ögesine (*devinimin yolu* ya da *devinimin tarzı*) daha çok dikkat ettiklerini tespit edebilmek için, önce bu devinim sahnelerinin her birinde hangi kısımların yol, hangilerinin tarz ifade ettiğinin belirlenmesi gerekmektedir. İlk olarak bu alanda yapılmış olan önceki çalışmalar tarandı (*örn.* Papafragou ve diğ., 2008; Holsanova ve diğ., 2010 ya da Flecken 2010), fakat anlamlı bir sonuca ulaşamadı. Bu konuda yapılan sınırlı sayıdaki çalışmanın her birinde farklı sınıflandırmalar yapılmıştı ve hiçbirisi de bizim amacımıza uygun değildi. Bu nedenle, bu çalışma için kendi alanlarımızı, kendi amacımız doğrultusunda, kendimiz belirlemeye karar verdik. Devinimin yolu belirlenirken, yazındaki genel yargı temel alındı: yol, kaynak zemin ile hedef zemin arasındaki aralığı ve aynı zamanda da bu iki noktayı kapsamaktadır. Böylece, her üç arka plan için de farklı birer “devinim yolu” alanı belirlendi. Bu alanlar; içeri ve dışarı yönlerini içeren görüntüler için apartman kapısı ile eylemin başladığı/bittiği kaldırım kenarı arasındaki bölgeyi içermektedir. Yukarı/aşağı yönlerini içeren görüntüler için bu alan, temelde eylemin başlangıç ve bitiş noktası arasındaki alanı kapsamaktaydı. Karşıdan karşıya geçme ögesi içeren görüntüler için devinimin yolu bölgesi ise, iki kaldırımın arasındaki bölge olarak belirlendi.

Devinimin tarzını ifade edecek olan bölgeleri belirlemek ise, çok daha zahmetli bir işti. Devinimin tarzı, bu eylemleri harekete döken kişinin vücut hareketleriyle belirlendiği için, bu bölgeyi “o kişinin vücudunun tamamı” olarak almayı uygun gördük. Fakat kişi devamlı hareket halinde olduğu için, bu belirleme işi çok daha uzun süre aldı. Bu işlem, kişinin her bir videonun, her bir karesinde ekranın neresinde olduğunun tek tek tespit edilmesini gerektirmekteydi. Bu amaçla, özel yazılımımıza yüklemek üzere, bu alanları tek tek kendimiz işaretledik. Ardından, her bir katılımcıdan elde edilen veri, bu sistem yardımıyla analiz edildi ve elde edilen “devinimin yolu”, “devinimin tarzı” ve “diğer” yüzdeleri listelendi.

Veri analizi için Tek-Değişkenli Varyans Analizi (*Univariate ANOVA*) yöntemi kullanılmıştır. Bağımsız değişken olarak *dil grubu*, bağımlı değişken olarak ise, devinimin tarzı yüzdesinin tarz ve yolun toplam yüzdesine oranı olarak ifade

edilebilecek olan *devinim tarzı oranı* değeri alınmıştır. Elde edilen sonuçlar göstermiştir ki, katılımcıların bu deneydeki göz izleme sonuçları, dolayısıyla dikkat örüntüleri, arasında anlamlı herhangi bir dillerarası farklılık gözlenmemektedir. Anadili Türkçe olan katılımcılar da, İngilizce olan katılımcılar da, Fransızca olan katılımcılar da, baskın bir yüzdeyle daha çok devinim tarzı ögesine bakmışlardır³⁸.

4.5. Betimleme Öncesi Göz-İzleme Deneyi

Önceki göz-izleme deneyinden farklı olarak bu deneyde katılımcılar, görüntüleri amaçsız bir şekilde değil, izledikten sonra sözlü olarak tanımlamak amacıyla izlemekteydi. Buradaki amaç, birinci ve üçüncü deneydeki göz hareketi verilerini kıyaslayarak, Slobin (1996a)'ın “konuşmak-için-düşünmek” varsayımını test etmektir.

Bu deneyden elde edilen göz hareketi verileri, aynı bir önceki deneyde olduğu gibi Tek-Değişkenli Varyans Analizi (*Univariate ANOVA*) yöntemiyle analiz edilmiştir ve yine bağımsız değişken olarak *dil grubu*, bağımlı değişken olarak *devinim tarzı oranı* alınmıştır. Sonuçlar, bu göz izleme deneyinin sözlü doğasının pek bir fark oluşturmadığını ortaya koymuş ve tüm katılımcılar, önceki deneyde de olduğu gibi, anadillerinin etkisi olmaksızın, devinim tarzına daha çok bakmışlardır.

5. TARTIŞMA VE SONUÇ

Bu çalışmanın amacı, dillerarası farklılıklardan yola çıkarak farklı anadil konuşurlarının devinim olaylarını ifade biçimleri ve kavramsallaştırma örüntüleri arasındaki benzerlik ve farklılıkları ortaya koymak olarak belirlenmişti. Bu amaçla, birbirini takip eden ve tamamlayan beş adet deney gerçekleştirildi.

Bu deneylerden *Görüntü Betimleme Deneyi* bir dil üretimi deneyiydi ve ilk temel hipotezimizi test ederek çalışmamızın bel kemiğini oluşturmaktaydı. Bu deneyde katılımcılar izledikleri videoları serbest bir şekilde ifade etmekteydiler ve U-

³⁸ Deneylerde kullanılan materyalin hareketli doğası gereği, *devinim tarzı bölgesi* olarak ifade edilen alanla *devinim yolu bölgesi* olarak nitelendirilen alan arasında yer yer çakışmalar olmaktadır. Okuyucuların çalışma boyunca bu önüne geçilemez durumu akıllarında tutmalarında fayda bulunmaktadır.

çerçevesi bir dil olarak nitelendirilen İngilizce ile E-çerçevesi diller olarak nitelendirilen Türkçe ve Fransızca'nın devinim olayı ifade biçimleri arasında kesin farklılıklar beklenmekteydi. Beklendiği gibi anadili İngilizce olan katılımcılarımız devinim tarzını çoğunlukla esas fiille vermeyi tercih ederken, anadili Türkçe ve Fransızca olan katılımcılarımız devinimin yolunu esas fiille vermeyi uygun görmüşlerdir. Hipotezimizi destekleyen bu sonuç, aynı zamanda Talmy'nin ikili sınıflandırmasını da desteklemektedir.

Yukarıda yorumlanan deneyin sonuçlarını tamamlayacağına inandığımız *Kabul Edilebilirlik Değerlendirme Deneyi*, bir dili anlama deneyiydi. Yani katılımcıların dil kullanımlarını değil, karşılarına çıkan kullanımları nasıl değerlendirdiklerini incelemekteydi. Bu deneyin sonucuna göz attığımızda, hipotezimizin bir kez daha desteklendiğini görmekteyiz. Çünkü bu veriler bize anadili İngilizce olan grubun devinim tarzını esas fiille veren cümlelere çok yüksek puan verdiğini, anadili Türkçe ve Fransızca olan grupların ise yüksek puanları devinim yolunun esas fiille verildiği cümlelere verdiklerini göstermektedir. Bu da yine Talmy'nin sınıflandırmasını destekler bir sonuçtur. Öte yandan, dil üretimi sonucundan farklı olarak, bu dili anlama deneyinde aynı tipolojik özelliğe sahip iki dil olan Türkçe ve Fransızca'nın anadil konuşurlarının performansları arasında da farklılıklar gözlemlenmiştir. Yani, tipoloji-içi bir farklılık tespit edilmiştir. Veriler yakından incelendiğinde, bu farklılığın devinim tarzı cümlelerine verilen puanlardan ziyade devinim yolu cümlelerine verilen puanlarla ilgili olduğu görülmüştür. Diğer bir deyişle, her iki grubun katılımcıları da (kendilerinden beklediği üzere) devinim tarzını esas fiille veren cümleleri düşük puanla değerlendirmişlerdir. Fakat devinim yolunun esas fiille verildiği, yani hedef dile uygun, cümlelerde Türkler Fransızlar'dan daha çekimser kalmıştır. Bütün Fransız katılımcılar bu cümlelere tam ya da tama yakın puan verirken; Türk katılımcılar, yine yüksek puanlar vermekle beraber, Fransız katılımcılar kadar kararlılık gösterememişlerdir. Bu durumun, biri biçimdizimsel diğeri sözcüksel olmak üzere iki açıklaması olabileceği düşünülmektedir. Türk katılımcıların devinimin yolunun esas fiille verildiği cümlelere Fransız katılımcılar kadar yüksek puan vermemesinin ilk nedeni bu cümlelerin dizimsel, yani sentaktik yapıları olabilir. Deneyde verilen tüm cümleler şu örnekteki dizimle verilmiştir: Adam + koşarak + binaya + girdi. Fakat sözcük dizimi esnek bir dil olan Türkçe'de aynı cümle "Adam binaya koşarak girdi" olarak da ifade edilebilir. Öte yandan,

biçimsel yani morfolojik olarak da zengin bir dil olan Türkçe’de aynı cümleyi “binaya girdi” yerine, “binanın içine girdi” ya da “binadan içeri girdi” şeklinde de ifade edebiliriz. Fransızca’da ise, biçimdizimsel anlamda bu kadar esneklik ve çeşitlilik mevcut değildir. Bu konuda yapılabilecek diğer bir açıklama da, Türkçe’nin devinim tarzı ifade eden belirteçlerinin zenginliğidir ve bu da –arak son ekinin biçimbilimsel üretkenliğinden kaynaklanmaktadır. Deneydeki cümlelerde sürekli “sendeleyerek” sözcüğüyle ifade edilen eylem “sallanarak” sözcüğüyle, “topallayarak” sözcüğüyle ifade edilen eylem “aksayarak” sözcüğüyle veya “hoplayarak” sözcüğüyle ifade edilen eylem “zıplayarak ya da sekerek” ifadeleriyle de verilebilmektedir. Türk katılımcıların bu cümlelere beklendiği kadar yüksek puan vermemesinde sözcük seçiminin de, yani bekledikleri sözcükleri orada bulamamalarının da, rol oynayabileceğini düşünmekteyiz.

Mevcut çalışmanın ikinci büyük araştırma sorusu da, anadilleri Türkçe, İngilizce ve Fransızca olan yetişkinlerin aynı devinim olayını kavramsallaştırma örüntülerinin birbirinden farklı olup olmadığıydı. Diğer bir deyişle, dilin düşünce üzerindeki etkisi üç sözsüz deneyle sınanmaktaydı. Bunlardan ilki *Benzerlik Değerlendirme Deneyi*’ydi ve sonucunda herhangi bir dillerarası farklılık bulunamamıştı. Bu da bize dilden bağımsız algıyı, daha genel ifadesiyle düşüncenin evrenselliği ilkesini hatırlatmaktaydı. Bu deney sonuçları hipotezimizi desteklemekteydi, fakat daha etkin bir yöntem olduğuna inandığımız göz izleme verilerinin analizi sonlanana kadar net bir sonuca varmaktan kaçınmıştık. Sözsüz göz izleme deneyimiz olan *Görüntü İnceleme Deneyi*’nin sonucunda da herhangi bir dil etkisi görünmemesi, hipotezimizi doğrulamış oldu.

Sözsüz deneyimizdeki göz izleme sonuçlarıyla, betimleme deneyi öncesi göz izleme sonuçlarını kıyaslayarak, dil kullanım amacıyla izlenen devinim olaylarının amaçsızca izlenenlerden farklı bir şekilde değerlendirilip değerlendirilmediğini de sorgulamıştık. Bu bir anlamda Slobin (1996a)’ın “konuşmak-için-düşünmek” varsayımını test etmektir. Kişinin betimlemeden önceki izlemelerinde, kendisine sunulan görüntüleri kuracağı cümlelerin yapısına uygun olarak inceleyeceğini öngörmekteydik. Diğer bir deyişle, anadili İngilizce olan katılımcıların betimleme deneyinde dillerindeki baskın örüntüye uygun olarak devinim tarzına daha fazla bakacakları; anadilleri Türkçe ve Fransızca olan katılımcılarınsa, yine dillerindeki

baskın örüntüye uygun olarak devinim yoluna daha fazla bakacakları varsayılmaktaydı. Fakat sonuçlar böyle bir etki vermedi. Betimleme öncesi göz-izleme verilerinde de, aynı diğer göz izleme deneyimizde olduğu gibi, herhangi bir dil etkisi tespit edilemedi. Slobin'in varsayımıyla çelişkili görünen bu sonuç, *evrensel yaklaşımla* (Jackendoff, 1990, 1996) uyumluluk göstermektedir.

Bu çalışma kapsamında gerçekleştirilen deneyler birçok yönüyle özgün değerlere sahiptir. Çalışmada, devinim eylemlerinin ifade biçimleri, hem üretim hem de anlama yönüyle ele alınmıştır. Yalnızca devinim olaylarının ifade biçimleri değil, aynı zamanda bu olayların insan zihninde nasıl kavramsallaştırıldıkları da incelenmiştir. Ayrıca mevcut çalışma; farklı dilleri kıyaslaması yönüyle bir *karşılaştırmalı dilbilim çalışması*, dil ve düşünce ilişkisi konusunu ele alması yönüyle bir *dil felsefesi çalışması* ve de kullandığı deneysel yöntemler açısından bir *psikodilbilim ve bilişsel bilimler çalışması* olarak nitelendirilebilecek çok yönlü bir çalışmadır. Yazına metodolojik olarak yapılan temel katkı da şöyledir: Alanyazındaki devinim olayı çalışmalarının çoğunda deney materyali olarak statik resim ya da animasyon kullanımı tercih edilirken, bu çalışmada devinimin eksiksiz yansıtılması ve gerçeklik olgusunun yitirilmemesi amacıyla gerçek videolar kullanılmıştır. Bildiğimiz kadarıyla, Türkçe devinim olayları için gerçek video çekimleriyle yapılmış olan başka bir çalışma bulunmamaktadır.