

This is to certify that the

dissertation entitled

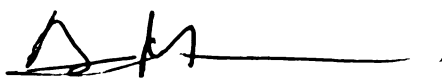
TENSE, ASPECT, AND EVENT REPRESENTATIONS
IN ENGLISH AND JAPANESE

presented by

Ayako Yamagata

has been accepted towards fulfillment
of the requirements for

Ph.D. degree in Linguistics



Major professor

Date 7 May 2000

LIBRARY
Michigan State
University

PLACE IN RETURN BOX to remove this checkout from your record.
 TO AVOID FINES return on or before date due.
 MAY BE RECALLED with earlier due date if requested.

DATE DUE	DATE DUE	DATE DUE
050602 JUL 13 2002		092404
OCT 13 2004		
083105 OCT 09 2005		

**TENSE, ASPECT, AND EVENT REPRESENTATIONS
IN ENGLISH AND JAPANESE**

By

Ayako Yamagata

A DISSERTATION

**Submitted to
Michigan State University
in partial fulfillment of the requirements
for the degree of**

DOCTOR OF PHILOSOPHY

Department of Linguistics

2000

ABSTRACT

TENSE, ASPECT, AND EVENT REPRESENTATIONS IN ENGLISH AND JAPANESE

By

Ayako Yamagata

This study attempts to defend a theory of temporal reference that unifies the notions of tense and aspect, and apply it to the tense and aspect systems of English and Japanese. It argues that adopting Klein's (1994) system of temporal reference allows for the crosslinguistic differences between the English and Japanese tense and aspect systems to be accounted for in a simple manner with no extra stipulations.

The present study shows that the complications with the interpretation of tenses in different types of embedded clauses can be ascribed to the presence or absence of the intensional context, a shift of the deictic center for the tense interpretations, and the aspect of the predicates that appear with the tense morphemes. As a result, the semantics of tenses can be maintained consistent across different constructions and the tense system can be much simpler than those that have been proposed by some existing theories of tense. Both the English past and the Japanese *-ta* encode $TT < TU$ as a tense marker, while both the English present and the Japanese *-ru* encode $TT \supset TU$ for the present tense, though the Japanese *-ru* as a non-past tense marker additionally encodes $TU < TT$. The seeming differences between the English and the Japanese tense systems are attributed to the fact that Japanese employs the relative tense system on top of the absolute tense system. The relative tense is obligatorily used in the verb-complement structure in Japanese, while it is optional in *toki*-clauses and relative clauses. This is

because only complement clauses may involve the indirect speech context. The presence of SOT in English and its absence in Japanese in the verb-complement constructions can be explained by the fact that the complement past in Japanese is always evaluated with respect to the matrix past, and therefore, may not be able to overlap with the matrix event. Thus, there is no need to posit a SOT rule for English nor different semantics for the past tenses in English and Japanese.

This study also accounts for the puzzling behavior of the Japanese *Verb-te-iru* form, which can have both a perfect and a progressive interpretation. While the English present participial encodes $TT \subset TSit$ and the English past participial encodes $TSit \subset TT$, *-te* of *-te-iru* simply introduces a topic time without specifying how it relates to the situation time. This allows the TT of the sentences with *-te-iru* to be placed either within the situation time or in the post time of the situation, depending on the context. The existential meaning of the present progressive and the present perfect of English and the Japanese *-te-iru* is provided by their auxiliary verbs. The tense components of the auxiliary verbs of these aspectual forms encode $TT \supset TU$, and provide the current relevance of the described eventuality. The study shows that Klein's system clarifies the respective contributions of tense and aspect components of these complex temporal morphemes in English and Japanese in a simple manner without extra stipulations.

Copyright by
AYAKO YAMAGATA
2000

To my parents

ACKNOWLEDGMENTS

I would like to show my deepest appreciation to all my committee members: Alan Munn, Barbara Abbott, Mutsuko Endo Hudson, Susan Gass, and Dennis Preston. I am also greatly indebted to Cristina Schmitt, who joined the Department of Linguistics at Michigan State University after I left the department to work for another institution, and helped me with the completion of this thesis. I have been away from school almost for three years, and I would not have been able to complete this thesis if it had not been for the warm encouragement of all the people I list here.

I have been very fortunate to be a student of many wonderful faculty members at the Department of Linguistics at Michigan State University. As I started teaching Japanese linguistics at other institutions before the completion of my thesis, I realized that I have been incorporating the styles of the different people from whom I learned linguistics. As I occasionally hear my students characterize me as being enthusiastic about the subjects I teach and open and flexible to individual student's levels of understanding and needs, I realize that this is exactly how I characterize Alan Munn in his syntax classes. I believe that his enthusiasm in the subject matter and sincere intellectual curiosity for every new piece of data or unsolved problem inspired and involved many students. Although there are innumerable other assets of his that my ability cannot match, I will strive to approximate my level of instruction to his both inside and outside class. It is an honor to find myself trying to follow what he is doing for his students.

I am also grateful for the fact that I took Introduction to Linguistics from Barbara Abbott, who, not only taught me how intriguing the study of language is, but also how rigorous it is. Although I cannot yet teach a course in as organized a manner as she does, I

surely endeavor to follow her as an example. In my first year in the MA program, I also took the Structure of Japanese from Mutsuko Endo Hudson. She taught me how to look at the data analytically and read others' arguments critically. Her course formed a foundation for the Japanese linguistic courses that I teach. I am also grateful to Susan Gass, from whom I took second language acquisition courses. The training I had in her SLA courses helps me bridge my research in linguistics and language teaching. I would also like to show my gratitude to Yen-Hwei Lin, who taught me Phonology, and Kazuhiko Fukushima, who is currently at Kansai Gaidai University in Osaka, Japan. Their courses were both fun and enlightening, and they did a great job in pushing me to work hard. I would also like to show my appreciation to Dennis Preston, who always gave me insightful comments from different perspectives.

My special thanks goes to Mutsuko Endo Hudson, who taught me Japanese Pedagogy at the Summer Institute of Columbia University and invited me to MSU to pursue my graduate degree in linguistics. If I hadn't met her there, I would have missed many great opportunities for learning and teaching.

My deepest and foremost gratitude goes to Alan Munn and Cristina Schmitt, who have supported me in my research since the time they have come to MSU. I greatly appreciate their guidance and encouragement over the past few years. I admire their enthusiasm and the integrity in their research, genuine intellectual curiosity to any language data, academic rigor, and how they care for their students. Their hard work and sincere attitude toward linguistics inspire me constantly.

Finally, I would like to express my gratitude to my parents who taught me how to love people and the world, and my dearest aunt, Emiko, who has supported me both

mentally and financially throughout my education. I would also like to thank my boyfriend, Iwao, who has given me strong moral support from across the Pacific Ocean, Dianne and Michael Alexanian, who offered me lodging every time I visited East Lansing, my uncle, Hiro and his wife, Misako, who also gave me warm moral support. My uncle Hiro passed away last fall. I wish he were alive to see me complete this dissertation. I am also grateful to all my friends inside and outside the US who cheered me on over the years, and the colleagues in the institutions I worked for in the past few years. Without their kind support and encouragement, I would never have come this far.

TABLE OF CONTENTS

List of tables.	xiii
List of abbreviations.	xiv
Chapter 1: Introduction	1
1.1 The aim of the study.	1
1.2 Theories of tense	3
1.3 Problems to be investigated	4
1.3.1 Sequence of tense	4
1.3.2 Tenses in adverbial clauses.	6
1.3.3 Tenses in relative clauses	8
1.3.4 SOT and aspect	10
1.4 Theoretical framework and hypotheses	16
1.4.1 The theoretical framework (Klein, 1994)	16
1.4.2 Hypothesis and more questions.	26
1.5 Summary	32
Chapter 2: Theories of Temporal Reference and Sequence of Tense	33
Introduction	33
2.1 Problems with the traditional analysis	34
2.2 Enç's (1989) 'Anchoring Conditions' for tense	36
2.3 Ogihara's (1989, 1995a,b) account of SOT	42
2.3.1 Referential vs. quantificational analyses of tense	42
2.3.2 Tense deletion and the difference between English and Japanese	44

2.3.3	Ogihara's account of DA reading of present-under-past and <i>de re</i> attitude	49
2.4	Abusch's theory of tense interpretation	54
2.4.1	Sequence of tense, intensionality and scope (Abusch 1988) .	54
2.4.2	Interpretation of <i>de re</i> belief (Abusch 1991, 1994, 1997) . . .	56
2.4.3	Upperlimit phenomenon and division of labor (Abusch 1994, 1997a)	59
2.5	The role of pragmatics in tense interpretations (Costa 1972)	64
2.6	Hornstein (1993)	68
2.7	Nakamura (1994)	73
2.8	Summary	77
Chapter 3:	Temporal Reference in English and Japanese	79
	Introduction	79
3.1	The SOT data revisited	80
3.1.1	Complement tenses in English and Japanese	80
3.1.2	Interpretation of complement tenses in English	83
3.1.2.1	Past under past	83
3.1.2.2	Present under past	88
3.1.2.3	Tense or aspect?	92
3.1.3	Interpretation of complement tenses in Japanese	99
3.1.4	Problematic cases and revision of future tense in English. .	104
3.2	Tenses in adverbial clauses	120
3.2.1	Semantics of <i>when</i> -clauses and <i>toki</i> -clauses	120
3.2.2	Interpretation of tenses in <i>toki</i> -clauses	121

3.2.2 The ordering of events: tense or pragmatics?	130
3.3 Tenses in relative clauses	135
3.3.1 Tenses in relative clauses in English	135
3.3.2 Tenses in relative clauses in Japanese	139
3.4 Division of labor and a note on the present tense	146
Chapter 4: Aspect and event representations	149
Introduction	149
4.1 Application of Klein's model to aspectual forms in English	150
4.1.1 The progressive in English and eventive interpretations . .	151
4.1.2 Peculiarity of present tense in English	156
4.1.3 Representation of the present perfect in English	158
4.1.4 Representation of sentences with <i>-te-iru</i> in Klein's model	161
4.1.5 Summary	166
4.2 Ambiguity of <i>-te-iru</i>	168
4.2.1 The stage-level and individual-level distinction (Ogihara, 1999)	168
4.2.2 <i>-Te-iru</i> as a stage-level operator (Yamagata, 1998)	176
4.2.2.1 The problems with the type 4 verbs	176
4.2.2.2 <i>-Te-iru</i> with stative verbs	181
4.2.2.3 <i>-Te-iru</i> with non-stative verbs	182
4.2.3 Existential readings and <i>-te-iru</i>	183
4.3 The semantics of <i>-te</i> in the aspectual composition	186
4.3.1 The present and past participials in English and <i>-te</i> of <i>-te-iru</i>	186

4.3.2 Underspecification of <i>-te</i> and theoretical consequences . . .	188
4.3.3 Remaining problems and the further modification on Klein (1994)	190
4.4 Aspectual system of Japanese	192
4.5 Summary	193
Chapter 5: Conclusion	195

LIST OF TABLES

Table 1: Basic tenses in English	26
Table 2: Basic tenses in Japanese	27
Table 3: Aspectual distinctions in English and Japanese	29
Table 4: Interpretations of complement tenses embedded under past in English and Japanese	81
Table 5: Semantics of components of aspectual forms	187
Table 6: Aspectul System in Japanese	193

KEY TO ABBREVIATIONS

ACC	accusative marker
ASP	aspectual marker
COMP	complementizer
COP	copula
DA	double access
DAT	dative case marker
DP	determiner phrase
GEN	genitive case marker
LOC	locative marker
NOM	nominative case marker
NON PAST	non past tense marker
NP	noun phrase
PAST	past tense marker
PRES	present tense marker
SOT	sequence of tenses
TOP	topic marker
TSit	time of situation
TSit _A	time of situation of the adverbial clause
TSit _C	time of situation of the complement clause
TSit _M	time of situation of the main clause
TT	topic time
TT _A	topic time of the adverbial clause
TT _C	topic time of the complement clause
TT _M	topic time of the main clause
TU _M	time of utterance
TU _A	time of utterance of the adverbial clause
TU _C	time of utterance of the complement clause
TU _M	time of utterance of the main clause
VP	verb phrase

CHAPTER 1

Introduction

1.1 The aim of the study

The function of temporal expressions in natural language is to relate events and situations that are described by linguistic means to the time in our world. Two major grammatical categories of temporality for this purpose, tense and aspect, have long been the objects of research in linguistics. However, while both tense and aspect are grammatical expressions of temporal reference, they are usually assumed to be mutually discrete by many modern linguists, and most studies on temporality within the framework of contemporary formal linguistics focused on one, excluding the other, and very few attempted unification of these two related, but distinct concepts.

The aim of this thesis is to defend a theory of temporal reference which unifies the notions of tense and aspect, and to account for the behavior of the primary tense/aspect morphemes in English and Japanese in event¹ representations in a principled manner. Specifically, I will apply the theory of temporal reference advanced by Klein (1994) to Japanese, and account for the behavior of the primary temporal morphemes of Japanese in contrast to the primary temporal morphemes of English, elucidating the properties of tense and aspect systems that are common across these two unrelated languages as a descriptive task. As a theoretical task, this study attempts to distinguish the contribution of the theory of tense and aspect from the contribution of other modules of the grammar, and defend the existence of the independent system of temporal reference in natural

¹ I will use the term ‘event’ as it appears in a phrase like ‘event representation’ or ‘representation of events’ interchangeably with ‘eventuality’. Eventuality is a cover term for states, processes, and events (cf. Bach, 1989). I will use a term ‘situation’ as a cover

language grammar. I will argue that adopting Klein's system in conjunction with Abusch's (1988, 1991, 1994, 1997a) theory of intensionality allows for the differences between English and Japanese tense and aspect systems to be accounted for in a simple way without extra stipulations. I will also show that Klein's system can be used to account for the puzzling behavior of the Japanese *Verb-te-iru* form, which can have both a perfect and a progressive interpretation. I will argue that, rather than treating *-te-iru* as lexically ambiguous, we can give it a single meaning from which both interpretations can be derived.

This thesis is organized as follows: In this chapter I will clarify the position of the present study, introduce the problems to be investigated, and present the basic proposals. Chapter 2 will review the proposals and analyses in the previous literature on tense, discuss their problems, and lay out the foundations for the analysis of the temporal morphemes that will follow. In Chapter 3 I apply Klein's theory of temporal reference to the analysis of basic tense morphemes of English and Japanese in various syntactic environments, incorporating the generalizations obtained in Chapter 2. In Chapter 4, I extend the analysis of temporal morphemes to cover the major aspectual forms of English and Japanese, and will investigate their functions in event representations in contrast to the functions of tense morphemes discussed in Chapter 3.

The remainder of this introductory chapter proceeds as follows: section 1.2 introduces various theories of tense in their bare outlines and clarifies the position of the present study. Section 1.3 presents the data and the problems to be investigated, and section 1.4 introduces the theoretical framework, discusses its advantages over other

term for states and processes, and the term 'eventuality' is used when I would like to

approaches to tense, and presents hypotheses on the nature and functions of temporal morphemes to be investigated. The last section will provide a summary of this chapter.

1.2 Theories of tense

While logicians treat tense as an operator that shifts the original evaluation time of non-tensed sentences (Prior 1967), and there are some studies by linguists that follow this tradition (Kamp 1971, Vlach 1973), most linguists nowadays represent natural language tense in different ways. There are two opposing theories in standard treatments of tense which do not introduce additional tense operators in the logical system: 1) the quantificational theory, and 2) the referential theory. The proponents of the quantificational theory (e.g., Dowty 1979; Partee 1984; Ogihara 1989, 1995) treat tense as quantifiers over times, which is much the same in spirit as treating tenses as operators, though the proponents of the quantificational theory represent such properties of tense using the ordinary predicate logic with the addition of a tense variable. In contrast, the proponents of the referential theory propose that natural language tense introduces free variables that receive their value as time intervals that are salient in the context (Partee 1973; Enç 1986, 1987; Abusch 1988, 1991, 1994). In fact, there is a third position, which treats tenses to be consisted as syntactic primitives. Theories of this sort include the theory of tense proposed by Hornstein (1990) and those proposed by Zagana (1995) and Stowell (1995, 1996). The last two contrast with all the others in that they take the tense phrase as a dyadic predicate that takes evaluation time and event time as syntactic arguments.

include both situations and events.

Although the split is usually considered to be either quantificational vs. referential or semantic vs. syntactic treatments of tense, I will propose a different division of the existing theories of natural language tense: 1) the theories of temporal reference which assume that tense and aspect are totally distinct objects of research (e.g., most of the above-mentioned theories), and 2) the theories of temporal reference in which tense and aspect are treated in unified manner (e.g., Reichenbach 1947², Klein 1994, Zagana 1997). This dissertation takes the second position and will show that natural language tense phenomena can only be explained by incorporating the theory of aspect.

1.3 The problems to be investigated

1.3.1 Sequence of tense

An adequate theory of temporal reference of natural language should be able to account for the following data.

- (1) John said that Mary **was** pregnant.
 - a. John said that Mary had been pregnant. (shifted reading³)
 - b. John said, "Mary is pregnant." (simultaneous reading)
- (2) John said that Mary **is** pregnant. ('double-access' reading)
- (3) John-wa Mary-ga ninsinsi-te-i-ta to it-ta
 John-TOP Mary-NOM pregnant-ASP-PAST COMP say-PAST
 'John said that Mary had been pregnant.' (shifted reading only)

² Although Reichenbach (1947) treats the progressive aspect as 'continuous tense' and also treats the perfect as representing tense on a par with the simple past, and Vlach (1993) explicitly argues against the Reichenbachian treatment of tense, I will include Reichenbach (1947) here, since his work did not exclude aspect from the theory of temporal reference, despite giving the wrong label.

³ The 'shifted-reading' refers to a reading of past-under-past constructions in which the time of embedded situation is taken to be prior to the time of the matrix past.

- (4) John-wa Mary-ga ninsinsi-te-i-**ru** to it-ta
 John-TOP Mary-NOM pregnant-ASP-NON PAST COMP say-PAST
 a. John said, "Mary is pregnant." (simultaneous reading)
 b. John said that Mary is pregnant. ('double-access reading')[optional])

An English sentence like (1) in which the past tense in the complement occurs under the matrix past may receive an interpretation in which the complement event precedes the main event in the time sequence as given in (1a). However, it may also receive an interpretation in which the evaluation time of the event in the complement clause is simultaneous with the time of the matrix event as shown in (1b).

The fact that we have this latter interpretation for the past-under-past construction like (1) has led some researchers to speculate that the past tense in the complement is in fact the present tense in disguise, and that English has a rule to transmit the matrix past tense to the embedded tense. This is called the sequence of tense (hereafter, SOT) rule. However, as shown in (2), in which the complement predicate under the matrix past is marked with the present tense, this rule seems to be optional. Moreover, the present-under past construction like (2) receives a peculiar reading, which is not shared by the simultaneous reading of (1).

It has been pointed out by Comrie (1985) and by many others (e.g., Costa 1972, Smith 1978, Enç 1987, Ogihara 1989, and Abusch 1991) that constructions like (2) with the present tense under the past tense show the current relevance of the situation expressed by the complement clause. Following Ogihara (1989), let us call this the 'double-access' reading. The 'double-access' (hereafter, DA) reading is so called because in order to interpret the complement of the present-under-past construction correctly, we

need to have access to both the evaluation time introduced by the matrix clause and the speech time.

When we turn to the surface equivalent of the past-under-past in Japanese, which is shown in (3), we only have the shifted reading like the English example in (1a). In order to have a simultaneous reading like (1b), the complement predicate has to be marked with the non-past tense morpheme in Japanese as shown in (4). What is interesting here is that the present-under-past construction in Japanese may also have the DA reading on top of the simultaneous reading, which is not obligatory unlike the English counterpart.

An adequate theory of temporal reference in natural language must be able to account for the presence of the SOT phenomenon in English and its absence in Japanese, and the obligatoriness of DA reading in English on the one hand, and its optional nature in Japanese on the other.

1.3.2 Tenses in adverbial-clauses

Japanese tense morphemes exhibit different behavior from English tense not only in the complement clauses but also in the adverbial clauses.

- (5) I met him when I was going to Tokyo.
- (6) Tokyo-ni ik-u toki kare-ni at-ta
Tokyo to go-NON PAST when he to meet-PAST
'I met him when I was going to Tokyo.'
- (7) I will meet him when I get to Tokyo.
- (8) Tokyo-ni it-ta toki kare-ni a-u
Tokyo to go-PAST when he to meet-NON PAST
'I will meet him when I get to Tokyo.'

- (9) Tokyo-ni it-**ta** toki hikooki-no naka-de kare-ni at-**ta**
 Tokyo to go-PAST when plane-GEN inside-at he to meet-PAST
 'I met him on the plane (to Tokyo) when I went to Tokyo.'
- (10) Kondo Tokyo-ni ik-**u** toki Shinjuku-de kare-ni a-**u**
 next time Tokyo to go-NON PAST when Shinjuku in he to meet-NON PAST
 'I will meet him in Shinjuku next time when I go to Tokyo.'

The Japanese counterpart of (5) as shown in (6) has non-past tense marking on the verb in the *when*-clause despite the fact that the sentence as a whole has evaluation time in the past. In contrast, the Japanese counterpart of (7) as given in (8) has past tense marking on the verb in the subordinate clause even though it in fact refers to the future time. Some researchers have attempted to account for the behavior of these temporal morphemes in Japanese by positing that they are not tense markers, but aspectual markers (e.g. Ando, 1986, among others). The proponents of this view argue that *-ru* is a marker of imperfective aspect that indicates incompleteness of the embedded event with respect to the matrix event, while *-ta* is a marker of perfective aspect that indicates completion of the embedded event with respect to the matrix event. However, this may not provide a full account of the behavior of these temporal morphemes, as the situations depicted in (6) and (8) can be rephrased as (9) and (10). In (9), despite the fact that the embedded event cannot be understood to have been completed with respect to the matrix event, *-ta* is used, and in (10), *-ru* is used to describe the event which must precede the matrix event. Thus, the position that *-ta* always indicates completion and *-ru*, incompleteness, of the embedded eventuality with respect to the matrix eventuality, cannot be maintained.

An adequate theory of tense and aspect must be able to account for the similarities and the differences between the pairs like (5) and (6), (7) and (8), (6) and (9), and (8) and (10). It also has to be able to provide an account of the behavior of the temporal

morphemes in Japanese, which do not seem to behave like tense markers in a strict sense, as observed in (4), (6), and (8).

1.3.3 Tenses in relative clauses

It has been pointed out by previous studies on tense that tenses in relative clauses exhibit different behavior from those in the complement clauses (cf. Enç 1987, Abusch 1988, Ogihara 1989). The following example from Ogihara (1989: 96 (27)) illustrates this point.

(11) John saw [NP a man [S' who was laughing]].

As pointed out by Ogihara, the above sentence allows any temporal relation between the matrix and the main events. That is, the time of the man's laughing can be prior to, simultaneous with, or subsequent to the time of John's seeing it. This observation has led some researchers to believe that tenses in relative clauses are assumed to allow an independent reading⁴ (Enç 1987, Abusch 1988, Ogihara 1989). The Japanese surface equivalent (i.e., the past-under-past construction) will allow the same possibilities for the interpretation.

(12) John-wa [NP warat-te i-ta otoko]-o mi-ta
John-TOP laughing be-PAST man -ACC see-PAST
'John saw a man who was laughing.'

(12) allows the same three temporal relationship between the time of John's seeing and the time of man's laughing: precedence of the embedded event to the main event, simultaneity of two events, and subsequence of the embedded event to the main event.

⁴ 'Independent' in a sense that the interpretation of the tense in the embedded clause does

However, the following sentence in (13), which is another variation of the Japanese translation of (11), has only two readings instead of three.

- (13) John-wa [NP warat-te i-ru otoko]-o mi-ta
 John-TOP laughing be-NON PAST man -ACC see-PAST
 a. John saw a man who was laughing. (simultaneous reading / *shifted reading)
 b. John saw a man who is laughing.

In contrast to (12), (13) can only be interpreted either as a past incident of John's seeing the laughing man or John's seeing in the past of the man who is laughing at the utterance time. (13) does not allow the interpretation in which the time of man's laughing is prior to the time of John's seeing it. Furthermore, as pointed out by Nakamura (1994b), the present-under-past of the relative-clause construction like (13) does not have a DA-reading for the embedded event, unlike the present-under-past of the complement construction (cf. (4) in section 2.1). That is, the man's laughing in (13) may not hold both at the time of seeing and at the time of utterance. In other words, the incident of the man's laughing and the incident of John's seeing the man are necessarily disjoint with the interpretation in which the man's laughing holds at the present time. This observation holds not only for the Japanese data, but also for the English data. The English translation of the Japanese sentence in (13), which is given in (13b), does not have a DA-reading like the sentence in (2) in section 1.3.1. Thus, in relative-clauses, the temporal morphemes in English and temporal morphemes in Japanese seem to behave similarly, compared with their different behaviors in complement clauses and in adverbial clauses.

So far, we have the following generalizations.

not have to be dependent on the tense in the main clause.

- i) Temporal morphemes may exhibit different behavior between English and Japanese (cf. (1) and (2) vs. (3) and (4) for complement clauses; (5) vs. (6) and (7) vs. (8) in adverbial clauses).
- ii) Temporal morphemes in both English and Japanese may be interpreted differently across different constructions (cf. (2) and (4) vs. (13))

At this point, we may ask ourselves a number of questions: what exactly are the constraints imposed by these temporal morphemes on the ordering relationship of events and situations? Are the functions of temporal morphemes of English and Japanese different or are they basically the same? If they are the same, what is it that makes two systems look different as we observed in case of interpretations of tenses in complement clauses and adverbial clauses? Since they also exhibit some similarities in relative clause constructions, we wonder if they are really as different as they seem. Do the constraints stay the same across different constructions (e.g., the complement structure, the relative clause constructions), or do they vary with syntactic structures? If the constraints imposed by the temporal morphemes on the interpretations of the temporal ordering are consistent in a language, what is it that differentiate the interpretation of the complement tense and the interpretation of the relative clause tense, for example? The present dissertation seeks the answers to these questions.

1.3.4 SOT and aspect

When we look at other languages in the world, we find that what we observed as the SOT phenomenon in English is rather a peculiar fact, and that there are languages other than

Japanese which do not exhibit this property. For example, Russian does not have the SOT rule, and tense in the Romance languages exhibits the SOT-like behavior only in certain environments. In Brazilian Portuguese, for example, SOT is possible only with imperfective verbs⁵.

With a closer look at some English data, we see that the SOT phenomenon not only has very limited cross-linguistic validity, but also has limited distribution even within English. SOT is typically observed when predicates in the complements are stative and the rule may not apply when the predicates in the complements are non-stative. (The following data is from Enç, 1987: 634)

- (14) a. Mary found out that John failed the test.
b. The gardener said that the roses died.
c. Sally thought that John drank the beer.

All the sentences in (14) clearly have the shifted reading. However, (14a) and (14b) do not have the simultaneous reading, and thus, are not considered to have undergone the SOT rule. (14c) may have the simultaneous reading if we interpret a verb *drank* to be referring to the past habit of John. However, we may not be able to interpret the incident of John's drinking the beer to be simultaneous with Sally's thinking of it with an eventive interpretation of *drink the beer* in the complement. What is interesting, furthermore, is that when the complement verbs are marked with the present tense, which is what we would have when the SOT rule does not apply to the underlying present tense in the complement, the sentences in (14) will become either semantically anomalous (in case of

⁵ Cristina Schmitt (p.c.).

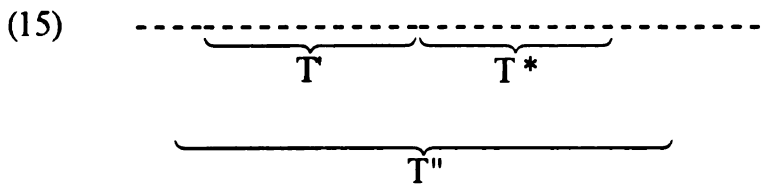
(14a) and (14b)) or will have different meanings from the meanings borne by the sentences with past tense marking on the complement verbs (in case of (14c)).

- (14') a. ?Mary found out that John fails the test.
b. ?The gardener said that the roses die.
c. Sally thought that John drinks the beer.

(14'a) sounds as though Mary found out that John always fails in a particular kind of test, and (14'b) sounds as though the gardener predicts that the certain type of roses being talked about are destined to die. The embedded clause in (14'c) does not refer to a particular 'drinking incident', but refers to John's habit of drinking a certain kind of beer as one of his attributes. None of these sentences with the present tense marking on the complement verb would allow an eventive reading. This observation is consistent with our generalizations on sentences in (14), and if we assume the existence of the SOT rule, we may be able to conclude that the rule may only apply when the complement predicate is interpreted to be non-eventive and that is why only (14'c) can undergo the rule with natural interpretation. However, such a generalization will make the SOT rule very unattractive to be part of the grammar. Not only is it optional and a language-specific rule, but its application is also restricted to a limited set of predicates in a language. While keeping the SOT rule despite its extremely limited distribution is one possibility, we may also speculate that there is no such rule as the SOT in the grammar. The sole motivation for positing the rule is that it is useful for explaining the simultaneous reading of the past-under-past construction. However, as we have observed above, what seems to be responsible for the simultaneous reading of the complement construction is in fact the

temporal overlapping of the complement situation and the matrix event. Such temporal overlapping between the situation described by the complement predicate and the matrix event seems to be determined by the nature of the complement predicate. Specifically in the above cases of (14) and (14'), the complement situation and the main-clause event can be interpreted to be overlapping in time only if the complement predicate refers to the individual-level property of the subject. If the simultaneous reading of the past-under-past constructions can be systematically explained without the SOT rule as in the case of (14), we may as well discard this rule in its entirety.

Enç (1987) provides the following diagram to explain the double-access reading of the present-under-past construction (Enç, *ibid.*: 637), which further supports our hypothesis on the alternative explanation for the seeming referential dependency of the complement tense on the matrix tense.



John said that Mary is pregnant.

As pointed out earlier, in order to correctly interpret the complement of the present-under-past construction like (2) in section 2.1, which is repeated under the diagram in (15), the time of the complement clause T'' must encompass both the time of the matrix event T' and the speech time T^* . Enç provides the following evidence to show this point. (Enç *ibid.*, ‘?’ mine.)

(16) ?John heard two years ago that Mary is pregnant.

Enç explains that the above sentence is anomalous because the time of Mary's pregnancy cannot encompass the two times, given the normal length of human pregnancies. Then, what is responsible for the DA interpretation of the present-under-past construction is the inclusion relation between the matrix event and the embedded eventuality, which also encompasses the speech time.

Enç's observation on DA reading is consistent with our earlier observation on the simultaneous reading of the past-under-past constructions. In order for us to obtain the referential dependency of the complement tense on the matrix tense, it seems that there has to be temporal overlapping or inclusion relationship between the complement situation and the matrix events/situations. Such overlapping or inclusion relationship may partly be determined by our knowledge of the world as in the case of the contrast between (15) and (16), but also crucially depends on the aspect of the situation in the complement clause as discussed in the case of (14) and (14'). The complement situations that allow the simultaneous reading that we have seen so far are either an individual-level property of the subject and/or a stative situation, neither of which are eventive. Both stative predicates and individual-level predicates refer to unbounded homogeneous state. If these properties of the complement predicate give rise to the seeming referential dependency of the complement tense on the matrix tense in the past-under-past constructions, we would expect that having any predicates with such properties in the complement clause may induce the SOT effect. This hypothesis will further be tested with various predicates in both English and Japanese. Before we leave this section, let us turn to more evidence that shows the validity of our hypothesis.

Zagona (1997) has noticed that apparent referential dependencies between times of two events may be found in complement structures which lack tenses. See the following examples from Zagona (1997: 44).

(17) Sue heard [that Mary was pregnant] (Inclusion/Precedence/*Subsequence)

(18) Sue heard about [Mary's pregnancy] (Inclusion/Precedence/*Subsequence)

In both (17) and (18), Mary's pregnancy may be interpreted to be temporally contemporaneous with or before Sue's hearing about it, but neither (17) or (18) has an interpretation in which the subordinate eventuality follows the matrix eventuality in time sequence. In other words, despite the fact that (18) lacks the subordinate tense, the same temporal dependency that obtains between the matrix 'hearing' and the complement eventuality in (17) holds in (18) as well. This clearly indicates that this dependency cannot be ascribed to the dependency of the finite tense in the complement clause.

Based on the above data, Zagona argues that there are two distinct types of dependencies: 1) event-event dependencies; and 2) event-evaluation-time dependencies⁶, and the former type of dependencies as exemplified by (17) and (18) above is not something that is to be accounted for by the theory of tense. Specifically, she argues that there is no tense-specific subtheory in the grammar, and that event-event dependencies in particular, should fall naturally from the aspectual nature of the events of the complement clause and the event-evaluation-time dependencies is explainable solely by the binding theory, which exists independently of the tense system.

⁶ This is represented by occurrences of *would* in complement clauses in English.

The present study shares Zagana's (1997) view on the independence of the event-event dependency from the tense-specific module of the grammar, and agrees with her position that we need to incorporate aspect to account for the natural language tense phenomena. This thesis, however, takes a position that we do need a system of temporal reference as an independent module of natural language.

One interesting fact about the interpretation of tense in (17) is the lack of reading in which the complement tense refers to the time subsequent to the time denoted by the main verb. If we take the function of past tense in English to simply locate the event or situation prior to the speech moment, this cannot be explained. However, we do not want to ascribe this fact to the nature of tense, either, because lack of such reading also applies to (18), which does not have finite tense in the complement. The present study will also investigate what is responsible for the lack of subsequence interpretation of the complement event in constructions like (17) and (18).

4 Theoretical frameworks and hypotheses

4.1 Theoretical framework (Klein, 1994)

For the purpose of description of the functions and behavior of temporal morphemes, this study adopts the basic assumptions of the theory of temporal reference advanced by Klein (1994). One motivation for adopting Klein's theory of temporal reference is that the system uses the same notational device for both representation of tense and representation of aspect as grammatical categories. Thus, it allows us to unify tense and aspect, which function complementarily in temporal representation of events and situations in natural language. Another reason for choosing Klein is that the theory

incorporates important facts about semantic nature of predicates that crucially affect the interpretation of temporal expressions.

Klein (1994) introduces three basic components for the system of natural language temporal representation: TU (time of utterance), TT (topic time), and TSit (time of situation). This tripartite system of temporal representation has apparent resemblance to the well-known system of tense representation advanced by Reichenbach (1947). TU corresponds to S, or speech time, and TSit roughly corresponds to E, or event time of the Reichenbachian and the neo-Reichenbachian theories of tense (Reichenbach 1947, Hornstein 1990). However, a closer examination of the two systems reveals many notable differences between them. First of all, the nature of TT or topic time in Klein's system is fundamentally different from that of the R point or reference time in Reichenbach's system. TT is defined by Klein as THE TIME SPAN TO WHICH THE SPEAKER'S CLAIM ON THIS OCCASION IS CONFINED (Klein 1994: 4, capital letters his). In contrast, the R point was never clearly defined by Reichenbach. Many researchers seem to interpret R point to be the secondary deictic center (with the primary one being the speech time) in the representation of the complex tenses such as the past perfect and the future, and is descriptively something that distinguishes the present perfect from the simple past in English (cf. Comrie, 1976, among others).

Adopting the notion of TT in place of the R point in the three-point temporal reference system will solve many recalcitrant problems that are associated with the unclear nature of the R point (cf. Binnick, 1991; Klein, 1994). For example, Reichenbach explains that the R point for a sentence like *Peter had gone* is some time between the time when Peter went and the speech time, and that it is determined by the context. However,

the R point for a simple-past sentence like *Peter went to the party* is surely not a time point between E and S. In fact, the R point in the representation of plain past sentence is placed contemporaneous with the event time, and therefore is somewhat superfluous. The R point for a present-perfect sentence like *Peter has gone*, on the other hand, is contemporaneous with the speech time. Although the distinction between the above three temporal expressions in English is made by different positioning of the R point in Reichenbach's system, it is not entirely clear how we determine the location of the R point for each and why it should be that way. If we replace the R point by TT in Klein's theory of temporal reference, we would have an account for why TT should be placed at the respective position in the temporal representation of each of the above sentences. For example, TT for the simple past is placed anterior to the time of utterance because the speaker's claim on the described situation should be confined with some time period before the speech moment. In contrast, TT for the present perfect is placed simultaneous with the utterance time because the speaker's claim on the described event is made on the time period including the utterance time. Thus, we have motivations for different placing of TT in Klein's theory, which was not the case with Reichenbach's R point.

The most significant advantage of Klein's theory of temporal representation over Reichenbach's system, however, is that the system enables us to unify the notions of tense and aspect, while clarifying their distinct natures and functions. In contrast to Reichenbach, who apparently assumed that all the temporal morphemes of English that indicate ordering relationship of events and situations are tense forms, Klein contends that they encode both tense and aspect. Klein arranges his three primitives: TU, TT, and TSit, to represent the distinct temporal relations expressed by tense and aspect

respectively, which will be introduced below. Although Reichenbach's analysis of temporal morphemes in English involves the analysis of the progressive and the perfect forms, aspectual distinctions among temporal morphemes were only implicit in Reichenbach's system, and neither the progressive nor the perfect were identified as encoding aspect in Reichenbach (1947).

In Klein's system of temporal reference, relative locations of events and situations described by sentences with respect to a given deictic center (which is usually the utterance time) are represented by the relation between TT and TU. This is a definition of tense. In contrast, the way in which the described events and situations are set against a given time frame, which we associate with the notion of aspect, is defined by the relation between TT and TSit. With this notational device, for example, the simple past-tense form in English encodes the relation TT BEFORE TU⁷ as a tense form, and TT AT TSit as an aspectual form. A complex tense like the present perfect encodes the relation TT INCLUDES TU⁸ as a tense form, and the relation TT AFTER TSit as an aspectual form.

Another major difference between Reichenbach's system and Klein's system is that the latter incorporates different aspectual nature of lexical content of a predicate to be part of the theory of temporal reference. Klein distinguishes three types of situations: 0-state, 1-state, and 2-state. What he calls 0-state refers to situations that do not entail any sort of change over time (e.g., *The book is in Russian*). In my understanding, this characterization of 0-state situations roughly corresponds to the nature of situations

⁷ Klein's actual representation for the past tense is TU AFTER TT. However, I will represent this relation as TT BEFORE TU throughout this dissertation, as positioning TT relative to TU is intuitively more appealing to me than the reverse.

depicted by individual-level predicate in Carlson's (1977) sense. Since there is no change in time implied in the nature of the predicate, sentences with this group of predicates may not involve any contrast within and outside of TT.

In contrast, 1-state refers to a situation that has some implication of change over time (e.g., *Mary is in the room*⁹). One crucial difference between 0-state and 1-state is that since the lexical content of the latter implies change in state over time, there may be contrast between inside and outside of TT. For example, sometime after one said, "Mary is in the room", Mary may still be in the room but it may also be the case that she is not there any longer. This is different with 0-state such as *The book is in Russian*. If the book is Russian, it was in Russian, and it will be in Russian as long as it exists. While TT of both 0-state and 1-state are unbounded, 1-state is different from 0-state in that TSit of 1-state can be confined within some definite time span. Thus, it is possible to conceive of contrast between inside and outside of TT for 1-state at least by inference, but not for 0-state.

Finally, 2-state refers to situations described by predicates that involve definite change of state in their lexical content. Thus, TT of 2-state involves a state contrast in itself. For example, a sentence *Mary left* involves the first state in which Mary hasn't yet left and the second state in which Mary has already left. Klein call the first state SOURCE

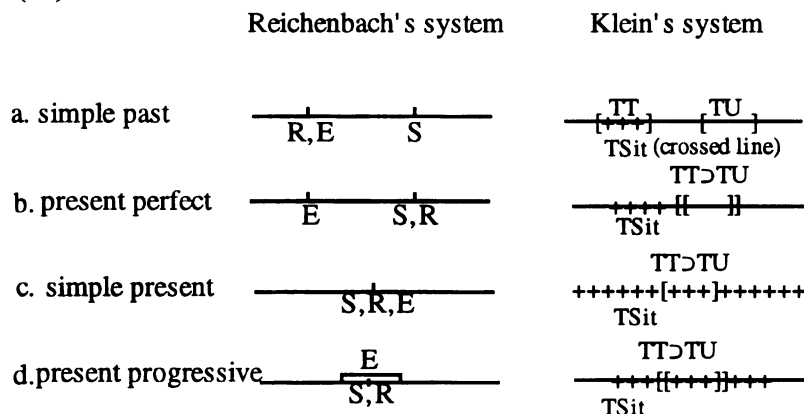
⁸ Klein represents this relation as TU INCLUeded in TT, but for the same reason as I stated above in the footnote 5, I will represent this as TT INCLuding TU throughout this dissertation.

⁹ The situation depicted by the lexical content of the sentence *Mary is in the room* does not entail change over time in itself, but has implication for change: Mary may be in the room at the time of utterance, but she may not necessarily be in the room after a while. Such implication does not hold with a 0-state situation like *The book is in Russian*.

STATE (SS) and the second TARGET STATE (TS). Then, 2-state refers to a situation that involves SS and TS in the lexical content.

Based on these assumptions, let us see how basic tenses and aspects in English are represented in Klein's system in contrast to Reichenbach's representations.

(19)



In Klein's representation on the right, topic time (TT) and utterance time (TU) are represented by square brackets, and situation time (TSit) is shown by line drawn with crosses. In (19b) and (19d) the outer square bracket indicates TT and the inner bracket indicates TU, and $TT \supset TU$ reads as Topic Time includes Utterance Time. We assume that the situation represented here is 1-state such as a situation described by a sentence *John smokes*. In the present tense form, the predicate of 1-state like *John smokes* refers to a timeless property, and therefore, TT is not confined and is equated with the lexical content of the proposition, and it properly includes TU, which is indicated by square bracket. In this system, simple past in English encodes the *perfective* aspect, and is defined as TT AT TSit, as shown by the diagram on the right of (19a). In contrast, the present perfect encodes the perfect aspect, which is defined as TT AFTER TSit as shown

on the right of (19b). According to Smith (1997), the term ‘perfective’ refers to a closed aspect, which encompasses the beginning point and the end point of a situation, while ‘perfect’ constructions generally convey that the situation precedes Reference Time and that they have a resultant stative value. If we replace Reference Time with Topic Time in the above description of the ‘perfect’ aspect, this characterization will fit Klein’s definition of ‘perfect’ in English.

Klein’s representations of the grammatical aspects of English as shown above capture our intuitions on aspectual distinctions, which have been discussed on numerous occasions in the previous work on aspect but never received anything more than metaphorical descriptions. As pointed out by Comrie (1976), the perfect provides a view of situations as if we are looking them from outside, while the progressive provides a view of situations as if we are looking them from inside. If we see how TT is placed with respect to TSit in the above diagram for the perfect and the progressive respectively, it seems to be obvious why we perceive them as such. In the case of perfect, TT, or the time span to which the speaker’s claim on the described situation is confined, is placed posterior to TSit, which is the time associated with the situation represented by the lexical content of the proposition. Thus, we necessarily represent a view of the situation described by a verb from outside of the situation when we have perfect marking on a verb. In contrast, with the progressive morpheme on the predicate, we place TT within the situation described by the lexical content of a verb. Thus, we inevitably obtain the view of the situation from within the time span associated with the lexical content of a verb. Klein’s characterization of simple past as encoding perfective aspect also has an intuitive

appeal in explaining why we seem to take events or situations as indivisible wholes when we use simple past.

Another strength of Klein's theory lies in that Klein's distinction of the situation types described by the lexical content of predicates accounts for some of the puzzles concerning the distribution of the progressive form in English. It is well attested that states in general resist being changed into the progressive form (Vendler 1967, Dowty 1979). As far as I am aware, stative predicates in English fall into either property-denoting generic predicates or predicates denoting stage-level property of a described entity. Take a sentence *The book is in Russian*, for example. This is what Klein calls 0-state situation, and the predicate *be in Russian* is a property-denoting term. The relation among TT, TU, and TSit for this sentence is schematized in (20).

(20) TT⊃TU (The square bracket represents TU)
 +++++++[++++]++++++
 TSit

The book is in Russian.

According to Klein, lexical contents of 0-state do not have any TT-contrast. He explains this as 'if they are linked to a particular TT, then they are automatically linked to any other TT' (Klein, *ibid.*: 101). TSit of 0-state extends over the entire time, and therefore, TT is always included in TSit. Since the progressive aspect encodes the relation TT INCLUDED in TSit, adding this aspectual form to any predicate of 0-state content will not change the aspectual nature of the proposition. Thus, the progressive counterpart of a sentence in (20), even if it were a grammatical sentence, would represent the situation in

exactly the same way, and hence, we have no reason to have variations with this aspectual form for 0-state situations.

Let us now turn to an eventuality described by a sentence with a stage-level predicate like *Mary is in the room*.

- (21)
$$\begin{array}{c} \text{TT} \supset \text{TU} \\ \text{-----} \text{+++} [[\text{++++}]] \text{+++} \text{-----} \\ \text{TSit} \end{array} \quad \begin{array}{l} \text{(The outer square bracket represents TT and} \\ \text{the inner square bracket represents TU)} \end{array}$$

Mary is in the room.

A stage-level predicate like *be in the room* corresponds to what Klein calls 1-state predicates. While both 0-state and 1-state do involve TT contrast, the situation described by the lexical content of a 1-state predicate has both the beginning and the end point. However, as shown by the diagram in (21), the proposition *Mary is in the room* already encodes the relation between TT and TSit that would be obtained by the progressive aspectual marking on the predicate. This is due to the nature of this predicate, which refers to a stage-level property of the subject. A stage-level property, as characterized by Carlson (1977), refers to a stage, or one of the realizations of a described individual. Since the progressive is the function to derive a developmental stage of a described individual, the modification of 1-state expressed by stative predicates that denotes stage-level properties with the progressive form would be redundant. In contrast, the 1-state expressed by non-stative verbs such as *sleep* will have a progressive counterpart, as a predicate like *sleep* does not encode the relation $\text{TT} \subset \text{TSit}$. 1-state situations described by non-stative verbs may indicate a stage-level property of an individual only when they are

in the progressive form, which modifies TT of the proposition so that it will be confined within TSit associated with the lexical content of the predicate in its bare form.

Simply representing the event time E as an extended line rather than a point as shown in Reichenbach's representation of the present progressive in (19d) does not provide any explanation for the incompatibility of statives with the progressive form.

As a final piece of support for Klein's system of temporal representation I would like to point out an interesting fact that is observed when 0-state sentences are expressed in the past tense form.

(22) TT TU
 +++++[++++]+++[++++]++++
 TSit

The book was in Russian.

As in the case of the present-tense counterpart in (20), situations described by 0-state sentences do not involve TT contrast. However, (22) differs from the present-tense counterparts in that it does not indicate the present relevance of the situation despite the fact that the situation expressed by the lexical content of the proposition may still hold at TU. This is due to the nature of past tense, which is defined as TT BEFORE TU. As this condition necessarily makes TT and TU disjoint, we have an effect that TT is confined within some time frame which necessarily precedes the time of utterance. This effect of the right boundedness of the situation in the past tense form is captured nicely in Klein's system of temporal representation.

Although the above points were not explicitly made in Klein (1994), I take these to be crucial arguments for choosing Klein's theory over Reichenbach's theory of tense.

The fact that the system can account for the effect of eventive reading for individual-level predicates in the past tense form along with the fact that it explains the restrictions on the occurrence of aspectual forms such as the progressive form without any further stipulations definitely points to the explanatory power and the soundness of the system.

1.4.2 Hypotheses and more questions

With the notational device introduced above, I will assume that temporal morphemes in English and Japanese encodes the following distinctions in tense and aspect respectively.

<u>English</u>	<u>TENSE</u>	<u>NOTATION</u>
<u>...s</u>	<u>PRESENT</u>	<u>TT INCLUDES TU</u>
<u>...ed</u>	<u>PAST</u>	<u>TT BEFORE TU</u>
<u>will</u>	<u>FUTURE</u>	<u>TT AFTER TU</u>

Table 1: Basic tenses in English

As shown in Table 1, I will assume that English exhibits three-way tense distinctions. To assume that an auxiliary *will* is a tense marker is controversial, and some researchers argue that it should rather be considered as a modal auxiliary than a tense marker (e.g., Enç 1989). While there may be many occasions when *will* expresses an intention of a subject, it is evident that *will* locates TT after TU with some systematicity. Thus, to the extent that modal-like occurrences of *will* are not incompatible with the function of this auxiliary to place TT posterior to TU, I will assume that *will* encodes the future tense, though it is most natural to believe that it may also encode some modal property. Let us now turn to the tense system in Japanese.

Japanese	TENSE	NOTATION
<i>-ru</i>	NON- PAST	TT NOT BEFORE TU
<i>-ta</i>	PAST	TT BEFORE TU

Table 2: Basic tenses in Japanese

Table 2 shows the tense distinctions in Japanese. In contrast to the English tense system, the Japanese tense system has only two-way distinctions. It does not have a grammatical device to uniquely indicate the future time: the morpheme *-ru* may indicate either PRESENT or FUTURE. That is to say, *-ru* may encode either the relation TT INCLUDES TU or the relation TT AFTER TU as a tense marker, and contrasts with *-ta*, which encodes the relation TT BEFORE TU as a past-tense marker.

As I mentioned earlier in section 1.3.2, some researchers claim that the properties of *-ru* and *-ta* are exclusively aspectual (e.g., Ando 1986). I have already presented some basis for disagreement to such a claim in section 1.3.2, where I stated that if *-ru* and *-ta* indicate incompleteness and completion of the embedded event with respect to the time of the main event respectively, we do not have account for the occurrences of these morphemes in sentences like (9) and (10), where *-ta* co-occurs with the event which has not yet happened at the time of the main event and *-ru* co-occurs with the event which precedes the main event. In fact, there are even more obvious counter-examples to the claim that the nature of these morphemes are exclusively aspectual.

(23) *Kinoo nihon-ni ik-u
yesterday Japan to go-RU
(I/he/she) was going to Japan yesterday.

(24) *Asita nihon-ni it-ta
tomorrow Japan to go-TA
(I/he/she) will have gone to Japan tomorrow.

If *-ru* is exclusively an imperfective marker and does not encode tense, how should we explain the fact that it cannot express an imperfective (or incomplete) situation at some time in the past as shown in (23)? Likewise, if *-ta* is exclusively a perfective marker, why is it that *-ta* cannot be used to describe a situation which is to be completed in the future time as in (24)? Although there are data which are inconsistent with the view that *-ru* and *-ta* are exclusively tense markers,¹⁰ since such is not the claim of the present study, the burden of proof seems to be on those who attribute only aspectual nature to these morphemes.

Given that *-ru* and *-ta* have properties to indicate relative location of TT of the sentence with respect to a given anchor (i.e., TU in Klein's system) as tense markers, the descriptive task of the present dissertation is to show how this system can explain the problems we have seen in section 1.3. All the data (1) through (10) will receive natural explanation if we employ a concept of shift of the deictic center for tense interpretation (=TU) of the embedded clause to TT of the matrix clause, which is obligatory to complement clauses and optional for adjuncts.

The shift of the deictic center for the tense interpretation have been discussed elsewhere in the previous literatures on tense. Some researchers call the system which employs such a shift as the 'relative tense system' in contrast to the absolute tense system in which the deictic center or the anchor for the tense interpretation always stays the same (e.g., Comrie 1985). Others labeled such shift as 'shift in viewpoint', but

¹⁰ For example, the occurrence of *-ru* in so-called historical past, and the occurrence of *-ta* in a sentence like *Aa, tukare-ta!* 'Oh, I'm tired!' The latter is considered to be a modal use of *-ta* (Kindaichi 1976).

without defining what the ‘viewpoint’ is (e.g. Soga 1983). The present study attempts to provide more explicit characterization and seeks for a motivation for such shift.

One problem in Table 1 is the treatment of the present tense. It is not at all obvious if English has the present tense as a grammatical category. In both English and Japanese, non-stative verbs typically do not indicate the situation which has the present-time relevance except as denoting an individual-level property, and only statives such as *know* in English and *iru* ‘exist’ in Japanese typically indicate the present time with the present-tense morpheme (for English) or non-past tense morpheme (for Japanese). Nevertheless, I will tentatively assume that the present tense morphemes in English are indeed tense markers, and investigate if such an assumption can adequately account for their distribution.

Moving onto the aspectual distinctions, we hypothesize the following aspectual distinctions among the temporal morphemes of English and Japanese, and will verify whether or not such classification is a plausible one.

English	Japanese	ASPECT	NOTATION
<u><i>be...ing</i></u>	<u><i>-te-i</i></u>	IMPERFECTIVE	TT INCLUDED in TSit
<u><i>...ed</i></u>	<u><i>-ta</i></u>	PERFECTIVE	TT AT TSit
<u><i>have...ed</i></u>	<u><i>-te-i</i></u>	PERFECT	TT AFTER TSit

Table 3: Aspectual distinctions in English and Japanese

Table 3 should be understood to represent the aspectual property encoded by each morpheme. The functions of each morpheme that may not be purely temporal are not being considered for that matter. For example, it is possible that these morphemes have modal properties or even other functions as well, but the present study does not concern

those functions. Since Table 3 is by no means to claim that there is a one-to-one correspondence between the morphemes and the given aspectual information, identifying *-ta* in Japanese to encode past tense in Table 2 and having the same *-ta* to encode perfective aspect in Table 3 are not contradictory. Also, it is not my intention to claim that the progressive in English is a grammaticized imperfective form. There is a well-established distinction between the progressive and the imperfective aspect, and I do respect the distinction. However, I follow Comrie (1976) in assuming that the progressive is a type of the imperfective aspect, and in that sense, I believe that progressives in English encode TT-TS_{it} relation associated with imperfective aspect in terms of their purely temporal property.

The Table 3 also shows that *-ta* in Japanese encodes perfective¹¹ aspect, and contrasts with *-te-i*, which indicates imperfective aspect. This radically departs from the traditional view in which *-ta* as a perfective marker contrasts with imperfective *-ru*. Among the previous studies on temporal morphemes in Japanese one that is somewhat similar to the present analysis of *-ru* and *-ta* was given by Machida (1989), though he assumes that meanings indicated by these morphemes diverge between stative and non-stative predicates; i.e., *-ru* and *-ta* indicate perfective aspect with non-stative predicates and imperfective aspect with stative predicates. I will argue that the fact that *-ru* and *-ta* seem to show imperfective aspect with stative predicates naturally follows from the nature of statives (i.e., incapable to have TT contrast), and therefore, the aspect born by these morphemes stays constant across different types of predicates they are attached to.

¹¹ The term ‘perfective’ and ‘perfect’ are defined by the notations given in Table 2, following Klein’s terminological definitions.

One major difference between English and Japanese in morphological realization of the aspectual distinction that is shown in Table 3 is that two totally different aspects; perfect and progressive (imperfective), which are realized as different morphemes in English, are encoded by a single morpheme in Japanese, namely, *-te-i*¹². This interesting fact has inspired many researchers to investigate this aspectual form *-te-iru* (e.g., Kindaichi 1976, Kunihiro 1982, Teramura 1984, Kudo 1989, Jacobsen 1991, McClure 1993, Ogiwara 1997, among others), but none has so far been successful in explaining why the form can express such distinct meanings, nor could they unify various different occurrences of this aspectual form. A proposal on the aspect encoded by *-te-iru* that is similar to the current proposal was given by Kudo (1989), in which she argued that *-te-iru* may encode either ‘durative’ aspect or ‘perfect’ aspect, and that it contrasts with *-ta*, which indicates that an event is an indivisible chunk. I basically agree with her intuition and the discourse functions of *-te-iru* presented in her study. However, Kudo (1989) does not have an account for why *-te-iru* may encode two different aspects. The present study attempts to provide an account for why this form may represent such distinct notions as perfect and imperfect, and proposes a single semantic property that unifies all the occurrences of this aspectual form. I will argue that *-te-iru* is a stage-level operator that takes an eventuality described by the proposition and convert it to a description of one of the realizations of an individual in Carlson’s (1977a) sense. I will argue that such unitary semantic nature that embraces all the occurrences of this form is naturally derived from the semantics of *-te* and the auxiliary *iru* of *-te-iru*. Specifically, I will propose that

¹² Since this form is usually referred to as *-te-iru* with the non-past tense morpheme *-ru* in most of the literature on this form, this study will also follow this tradition and will use

the ambiguity of the sentences with this form between a perfect and imperfective reading can be attributed to the underspecified temporal nature of the morpheme *-te* of *-te-iru*, and that the existential or stage-level property is provided by the auxiliary *iru*.

1.5 Summary

In this chapter, I have shown that in spite of the fact that research on tense in both English and Japanese is abundant, the rules governing the interpretations of temporal morphemes in English and Japanese are not at all clear. It seems to be the case that one needs to incorporate the theory of aspect in order to obtain more complete picture of the problematic tense phenomena such as a sequence of tenses. Furthermore, tense and aspect, which, I believe, are complementary in the system of temporal reference and representation of events and situations in natural language discourse, must be treated in one whole picture for us to truly understand their respective role in the grammar. The theory of temporal reference advanced by Klein (1994) seems to be a strong tool for description of functions of the temporal morphemes in English and Japanese by allowing us to unify the notions of tense and aspect. In the next chapter, we will examine the previous accounts of the sequence of tense phenomena, and summarize their findings and the problems in order to form theoretically and descriptively more sound foundation for our analysis of the functions of temporal morphemes that will follow.

-te-iru, unless we specifically talk about the property of this complex morpheme independent of *-ru* or the past-tense variation of this form, *-te-i-ta*.

CHAPTER 2

Theories of Temporal Reference and Sequence of Tense

Introduction

This chapter examines the existing theories of tense and their accounts of the SOT data, and summarizes their findings and the problems. Part of the aim of this chapter is to sort out what is and what is not to be included in the theory of temporal reference. By separating the principles and the constraints that interact with the system of temporal reference, but belong to other modules of the grammar, the system of temporal reference will become considerably simpler than what some of the existing theories claim it to be. Another aim of this chapter is to provide sound foundations for the account of the behavior of temporal morphemes by eliminating some of the inadequate generalizations of the previous studies and by verifying others that are both empirically and theoretically well grounded.

Section 2.1 will go over the problems of the traditional analysis of SOT, and section 2.2 examines Enç's solution to some of those problems. Section 2.3 and 2.4 will review the approaches to SOT and DA reading of present-under-past by Ogihara (1989, 1995a,b, 1997) and Abusch (1988, 1991, 1994, 1997a,b), which incorporate a theory of intensionality into the theory of tense. These sections will show that interpretation of tense may involve a semantic concept of intensionality, which is independent from the tense system but interacts with it in an important way. Section 2.5 will review Costa's (1972) account of SOT to draw our attention to pragmatic factors to be considered in interpretation of temporal expressions. Section 2.6 discusses Hornstein's (1990) theory, which contains a number of important insights on the nature of tense. Hornstein presents

properties of tense which are not those of operators but more like those of adverbs, and this thesis shares his view on the nature of tense. Section 2.7 will examine Nakamura's (1994) account of the different behavior of tense morphemes in Japanese and English, which is based on the syntactic theories of tense proposed by Stowell (1993, 1996) and Zagana (1990). While this study has nothing to say against Stowell's and Zagana's syntactic approaches to tense, it will be shown that Nakamura's analysis of the nature of tense morphemes of Japanese turns out to be problematic in many respects. Finally, section 2.8 will summarize the generalizations obtained in the chapter.

2.1 Problems with the traditional analysis

One of the problems of the traditional analysis of SOT is that if the SOT rule is a syntactic rule that simply transmits the matrix past to the embedded past, one would expect that it should apply uniformly to different constructions with embedded clauses, but it does not. As pointed out in section 1.3.1, there are some differences in the interpretation of embedded past in the complement construction and the relative-clause construction. Another problem is that the rule has only very limited applicability, which makes it hard for us to believe that it is part of the grammar. The most serious problem may be the one that has been pointed out by Enç (1987): there is a difference in the interpretation between the past-under-past and the present-under-past constructions, which should not be expected if the assumptions of the traditional analysis were correct. Below, we will briefly go over this problem to provide a basis for the later discussions on this issue.

The traditional analysis of SOT assumes that the simultaneous readings of the past-under-past construction, which is repeated here as (1) from section 1.3.1, has the underlying structure as shown in (1') before the application of the rule.

- (1) John said that Mary was pregnant.
 John said, "Mary is pregnant." (simultaneous reading)
- (1') PAST [John say [PRES [Mary be pregnant]]]

According to the traditional analysis, the SOT rule applies to the structure in (1'), which has present tense in the complement clause, copies the past tense morphology of the matrix verb onto the complement verb, and we will have the surface form as shown in (1). In contrast, an output of the non-application of the SOT rule will generate a sentence with the present tense in the complement as a surface form, which is repeated here as (2) below from section 1.3.1.

- (2) John said that Mary is pregnant.

Since (1) and (2) share the same underlying form (1'), we expect that (1) and (2) would have the same interpretation. As pointed out earlier in section 1.3.1, however, there is a difference between (1) and (2) in their interpretations. The present-under-past construction like (2) has an interpretation which is not shared by (1), namely, the double-access reading. If we assume that (1) and (2) have the same underlying structure shown in (1') as claimed by the proponents of the traditional analysis of SOT, this is not explained.

2.2 Enç's (1987) 'Anchoring Conditions' for tense

As a solution to the above mentioned problem with the traditional approach to the SOT phenomenon, Enç (1987) proposes a syntactic account of SOT, drawing a basic insight from the referential treatment of tense. She takes the standard view of tense in the framework of the Government and Binding theory (Chomsky 1981, 1986), and assumes that tense is in Infl in the structural configuration. Based on the observation that there are close connections between Comp and Infl (cf. Stowell, 1981), Enç assumes that the specifier of tense is located in its Comp, which governs Infl (hence tense), and proposes the following 'Anchoring Conditions' for tense as part of the grammar of natural language, which applies at D-Structures in the GB model (originally (27) in Enç 1987: 643).

(3) *Anchoring Conditions*

- a. Tense is anchored if it is bound in its governing category, or if its local Comp is anchored. Otherwise, it is unanchored.
- b. If Comp has a governing category, it is anchored if and only if it is bound within its governing category.
- c. If Comp does not have a governing category, it is anchored if and only if it denotes the speech time.

In this system, the two readings of the sentence, *John said that Mary was pregnant*, can be represented as follows.

(4) John said that Mary was pregnant.

- a. [S' Comp₀ [S NP [_T PAST_i [V [Comp_i [NP PAST_j ...]]]]]] (shifted)
- b. [S' Comp₀ [S NP [_T PAST_i [V [Comp [NP PAST_i ...]]]]]] (simultaneous)

In both (4a) and (4b) the matrix Comp, which is not governed, is anchored by satisfying the anchoring condition given in (3c) by denoting the speech time. The matrix tense is

bound by the matrix Comp, and since it has PAST, it is anchored to the time before the time of the matrix Comp, which is the time before the speech time. The Comp of the complement clause in (4a) is governed by the matrix verb, and its governing category is the matrix S. Since the matrix tense, as a proper antecedent, binds the complement Comp, both the complement Comp and the complement tense are anchored. Because of the semantic nature of PAST in the complement, the time of Mary's pregnancy is placed prior to the matrix PAST, which is the time of John's saying it, and the shifted reading is derived.

In (4b), on the other hand, the Comp in the complement clause does not have an index. This means that the anchoring of the complement tense is done by binding rather than through its Comp. Since the minimal domain that contains a subject c-commanding the governor Comp is the matrix S, according to Enç, the governing category of the complement Comp becomes the matrix clause. Thus, an embedded PAST is co-indexed with the matrix PAST through binding, and they denote the same time, which is the time prior to the time of the matrix Comp, the speech time. The simultaneous reading is hence derived.

In contrast to past-under-past constructions, the account of present-under-past constructions is not so straightforward. Since the Anchoring Conditions as they are formulated in (3) will yield a wrong result for the interpretation of a present-under-past sentence like *John said that Mary is pregnant*, Enç suggests that we should revise the analysis of present tense rather than modifying the Anchoring Conditions. To justify this line of revision, Enç points out that in languages such as Russian, the present tense under the matrix past receives the same indexing as the matrix past, and argues that present

tense is not inherently related to the speech time in such languages, while it always denotes the speech time in English. Then, in order for the present tense in the complement clause to have access to the speech time, she introduces the following reindexing rule for the languages like English, in which the complement present tense denotes the speech time.

- (5) At LF, change the referential index of the present tense and its Comp to 0.
(originally (40) in Enç, 1987)

The addition of the condition in (5) to the system will allow the complement present tense to receive the speech time as its referential index. This solves one problem, but the system still needs to account for DA reading of the present-under-past. Enç argues that the inclusion relation between the complement present and the time denoted by the matrix past is to be determined not by the Anchoring Conditions in syntax, but by an interpretive rule in the semantics, which determines the antecedent of temporal expressions. She stipulates that “all temporal expressions carry a pair of indices, the first one identifying their referent, and the second establishing their link to other referents” (1987: 651), and argues that “when two temporal expressions in a sentence share a second index, the denotation of the lower one will be included in the denotation of the higher one” (ibid.). Now let us see how a present-under-past sentence, *John said that Mary is pregnant*, can receive DA reading in Enç’s model.

- (6) $[_{s'} \text{Comp}\langle 0, i \rangle [_{s'} \text{PAST}\langle j, k \rangle [_{s'} \text{Comp}\langle 0, k \rangle [_{s'} \text{PRES}\langle 0, k \rangle]]]]$
(originally (44) in Enç, 1987: 652)

(6) is a schematic LF representation of the present-under-past construction. By application of the Anchoring Conditions, the complement present is bound by the matrix past, but the referential index of the present tense was rewritten to 0 at LF by (5). Here, by stipulation, only the first index (which determines the referent of the time expression, according to Enç) is rewritten, thereby leaving the second one still be the same as that of the matrix tense (which was derived by the application of the Anchoring Conditions). In (6), the time denoted by the complement present is included by the time denoted by the matrix past because they share the same second index. However, this would yield a wrong result: PRES cannot be included in PAST to denote a time interval that includes the moment of utterance. In order to avoid this problem, Enç further stipulates that the complement present must scope out at LF as shown below.

$$(7) \quad [s \cdot \text{Comp}_{\langle 0, j \rangle} [s [s \cdot \text{Comp}_{\langle 0, k \rangle} [\dots \text{PRES}_{\langle 0, k \rangle} \dots]] [s \text{ NP } [\text{PAST}_{\langle j, k \rangle} [V \ e]]]]]$$

(originally (45) in Enç, 1987: 653)

The complement S' is adjoined to the matrix S in the above structure. Now that the complement present is out of the scope of the matrix past, it receives the time interval including the moment of speech as its referent, according to its first index 0. The matrix past, which now appears after the complement present in its sister node in linear order, is included in the denotation of the present tense. Hence, we successfully derive the DA reading for the present-under-past construction.

Aside from the many stipulations that she had to make to get her Anchoring Conditions to work, Enç's theory seems to suffer some serious problems. My first objection to her proposal is an intuitive one. In Enç's theory of temporal reference,

indexing of tense and Comp (aside from the re-indexing mechanism for the interpretation of present-under-past we discussed above, which is done at LF) is done in the semantics by a set of semantic rules provided specifically for tense. The interpretation of tense is determined by the Anchoring Conditions together with this set of semantic rules. Assuming that the conditions apply at the D-Structure level, as Enç proposes, the semantic rules for indexing, as an interpretive rule, should apply to the output of application of the Anchoring Conditions at LF in the model provided in the GB framework. However, her proposal seems to suggest that different interpretations of the complement tense provide different indices to its Comp, and different syntactic treatment is required accordingly. Take (4) for example; for the Anchoring Condition to correctly predict the desired interpretations, we need an index i for the lower Comp and another index j for the embedded past for a shifted reading, while giving the lower Comp 0-index for a simultaneous reading. This must be done by semantic rules before the Anchoring Condition applies. This derivational process goes against the basic assumptions of the Government and Binding theory of grammar which her study takes as a framework.

Second point of disagreement comes from the articulation of the theoretical details of her system. Enç's theory explains the tense dependencies in complement clauses as anaphoric co-reference, which is, at least partially, determined by the binding theory. Since we have so much evidence to believe that the binding theory exists as part of the natural language grammar, if we can explain the interpretation of tense by the binding theory, which is supported independently of the tense system, it is a definite theoretical advantage that Enç's theory can enjoy over the traditional analysis. However, the definition of the governing category, which is of a central importance in her system is not

necessarily clear. Zagona (1997) points out that assuming that the definition of governing category is a Complete Functional Complex in Chomsky (1986), governing category for the embedded tense in (4b) above must be the lower clause (IP or CP), which has subject NP and the Comp, whichever is taken as the argument to close off the functional complex, rather than the matrix clause. I do not quite agree with Zagona in this respect, since taking the matrix clause as the governing category for the embedded tense in (4b) could be maintained with the notion of *BT*-compatible indexing, which is also defined in Chomsky (1986). As I pointed out above, however, what she takes as governing category is not explicitly laid out in Enç (1987), and therefore, the problem pointed out by Zagona remains unsolved.

The most serious problem with Enç's theory, however, is an empirical one. As pointed out by Ogihara (1989), the theory cannot deal with the occurrence of future tense embedded under past tense. Ogihara presents the following example (1989: 157, originally (97)) to illustrate this point.

- (8) I told Bill that you would say that you *only had three magic tricks to do*, but it looks as if you have brought enough equipment to do six or seven.

Although Enç (1987) does not treat the English future auxiliary *will* as a tense morpheme, it is obvious that the future time denoted by the auxiliary *will* also subjects to SOT when embedded under past tense, and thus, its behavior needs to be accounted for. In the above example, the time of 'your saying' is subsequent to the time of 'my telling this to Bill', and the time of 'your having three magic tricks to do' is taken as simultaneous with the time of 'your saying'. Ogihara applies Enç's theory to this example, and provides the following two options for indexing (*ibid.*: 159).

- (9) a. I PAST_t tell Bill that you PAST_t will say that you only PAST_t have three...
- b. I PAST_t tell Bill that you PAST_{t'} will say that you only PAST_{t'} have three...
where $t' < t$

Whichever option we take, the fact that the embedded *saying* is subsequent to the matrix *telling* cannot be derived.

In sum, although Enç's theory of tense has definite advantages over the traditional account of the SOT phenomenon in that it eliminated the undesirable SOT rule and is attractive in a sense that it reduces the account of anaphoric nature of tense morphemes to generalizations obtained by the binding theory, it has both theoretical and empirical problems, and cannot be maintained to be part of the natural language grammar.

2.3 Ogihara's (1989, 1995a,b) account of SOT

2.3.1 Referential vs. quantificational analyses of tense

Based on the data from Japanese, which is a non-SOT language, Ogihara (1989) proposed a new theory which accounts for the SOT phenomenon in English and its absence in Japanese.

Unlike Enç (1987), who supports the referential theory of tense, Ogihara (1989, 1995a,b) takes the position of the quantificational theory of tense, which assumes that interpretation of tense in natural language involves existential quantification over times. With some apparent counter-arguments to the quantificational theory of tense by the proponents of the referential theory in mind, Ogihara provides a revision of the quantificational analysis of tense of the type developed by Dowty (1979) with the notion of *de re* attitude reading, which was first discussed by Quine (1956), and formalized for

the interpretation of tense by Abusch (1988, 1991, 1994) about the same time with Ogihara's work. Details aside, what Ogihara attempted to do is to combine the explanatory forces of the quantificational theory with those of the referential theory to cover the data in which the occurrences of tense seem to be referring to particular time intervals which are contextually salient rather than the unrestricted moments of time. Such a context-dependent property of tense may be most cogently represented by a famous example given by Partee (1973), with which she argued for the anaphoric nature of tense analogous to pronouns. Partee pointed out that the sentence in (10), when uttered halfway down the turnpike, does not have either of the possible interpretations provided by the standard quantificational analysis of tense. (The formal semantic representations provided below for Partee's example is from Ogihara (1989: 40).)

- (10) I didn't turn off the stove.
 a. $\neg \exists t [\text{PAST}(t) \ \& \ \text{AT}(t, \text{I-turn-off-the-stove})]$
 b. $\exists t [\text{PAST}(t) \ \& \ \text{AT}(t, \neg \text{I-turn-off-the-stove})]$

(10a) reads that there exists no time in the past at which I turned off the stove, and (10b) reads that there exists some time in the past at which I did not turn off the stove. Neither of these can be an appropriate interpretation for the sentence (10) in the above circumstances, as discussed by Partee. The sentence clearly refers to a particular time interval in the past, which must be clear in the context in which (10) was uttered.

Ogihara realized this problem with a purely quantificational treatment of tense. However, he does not support the referential theory of tense based on some empirical evidence against it, and hence, chooses to revise the quantificational theory with the addition of a theoretical device to incorporate context-dependent properties of tense to

the system. He illustrates the difference between the referential analysis and his own with formal representations of a simple past sentence *I saw Mary* (as an answer to a question by someone who is looking for a Mary) in two approaches respectively as follows (Ogihara 1995b: 667).

- (11) a. $\exists t [t < s^* \ \& \ t \subseteq t_{RI} \ \& \ \text{see}'(t, I, m)]$ (quantificational analysis + contextual restriction)
 b. $t < s^* \ \& \ \text{see}'(t, I, m)$ (referential analysis)

In order to restrict the quantificational force of the existential quantifier, which would lead to a wrong prediction if unrestricted, Ogihara introduces the function t_{RI} , which is ‘a free time variable whose value is the time interval that is salient in the given context’ (ibid.) as shown in (11a), and restrict the time when the speaker saw Mary to fall within some contextually salient interval, thereby capturing the basic insight of Partee’s (1973) referential analysis. (11a) thus reads as “I saw Mary at some time in the past (i.e., a time interval that precedes the speech time) that falls within some relevant time interval salient in a given context.” In contrast, (11b), which is a formalization of a reading of this sentence in a referential analysis, would read, according to Ogihara, as “I saw Mary *throughout* some contextually salient past interval.” (Ogihara, 1995b: 667; *italic his*). Evidently, the latter will not serve as a felicitous answer to someone who is trying to find Mary.

Ogihara’s theory of tense seems to be quite attractive so far, but it becomes somewhat quaky when it comes to an explanation of the SOT data. In the next section, we will examine how this theory accounts for the interpretation of complement tenses.

2.3.2 Tense deletion and the difference between English and Japanese

Ogihara (1989, 1995a,b) assumes that there must be a special syntactic rule in the grammar to account for the SOT phenomenon, which applies at LF after Quantifier Raising and before the semantic interpretation of tense. This is different from the traditional view of the SOT rule, as Ogihara characterizes it, since the SOT rule in the traditional analysis is assumed to apply somewhere between S-structure and PF if it is recast in the GB framework. In place of a traditional tense-copying SOT rule, Ogihara posits a “tense deletion” rule which deletes an embedded tense under identity with the immediately higher tense. This rule is formulated as follows (originally (18) in Ogihara, 1995b: 673).

- (12) A tense morpheme α can be deleted if and only if α is locally c-commanded by a tense morpheme β (i.e., there is no intervening tense morpheme between α and β), and α and β are occurrences of the past tense morpheme.

Application of this newly defined SOT rule to a past-under-past sentence *John said that Mary was sick* with a simultaneous interpretation changes the LF representation of this sentence from (13a) to (13b). (originally in Ogihara, *ibid.*: 674 (19a, b))

- (13) a. John PAST say that Mary PAST be sick
b. John PAST say that Mary $\emptyset$ be sick

Now (13) translates into (14), which is an application of Ogihara’s quantificational analysis of tense with contextual restriction.

- (14) $\exists t_1 [t_1 < s^* \ \& \ t_1 \subseteq t_{RI} \ \& \ \text{say}'(t_1, j, \wedge[\text{be-sick}'(m)])]$
(originally (20) in Ogihara, 1995: 674)

As the translation shows, the predicate in the complement clause ‘be sick’ in this sentence does not have a time argument unlike the higher predicate, and therefore is taken to be in the same set of world-time pairs as the proposition as a whole represents, and hence, the simultaneous reading is derived.

The shifted reading of the same sentence, on the other hand, is derived by adopting Lewis’s (1979) *de se* analysis of propositional attitudes, which argues that the object of an attitude verb should be understood as “self-ascription of properties” (Lewis, *ibid.*: 521) rather than a proposition (i.e., a set of worlds). Ogiwara applies this idea to his analysis of tense by assuming that the object of an attitude to be a “property of times” and that indirect discourse verbs (e.g., *say*) constitute a subtype of propositional attitude verbs (e.g., *believe*). With these assumptions, the translation of the shifted reading of *John said that Mary was sick* is represented as follows.

$$(15) \exists t_1 [t_1 < s^* \ \& \ t_1 \subseteq t_{R1} \ \& \ \text{say}'(t_1, j, \sim \lambda t_2 [\exists t_3 [t_3 < t_2 \ \& \ t_3 \subseteq t_{R3} \ \text{be-sick}'(t_3, m)])]] \\ \text{(originally (22) in Ogiwara, 1995: 675)}$$

Unlike (14), the embedded predicate in this translation has a time argument, which is placed before the ‘saying’ time and is contextually restricted. Since the time of John’s saying itself is placed earlier than the speech time, we obtain the shifted reading for the embedded tense.

Ogiwara argues that Japanese counterparts of the two interpretations of the above English sentence, *John said that Mary was sick*, the one with a simultaneous interpretation, which has a present-under-past surface form in Japanese, and the other with a shifted interpretation, which has a past-under-past surface form in Japanese, as

shown in (16a,b) below, would have exactly the same semantic representation as their respective English counterpart.

- (16) a. John-wa [Mary-ga byooki-da] to it-ta
 John-TOP Mary-NOM be sick-PRES that say-PAST
 ‘John said that Mary was sick.’ (simultaneous reading)
- b. John-wa [Mary-ga byooki-dat-ta] to it-ta
 John-TOP Mary-NOM be sick-PAST that say-PAST
 ‘John said that Mary was sick.’ (shifted reading)

According to Ogihara, (16a) translates exactly the same as (14), and (16b) as (15). He argues that ‘Japanese sentences in the present tense are tenseless sentences’ and that ‘(unlike English) Japanese lacks a tense deletion rule.’ (Ogihara 1995: 676)

Before we move onto his analysis of present-under-past constructions and the DA reading in English, let us go over the syntactic part of his analysis of past-under-past constructions. Ogihara argues that the sentence *John said that Mary was sick* undergoes a tense-deletion rule at LF, which applies only when a certain condition is met, namely, when the complement past tense is embedded under the matrix past tense. This is highly stipulative, and as far as I see, there is no independent motivation for positing such a tense deletion rule. We do not have any account for why this tense deletion does not apply to the past-under-present or the-past-under-future constructions, for example. In this respect, Ogihara’s tense deletion rule has the same empirical problem as the traditional SOT rule.

Another point of disagreement concerns his treatment of cross-linguistic difference between English tense and Japanese tense. In order to explain the absence of the SOT phenomenon in Japanese, he argues that embedded non-past tense morpheme *-ru* in

Japanese is 'tenseless', and hence Japanese lacks the tense-deletion rule (simply because it does not need it), while English present tense always denotes the speech time. This reminds us of Enç's analysis of the difference between present tense in an SOT language like English and present tense in a non-SOT language like Russian, which also claimed that English present tense is inherently present-time denoting. Ogihara (1995b) argues that this assumption is needed independently of the interpretation of tense in complement constructions, referring to the difference between the behavior of present tense of English and that of Japanese in relative clauses. His argument proceeds as follows: since present tense in relative clauses embedded in the matrix past tense in Japanese can be interpreted to be simultaneous with the matrix past while present tense in English in the same environment always denotes the speech time, we need a special treatment of English present tense, and hence, we must posit a condition which dictates that present tense in English is always linked directly to the speech time. If I recapitulate his argument in a syllogistic form, it seems to go as follows: (i) English PRES and Japanese PRES behave differently in relative clauses.; (ii) Thus, it must be the case that English PRES always denotes the speech time while Japanese PRES is tenseless.; (iii) Therefore, English PRES and Japanese PRES behave differently in complement clauses. If we intend to account for the functions of tense morphemes in both languages across different constructions, this is completely circular and can be of no explanation. Since we do not have a systematic account of their behavior in any type of clauses yet, a speculation formed on the behavior of tense in one clause type without an independent motivation or other plausible supports cannot serve to account for its behavior in another type of clause. Ogihara's claim that Japanese present tense is tenseless while English present tense is inherently

present-time denoting needs to be supported by more empirical data, including their occurrences in simple sentences, independently from their behavior in embedded clauses.

2.3.3 Ogihara's account of DA reading of present-under-past and *de re* attitude

Ogihara (1989) noticed that for a present-under-past sentence to be true, the state expressed in the complement does not have to obtain at the speech time. He illustrates this point with the following example (1989: 287).

(17) At 10 A.M.:

John and Bill are peeping into a room. Sue is in the room.

(a) John: (near-sighted) Look! Mary is in the room.

(b) Bill: What are you talking about? That's Sue, not Mary.

(c) John: I'm sure that's Mary.

1 minute later (Kent joins them); Sue is still in the room.

(d) Bill: (to Kent) John said that Mary is in the room. But that's not true. The one that is in the room is Sue.

In (17), despite the fact that the person who is actually in the room is not Mary, but Sue, and Bill believes that it's not Mary, Bill's report about John's statement has present tense for the complement *Mary is in the room*. Based on this observation, Ogihara concludes that what should be true at the speech time is not the content of the subject's belief itself, but a state to which the subject of the sentence ascribes the property that is denoted by the content of the belief.

Incorporating the basic insight of an eventuality-based semantics (cf. Davidson, 1976; Bach 1986) and Lewis's (1979) theory of *de se* propositional attitude, Ogihara (1995b) proposes an analysis of double-access reading of present-under-past in English, which is built on Cresswell and von Stechow's (1982) formalization of *de re* attitude report, by reformulating *de re* attitude to be a report about a state of an individual. He argues that interpretation of present-under-past involves a *de re* reading of present tense, and that present tense in English should be understood as a generalized quantifier of states involving the speech time. Under this view, a present-under-past sentence in (18a) with its LF representation after tense movement in syntax in (18b) translates as (19).

- (18) a. John said that Mary is in the room
 b. [_S Past₀ [_S John e₀ say that [_S Pres₁ [_S Mary S₁ be in the room]]]]

(19) $\exists e[e < s^* \ \& \ \text{say}'(e, j, \lambda t \lambda x \exists s[\text{exist}'(s^*, s) \ \& \ \text{be-in-the-room}'(s, m)])]$

However, Ogihara points out that this translation wrongly predicts that a present-under-past sentence in (18a) is synonymous with a sentence like (20) below.

(20) John said that Mary would be in the room.

Clearly, (18a) does not entail John's prediction in the past about Mary's future location, which is one interpretation of (20), and thus, these two must be distinguished.

In order to solve this problem, Ogihara proposes a constraint that dictates: "any attitude report must be made in such a way that the temporal directionality of the original attitude as reported by the sentence agrees with the temporal directionality of the tense morpheme used in the verb complement clause" (1995b: 204). He further defines the

temporal directionalities of tenses as: “simple past tense is previous-time-oriented; simple present tense is current-time-oriented, and future auxiliary (*will* or *would*) is future-time-oriented” (ibid.). This constraint, what he calls “temporal directionality isomorphism”, which was first introduced in Ogihara (1989), is highly stipulative, and does not have any other motivation than rescuing his theory from the unfortunate outcome of its application to crucial data, and therefore cannot be supported to be part of the grammar. However, let us suppose for now that this constraint is at work for an expository purpose of Ogihara’s account of present-under-past.

With the proposed constraint on the temporal directionality, the translation in (19) turns out to be illicit, as it violates the temporal directionality of the tense morphemes: the future-orientation of the lower tense as it is formulated in (19) does not agree with the present tense in the surface form in (18a). In order to derive a correct interpretation, Ogihara suggests that the embedded present tense must move to a higher position, adjoining to the matrix S as shown in (21).

(21) [_S Pres₂ [_S Past₀ [_S John e₀ say that [_S S₂ [_S Mary S₁ be in the room]]]]]]

This LF representation translates to (22).

(22) $\exists s[\text{exist}'(s^*, s) \ \& \ \exists e[e < s^* \ \& \ \text{say}'(e, j, s, \lambda t \lambda s_1 [\text{be-in-the-room}'(s_1, m)])]]$

According to Ogihara, (22) reads as “there exists a state S now such that John talks in the past as if he ascribes to S the property of being a state of Mary’s being in the room” (1995b: 205), and we finally derive an appropriate *de re* attitude interpretation for a complement present tense.

Returning to the aforementioned problem of temporal directionality isomorphism, a review of Ogihara's (1995b) analysis of present-under-past by Abusch (1997b) also points out the problem in this alleged constraint. She argues that because this constraint is only stated in descriptive terms and also because key notions such as "temporal directionality of the original attitude" and of an attitude being "future-oriented" are not defined, it cannot be considered to be included in a grammar. Since this constraint is central to Ogihara's analysis of DA reading of present-under-past constructions, as we have seen in the preceding discussions, the theory cannot survive if this constraint is not supported. Furthermore, a careful review of Ogihara's syntactic formulations to derive DA reading of present-under-past by Kusumoto (1996) revealed that the derivation involves illegitimate syntactic movement. Although we do not go into the detail of this problem, it suffices to say that the theory cannot be maintained as it is currently formulated.

Aside from the above mentioned theoretical problems in technical details, we also have empirical evidence against Ogihara's state formulation of *de re* attitude report of present-under-past. Ogihara provides the following dialogue to show that whether or not the state referred to in the subject's report holds in the actual world is a crucial factor for the use of present tense in the content of the report. (Ogihara, 1995b: 188; originally (22))

(23) John and Bill are peeping into a room. Sue is in the room.

(a) John: (near-sighted) Look! Mary is in the room.

(b) Bill: What are you talking about? That's Sue, not Mary.

(c) John: I'm sure that's Mary.

Sue leaves the room. One minute later, Kent joins them.

(d) Bill: (to Kent) #John said that Mary is in the room.

Ogihara notes that the use of present tense in (d) is inappropriate after Sue has left the room. However, if we assume that none of Bill, John, or Kent has noticed that Sue left the room, I think it is perfectly fine for Bill to say (d). When he finds out that Sue is no longer in the room, he would probably re-state it as *John said that Mary was in the room*. However, at the time when Bill uttered (d) without knowing that Sue has left, (d) should not be semantically anomalous as Ogihara assumes. This seems to be problematic to his state formulation of *de re* attitude belief because the property of the state of an individual who is the target of the belief of the subject of *de re* report does not have to be present for the utterance (d) to be acceptable. The acceptability of this sentence here seems to subject to the rules of pragmatics, which may be important for us to derive appropriate temporal interpretation, but which, I believe, should constitute separate module from the theory of temporal reference. We will come back to this issue in Chapter 3.

The use of present tense under past in the above situation is perfectly consistent with Klein's assumption on the nature of TT (topic time), which is defined as "the time span to which the speaker's claim on this occasion is confined" (1994: 4). If the TT of the embedded situation is taken to be the speaker's claim, the use of present tense, which merely introduces the relation $TU \subset TT$, does not have to do anything with the reality as

long as there exists a pragmatically felicitous context for the speaker to use the present tense (i.e., that Bill doesn't know that Sue has left the room in the above case).

Despite the above mentioned problems, however, the fact that Ogihara noticed the crucial role of the intensional context for the interpretation of tenses is an important step forward for the analysis of natural language temporal expressions. A very similar and yet different analysis of tense with the theory of intensionality was proposed independently by Abusch (1988, 1991, 1994, 1997a), which we will examine in the next section.

2.4 Abusch's theory of tense interpretation

2.4.1 Sequence of tense, intensionality and scope (Abusch 1988)

Slightly preceding Ogihara's (1989) work, Abusch (1988) unfolded a similar analysis of simultaneous reading of past-under-past and of DA reading of present-under-past as Ogihara's (1989). Her analysis is based on the observations on the behavior of tenses embedded in an intensional context. With the following data (Abusch, 1988: 2), Abusch convincingly argues against non-SOT theories such as the one proposed by Dowty (1982b), which she calls INDEPENDENT THEORY OF TENSE (capital hers).

- (24) John decided a week ago that in ten days at breakfast he would say to his mother that they were having their last meal together.

The independent theory of tense assumes that each tense should be interpreted with respect to the time of utterance independently from the matrix verb. However, if the utterance time is taken as an evaluation time in the interpretation of the past tense embedded under *would* in (24), we would have an interpretation in which the time of having their last meal together precedes the utterance time, while in fact it follows it.

Based on this observation, Abusch argues that any theory which claims that interpretation of embedded tense is independent of the matrix tense must be false, and that some kind of SOT rule is necessary to explain behavior of the past tense morpheme in English as observed in (24).

Abusch considers an English verb ‘say’ to be of a kind of intensional attitude verbs along with the standard intensional attitude verbs such as ‘believe’ or ‘know’. Based on this assumption on the semantics of ‘say’, Abusch argues that the reason that the complement of saying in (24) is temporally dependent on the time of saying is the intensionality imposed on the complement by the intensional attitude verb ‘say’.

In order to show an importance of the presence of an intensional context for the interpretation of tense, Abusch further provides an example of tense in a relative clause, whose head noun is an argument of an intensional verb and therefore subject to the same restriction as tenses in the complements of intensional verbs¹.

(25) John looked for a woman who married him.

Abusch points out that an intensional (i.e., de-dicto) reading of (25), in which John looked for a woman, but not a particular one, who married him, does not have a so-called ‘forward shifted reading’, in which the embedded tense follows the matrix tense in time sequence. An independent reading is possible only with an extensional (i.e., de-re) reading, in which John looked for a particular woman who married him. Abusch argues that this, together with the data like (24), constitutes evidence against the independent

¹ Abusch classifies English verbs like ‘look for’ and ‘need’ as intensional transitive verbs, which also create an intensional context for their arguments just like intensional attitude verbs like ‘believe’ and ‘say’.

theory of tense. Minimally, these data show that some tense interpretations are sensitive to scope relations and that intensionality interacts with tense phenomena in an interesting and important way.

Abusch's argument for the importance of intensionality for the interpretation of temporal dependency brings us back to our earlier question regarding the data presented in section 1.3.4 (originally in Zagana (1997: 44)), which are repeated below in (26).

(26) Sue heard [that Mary was pregnant] (Inclusion/Precedence/*Subsequence)

(27) Sue heard about [Mary's pregnancy] (Inclusion/Precedence/*Subsequence)

Both (26) and (27) lack a reading in which the time of the embedded situation is interpreted to be subsequent to the time of the matrix event. If we apply Abusch's theory of intensionality, lack of such readings can be explained by the presence of intensional context imposed by the matrix verb *hear* in both (26) and (27). However, as pointed out in section 1.3.4, the fact that there is exactly the same effect on both the tensed complement in (26) and a nominal complement in (27) suggests that the influence of intensional context is independent of tense. Therefore, I will argue that the theory of intensionality, as important to interpretation of tense as it may be, should be separated from the theory of temporal reference. It is of our concern, however, to further elucidate how the theory of intensionality interacts with the interpretation of tenses in embedded context, as it will help us to identify the unique function of tenses separated from what should be ascribed to other modules of the grammar.

The following two sections will examine Abusch's theory in more detail, and abstract what is relevant for our later discussions.

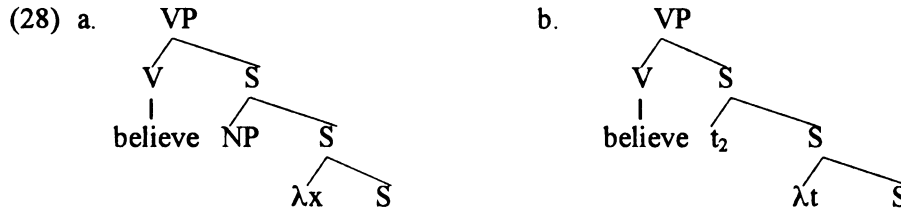
2.4.2 Interpretation of *de re* belief (Abusch 1991)

Abusch's theory of temporal interpretation is based on the analysis of *de re* attitude belief by Lewis (1979).

De re belief is best exemplified by Quine's (1956: 185) well known example of Ralph's belief about Ortcutt. One day, Ralph had a glimpse of Ortcutt in a brown hat and believes that the man is a spy. On another occasion, however, Ralph saw Ortcutt in a gray hat and believes that the man he has so glimpsed is not a spy. If we assume that the object of Ralph's belief is a proposition (i.e., a set of worlds), it would mean that Ralph has a contradictory belief about Ortcutt, but this is not what we ascribe to Ralph's belief. In order to solve this puzzle, Lewis (1979) introduced a notion of centered worlds, which he defined as pairs of a world and a designated inhabitant. He proposed that the object of belief should be a set of centered worlds, and that the object of *de re* belief (which he calls *res*) must be related to the one who possesses the belief by a suitable relation of acquaintance. By introducing a set of centered worlds and a suitable acquaintance relation (i.e., that Ralph saw Ortcutt in a certain outfit to form a certain belief in a given centered world in the above case), we can ascribe two different beliefs about Ortcutt to Ralph that are non-contradictory, despite the fact that the objects of Ralph's belief refer to the same individual in the real world.

Cresswell and Von Stechow (1982) generalized this *de re* belief to constituents other than NPs, and Abusch (1991) applied this *de re* attitude belief to interpretation of

tense in intensional contexts by replacing *de re* belief about individuals by *de re* belief about time intervals. See her representations of VP of *de re* belief for NP (1991: 6) and for tense (1991: 8) respectively.



Just as the NP complement is scoped out as *de re* belief in (28a), the embedded tense is scoped out of the complement clause as *de re* in (28b). According to Abusch, in a present-under-past sentence like *John believed that Mary is pregnant*, a tense variable t_2 in (28b) picks up the interval that subsumes John's believing time and the utterance time, and this time interval must also be some interval at which Mary is pregnant in John's centered world. However, as we have already seen in section 2.3.3 (example (22)) with Ogihara's example, the actual state of an individual based on which the subject's belief was formed may not necessarily have to continue up to the utterance time. Furthermore, as pointed out by Abusch (1997b) herself, what forces the *de re* present tense to overlap the original attitude time is not entirely clear in Abusch (1991). In her 1991 paper, she claims that the overlap of the denotation of the present tense with the subject's believing or saying time and with the utterance time in case of present under past constructions follows from rules of grammar. Abusch (1994/1997a) takes this position to be implausible and stipulative, and reformulates her theory so that this effect of overlap is derived by independently motivated factors: the semantics of present tense, the semantics of *de re* interpretation, and what she calls 'the upper limit constraint' on the reference of

tense node. The account of present under past constructions that I will present in the next chapter agrees with Abusch's position: the peculiar interpretation of present under past must be derived by independently motivated modules of the grammar. The theory of temporal reference by Klein (1994), by incorporating the theory of intensionality of Lewis's (1979) tradition and aspect will do this task. Below, I will briefly go over the revised version of Abusch's theory of *de re* present tense, mainly for the purpose of exhibiting problematic data for our later discussions.

2.4.3 Upper limit constraint and division of labor (Abusch, 1994, 1997a)

In Abusch (1994, 1997a), which are improved versions of her earlier analysis of *de re* attitude belief, she argues that simply incorporating intensional theory of *de re* interpretation to the independent theory of tense is not sufficient to deal with the full range of data. One of the empirical problems of the independent theory with *de re* attitude belief (which Abusch labels as "extensional *de re* theory") is what she calls "upper limit phenomena," and refers to the fact that the range of denotation of tenses is limited by local evaluation time. The following examples illustrate this problem (Abusch, 1997a: 16, originally (27)).

(29) Last Monday John Past₂ believed that he Past₃ was in Paris on Tuesday₃.

The past tense on the embedded copula (i.e., Past₃), cannot denote a time posterior to the time of believing. In other words, *Tuesday* in (27) cannot be the Tuesday following last Monday, but must be the Tuesday of the previous week. Abusch points out that the extensional *de re* theory (i.e., independent theory of tense + theory of intensionality)

cannot account for the lack of reading in which the lowermost past tense is after the time of believing, unless we have a stipulation that an acquaintance with future times cannot constitute a suitable acquaintance relation.

Another problem with the extensional *de re* theory is illustrated with the following example (Abusch, 1997a: 17).

(30) Sue Past₃ believed that she Past₃ would marry₂ a man who Past₂ loved her.

The problem that Abusch points out with (30) is the fact that the time of loving is taken to be simultaneous with the marrying time despite the use of past tense. Abusch argues that since the loving time, in contrast to the marrying time (which is in the belief context), is outside the belief context, its simultaneity with the marrying time cannot be accounted for by *de re* theory.

While Abusch provides her own solution to these problems, we will see that the theory of temporal reference that is to be defended in the present study is also capable of accounting for the constraints on the interpretations in these sentences. We will come back to these data in chapter 3, but let us further examine the problems under discussion here. Abusch explains that the lack of forward shifted reading of the embedded past in (29) can be explained by what she calls “the upper limit constraint (ULC),” which dictates that the local evaluation time is an upper limit for the denotation of tenses. First, she examines the behavior of modal auxiliaries such as *might* and *ought*, and argues that these modals are ‘semantically tenseless, and directly pick up the local evaluation time as a modal perspective’ (Abusch, 1997a: 23). Then, based on the following data, she shows

that the evaluation time of a modal in an intensional context is bound by the time of the matrix verb.

- (31) a. John married a woman [who might_0 become rich].
b. John believed λt_0 [his bride might_0 become rich].

The evaluation time for the embedded modal auxiliary *might*₀ in (31a) is the utterance time because it is free (i.e., not bound by anything), while the evaluation time for the embedded *might*₀ in (31b), which is bound by lambda, is believer's now. Abusch takes this to be evidence for the presence of an evaluation time abstractor on intensional arguments, and further argues that since future is always indeterminate, 'the now of an epistemic alternative is an upper limit for the denotation of tenses' (Abusch 1997a: 24). Thus, the forward shifted report as given in (32b) below (ibid., originally (46)) is prohibited.

- (32) a. Mary believed that John was afraid during the last thunderstorm.
b. *Mary believed that John was afraid during the next thunderstorm. (* mine)

I do not have any argument against Abusch's proposal on the existence of ULC (upper limit constraint), and I believe that it is a perfectly reasonable constraint to have in order to derive desired interpretations in the above cases of (29) and (32b). However, the fact that the time of the embedded eventuality cannot follow the time of believing or saying seems to be a natural consequence of interaction between what is available as linguistic form and our conceptualization of the semantic relationship between the act of saying or believing and the content of believing or saying. If the content of believing or saying is the future event with respect to the time of believing or saying, the grammar would not allow us to present them as if it was a fact, but would require us to present it as a prospect.

This is exactly what the auxiliary *will*, or in this case, its past-tense counterpart, *would*, will do. In fact, both (29) and (32b) will have a future shifted reading of embedded past if we replace the simple past tense of a copula in these sentences with *would* plus copula.

(29') Last Monday John Past₂ believed that he Past₃ would be in Paris on Tuesday₃.

(32') b. Mary believed that John would be afraid during the next thunderstorm.

I see no reason for the grammar to endow the embedded tense in (29) and (32) to express the forward shifted reference, when it already has an unambiguous means to express a desired ordering relationship as shown in (29') and (32'). The fact that the complement of believing or saying cannot refer to the fact that is posterior to the time of believing or saying is a matter of pragmatics and of our knowledge of the world. I believe that we are not encouraged to express what is unknown to us as if it was known, since it will violate the Cooperative Principles of conversation in Gricean (1968) sense. As we saw in (29') and (32'b), it is perfectly fine for us to report the content of belief as future prospect with a relevant linguistic expression for this purpose. Therefore, I would consider this constraint to be independent of the theory of tense on a different ground from Abusch.

Let us now return to another problem, which is raised by (30), in which the lower most past tense indicates the simultaneity with the intermediate past tense.

(30) Sue Past₃ believed that she Past₃ would marry₂ a man who Past₂ loved her.

Abusch explains that because of the semantics of *would*, the marrying time is interpreted to be posterior to the believing time. The time of loving, on the other hand, cannot be interpreted to be after the marrying time, but it should not be restricted to the time prior

to the marrying time, as it can also be interpreted to be simultaneous with it. Abusch (1997a) states that the extensional de re theory cannot explain this, but I would argue otherwise. As we saw in chapter 1 (section 1.3.4), when the situation embedded under past is stative, we have simultaneous reading of embedded tense with the matrix tense due to unbounded nature of stative situations. If we incorporate this idea to the account of (30), the simultaneity of the loving time with the marrying time naturally follows from the aspectual nature of the lower most predicate *love*, which is stative. Therefore, this data is not incompatible with the assumptions of the extensional de re theory, and I consider that the independent theory of tense with some version of de re analysis is still dependable if it incorporates the theory of aspect.

In the preceding sections, we have seen that the theory of intensionality proposed by Lewis (1979) and was applied by Abusch (1988, 1991, 1994, 1997) and Ógihara (1989, 1995a,b, 1997) to their theories of tense, has a crucial role in the interpretation of natural language temporal expressions. However, the fact that the intensionality plays the same role in the determination of ordering relationship of events and situations in an environment without finite tense (section 2.4.1, (26), (27)) suggests that it should belong to an independent module of grammar. Thus, it should be separated from the theory of tense. Some other constraints, which interact with the interpretation of tense, seem to belong to pragmatics rather than to the theory of temporal reference (section 2.3.3, (22); section 2.4.3, (29), (32)). Furthermore, we also observed that a problem of simultaneous interpretation of past under past in a certain environment (section 2.4.3, (30)) may be explained by incorporating theory of aspect as suggested in Chapter 1.

The next section reviews Costa's (1972) analysis of SOT, which provides us with an important insight into pragmatic/situational factors that interact interestingly with tense interpretations, namely, the speaker's attitudes towards the content of the report and presuppositions.

2.5 The role of pragmatics in tense interpretations (Costa 1972)

With the following example, Costa argues that whether the complement describes a condition that is still true at the time of utterance or not is crucial to the choice of present tense in the complement, and that the speaker has a choice between using present or past for the complement verb, depending on his attitude towards its content. (The following sentence is from Costa (1972: 44), originally (23), underlines hers.)

- (33) John didn't realize that you had to declare that you weren't a Communist to get
a US visa. have aren't

Although Costa does not provide an explanation as to how the choice between present and past is made in (33), it is obvious that if the speaker chooses to present what is stated in the complement to have a current effect, he or she would use present tense in the complement. In this sense, it may not necessarily be accurate for us to say that what is crucial for the choice of tense in the complement is whether or not the complement describes a condition that holds true at the time of utterance. Because the speaker's belief, i.e., that you have to declare that you aren't a Communist to get a US visa in (33), may turn out to be false in the actuality. And yet, as long as the speaker believes that the condition is currently at work, (33) is felicitous with present tense in the complement.

What matters for the choice of the present tense in the complement, then, is the relevance of the speaker's claim as to the content of the complement to a particular context where the utterance is made.

In order to show the difference in nuance between past under past (with simultaneous interpretation) and present under past, Costa further provides the following example in which the complement describes what she calls ‘a universal of behavior’ (Costa, 1972: 45, originally (26), underline hers).

- (34) I'm afraid Bill and I haven't found any more cool facts for your book on How People Behave, but John asked us to tell you he noticed that people can't help touching walls that carry the sign: "Don't Touch - Wet Paint". ?couldn't help

In (34) the use of present tense is perfectly fine in the given situation, while the use of past tense sounds strange. Costa explains that the crucial factor here is what she calls the notion of “presupposed relevance” (ibid.), which can be defined as what the speaker considers to be relevant to the conversation he is engaged in. The point is further enforced by the following examples (ibid., originally (27) and (28) respectively).

- (35) Did Sarah have any ideas about what might be wrong with my marriage?
- (36) Well, she mentioned that married couples often discover that they
wrongly think that their sex-life is perfect.
*thought *was *discovered

As shown above, when the past tense is used in the complement in (36) as an answer to (35), the utterance would not make a relevant answer to the question.

All the above observations by Costa turn out to be consistent with Klein's definition of TT (topic time) and of present tense and past tense. According to Klein, present tense introduces a relation $TU \subset TT$ (utterance time is included in topic time) and past tense $TT < TU$ (topic time is before utterance time). The fact that the content of the utterance is asserted to have a present time relevance (at the speech moment) naturally follows from the definition of present tense, that is, TT includes the speech time. Likewise, the fact that the content of the utterance introduced with past tense is not interpreted to have a present-time relevance also naturally follows from the definition of past tense, that is, TT is confined to the time span before TU. The crucial point in both cases is the nature of Klein's TT, which is defined as the time span to which the speaker's claim is confined, and is distinguished from the actual time of the situation described by the utterance.

Another important point in Costa's work on SOT is her classification of verbs in terms of their behavior with respect to SOT. She argues that English verbs can be classified into two groups: A-verbs, which optionally undergo the SOT rule, and B-verbs, which impose obligatory SOT. Examples of each type are as follows (Costa, *ibid.*: 46, originally (30) and (31) underlines hers).

(37) Bill forgot (mentioned /regretted /realized /discovered /showed /noticed /was
A- amazed /was concerned /said /reported) that coconuts grew high up on trees.
VERBS grow

(38) Bill knew (was aware /thought /believed /imagined /figured /dreamed /wished
B- hoped /asserted /alleged /insisted /quipped /snorted /whispered) that the new
VERBS President of Chorea was really a Chai CIA agent.
*is

Costa classifies A-verbs further into two classes: factives and entailment verbs on the one hand, and a few non-factive reportative verbs such as *say* and *report*, which, she assumes, are capable of being used factively. B-verbs, on the other hand, consist of what she calls manner verbs of saying, non-factive verbs, and ill-behaved factive verbs such as *know*, whose factivity can be canceled. Costa makes an interesting remark on manner verbs of saying by stating that ‘they are always associated with a point of time in narrative discourse, thereby establishing a distance between speaker and sentence so that the speaker cannot identify with the complement’ (ibid.). This sounds somewhat similar to characterizing them as creating an intensional context for a complement. It seems that two classes above basically contrast in whether they are optionally or obligatorily intensional. A-verbs seem to undergo SOT when they are used intensionally, and allow present marking of the complement when they are used factively. B-verbs must undergo SOT because they are obligatorily intensional. It would be interesting to test how different types of verbs behave in SOT. However, since such an investigation is beyond the scope of the present study, here we confine ourselves to a finding that a verb in the matrix clause may affect the interpretation of the complement tenses due to the presence or absence of the intensionality imposed by the matrix verb.

In this section, we briefly reviewed Costa’s study on SOT, and found that the speaker’s attitude and presuppositions in a context of the utterance play an important role in the interpretation of and the choice of the complement tense. We also learned that the type of the matrix verb may affect SOT. These points will be incorporated into our investigation of the tense systems of English and Japanese in Chapter 3. Before we move

on to our analyses, however, I will discuss two other approaches to SOT, and show how the present study is different from them.

The next two sections examine two syntactic theories of tense interpretations: one by Hornstein's (1990) Neo-Reichenbachian approach, and the other by Nakamura's (1994) application of syntactic theory of tense of the type proposed by Stowell (1993) and Zagana (1990) to Japanese data.

2.6 Hornstein (1990)

Hornstein (1990) advanced a syntactic theory of tense which takes Reichenbach's (1947) three time coordinates to be syntactic primitives. He proposed that a tense system constitutes an independent linguistic level, which comprises of three points: S, R, and E that are linearly ordered. The ordering relationship of three time primitives in Hornstein's model is represented as being separated by either a line or a comma. When two points are separated by a line, the leftmost point is interpreted as temporally earlier than the other, and when two points are separated by a comma, they are interpreted to be contemporaneous. Based on the above assumptions, he proposes that the following six tenses constitute the basic tense structures (BTSs) for English, from which all the other complex tenses are derived.

(39)	S,R,E	present
	E,R _ S	past
	S _ R, E	future
	E _ S,R	present perfect
	E _ R _ S	past perfect
	S _ E _ R	future perfect

Hornstein convincingly argues that the nature of tenses is fundamentally different from that of operators, drawing our attention to the following facts (1990: 166):

(40)

- (i) A tense in an embedded clause can only be temporally dependent on the tense of the clause under which it is immediately embedded.
- (ii) A tense within a relative clause, regardless of how deeply embedded the relative clause is, is never temporally dependent on any other tense. In other words, it is always temporally interpreted relative to the moment of speech.

Generalizations given above amount to saying that tenses can never allow intermediate scope positions despite the fact that they allow widest possible scope. If tenses are operators, this cannot be explained. Hornstein argues that the property of tense as presented in (ii) above can be explained if we assume that all tenses have an S point and that the tense interpretation is subject to “the principle of full interpretation, which maps S onto the utterance time if S is otherwise unanchored” (1990: 167). As this wording suggests, Hornstein’s definition of the S point is not quite the same as the S point in Reichenbachian treatment of tense. In order to incorporate the temporal shifting characteristic of finite embedded clauses into the theory, Hornstein redefines the S point so that it can be anchored to times other than the speech moment, while preserving its default temporal value to be the utterance time. This redefinition allows S to serve as a deictic center for interpretation of tense, and at the same time, endows the Reichenbachian theory a power to explain the behavior of tenses in embedded clauses.

Hornstein argues that complex tense structures arise by modifying BTSs (basic tense structures) in (39) either by temporal adverbs or by a syntactic rule such as SOT, and that there is a constraint on the derivation of complex tense structures (CDTS), which

dictates that a derived tense structure (DTS) must preserve BTS. According to Hornstein, BTSs are preserved under the following conditions (1990: 15, originally (13)).

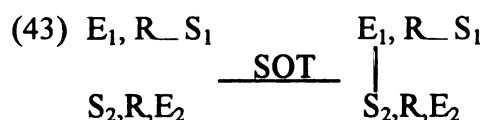
(41) BTSs preserved iff

- a. No points are associated in DTS that are not associated in BTS.
- b. The linear order of points in DTS must preserve BTS.

Let us now examine how this Neo-Reichenbachian theory of tense accounts for our SOT data.

- (42) a. John heard that Mary was pregnant.
 b. John heard that Mary is pregnant.

(42a) with simultaneous reading (i.e., the SOT version) will receive the following tense structure in Hornstein's model.



What appears on the left side of (43) is the tense structure of (42a) before the application of the SOT rule. The upper diagram shows the tense structure of the matrix clause and the lower one that of the embedded clause. The SOT rule shifts S_2 and associates it with the matrix E_1 . The derived tense structure that appears on the right side in (43) preserves BTS, and thus the structure is well-formed. The association line between E_1 and S_2 makes the evaluation time of the complement to be contemporaneous with the time of hearing, and simultaneous reading is hence obtained. In contrast, DTS of (42b), which is what we obtain by non-application of SOT, is identical to BTS as shown below.

(44) $E_1, R _ S_1$

S_2, R, E_2

Since there is no association line between time coordinates of the matrix tense structure and those of complement tense structure, S_2 denotes the time of utterance by default, and Mary's pregnancy is interpreted to be contemporaneous with the utterance time.

The problem with Hornstein's account of SOT, as I see in the mechanism introduced above, is that the system cannot account for a DA reading of present-under-past sentence like (42b). Since S_1 and E_1 are separated by a line, E_1 is interpreted to be anterior to S_1 . However, E_2 , which is supposed to be contemporaneous with the utterance time, cannot be interpreted to overlap with E_1 , because E_1 is necessarily disjointed from the time span associated with E_2 . Thus, Hornstein's theory of tense, as simple and attractive as it is, cannot account for the full range of data as it is currently formulated. Hornstein's theory of tense may be revised to incorporate the theory of intensionality to cope with this problem of the DA reading. The present analysis of tense take this line of integration of what seems to be well-motivated and descriptively accurate in each theory to Klein's model of temporal reference. This thesis also incorporates many of the generalizations made by Hornstein (1990) to the explanation of the behavior of the temporal morphemes of English and Japanese in Chapter 3. In this regard, I will point out another important contribution of Hornstein's theory to the study of natural language temporal reference.

An important point of departure from earlier analyses of tense in Hornstein's theory of tense lies in that he treats tenses to be abstract entities distinguished from the actual morphemes that encode them. Most of the previous analyses treat tenses to be

identical to their carriers (i.e., temporal morphemes), which runs into a number of problems including an unconventional use of tenses such as historical present and a modal-like property of some temporal auxiliaries. If tenses and their carriers are distinguished, there need not be one-to-one correspondence between tenses and the meanings borne by the morphemes that encode tenses. Such a view is consistent with the idea that some temporal morphemes exhibit modal functions, and also with the position of the present study that temporal morphemes of English and Japanese encode both tense and aspect.

Another important point in Hornstein's analysis of tense which is relevant to the present study is found in his treatment of perfect of English. Although Hornstein explicitly states that tense and aspect form separate modules of the grammar and that his theory only concerns the system of tense, his analysis of English perfect shares important insights with Klein's theory of temporal reference. Hornstein argues that tense represents the SR relation and that the perfect morpheme *have* determines the RE relationship. If we replace S with TU (time of utterance), R with TT (topic time), and E with TSit (time of situation) in Klein's model, the above relations defined by Hornstein (i.e., SR relation that is determined by tense and RE relation that is determined by the perfect morpheme *have*) are exactly the same as the relations that are determined by tense and aspect respectively in Klein's theory, namely, that tense defines the relation between TU and TT and that aspect defines the relation between TT and TSit.

In sum, Hornstein's theory, though it cannot account for a peculiar reading of present under past, has many important implications to the analysis of natural language temporal expressions. In particular, his position that tenses are more like adverbs than

operators in nature is shared by the present study. The assumption that tenses and their carriers should be separated is also a significant breakthrough in the theory of temporal reference, and this position is shared by others who proposed different kind of syntactic theories of tense, namely Zagana (1990) and Stowell (1993, 1996). The next section will examine Nakamura's (1994) analysis of tense morphemes as contrasted with tense morphemes in English, which applied Stowell's (1993) theory of tense.

2.7 Nakamura (1994)

Based on the theories of tense advanced by Stowell (1993) and Zagana (1990), which argue that tense is 'a dyadic predicate of temporal ordering that takes two time-denoting phrases as its arguments' (Stowell 1993: 1), Nakamura (1994) accounts for the difference between English and Japanese in the behavior of tense morphemes in embedded clauses in terms of their difference in semantic/syntactic nature.

In order to refresh our memory regarding the difference between English and Japanese, the past-under-past sentence of English and its surface equivalent in Japanese are repeated here from section 1.3.1.

- (45) John said that Mary was pregnant.
 a. John said that Mary had been pregnant. (shifted reading)
 b. John said, "Mary is pregnant." (simultaneous reading)
- (46) John-wa Mary-ga ninsinsi-te-i-ta to it-ta
 John-TOP Mary-NOM pregnant-ASP-PAST COMP say-PAST
 'John said that Mary had been pregnant.' (shifted reading only)

Based on the fact that past tense morpheme *-ta* in complements does not allow simultaneous reading, Nakamura argues that unlike English past tense morpheme *-ed*,

which is a PAST POLARITY ITEM (PPI) that is semantically empty and must be licensed by a real Past tense in the head of TP, Japanese *-ta* is an actual spell-out of semantic Past in the head of TP. According to Nakamura, this is specifically why the simultaneous reading is unavailable for past under past in Japanese as shown in (46). Japanese non-past tense morpheme *-ru* is also analyzed as a non-polarity item, but in a sense that it is not an ANTI-PAST POLARITY ITEM (anti-PPI) like English present tense morpheme. Anti-PPI items, as defined by Stowell, must not be in the scope of a semantic Past. Since Japanese *-ru* is not subject to this restriction, it appears inside the scope of the semantic Past, and hence, we obtain a simultaneous reading for the present embedded under past as shown in (47) below.

- (47) John-wa Mary-ga ninsinsi-te-i-ru to it-ta
 John-TOP Mary-NOM pregnant-ASP-PRES COMP say-PAST
 ‘John said that Mary was pregnant.’ (simultaneous reading)

As mentioned at the beginning of this section, these analyses are based on the assumptions made by Stowell (ibid.) that tense is a dyadic predicate of temporal ordering which realizes as a head of TP and takes reference time as its external argument and event time as its internal argument, and that the tense morphemes of English are polarity items that must be licensed by abstract tense under TP. Analyzing tense morphemes to be distinguished from the abstract notion of tense is the same in spirit with Hornstein’s analysis of tense and temporal morphemes, which is empirically well-motivated, considering some of the unconventional use of tense morphemes in natural language. A unique contribution of Stowell’s theory of tense is that it allow us to present the derivation of temporal ordering with the existing model of the structure of language.

However, Nakamura's application of Stowell's model to the account of the behavior of temporal morphemes of Japanese seems to have a number of empirical problems.

First, Nakamura's account of non-past tense morpheme *-ru* of Japanese does not account for the optional DA reading of present-under-past Japanese sentence like (47). DA reading of present under past in English is ascribed to the anti-polarity nature of its present tense morpheme. Due to its nature as an anti-past polarity item, a complement present tense under matrix past in English must scope out of the matrix past, leaving its copy behind. Then, by assumption, the event time of the complement is computed in a way that it encompasses both the event time evaluated in the in-situ position and the event time at the scoped-out position, yielding DA reading. If we assume that Japanese non-past tense morpheme *-ru* is fundamentally different in nature from English present in that it is not an anti-past polarity item, the computation of DA reading that we have just reviewed for English present tense is not available for *-ru*. Then, we cannot account for the optional DA reading of (47).

Another problem with Nakamura's analysis lies in his treatment of difference between the behavior of tenses embedded under the matrix intensional verbs and those embedded under the matrix factive verbs. Nakamura argues that while independent reading is unavailable for the complement of an intensional verb like *-iw* 'say', the complement of a factive verb like *-sir* 'know' allows such reading when it appears in the future tense. He provides the following examples to illustrate this point (1994: 366, originally (7) and (8)).

- (48) Hanako-wa [Taroo-ga gakusei da to] yu-u daroo
 Hanako-TOP Taroo-NOM student be COMP say-PRES probably
 ‘Hanako will say that Taro is a student.’
- (49) Hanako-wa [Taroo-ga gakusei de ar-u koto]-o sir-u daroo
 Hanako-TOP Taroo-NOM student be-PRES fact ACC know-PRES probably
 ‘Hanako will learn (the fact) that Taro is a student.’

Nakamura argues that (49) can either be interpreted to mean that Taroo’s being a student obtains at the speech moment independently of the futurity of the matrix verb, or that Taroo will be a student at the future time of Hanako’s learning of it. In contrast, (48) does not have this ambiguity, and only has a simultaneous reading. Nakamura ascribes this difference to the syntactic nature of the complement in respective case: while *sir-* ‘know’ takes a complex NP/DP complement, which is assumed to undergo optional LF movement outside the scope of the matrix Past, *iw-* ‘say’ takes a CP complement and may not allow this option. However, as Nakamura points out himself, this cannot explain the fact that independent reading is unavailable when matrix factive verbs are in the past tense. See the following example from Nakamura (1994: 369, originally (13)).

- (50) Taroo-wa [Hanako-ga ninsin-si-te-i-ru koto]-o sit-te i-ta
 Taroo-TOP Hanako-NOM be-pregnant-PRES fact ACC know-NF be-PAST
 ‘Taroo knew that Hanako was (*lit.* is) pregnant.’ (simultaneous reading)

Not only Nakamura’s syntactic account of the difference between (48) and (49) cannot explain the lack of independent reading for (50), but it also fails to account for the fact that (50) in fact has DA reading. Furthermore, Nakamura does not have an account for the lack of DA reading for (49), either. It should also be pointed out that an intentional verb like *iw-* ‘say’ can possibly take NP/DP complement just as a factive verb like *sir-*

‘know’ does, which further undermines the argument that the difference in temporal interpretation between the complement of intensional verbs and that of factive verbs should be ascribed to the syntactic property of their respective complement type. Consider the following example.

- (51) Mary-wa [John-ga gakusee da to yu-u koto]-o it-te-i-ta
Mary-TOP [John-NOM student be COMP say-PRES fact]-ACC say-ASP-PAST

As shown in (51), since taking NP/DP complement is not a unique syntactic property of factive verbs, ascribing the difference between (48) and (49) to the difference between in their complementation is not empirically well-motivated. Furthermore, the fact that (51), with their NP/DP complement, still does not have an independent reading, is a conclusive evidence against Nakamura’s position that the independent reading is attributed to a syntactic property of the complement. The solution should reside in somewhere else.

Owing to Costa’s (1972) work on SOT, we know in fact that factive verbs (at least those in English) may optionally allow SOT but not always, while SOT is obligatory for intensional verbs. We also know from Abusch (1997a) that futurity of the matrix verb interacts with the interpretation of embedded tense in an interesting way. Thus, we may possibly have an alternative account for the puzzles with Japanese data we observed in this section. Chapter 3 will provide an alternative analysis.

2.8 Summary

This section examined various theories of tense that have been proposed to date and their approaches to SOT data. Each theory that we reviewed has its own strength and weakness, but all of them had important implications to our analysis of natural language

temporal expressions. The theory of temporal reference that will be tested in the present study will incorporate the generalizations obtained in the predecessor's work on temporal expressions, namely, evaluation time shifting property of tense, referential nature of tense, the role of intensional context imposed by certain types of matrix verbs on the interpretation of embedded tense, and role of speaker's attitude and presuppositions on the interpretation of tense. In this sense, the present study is an integrated approach to tense which incorporates all the benefits of the predecessors' painstaking work. This thesis, however, is an attempt to make its own contribution to the field in a sense that it tries to separate what should be considered to be part of the theory of temporal expressions (which, I believe, to constitute an independent module and may possibly have universal applicability) from those that belong to other modules of grammar.

One thing that is lacking in all the previous accounts of SOT and other behavior of tense is consideration of aspect of the predicates involved. As we saw in chapter 1, the aspect of the predicate in the complement has a crucial role in the interpretation of past under past constructions in English. We have also seen in this chapter that some of the problematic data (section 2.4.3, (30)) may receive an explanation if we incorporate the notion of aspect in interpretation of tenses. The subsequent chapters will further examine how tense and aspect interact while making respective contributions to the representation of events and situations in natural language, and attempt to provide a unified analyses of the behavior of temporal morphemes of English and Japanese.

CHAPTER 3

Temporal Reference in English and Japanese

Introduction

This chapter examines the behavior of temporal morphemes of English and Japanese in the light of the generalizations obtained in Chapter 2 by applying Klein's three temporal primitives to the representation of various complex tense structures in both languages. The aim of this chapter is to show that complications with the interpretation of tenses in different types of embedded clauses can be ascribed to the presence or absence of an intensional context and/or the aspect of the predicates that appear with the tense morphemes. As a result, the semantics of tenses can be maintained consistent across different constructions and the tense system can be much simpler than has been claimed by some existing theories of tense.

I will argue that the fundamental difference between temporal morphemes in English and temporal morphemes in Japanese in their behavior in complement clauses can be explained in terms of the nature of the shift of deictic center for the interpretations of embedded tenses in the respective tense systems. Aside from that, the functions of the basic tenses can be kept constant between these two totally unrelated languages, and the apparent complexity in the behavior of their tense morphemes result from the interaction of the system of temporal reference with other modules of the grammar.

The organization of this chapter is as follows: section 3.1 examines the interpretation of tenses in complements in Klein's model of temporal representation incorporating the theory of intensionality and the theory of aspect. Section 3.2 extends the analysis of temporal interpretation presented in section 3.1 to tenses in adverbial clauses, namely, *when*-clauses in English and *toki*-clauses in Japanese. Section 3.3 further

examines the occurrences of tense morphemes in relative clauses in both languages. Finally, section 3.4 summarizes the generalizations obtained in this chapter.

3.1 The SOT data revisited

3.1.1 Complement tenses in English and Japanese

This section examines the interpretation of complement tenses in English and Japanese using Klein's model of temporal representation by incorporating the basic insight of the theory of intensionality (cf. Lewis 1979; Abusch 1991, 1994, and 1997a). I will argue that the fundamental difference between English and Japanese in the behavior of their tense morphemes in the verb-complement clauses lies in the fact that Japanese requires a shift of deictic center for the evaluation of the complement tenses while English does not. This subsection provides an empirical support for such a proposal.

The following Table recapitulates the difference between English and Japanese in the interpretations of complement tenses embedded under past that we observed in Chapters 1 and Chapter 2. PAST, PRESENT, NON PAST in the Table 1 represent the temporal morphemes used for the respective tenses, and 'past shifted', 'simultaneous', and 'DA (=double access)' represent respective interpretations.

English	Interpretations	Japanese
PAST	past shifted	PAST
	simultaneous	
PRES	DA	NON PAST

Table 4. Interpretations of complement tenses embedded under past in English and Japanese

As we see in the Table 1, the simultaneous reading, which is available for past under past in English is only available with an embedded non-past in Japanese. Past under past in both languages has a past shifted reading, while present under past in English and non-past under past in Japanese both have a DA reading. When we examine the examples in (1), we find that the difference between the indirect and the direct speech in English in the available interpretations is quite similar to the difference between English and Japanese that we see in Table 1.

- (1) a. John said that Mary was pregnant. (past shifted/simultaneous)
 John said that Mary is pregnant. (DA)
- b. John said, "Mary was pregnant." (past shifted)
 John said, "Mary is pregnant." (simultaneous)

Notice that the pattern that we see for the interpretation of tenses in the quoted speech in English shown in (1b) is similar to the one we find for the interpretation of complement tenses in Japanese given on the right of Table 1. That is, the past under past in both Japanese and quoted speech in English only allows a past-shifted interpretation, while

present under past in both Japanese and quoted speech in English allows a simultaneous reading.

We know that switching from the direct speech to the indirect speech involves a shift of deictic center for the interpretation of deictic expressions, which is illustrated by the following examples with personal pronouns in English.

- (2) a. John said, “I am sick.” (‘I’ = John)
- b. John said that I am sick. (‘I’ = the speaker; ‘I’ ≠ John)
- c. John said that he is sick. (‘he’ = John)

In the direct speech in (2a), ‘I’ refers to the matrix subject ‘John’, while in the indirect speech in (2b), ‘I’ refers to the speaker but not ‘John’. In order for the subject in the content of what is said to be the matrix subject in the indirect speech, a personal pronoun in the embedded subject position must be changed from ‘I’ to ‘he’ as in (2c). This observation leads us to believe that the difference in interpretations between (1a) and (1b) involves a shift of deictic center for the complement tenses in the indirect speech.

Based on the above observations concerning the direct and the indirect speech distinctions, I would argue that the interpretation of tenses in verb complements in Japanese, which patterns with the interpretation of tenses in the indirect speech in English, involves a shift of deictic center for the embedded tenses. The rest of this section will examine how this proposal, together with the theory of intensionality and the generalizations obtained in the study of aspect, helps us explain the differences in behavior of tense morphemes in English and Japanese that have been discussed in the previous literature on tense.

3.1.2 Interpretation of complement tenses in English

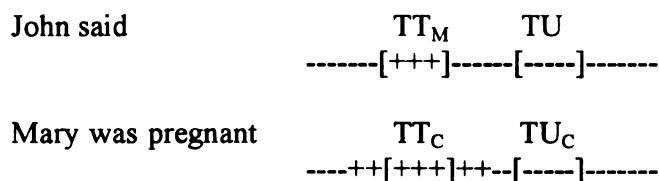
3.1.2.1 Past under past

This subsection examines the interpretation of complement tenses in English. Below we will investigate how the model proposed in this study accounts for the data of complement tenses embedded under past in English. The following data are repeated from earlier chapters.

- (3) John said that Mary was pregnant.
a. John said that Mary had been pregnant. (shifted reading)
b. John said, “Mary is pregnant.” (simultaneous reading)
- (4) John said that Mary is pregnant. (‘double-access’ reading)

Let us tentatively assume that tenses in any environment may have independent readings unless otherwise designated by other modules of the grammar. That is, complement tenses, tenses in adverbial clauses, and tenses in relative clauses may all be evaluated with respect to the utterance time just like tenses in simple sentences unless there is an additional operation required for the interpretation outside the tense system. Based on this assumption, a first approximation of the temporal representations for past under past like (3) may be schematized as below using Klein’s three temporal primitives.

- (5) John said that Mary was pregnant.



Since both the main and the complement clauses have past tense, TT of both clauses are placed before TU or the utterance time. Here, the temporal orderings between TT_M and TT_C are underdetermined by the tense system: their relationship could be 1) precedence, 2) simultaneity (or temporal overlapping or inclusion), or 3) subsequence. However, one option must be excluded by the requirements of the theory of intensionality. Recall from Chapter 2 that Abusch (1988, 1991, 1994, 1997a) considers ‘say’ in English is a type of ‘intensional attitude verbs’ like the standard intensional attitude verbs such as ‘believe’ or ‘know’, and creates an intensional context for the interpretation of its complement. The intensional context created by the semantics of the matrix verb ‘say’ makes it impossible for the content of saying to be temporally subsequent to the act of saying. Thus, the interpretation in which Mary’s pregnancy is subsequent to John’s saying is excluded. In the spirit of Abusch (1997a), let us assume that there is ‘an evaluation time abstractor’ (1997: 23) that is adjoined to intensional complements and requires the time of the complement situation to be dependent on the time of matrix event when a matrix verb introduces an intensional context. See the following LF representations of (5) as shown in (6a), in contrast to the LF representation of a sentence with a non-intensional context that is shown in (6b).

- (6) a. John said t_1 that λt_2 [Mary was t_2 pregnant].
b. John met t_1 a woman [who was t_2 pregnant].

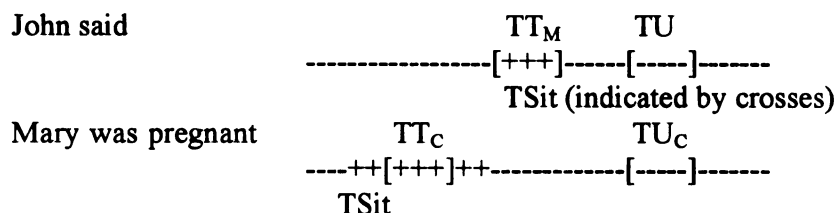
In (6a) the time of the complement situation is bound by an evaluation time abstractor represented by λt_2 , which was introduced to the LF by the intensionality imposed by the matrix verb ‘say’. The time of the complement situation, therefore, should be interpreted to be either simultaneous or anterior to the time of the matrix event. On the other hand,

the time of the relative clause in (6b) is not bound due to the absence of an intensional context, and hence, any time ordering relation becomes possible between the matrix and the main situations.

It should also be pointed out that if the sentence in (5) were John's prediction about Mary's future state, English has an independent grammatical device for expressing such a relation between the main event and the complement situation; namely, *John said that Mary would be pregnant*. Under our assumption, whether the auxiliary *will* is considered to be a future tense marker or a modal expression, it systematically functions to confine the situation expressed by the sentence to the time span that is posterior to the utterance time (i.e., the default value of TU). When there is a unique way to express the believer's future, it would be most natural for us to assume that the language may use the device to represent this relation between the matrix and the complement situations over the devices to indicate the believer's now or the believer's past.

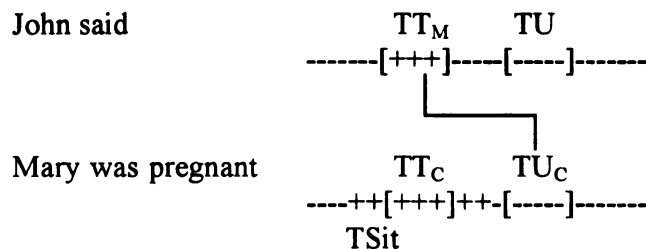
After excluding one out of three possible interpretations for temporal ordering of matrix and complement situations, we are now left with two options: the simultaneous interpretation and the shifted interpretation. However, I would argue that the model represented in (5) will only give us the simultaneous interpretation, and that it is logically impossible for us to derive a shifted reading with the temporal structure in (5). Let us suppose that the distance between TT and TU on the diagram represents the time length. The relation we need for the shifted reading is $TT_C < TT_M$. Assuming that the default values of both TU and TU_C are the utterance time, (5') below shows the relation between TT_M and TT_C that we would need in order to get the shifted reading.

(5') John said that Mary was pregnant.



In (5') TT_C is placed anterior to TT_M, and so we would expect to obtain a shifted reading. However, the shifted reading that we would have with the temporal structure given in (5') necessarily entails the simultaneous reading as well due to the unbounded aspectual nature of the complement situation. That is, *John said that Mary was pregnant* in its shifted reading with the temporal structure in (5') should mean that John said that Mary had been pregnant at some earlier time and that she was still pregnant when he said that. Obviously, we do not have such an interpretation, and the model presented in (5') makes the wrong prediction. I will propose that English past-under-past construction may optionally undergo a shift of deictic center for the evaluation of the complement tense in order to allow a shifted reading. Then, the correct temporal structure for the shifted reading of (5) would be the one shown in (7) below.

(7) John said that Mary was (=had been) pregnant. (shifted interpretation)



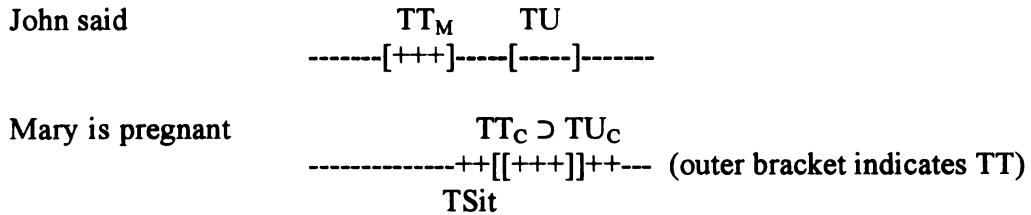
There is an association line between TT_M and TU_C , which indicates that there is shift in the deictic center for the interpretation of the complement tense. The shift in the deictic center means that the default anchoring point for the interpretation of tense, with respect to which the relative location of the described event or situation is indicated, changes its value from the utterance time to some other time. In the above case, the value of TU_C has been shifted from the utterance time to the topic time of the main clause as indicated by the association line between TU_C and TT_M . Tense in the complement clause, which is evaluated with respect to TU_C , now finds its relative location with respect to TT_M , which is evaluated with respect to TU of the matrix clause whose default value is the speech moment. A complement situation may still continue up to the left boundary of TU_C . However, TT_C is now confined within some time span that is necessarily disjoint from TU_C due to the definition of past tense. Thus, TT_C is also disjoint from TT_M due to the association line between the complement and the matrix clauses that makes TU_C and TT_M contemporaneous. Hence, by assuming the shift in deictic center for the evaluation of the complement tense, we can exclude an undesirable temporal overlap between TT_C and TT_M that subsisted with the model represented in (5'), and correctly predict the shifted reading for (3). We still assume the existence of an evaluation time abstractor in (7) as we did for (5) just because we assume that there is no difference between (5) and (7) in terms of the intensionality imposed on the complement. However, the abstractor has no unique contribution for the evaluation of the complement tense in (7), as the shift in anchor for the interpretation of the complement tense makes the complement time dependent on the TT of the matrix clause independently.

To summarize the discussion of the past under past construction in English, two different readings, the simultaneous reading and the shifted reading, can be explained by assuming the optional shift in deictic center for the interpretation of the complement tense. While an evaluation time abstractor proposed by Abusch (1997a) is crucial to obtain the simultaneous reading, it is not considered to be part of the tense system, as has been presented in Chapter 2. The evaluation time abstractor is imposed by the theory of intensionality, which constitutes a separate module of the grammar from the tense system. The intensionality provides further constraint on the temporal ordering of events and situations that may be underdetermined by the tense system. The function of the tense system by itself, however, is only to provide the relative location (i.e., posterior or anterior relation) of TT with respect to TU.

3.1.2.2 Present under past

Let us now move on to the present under past construction. Under the present model, present under past like (4) receives a temporal representation as shown in (8).

(8) John said that Mary is pregnant.



Despite the fact that the topic time of the matrix event is placed prior to the utterance time and that the topic time of the complement situation is contemporaneous with the utterance time, we take the topic time of Mary's being in the room to be temporally

overlapping with John's saying time. This is ascribed to the intensional context created by the matrix verb 'say' as in the case of (5) above. The existence of an evaluation time abstractor for the interpretation of the complement tense makes it temporally dependent on the matrix tense. However, as noted earlier, there is a crucial difference between past under past like (3) and present under past like (4) in their complement tense interpretations. While the complement tenses in both constructions have something in common in that they are both temporally dependent on the matrix tense, the complement in present under past is interpreted to have current relevance, which is not shared by the complement in past under past. The fact that both constructions exhibit the temporal dependency of the complement situation on the 'saying' time should be attributed to the intensional context imposed by the matrix verb. The difference in the interpretation between the two, under the present model, is explained by the difference between their complement tenses in their respective function in temporal ordering of the described eventualities. If we assume that tenses in any environment can receive an independent interpretation when there is no further constraint imposed by other modules of the grammar, the current relevance of the complement situation conveyed by the present under past can be ascribed to the nature of its present tense. The fact that TT_C in (8) is interpreted to have current relevance is due to the definition of present tense, which specifies the relation $TT \supset TU$ (i.e., topic time includes utterance time). Since there is no association line between the matrix and complement tenses, TU_C receives the same value as the value of the matrix TU , which is the utterance time. Since TT_C includes the utterance time by the definition of present tense, the claim made by the complement is naturally understood to have current relevance at the utterance time.

Under the current model, the time of Mary's being pregnant does not necessarily have to be the reflection of an actual situation that obtains at the speech moment. This is because TT or topic time as defined in the Klein's theory is only the time span to which the 'speaker's claim' on the described eventuality is confined, and it has no direct effect on the truth of the proposition in the model-theoretical sense. As long as there exists some situational factors that make the use of present tense in the complement desirable, whether or not Mary is actually pregnant does not matter. In this respect, the application of Klein's theory of tense to the analysis of the SOT phenomenon has another advantage. It provides an alternative explanation for another recalcitrant problem with SOT. The following sentence was pointed out by Abusch (1991) to be problematic to the analysis that assumes that the use of present tense embedded under past requires the complement situation to be actually true at the utterance time.

- (9) John said last month that Mary is pregnant but actually she has just been overeating for the last three months. (Abusch, 1991: 2, originally (5))

In (9), the use of present tense in the complement clause turns out to be grammatical, despite the fact that Mary is *not* pregnant, but has just been overeating. This constitutes a serious problem to the position that the use of the present tense under the matrix past requires the complement situation to be true at the speech moment.

Abusch (1991) explains that the use of present tense in the complement in the above example 'may only seem odd when Mary's symptoms (i.e., her big belly) cease sometime between John's believing time and the utterance time' (1991: 2). This is tantamount to saying that as long as the state that gives rise to the subject's belief (e.g.,

Mary's big belly in the above case) holds at the utterance time, one can safely use the present tense in the complement. The same fact was pointed out by Ogihara (1995b) with the following dialogue (1991: 188).

(10) John and Bill are looking into a room. Sue is in the room.

(a) John (nearsighted): Look! Mary is in the room.

(b) Bill: What are you talking about? That's Sue, not Mary.

(c) John: I'm sure that's Mary.

Sue leaves the room. One minutes later, Kent joins them.

(d) Bill (to Kent): # John said that Mary is in the room.

With example (10) Ogihara argues that the utterance *John said that Mary is in the room* is odd in a situation where the person who John mistook to be Mary (i.e., Sue) has already left the room. However, as I pointed out in Chapter 2 (section 2.3.3, (25)), the statement by Bill in (10d) turns out to be fine if we assume that Bill did not realize the fact that Sue left the room when he uttered (d).

Abusch's and Ogihara's examples, and the observation that (10d) can be an appropriate expression, are compatible with Klein's proposal on the nature of topic time. Klein defines topic time to be the time frame within which the speaker's CLAIM on the described eventuality is confined. That is to say, whether or not it reflects truth in the actual world is not necessarily crucial for the choice of tense. Under this view, the use of present tense in the complement in (9) and (10) is only to constrain the CLAIM made by the speaker to the time span that properly includes the utterance time. It does not necessarily have any implication to the truth of the proposition at the utterance time in the model-theoretical sense. As long as there exists a situation that makes the use of

present tense by the speaker relevant, it can be used within the limitation of the TT-TU relation encoded by the present tense. In other words, whether the use of present tense in the complement is felicitous or not may subject to an additional constraint imposed by situational factors or pragmatics. The system of tense, however, is independent of those pragmatic factors that may or may not make the utterance relevant for a given situation.

To summarize this section, the present-under-past construction in English, unlike the cases of past under past, does not involve the shift of anchor for the interpretation of the complement tense. The DA reading of present under past can be explained by the intensionality created by the matrix verb and the definition of the present tense in Klein's theory of tense. The TT (topic time) as defined by Klein turns out to be a strong tool for the explanation of the SOT phenomenon. Not only it is compatible with the interpretations of more common cases of past under past and present under past in English, but it also provides a straightforward explanation for the problematic cases that have been discussed in Abusch (1991) and Ogihara (1989, 1995).

3.1.2.3 Tense or aspect?

So far, it seems that the temporal overlapping between the topic time of the complement with the topic time of the matrix clause can straightforwardly be explained by an extra condition on the temporal interpretation imposed by the theory of intensionality. However, the following data reveal that an additional consideration is needed in order to embrace wider range of data.

(11) John said that Mary ran a mile.

(shifted interpretation only)

In (11), despite the existence of the intensional context imposed on the complement by the matrix verb ‘say’, we do not obtain a simultaneous interpretation between the matrix and the complement events. Intensionality does play a role in excluding an interpretation in which the complement event is temporally subsequent to the matrix ‘saying’. We are, however, left with only one possibility of temporal ordering of two events rather than the two possibilities that were available for the past under past construction in (3). While (3) allows both simultaneous and shifted readings, (11) allows only a shifted reading. It seems that the lack of a simultaneous interpretation in (11) comes from the difference between (3) and (11) in the complement eventuality. What is the difference between *Mary was pregnant* in (3) and *Mary ran a mile* in (11)? Zagana (1997) provides an answer to this puzzle.

Arguing against theories of tense that assumes tense-specific licensing conditions to account for the dependency of the complement eventuality on the matrix eventuality, Zagana (1997) revealed that relevant restrictions are observed in complements of nominals as well. Based on this fact, she argued that event-event dependencies should be ascribed to the aspect of the complement eventualities. Zagana points out that there is a significant contrast between culminating transitions and non-culminating transitions in the complement in terms of their temporal dependencies on the matrix tense. Her argument is twofold: i) while a non-culminating complement eventuality allows temporal dependency on the matrix eventuality, a culminating complement eventuality does not allow such dependency and is construed as temporally disjoint, and ii) the distinction between culminating and non-culminating eventualities holds across clausal and nominal complement categories and hence should be independent of the tense-licensing conditions.

What is relevant for the discussion at hand is the first part of her argument, and hence we contrast her data of non-culminating and culminating eventualities in finite clauses embedded under past. (Zagona, 1997: 49; originally (55) and (57) respectively)

(12) Non-culminating complement eventualities (Precedence / Inclusion)

- a. Sue heard that Mary designed bridges.
- b. Sue heard that Mary jogged at Greenlake.
- c. Sue heard that Mary baked bread.

(13) Culminating complement eventualities (Precedence / *Inclusion)

- a. Sue heard that Mary announced it at the meeting.
- b. Sue heard that Mary discovered a new planet.
- c. Sue heard that Mary designated some new members.

Between (12) and (13) only the sentences in (12) allow the simultaneous reading. The complement eventualities in (13) are necessarily disjoint from the time of hearing. This is because culminating transitions as in the complements of (13) are ‘bounded’ while non-culminating transitions as in the complements of (12) are ‘durative’.

The predicates in (12a-c) are property-denoting, and so they may occur in the present tense. On the other hand, the predicates in (13a-c) can only be eventive. Thus, the complement eventualities in (13) may never be expressed in the present tense. See the following present-under-past sentences, which are the non-SOT version of the sentences in (12) and (13).

(12') Non-culminating complement eventualities

- a. Sue heard that Mary designs bridges.
- b. Sue heard that Mary jogs at Greenlake.
- c. Sue heard that Mary bakes bread.

(13') Culminating complement eventualities

- a. *Sue heard that Mary announces it at the meeting.
- b. *Sue heard that Mary discovers a new planet.
- c. *Sue heard that Mary designates some new members.

The ungrammaticality of (13' a-c) confirms that the reason why the sentences in (13) lack simultaneous interpretation is not because of the tense system, but simply because the complement eventualities in (13) may not be interpreted to be durative under any circumstances, or at least in a pragmatically plausible sense.

Zagona's distinction between culminating and non-culminating eventualities can be applied to explain the difference between *Mary was pregnant* in (3) and *Mary ran a mile* in (11). States like *Mary was pregnant* corresponds to Zagona's non-culminating eventualities while *Mary ran a mile* corresponds to Zagona's culminating eventualities. Non-culminating eventualities are durative, and have an implication of possible continuation outside the time frame of its topic time. Thus, when we say that *Mary was pregnant*, we do not necessarily encompass the entire period of pregnancy with its beginning and end, nor do we claim that the situation does no longer obtain at the speech moment. For example, one can say, without any contradiction, something like: *Mary was pregnant when I saw her at the party three months ago, and she told me that she was expecting her baby in a couple of months, but I don't know if she has already delivered her baby or not*. Thus, despite the definition of past tense, which places TT anterior to TU, unbounded nature of stative predicates gives rise to an effect of temporal continuation of a described situation outside the temporal frame of its TT.

In contrast, *Mary ran a mile* in (11) is most likely taken to be an event, which is a temporally bounded occurrence. Description of an event may encompass its beginning and end, and there is no implication that it may obtain outside of its topic time. Thus, a statement like: *Mary ran a mile yesterday, but I don't know if she ran a mile today* refers

to two separate events, not one. Thus, the crucial difference in temporal interpretations of *John said that Mary was pregnant* and *John said that Mary ran a mile* is ascribed to the aspect of the complement eventuality.

However, as we see in Zagana's data presented in (12), it is not the case that non-stative verbs like 'run' in past tense always render eventive readings. See the example in (14) to illustrate this point.

(14) John said that Mary ran a mile every Sunday.

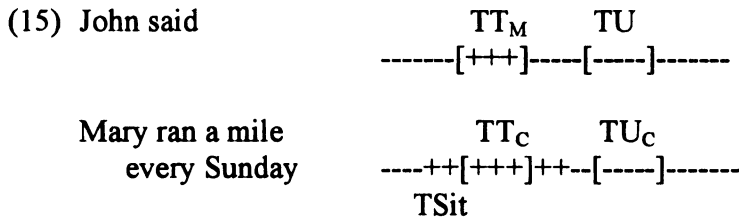
- a. John said, "Mary runs a mile every Sunday." (simultaneous interpretation)
- b. John said, "Mary used to run a mile every Sunday." (shifted interpretation)

(14) may have two different readings: one in which the embedded eventuality is taken to be co-temporal with the matrix saying time and another reading in which the embedded eventuality is considered to refer to the time interval prior to the matrix saying time.

The difference between (11), which receives an eventive reading on the one hand, and (14), which receives a durative reading on the other, is that the latter has an adverbial expression *every Sunday*. Apparently, this adverbial expression quantifies over events and provides a habitual interpretation for the sentence *Mary ran a mile*. Based on Mourelatos' (1978) distinction among events, processes, and states, we assume that events are countable occurrences while processes are mass or uncountable occurrences. However, as pointed out by Vlach (1993), frequency adverbials such as *often*, *usually*, *every week*, and *whenever Larry sneezed*, create process sentences. They convert a single-event denoting sentence into a sentence which refers to a pattern of events. In this sense, we can assume that *every Sunday* in (14) changed *Mary ran a mile* in (11) from an

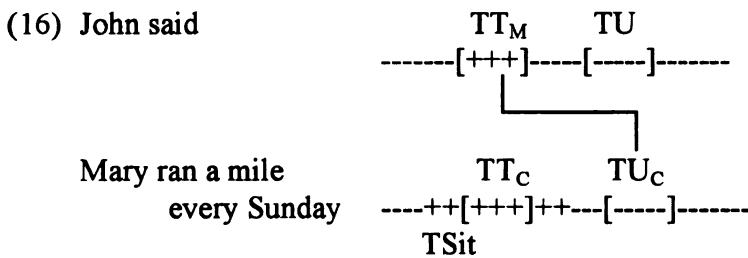
event into a process. A frequency adverbial expression changes the aspect of the sentence from terminative into durative.

Going back to the two possible readings of (14), the current framework explains each reading as follows. (15) is the temporal structure for the simultaneous reading (14a).



In (15) there is no association line that indicates a shift of the deictic center for the interpretation of the complement tense. The temporal dependency of the complement situation on John's saying time is ascribed to the intensional context created by the matrix verb just as in the case of (5) that we discussed earlier.

The shifted reading in (14b), on the other hand, is represented with an association line between the matrix TT and the complement TU.



The association line between TT_M and TU_C in (16) makes complement $TSit$ and the matrix TU necessarily disjoint. This will provides the shifted reading. However, since the complement eventuality is durative, there is an implication that the situation may obtain outside TT_C . Therefore, despite that fact that the content of what John said is

confined to some past interval prior to the time of saying (which is marked by TT_C), this does not exclude the possibility that Mary still ran a mile when John reported that she did, even though it may not have been known to John. Thus, with the temporal structure given in (16) we predict the correct reading of (14).

The two interpretations of (14) partterns with the interpretations of (3) in terms of the possible temporal relation exhibited between the matrix and the complement eventualities. What is common to (3) and (14) is that the compliment eventualities are DURATIVE. The situation described in the complement in (3) refers to a state and the eventuality described in the complement in (14) refers to an individual-level property of the subject, both of which is durative, or ‘non-culminating’ in Zagana’s term.

Thus, we may be able to conclude that the complication of the interpretation of past-under-past constructions in English is at least partially ascribable to the aspect of the complement eventuality, and that DURATIVE aspect of complement eventualities is a necessary condition for the simultaneous interpretation of past-under-past constructions. Thus, in addition to the intensionality imposed by the matrix verb, one needs to take into consideration the aspect of the complement eventuality in order to correctly predict the interpretation of the temporal relationship between the matrix and complement eventualities. In the remaining parts of section 3.1, we will focus our attention on well-discussed SOT data with stative complements, and will come back to the issue of the role of aspect and event structures in section 3.2 and section 3.3.

3.1.3 Interpretation of complement tenses in Japanese

This section attempts to account for the interpretation of complement tenses in Japanese in contrast to the interpretation of complement tenses in English discussed in the previous section. Let us first observe the Japanese data repeated from earlier chapters as (19) and (20) below, in comparison with the English data repeated here as (17) and (18).

(17) John said that Mary was pregnant.

a. John said that Mary had been pregnant.

(shifted reading)

b. John said, “Mary is pregnant.”

(simultaneous reading)

(18) John said that Mary is pregnant.

(DA reading)

(19) John-wa Mary-ga ninsinsi-te-i-ta to it-ta

John-TOP Mary-NOM pregnant-ASP-PAST COMP say-PAST

‘John said that Mary had been pregnant.’

(shifted reading only)

(20) John-wa Mary-ga ninsinsi-te-i-ru to it-ta

John-TOP Mary-NOM pregnant-ASP-NON PAST COMP say-PAST

a. John said, “Mary is pregnant.”

(simultaneous reading)

b. John said that Mary is pregnant.

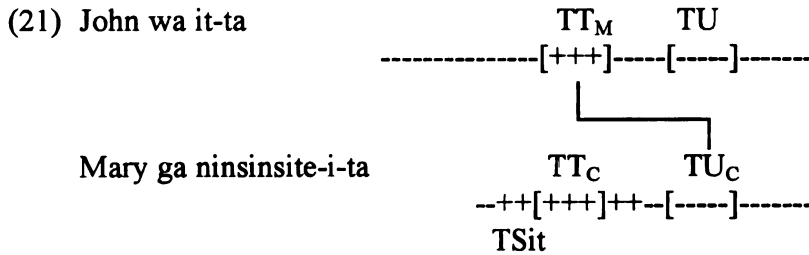
(DA reading)

Recall that in order to have a simultaneous interpretation, the Japanese counterpart of (17) must mark the complement verb with the non-past tense morpheme *-ru* as shown in (20). The surface equivalent (i.e., the past-under-past structure) of (17) in Japanese shown in (19) does not have a simultaneous reading. If we assume that past tenses in both languages encode the same TT-TU relation (namely, $TT < TU$), how should we explain the difference between English and Japanese as we observe here?

In order to account for the different behavior of past tense morphemes in English and Japanese as we see in the above data, I propose that the tense system in Japanese, unlike the tense system in English, involves an obligatory shift of deictic center for the interpretation of complement tense. One piece of empirical support for such a claim

comes from the fact that the SOT data in Japanese patterns the same as the data of the direct speech in English (cf. section 3.1.1, Table 1). Since the direct speech involves an obligatory shift of the deictic center, as we observed in section 3.1.1, it would not be implausible for us to expect that the SOT data in Japanese, which patterns the same as in the indirect-speech data in English except for DA reading of (20), may also involve the obligatory shift of the deictic center for the interpretation of the complement tenses.

Assuming that the deictic center for the interpretation of complement tenses in Japanese shifts from the time of utterance to the TT of the main clause, the structure for the interpretation of past under past like (19) is just like the one for its English counterpart in (17) with shifted reading. See the temporal structure in (21) for the Japanese sentence in (19).



*Subscript M and C stand for ‘matrix clause’ and ‘complement clause’ respectively.

The temporal structure in (21) is exactly the same as the temporal structure given earlier for the shifted reading of the past-under-past construction in English (section 3.1.1, (7)). Since the association line in (21) makes TT_C and TT_M disjoint, it is natural that the sentence receives only the shifted reading but not the simultaneous reading. Thus, the above structure in (21) correctly predicts the interpretation of (19) without positing anything special about the nature of the past tense in Japanese that is different from the

The temporal structure for the present-under-past construction such as (20), on the other hand, should look like (22) below.

John wa it-ta

TT_M TU
-----[+++]-----[-]-----
 └──────────┘

Mary ga ninsinsite-i-ru

TT_C ⊃ TU_C

-----++[[+++]]+-----
 TSit (TT is indicated by outer brackets)

101

(23) John said

$$\begin{array}{c} TT_M \quad TU \\ \text{-----} -[+++]\text{-----} [-]\text{-----} \end{array}$$

Mary is pregnant

$$\begin{array}{c} TT_C \supset TU_C \\ \text{-----} ++[[+++]]+-\text{-----} \\ TSit \end{array}$$

There is no shift in deictic center for the interpretation of the complement clause involved in (23) as indicated by the absence of an association line. Since the default anchor for the interpretation of any tense is the time of utterance, the value of TU_C , as well as for the value of the TU_M , is the utterance time. Then, TT_C , which properly includes TU_C , is understood to be contemporaneous with the utterance time. Hence, the interpretation of the complement tense in (23) is always the present time. As we have discussed earlier in section 3.1.2.2, the intensionality of the matrix verb makes TT_C dependent on TT_M independently of the tense system. Thus, TT_C now encompasses both the time of saying and the utterance time, and we obtain the obligatory DA reading for (23). In the case of Japanese present under past in (22), however, the temporal overlap between the complement eventuality and the matrix TU is only implied by the durative nature of the complement $TSit$. The inclusion of TU by TT_C is not obligatorily imposed by the tense system, unlike the English case.

The above explanation for the difference between the obligatory nature of the DA reading of present under past in English and the optional nature of the DA reading in Japanese is preferable to Ogiwara's account in a few respects. First, the account presented above does not have to stipulate a rule in the grammar whose sole existence is to explain the above difference between English and Japanese. As we reviewed in

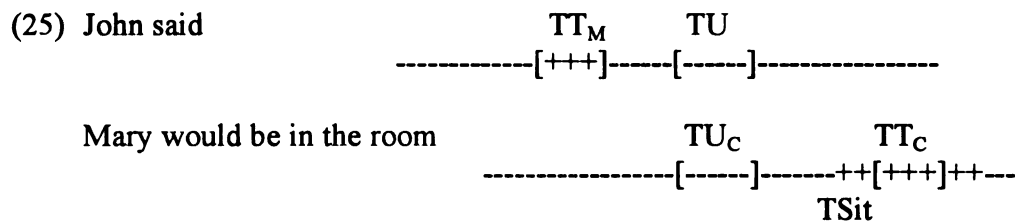
Chapter 2 (section 2.3.3), Ogihara proposed so-called ‘tense-deletion’ rule, which obligatorily applies to English but only optionally applies to Japanese. This rule, which lacks an independent motivation, also involves an illicit syntactic movement as pointed out by Kusumoto (1996), and therefore is not a desirable candidate to be part of the grammar of natural language. In contrast, the proposal made here posits the shift of deictic center for the interpretation of complement tenses in Japanese, but no other special rules in the grammar. Although assuming the obligatory shift of deictic center for the interpretation of tenses in Japanese but not for English is merely a stipulation at this point, an operation that shifts the deictic center for deictic expressions itself is not uncommon in natural language (e.g., use of present tenses for narration of past events in many languages). Second, we have empirical data in support of our analysis, namely, the fact that past tense in complements in Japanese behaves very similar to past tense in the direct speech, which involves the shift of anchor for deictic expressions (section 3.1.1). Third, by assuming that the difference between English and Japanese in tense interpretations lies in the nature of the shift of the deictic center, we can keep the definition of PAST in these two languages the same, namely $TT < TU$. Lastly, so far we have no counterexamples that cannot be accounted for by our analysis. The following section, however, presents the data that requires some revision in the functions of temporal expressions in English that we presented as hypotheses in Chapter 1 (section 1.4.2, Table 1).

3.1.4 Problematic cases and revision of future tense in English

Recall that in Chapter 1 we tentatively assumed that an auxiliary *will* in English encodes the $TT > TU$ (TT AFTER TU) relation as a future-tense marker. However, the following sentence raises a problem for such a position.

(24) John said that Mary would be in the room.

Let us suppose that (24) has two possible readings just like (3): the one with a shift of deictic center for the complement-tense interpretations and the one without it. In the latter reading, we have the following temporal structure.

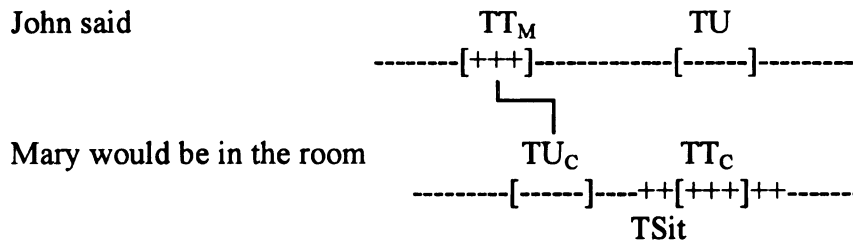


Since the temporal dependency of the complement situation to the matrix event can be ascribed to the intensional context created by the intensional verb ‘say’ in the matrix clause, we can explain the dependency of the time of the complement situation to the matrix ‘saying’ without an association line between TT_M and TU_C in (25). However, when there is no association line, the default value of TU_C is the utterance time. If we assume that *would* is the future-tense marker, we would necessarily place TT_C posterior to the utterance time as you see in (25). This goes against our intuition that the relative location of the state of *Mary’s being in the room* in (24) is indeterminant with respect to the utterance time.

Another problem with positing the temporal structure (25) for the sentence in (24) lies in the fact that a sentence like *John said that Mary will be in the room* would have the same temporal structure. If a single temporal structure corresponds to two surface forms, we would have the same problem as we had with the traditional analysis of SOT. Namely, we cannot explain the difference in interpretation between a sentence with *will* and a sentence with *would*, and we cannot explain why they have different morphological realizations on the surface forms.

What happens if we assume that (24) involves a shift of deictic center for the interpretation of the complement tense? The temporal structure would then be like (26).

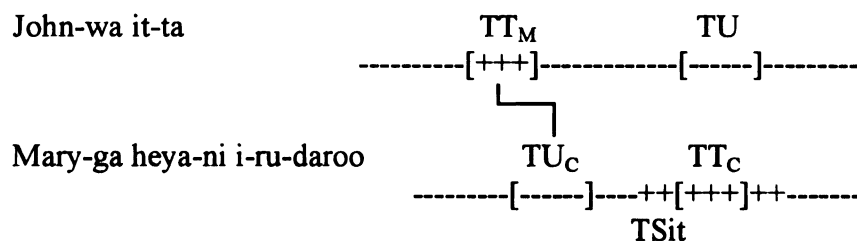
(26) John said



The temporal structure in (26) seems to be consistent with our intuition that the relative temporal location of Mary's being in the room is indeterminant with respect to the utterance time. However, this is not necessarily an ideal choice for us, either. The most serious drawback would be the lack of an independent motivation for making the shift of anchor obligatory only in cases when the complement predicate has a future-tense marker. Furthermore, if we posit the shift of the deictic center for the interpretation of the complement tense for an English sentence like (24), we need to explain why the complement verb has a past-tense morphological counterpart of *will* (i.e., *would*) in (24),

while a Japanese counterpart of (24) would have a non-past surface form for the same temporal structure. Observe a Japanese sentence and its temporal structure in (27).

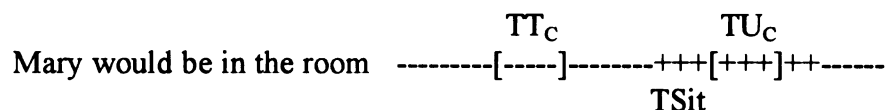
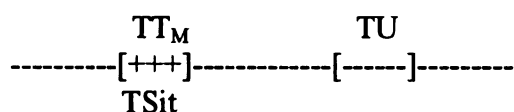
- (27) John-wa Mary-ga heya-ni i-ru-daroo to it-ta
 John TOP Mary NOM room in be-NON PAST-MODAL COMP say-PAST
 ‘John said that Mary would be in the room.’



Since Japanese lacks a grammatical device that uniquely marks the TU < TT (TT AFTER TU) relationship, the Japanese counterpart of (24) shown in (27) has a modal *daroo*, which represents the speaker’s conjecture on the possibility of some future eventuality. When we compare (26) and (28), we find that *will* is realized as *would* in (26), while *daroo* appears in a non-past tense form in (27), despite that (26) and (27) have the same temporal structure. Then, again, we face the same old problem: why is there a sequence of tense in English while Japanese doesn’t have one? Thus, even if we posit the shift of deictic center for the complement tense interpretation, our problems with the temporal structure for a sentence like (24) remain, given our current assumptions.

Based on the above facts, I propose a revision of the nature of an auxiliary *will* in English. Instead of positing that it is a grammaticized future-tense marker, I argue that it is better understood as encoding *prospective aspect*, which places TSit after TT. Let us examine the revised temporal structure for (24), which is shown in (28).

(28) John said



As we defined in chapter 1, the grammaticalized category of aspect encodes the relation between TT and TSit. TSit represents the time duration that is associated with the lexical content of the predicate (whether it is taken as a point in time or an interval). We now re-define the nature of *will* in English to be encoding an aspectual property of $\text{TT} < \text{TSit}$ (situation time is after the topic time). When we say, *we will go to the party*, the claim made with this utterance is confined to the present time, but it is a claim about our prospect of some future eventuality. Thus, when *will* occurs in its past-tense form *would*, the claim is confined within some time span before the utterance time, but the claim is, again, about the prospect for the eventuality that is placed posterior to the time of the claim.

Based on the above revision on the nature of *will* and *would*, the temporal structure in (25) illustrates the following generalizations:

- i) Past tense on both the matrix and the complement predicates introduce the relation $\text{TT} < \text{TU}$ (the topic time precedes the utterance time), and both the time of saying in main clause and the time of prediction in the complement clause receive past-time denotation.

- ii) The simultaneous reading (i.e., temporal overlap between the time of saying and the time of prediction) can be attributed to the effect of the intensional context imposed on the complement clause by the intensional verb 'say' in the matrix clause.
- iii) Posteriority of the complement situation with respect to the 'saying' time is explained by the aspect of *would* in the complement (which encodes the relation $TT < TSit$) and the temporal overlap between the matrix TT and the complement TT (which is constrained by the theory of intensionality).
- iv) Indeterminacy of the relative location of the complement situation with respect to the utterance time can be explained by the fact that the complement $TSit$ can be either before/on/after the time of utterance, given that the aspect of the complement predicate only specifies the relation $TT > TSit$, and does not tell where $TSit$ should be located with respect to TU_C .

With this revision of the temporal property of an auxiliary *will/would*, we can successfully explain the future-under-past construction in English, which have been problematic to the analyses of tense that have been proposed to date.

The above proposal for the nature of *will/would*, however, should not be considered to cover all the occurrences of these auxiliaries in natural language discourse. Both *will* and *would* as modal expressions may be used to represent the speaker's attitude toward a proposition¹. Under the assumptions of the present study, the fact that these auxiliaries can be used as modals should not be incompatible with the assumption that they also encode aspect. Since the system of temporal representation is assumed to exist independently from other modules of the grammar and a single morpheme can be

understood to encode information from different modules, it is perfectly reasonable for us to assume that *will* and *would* may encode both modal and aspectual properties. However, since it is not the purpose of the present study, we confine ourselves to exploring the purely temporal properties of these expressions.

The revision made on the nature of auxiliaries *will/would* also explains the following data with present and past embedded under the main predicate with *will* as well. Before we examine these data in the model adopted in this study, however, let me introduce how they are accounted for by Hornstein's (1990) theory of tense, which is similar to Klein's system of temporal reference.

- (29) a. John will believe that Mary was sick.
b. John will believe that Mary is sick.
c. John will believe that Mary will be sick.

Hornstein (1990) proposed an analysis of the SOT phenomena with the neo-Reichenbachian theory of tense, in which he uses Reichenbach's three temporal primitives: S (speech time), R (reference time), and E (event time) for representation of tense structures (cf. Chapter 2, section 2.6). In Hornstein's theory, S, R, and E are syntactic primitives and the SOT rule is a syntactic rule, which applies to the basic tense structure (BTS) and gives rise to a derived tense structure (DTS). The above sentences receive the following structures in Hornstein's theory (1990: 130, (25)).

¹ According to Lyons (1995), subjective modals express the speaker's belief or attitude toward the proposition, or his/her will or authority. I consider English *will* to be an example of this type of modal.

- (30) a.
$$\begin{array}{ccc} S_R, E_1 & & S_R, E_1 \\ E_2, R_S_2 & \xrightarrow{\text{SOT}} & \begin{array}{c} | \\ E_2, R_S_2 \end{array} \end{array}$$
- b.
$$\begin{array}{ccc} S_R, E_1 & & S_R, E_1 \\ S_2 R, E_2 & \xrightarrow{\text{SOT}} & \begin{array}{c} | \\ S_2 R, E_2 \end{array} \end{array}$$
- c.
$$\begin{array}{ccc} S_R, E_1 & & S_R, E_1 \\ S_2_R, E_2 & \xrightarrow{\text{SOT}} & \begin{array}{c} | \\ S_2_R, E_2 \end{array} \end{array}$$

The structures on the left are the temporal representations of (30a-c) before the application of the SOT rule (i.e., BTS), and the structures on the right are those after the application of the SOT rule (i.e., DTS). According to Hornstein, the sentences in (29a-c) are ambiguous between two readings and his theory predicts this ambiguity by the temporal structures before and after the application of the SOT rule. Before its application, independent readings are available for the complement tenses, while after the applications of the SOT rule, only the interpretations in which the complement eventualities are dependent on the matrix event time is available. This is, at least on the surface, very similar to what our theory would propose. However, two important facts cannot be explained in Hornstein's theory. One is the dependency of the complement eventuality on the matrix 'believing' time when the SOT rule does not apply. Another problem is that we do not have an account for why the morphological change on the complement predicate only applies when tenses are embedded under past and not when tenses are embedded under future.

The sentences in (29a-c) all exhibit the dependency of the complement eventuality with the believing time of the subject of the matrix clause, regardless of the application of the SOT rule. That is, if John believes something at time t_1 , then the content of his belief

must have some relevance to his belief at time t_1 independently of the SOT rule. In other words, there has to be some kind of mechanism that will connect the content of belief and the time at which the belief holds. Hornstein explains this temporal dependency between the complement and the matrix clauses to be the outcome of the application of the SOT rule, which associates E_1 and S_2 as we see on the right in (30). However, the structures on the left in (30) cannot represent the temporal dependency of the complement eventuality on the believing time. This temporal dependency for both of the two possible readings of the sentences in (29a-c) can straightforwardly be explained by the intensional context imposed on the complement eventuality by the matrix intensional verb *believe*, which applies uniformly regardless of the presence or absence of the shift of deictic center for the interpretation of the complement tense. Since our theory assumes that the intensional context is imposed on the complements by the matrix intensional verbs independently from the tense system, Hornstein's theory of tense may incorporate this idea to account for the temporal dependency of the embedded tense on the matrix tense that obtains independently of the SOT rule. However, as far as I understand, that is not the position that he takes.

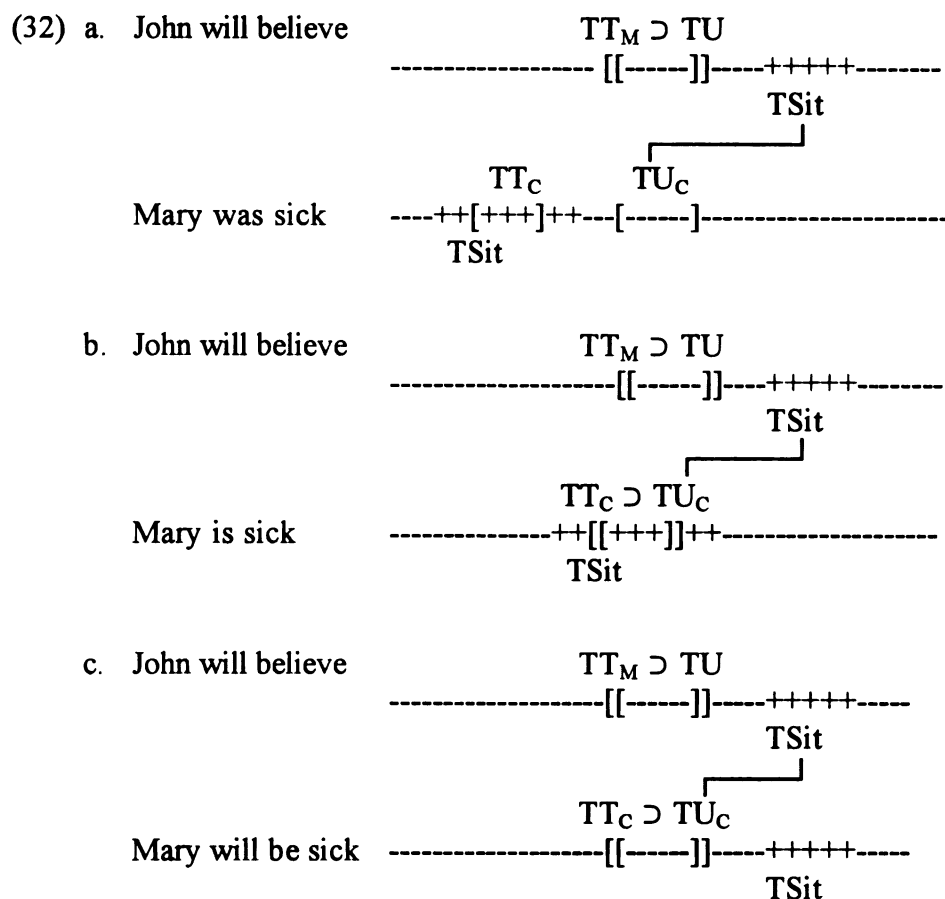
Another problem of Hornstein's theory lies in the assumption that a morphological change on the complement predicate that accompanies application of the SOT rule only applies when the complement tenses are embedded under the past tense. Since he assumes that a morphological change does not apply when tenses are embedded under the future as in the case of (29a-c), we do not see any change in the surface representations of the complement verbs after the application of the SOT rule. This is the same problem as the one that we faced with the traditional SOT theory which

assumed transmission of the morphology of the matrix predicate to the complement predicate only when tenses are embedded under past. The problem for both the traditional analysis and Hornstein's theory is that there is no explanation as to why this syntactic rule only applies when the matrix verb has past tense morphology. Furthermore, Hornstein's account of SOT would not be able to explain why the Japanese counterparts of the English SOT sentences never undergo morphological change of the complement predicate either. Such problems do not arise under the present analysis of SOT. Let us examine how (29a-c) can be explained in the theory of temporal reference proposed in this study.

- (31) a. John will believe $TT_M \supset TU$
 -----[[-----]]-----+++++-----
 TSit
- Mary was sick TT_C TU_C
 ----++[[+++]]+-[-]-----
 TSit
- b. John will believe $TT_M \supset TU$
 -----[[-----]]-----+++++-----
 TSit
- Mary is sick $TT_C \supset TU_C$
 -----++[[++++]]+------
 TSit
- c. John will believe $TT_M \supset TU$
 -----[[-----]]-----+++++-----
 TSit
- Mary will be sick $TT_C \supset TU_C$
 -----[[-----]]-----+++++-----
 TSit

(31a-c) represent the temporal structures of (29a-c) when there is no shift of the deictic center for the interpretation of the complement tense. Since the temporal dependency of the complement eventuality on the matrix believing time is the natural consequence of the requirement of the theory of intensionality, we do not need the SOT rule to account for this kind of dependency. As there is no shift of the deictic center for the complement tense interpretation (as is clear from the structural representations in (31a-c) that do not have an association line), tenses in the complement clauses in (31a-c) all receive an independent interpretation, capturing one of the two possible interpretations pointed out by Hornstein. Yet, due to the intensional context imposed by the matrix intensional verb *believe*, the temporal overlap between the matrix believing and the content of the subject's belief is captured, reflecting our intuitions as to the temporal relationship between the two eventualities involved in the description. Thus, our theory resolves one of the problems raised by Hornstein's account of (29a-c), namely, lack of the explanation for the temporal overlap between the matrix believing time and the complement eventuality before the application of the SOT rule, without any additional stipulation, which is a welcome result.

The other interpretations in which the relative temporal locations of the complement eventualities are evaluated with respect to the matrix believing time are derived when we shift the deictic center for the complement tense interpretation to the time span within which the lexical content of the matrix verb *believe* is associated with.



In one of the two interpretations for (29a), the time of Mary's sickness is indeterminant with respect to the matrix TU (or the utterance time). This is captured nicely in the temporal structure shown in (32a) which does not tell us relative location of TT_C with respect to the matrix TU. Since the aspectual property of *will* merely states that TSit is posterior to TT (which temporally overlaps with TU), and the complement past tense simply introduces the relation $\text{TT}_C < \text{TU}_C$, if we associate the matrix TSit with the complement TU_C , it may well be the case that TT_C is somewhere between TU and TSit of the matrix clause, but it can also be anterior to the matrix TU.

In one of the two possible interpretations for (29b), the present relevance of Mary's being sick is lost, and the complement eventuality is understood to be

simultaneous with the time of believing. This is also captured in (32b) with the same mechanism and assumptions. By associating TU_C to the matrix $TSit$, the deictic center for the interpretation of the complement eventuality shifts to the time of the matrix eventuality, and therefore, TU_C may not receive the utterance time as its value anymore. Then, the complement eventuality expressed by *Mary is sick* cannot be understood to be overlapping with the utterance time represented by the matrix TU , which explains why the current relevance of the complement eventuality is lost in this interpretation.

Finally, the fact that (29c) has an interpretation in which the complement eventuality is placed posterior to the matrix believing time is also captured by the present proposal without any further assumptions. In (32c) where an association line is drawn between the matrix $TSit$ and the complement TU_C , which indicates the shift of deictic center for the interpretation of the complement tense. Accordingly, the eventuality of Mary's being sick, which is placed posterior to the TT_C by the aspectual property of *will* in the complement clause, finds its location after the matrix $TSit$ (or believing time). This captures our intuition as to one of the available temporal interpretations of this sentence. Thus, not only does our theory solve one of the problems with Hornstien's theory, but it also provides accurate descriptions of the same set of data.

Another problem with Hornstein's account, namely, the lack of explanation for cross-linguistic facts, can also receive a natural account in the present study. Recall that the Japanese counterpart of a sentence like *John said that Mary would be in the room* (cf. the temporal structure given in (28)) does not have a morphological transmission of the matrix past onto the complement predicate. This is partly because Japanese does not have a lexical correlate of the auxiliary *will*, and a semantically closer counterpart, the

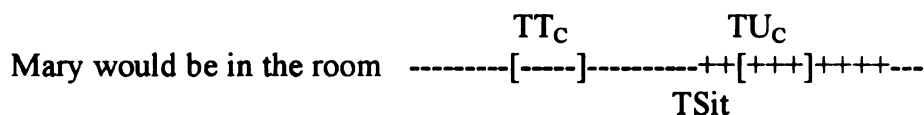
modal expression *daroo* (which is a derivative of a copula *da*) does not have a past-tense counterpart². And yet, it is possible for Japanese to express the same semantic content with non-past tense marking on the complement predicate as you see below.

- (33) John-wa Mary-ga heya-ni i-ru-daroo to it-ta
 John TOP Mary NOM room in be-NON PAST-MODAL COMP say-PAST
 'John said that Mary would be in the room.'

Hornstein only deals with English data and so it may not do justice to him if I simply say that he doesn't provide explanation of our Japanese data. However, since he assumes that tenses are abstract concepts and should be separated from their carriers (i.e., the tense morphemes), we should take what he proposes to be the properties of the tense system should apply uniformly to the tense system of other languages. Thus, assuming that Japanese also has past tense just like English, the property of past tense that is applicable to English should also apply to Japanese. For the same reason, the SOT rule as a syntactic rule should apply to Japanese as well, if nothing prevents its application. However, we have no account of why Japanese lacks a SOT rule and the morphological transmission of the matrix tense onto the complement tense. Under the present analysis, we do not need to assume a syntactic rule such as SOT or the transmission of the past tense morphology from the matrix predicate to the complement predicate in order to provide a unitary and consistent account of both the English and Japanese data. All we need to assume is that the shift of anchor is obligatory in the complement structure in Japanese while it is optional in English. Since Japanese does not have a past-tense counterpart of *daroo* for representing the future-under-past construction, the system

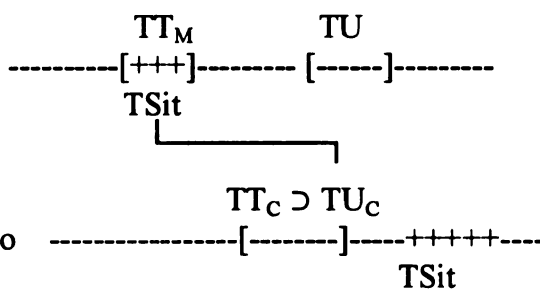
² Japanese has a form *dat-ta-roo*, which consists of a past-tense counterpart of a copula

(34) John said



In contrast, see the following temporal structure for a Japanese counterpart.

(35) John-wa it-ta



117

Since Japanese lacks a morphological counterpart of *daroo* with the same function as *would* in English, the anteriority of the time of the embedded TT with respect to the matrix TU must be encoded by some other means. The obligatory shift of the deictic center for the complement tense interpretation (which is indicated by the association line between TU_C and $TT_M/TSit$) will do this task. The indeterminacy of the location of the complement situation with respect to the utterance time is captured in this case as well just as in the case of (34), as there is nothing that tells us the relative location of the complement TSit with respect to the matrix TU in the above temporal structure. Thus, the theory of temporal reference presented in this study can predict these hitherto problematic cases with no further assumptions than those that have been supported independently.

The present proposal also solves another problem which was unexplained in the traditional analysis and in Hornstein's theory, namely, the fact that the SOT rule only applies when the matrix verb has a past morphology. Under the current proposal, the past tense morphology on the complement predicate is the realization of past tense. No morphology transmission rule is assumed. Thus, the fact that past tense embedded under future retains the past tense morpheme does not need to be explained as failure of the application of the SOT rule.

In our earlier discussion on the dependency of the time of the complement eventuality on the matrix eventuality (which is ascribed to the theory of intensionality), we did not make it clear which of the three temporal primitives of the description of the main eventuality the embedded eventuality should be associated with. As is obvious in the discussion above, the embedded eventuality depends on TSit of the matrix

eventuality. This is specifically because the content of ‘saying’ should be understood to be dependent on the time associated with the act of saying. In description of the past-under-past constructions, as long as the matrix verb is in the simple past, TT and TSit overlap completely due to the aspect of the simple past-tense form (TT AT TSit) and the verbal aspect of ‘say’ (which is NON-DURATIVE). This would be different if the main predicate represents some other aspect. In the present study, however, I do not intend to investigate all the possible variations of the combinations of tense and aspect in the complement constructions. Rather, I will focus on the kind of data that have received much attention from the previous studies on tense.

In the preceding sections we have seen how Klein’s system together with the notion of shift of anchor can explain the apparent difference in behavior of tense morphemes of English and Japanese in complement clauses. The apparent difference in behavior between tense morphemes of English and tense morphemes of Japanese is ascribed to the obligatoriness of the shift of deictic center for the interpretation of complement tenses in Japanese and the optional nature of such shift in English. The model presented here allows the function of past tense to stay constant across English and Japanese (i.e., to encode the relation $TT < TU$). The next section examines the behavior of temporal morphemes in both languages in adverbial clauses, and will see if the functions of the primary tenses can also be maintained across different constructions.

3.2 Tenses in adverbial clauses

3.2.1 Semantics of *when*-clauses and *toki*-clauses

This section examines the difference in behavior of tense morphemes of English and Japanese in adverbial clauses. Among various different types of adverbial clauses, the present study focuses on the occurrence of tenses in *when*-clauses in English and *toki*-clauses in Japanese. Before we move onto our analyses, however, some assumptions on the nature of *when*-clauses and *toki*-clauses need to be introduced.

Extending the temporal representation models developed by Hinrichs (1981, 1986) and Partee (1984) to description of temporal subordination structures in English, Spejewski and Carlson (1991) provide an analysis of *when*, arguing that *when* introduces a temporal subordination structure without specifying particular ordering relationship between the eventuality described by the *when*-clause and the one described by the main clause. For example, the following sentences taken from Spejewski and Carlson show that we can have any temporal ordering between the eventuality of the *when*-clause and the one of the main clause (332, originally (14)).

- (36) a. When Pam went to Chicago (e_1), she put her dog up in a kennel (e_2).
b. When Jean made the pancakes (e_1), she used molasses in the batter (e_2).
c. When Phil came into the house (e_1), he took his coat off (e_2).

In (36a) e_2 is understood to have happened before e_1 , in (36b) e_2 is understood to co-occur with e_1 , and in (36c) e_2 is understood to follow e_1 .

We assume that the nature of *toki* in *toki*-clauses in Japanese bears a close affinity to the nature of *when* in *when*-clause in that it introduces a temporal subordination structure by creating an eventuality frame with which the eventuality expressed by the

main clause coincides. Thus, when we say “ e_1 -*toki*, e_2 ”, the eventuality expressed in the main clause e_2 is understood to be co-temporal with the occasion of the eventuality e_1 introduced by the *toki*-clause, and *toki* itself does not dictate any particular ordering relation between the two eventualities.

Assuming the above mentioned nature for both *when* and *toki*, let us examine the behavior of *-ru* and *-ta* in *toki*-clauses in contrast to tenses in *when*-clauses in English.

3.2.2 Interpretation of tenses in *toki*-clauses

As we observed earlier in Chapter 1, both *-ru* and *-ta* seem to behave differently in *toki*-clauses from complements. The following examples are repeated from section 1.3.2.

- (37) Tokyo-ni ik-u toki kare-ni at-ta
 Tokyo to go-NON PAST when he to meet-PAST
 ‘I met him when I was going to Tokyo.’
 $at-ta (TT_M) < ik-u (TT_A) < TU$
- (38) Tokyo-ni it-ta toki kare-ni at-ta
 Tokyo to go-PAST when he to meet-PAST
 ‘I met him when I went to Tokyo.’
 a) $it-ta (TT_A) < at-ta (TT_M) < TU$
 b) $at-ta (TT_M) < it-ta (TT_A) < TU$
- (39) Tokyo-ni it-ta toki hikooki-no naka-de kare-ni at-ta
 Tokyo to go-PAST when plane-GEN inside-at he to meet-PAST
 ‘I met him on the plane (to Tokyo) when I went to Tokyo.’
 $at-ta (TT_M) < it-ta (TT_A) < TU$
- (40) Tokyo-ni it-ta toki kare-ni a-u
 Tokyo to go-PAST when he to meet-NON PAST
 ‘I will meet him when I get to Tokyo.’
 $TU < it-ta (TT_A) < a-u (TT_M)$

- (41) Tokyo-ni ik-u toki kare-ni a-u
 Tokyo to go-NON PAST when he to meet-NON PAST
 'I will meet him when I get to Tokyo.'
 a) $TU < a-u (TT_M) < ik-u (TT_A)$
 b) $TU < ik-u (TT_A) < a-u (TT_M)$
- (42) Kondo Tokyo-ni ik-u toki Shinjuku-de kare-ni a-u
 next time Tokyo to go-NON PAST when Shinjuku in he to meet-NON PAST
 'I will meet him in Shinjuku next time when I go to Tokyo.'
 $TU < ik-u (TT_A) < a-u (TT_M)$

As we observed earlier in section 1.3.2, tenses in *toki*-clause in Japanese, unlike tenses in *when*-clause in English, does not need to show tense agreement with the matrix tense. Thus, in (37), despite the fact that the evaluation time for the sentence as a whole is in the past, the tense in the *toki*-clause is marked with NON-PAST tense marker, and in (40), even though the statement is about some future eventuality, the predicate in the *toki*-clause has PAST tense marking. Based on these facts, some Japanese linguists (e.g. Ando 1986, among others) claim that *-ru* and *-ta* in Japanese are not tense markers but aspectual markers. The proponents of this view claim that *-ru* encodes imperfective and *-ta* encodes perfective or perfect³ respectively as aspectual forms but do not encode tense. They would argue that the use of *-ru* in (37) indicates incompleteness of the subordinate eventuality with respect to the time of the matrix eventuality, while the use of *-ta* in (40) indicates completion of the subordinate eventuality with respect to the matrix eventuality. However, this amounts to saying that *-ru* in the subordinate clause always indicates the simultaneity or posteriority of the subordinate eventuality to the matrix eventuality, and *-ta* in the subordinate clause always indicates the anteriority of the

³ The proponents of the aspectual view of the form claim that *-ta* encodes *kanryoo* in Japanese, which may be translated as either 'perfect' or 'perfective'. However, since

subordinate eventuality to the matrix eventuality. The behavior of *-ru* and *-ta* in the subordinate clauses as described here is not incompatible with our view that *-ru* encodes $TT \supset TU$ or $TT > TU$ and that *-ta* encodes $TT < TU$ as tense markers.

What is expressed by *-ru* in the *toki*-clause in (37) and by *-ta* in the *toki*-clause in (41) is the relative location of the subordinate eventuality with respect to the matrix eventuality. The non-past tense marker *-ru* in the *toki*-clause in (37) indicates that the subordinate event *Tokyo-ni iku* ‘to go to Tokyo’ is either contemporaneous or posterior to the matrix event *kare-ni au* ‘to meet him’, and the past tense marker *-ta* in the *toki*-clause in (40) indicates that the subordinate event is placed before the matrix event. In other words, the deictic center for the evaluation of the subordinate tenses in these examples is shifted from the utterance time to the matrix TT.

To view *-ru* and *-ta* as a relative tense marker is not a new proposal. In support of Soga (1983), who analyzed their use in the relative tense system, Ogihara (1999) recently argued against the aspectual view of *-ru* and *-ta*, saying that they are the tense markers which are not speech time oriented, but are evaluated with respect to the time of the subordinate events. However, the present study cannot commit to this position, ‘either, since the use of *-ru* and *-ta* in embedded clauses is not necessarily always event-time oriented. For example, (41) indicates that *-ru* has a speech-time oriented usage, in which the subordinate eventuality described with *-ru* may be placed prior the matrix eventuality (i.e., (41b)). Also, (38) indicates that *-ta* does not always show the completion or the anteriority of the described eventuality with respect to the matrix

they do not provide definitions of perfect and perfective, it is not clear which aspect they mean by ‘kanryoo’.

eventuality. (38b) clearly indicates that *-ta* in the subordinate clause may also have a speech-time oriented use.

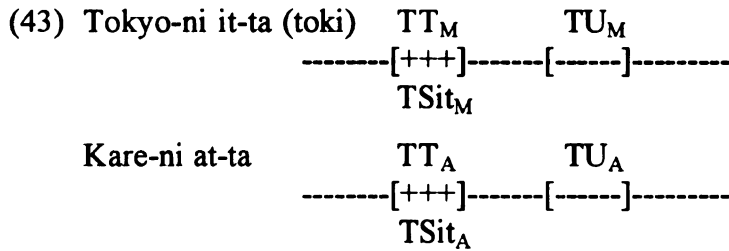
The optionality in the tense marking of the subordinate eventuality as we see in (37) and (38b), and (40) and (41b), reveals another important fact about tenses in *toki*-clauses in Japanese. That is, the shift of the deictic center for tenses in *toki*-clauses in Japanese, unlike in the case of complement tenses, is not obligatory.

One may argue that the preferred readings for (38) and (41) are (38a) and (41a) respectively. However, the mere fact that they are preferred does not mean that the grammar may not allow other readings. As shown in (39) and (42), the preferred readings are cancelable with a pragmatic bias towards readings with the opposite ordering between the matrix and the embedded eventualities. This is a strong piece of evidence in favor of the position of the present study that the temporal morphemes *-ru* and *-ta* in Japanese encode NON PAST ($TT \supset TU$ or $TT > TU$) and PAST ($TT < TU$) respectively as tense markers, and that the apparent complexity in their behaviors and the major difference between the English and the Japanese systems can be ascribed to the difference in the nature of the shift of the deictic center between the two systems.

Having said that the shift of the deictic center for the interpretation of tenses in *toki*-clauses is optional, we need to explain why it is so, compared to the obligatoriness of the shift of the deictic center for the interpretation of tenses in complement clauses. Apparently, this asymmetry comes from the fact that only the verb-complement construction, but not the adverbial construction, provides an indirect speech environment. Assuming that the shift of the deictic center for the tense interpretation is required in the

indirect speech (cf. section 3.1.1), we do not expect the adverbial construction to obligatorily shift the deictic center for the interpretation of the subordinate tense.

Based on the above generalization, the temporal structure for a sentence like (38) with past-tense marking on both the matrix and the subordinate clauses look like the one in (43).



*Subscript M and A stand for ‘matrix clause’ and ‘adverbial clause’ respectively.

In (43), tenses in both the matrix and the subordinate clauses will receive an independent reading. In other words, the temporal ordering between TT_M and TT_A is free. This explains why the preferred ordering $e_1 < e_2$ in “ e_1 -*toki*, e_2 ” is cancelable by a pragmatic bias towards the opposite ordering $e_2 < e_1$ as given in (39).

As we observed with the examples from Spejewski and Carlson (1991) in section 3.2.1, *when*-clauses in English with the past tense marking on both the matrix and the subordinate clauses may represent any temporal ordering between the matrix and the subordinate eventualities. While each sentence in (36) represent one particular ordering (i.e., either $e_1 < e_2$, e AT e , or $e_2 < e_1$), it may also be the case that a single sentence indicate more than one possibility of temporal ordering between the matrix and the subordinate eventualities. The following sentence, which is an English counterpart of (38), exhibits

the same ambiguity in terms of the temporal ordering of the matrix and the subordinate eventualities as (38) does.

(44) I met him when I went to Tokyo met < went or went < met

In (44), the meeting time can be either before or after the speaker's arrival in Tokyo: it is possible that meeting took place at the airport at the departure gate or on the plane, but it is also possible that it happened in Tokyo. The fact that there is freedom in temporal ordering of the matrix and the subordinate eventualities in (44) suggests that an eventuality described in *when*-clause may receive an independent reading, and that it has the same temporal structure as its Japanese counterpart shown in (43).

So far, there seems to be no major differences between the English and the Japanese tense systems except for the fact that the latter may involve a shift of the deictic center for the interpretation of the tenses in *toki*-clauses, while tenses in *when*-clauses in English are always evaluated with respect to the speech time. In other words, tenses in *when*-clauses in English must always be interpreted under 'the absolute tense system' while tenses in *toki*-clauses in Japanese may occur either in the absolute or 'the relative tense system' in Comrie's sense (1985). Thus, contrary to the commonly held view among Japanese linguists that there is a significant difference between tenses in English and tenses in Japanese, the semantics of basic tenses in both languages are revealed to be basically the same: the past tense in both languages encodes $TT < TU$, and the present tense in English and the non-past tense in Japanese both encodes the relation $TT \supset TU$ (though the latter also encodes the relation $TT > TU$ on top of $TT \supset TU$). Once again, this is due to the fact that the Japanese tense system may shift the deictic center for the

subordinate-tense interpretation from the speech time to the matrix topic time that the tenses in Japanese seems to behave differently from the tenses in English. The following data further support these generalizations.

- (45) a. Ne-ru toki ni denki-o kesi-ta
 go to bed NON PAST when light ACC turn off PAST ($e_1 = e_2$, OR $e_2 < e_1$)
 'I turned off the light when I was going to bed.'
- b. Ne-ta toki ni denki-o kesi-ta
 go to bed PAST when light ACC turn off PAST ($e_1 < e_2$, $e_1 = e_2$, OR $e_2 < e_1$)
 'I turned off the light when I went to bed.'
- c. *Ne-ta toki ni denki-o kesi-ta
 sleep PAST when light ACC turn off PAST
 '?I turned off the light when I slept.'

Let us suppose that e_1 represents the first eventualilty on the left (i.e., the eventuality in the *toki*-clause) and e_2 represents the matrix events on the right of the sentence string. In (45a) the non-past tense in *toki*-clause is evaluated with respect to the matrix topic time. According to our generalizations, we would expect that (45a) can be rephrasable with the past tense marking on the e_1 . As shown by (45b), our prediction is borne out: the temporal ordering of two eventualities that are allowed in (45a) can also be expressed with past-tense marking on the subordinate predicate as shown in (45b). This conforms to our analysis that we may have either the absolute or the relative tense system for tenses in *toki*-clauses. The ordering relation between the two eventualities are free in (45b), supporting our analysis that tenses in *toki*-clauses may receive an independent reading when both the matrix and subordinate predicates have the same tense. The English translation of (45b) also exhibits a free ordering of the two eventualities. Both in English and in Japanese, if e_1 and e_2 are interpreted to be two separate eventualities, we would

have the ordering $e_1 < e_2$. However, if e_2 is interpreted to have taken place during the time frame introduced by e_1 , either $e_1 = e_2$, or $e_2 < e_1$ becomes possible. This seems to conform to the generalizations made by Dowty (1986) that eventualities in narratives are interpreted to have occurred in the order they appear in the discourse unless such readings are cancelled by pragmatics or our knowledge of the world.

In contrast to (45a,b), (45c) in which the subordinate verb *neru* in Japanese is translated as ‘to fall asleep’ rather than ‘to go to bed’ turns out to be ill-formed. This is because e_1 cannot subsume e_2 within its time frame with a given interpretation, based on our knowledge of the world (i.e., turning off the light cannot be taken as a sub-event of falling asleep). Then, we are left with an option to interpret this sentence to represent two separate eventualities that happened in sequence. In this reading, however, we interpret (45c) to represent the temporal ordering of e_1 and e_2 in which e_1 (i.e., ‘to fall asleep’) precedes e_2 (i.e., ‘to turn off the light’). Since we do not normally expect one to be able to turn off the light in sleep, (45c) is ruled out by our pragmatic knowledge.

Thus, we may conclude that contrary to a commonly held view, the tense systems of English and Japanese are essentially the same except for the fact that the Japanese system may employ the relative tense system on top of the absolute tense system. The past tense in both languages encodes the relation $TT < TU$, and the present tense in both languages encodes the relation $TT \supset TU$ (with an additional $TT > TU$ for Japanese). The tense interpretations allowed with the use of respective tense morphemes in both languages subject to the same constraints imposed by these TT-TU relations regulated by the tense system. This is not surprising if we assume, as argued by Hornstein, that ‘the tense system constitutes an independent linguistic level’ and that ‘the mappings from

tense morphemes to temporal interpretations respect the formal constraints imposed by this level' (1990: 9).

Following Comrie (1985), the relative tense system uses a reference point other than the speech moment as the deictic center for the interpretation of tenses. In case of tenses in *toki*-clauses, the reference point shifts from the speech time to the time of the matrix eventuality. Unlike in the case of the complement construction, however, this shift of the reference point is optional in *toki*-clauses due to the lack of the indirect-speech environment. The optionality of the shift of the deictic center for the interpretation of tenses in *toki*-clause is supported by the empirical data in which the present tense under the matrix past tense can be replaced by the past tense with no change in temporal ordering of the two eventualities (cf. (37) and (38); (40) and (41)). The question is: what drives such a shift of the deictic center in *toki*-clause tenses, if it is optional? At this point, the only motivation that we could posit is to avoid the ambiguity that arises with the use of the same tense in both the matrix and the *toki*-clause. Recall that while both (37) and (38) encode $e_1 > e_2$, (38) with the past-tense on both clauses may also encode the opposite ordering: $e_1 < e_2$. (37) with the relative tense in the *toki*-clause, on the other hand, unambiguously represents the temporal ordering $e_1 > e_2$. In the same way, while (41) with the non-past tense on both clauses are ambiguous between $e_1 < e_2$ and $e_1 > e_2$, (40) with the relative tense in *toki*-clause unambiguously represent $e_1 < e_2$. In both of the above cases, when there is some other information that helps us determine the temporal ordering as in the case of (39) and (42), we do not necessarily have preferences for marking the tenses in the *toki*-clauses with the relative tense. Thus, I am inclined to believe that the optional shift of the deictic center for the interpretation of tenses in *toki*-

clauses is to avoid ambiguity in the interpretation. Though this is not a strong argument, it seems to be that all the data that we have seen so far are in support of this analysis.

3.2.2 The ordering of events: tense or pragmatics?

It seems to be the case that interpretations of tenses in both English and Japanese subject to the same pragmatic constraints. While the tense system by itself may allow any ordering of the main and the subordinates eventualities when they are marked with the same tense morphemes, our knowledge of the world may place an extra constraint on the ordering relationship between the two eventualities. Hence, we may be left with one particular ordering of two eventualities rather than having an ambiguous sentence. Then, the fact that bi-clausal sentences allow different possibilities for the ordering of the multiple eventualities while others have a single possible reading should not be ascribed to the complexity of the tense system. It is the complication that arises by interaction of the tense system with things outside the tense system.

The examples from Spejewski and Carlson (1991) presented in (36) in section 3.2.1 also conform to the above generalizations. Their sentences are repeated in (46).

- (46) a. When Pam went to Chicago (e_1), she put her dog up in a kennel (e_2).
b. When Jean made the pancakes (e_1), she used molasses in the batter (e_2).
c. When Phil came into the house (e_1), he took his coat off (e_2).

While a preferred ordering of the two eventualities in (46a) may be $e_2 < e_1$, the reverse ordering (i.e., $e_1 < e_2$) is also possible. If e_1 is interpreted to mean an occasion of Pam's visit to Chicago with all the preparations, e_2 is taken to be part of e_1 due to our knowledge of the world, and thus, can be interpreted to have taken place before Pam's going to

Chicago. However, if e_1 is interpreted to be Pam's action of going to Chicago, e_2 cannot be part of e_1 and therefore, $e_1 < e_2$ becomes the only possible ordering. In (46b), two eventualities are interpreted to be co-temporal because e_2 is taken to be part of e_1 due to our knowledge of the world. In contrast, (46c) is interpreted to represent the temporal ordering of $e_1 < e_2$, as it is a description of two separate actions in time sequence. In cases like (46c), where two eventualities can only be interpreted to represent two separate events that took place in the time sequence, one may not have the ordering $e_2 < e_1$ with simple past-tense making on both the matrix and the subordinate predicates. In Japanese, however, one may employ the relative tense system to represent the posteriority of the subordinate eventuality to the matrix eventuality. Compare the following two sentences.

- (47) a. Heya-o de-ta toki denwa-ga nat-ta
 room ACC go out PAST when phone NOM ring PAST
 'When I left the room, the phone rang.'
- b. Heya-o de-ru toki denwa-ga nat-ta
 room ACC go out NON PAST when phone NOM ring PAST
 'When I was about to leave the room, the phone rang.'

In (47a) with past tense in *toki*-clause gives us an interpretation in which the phone rang after the speaker went out of the room (i.e., $e_1 < e_2$). This is the only available reading for the ordering of the two culminating eventualities, both of which being marked by the past tense form, since the eventuality in the *toki*-clause cannot be interpreted to subsume the matrix eventuality in its time frame (i.e., *denwa-ga nat-ta* 'the phone rang' cannot be taken as subevent of *heya-o de-ta* 'went out of the room'). In contrast, (47b) with non-past tense in the *toki*-clause gives us an interpretation in which the speaker heard the phone ring when he or she is still in the room or at the door. In the second case, non-past

tense in *toki*-clause is evaluated with respect to the matrix TT and shows posteriority to the matrix eventuality. The English translation of (47b), on the other hand, keeps the past-tense marking on the predicate in *when*-clause, but employs the progressive form to indicate the imperfectivity of the subordinate eventuality with respect to the matrix eventuality. This is because English does not employ the relative tense system in *when*-clauses unlike *toki*-clauses in Japanese.

As mentioned earlier, the data like (47a,b) are often used as a piece of evidence for the claim that these morphemes are not tense markers, but aspectual markers. The proponents of such aspectual view of these temporal morphemes would argue that *-ta* in (47a) does not have a pure past-tense reading, but only indicates the anteriority of the described eventuality with respect to the matrix eventuality, and therefore, should be considered as a perfective marker. However, the fact that (47a) may exhibit only one possible ordering of the matrix and the subordinate eventualities cannot be ascribed to the peculiarity of *-ta* in Japanese. In fact, the English counterpart of (47a) also displays the same limitation on the temporal ordering of the two eventualities. This is the effect of the interaction of the aspect of the subordinate eventuality (which is ‘culminating’ or ‘bounded’) and the pragmatics (i.e., that we cannot take the subordinate eventuality in (47a) to subsume the matrix eventuality), which applies to both English and Japanese in the same manner. The proponent of the aspectual view of *-ru* and *-ta* may also argue that the fact that the subordinate eventuality in a sentence like (47b) has to be marked by *-ru* for the given interpretation indicates that *-ru* cannot be a tense marker, since the sentence as a whole refers to the past eventuality. However, we already know that such an argument may be valid only in the view of the absolute tense system, and that the data is

not incompatible with the definition of the non-past tense as encoding the temporal relation $TT \supset TU$ or $TT > TU$, in which TU is shifted to the time of the matrix eventuality. Furthermore, the fact that (47b) cannot be rephrased by replacing *-ru* with *-ta* is not because of the peculiarity of the Japanese tense system, either. It is because the resulting surface form may only allow the ordering of the two eventualities as given in (47a), but not the one in (47b), which is the effect of the pragmatics of the described eventualities. These facts have not been noticed by the proponents of the aspectual view of *-ru* and *-ta* nor those who consider that the data like (47a,b) are problematic to the position to view these morphemes as tense markers.

To recapitulate the observations provided above, the following generalizations can be made. The data that allow multiple possibilities for temporal ordering and those that allow only one interpretation differ in the relation between the matrix and the subordinate eventualities they exhibit. In the case of those that allow multiple possibilities for the ordering, the *when*-clause or *toki*-clause provides a temporal frame within which the matrix eventuality falls. In the case of those which allow only one interpretation, two eventualities that are represented as co-temporal by the semantic contribution of *when* or *toki* are assumed to occur in a certain time sequence in the real world. Observe the following examples for the former case, in which eventualities described in the *toki*-clauses provide a temporal frame for the description of the subordinate events.

- (48) Nihon-ni it-ta toki Deteroito-no kuukoo-de kamera-o kat-ta
 Japan to go PAST when Detroit GEN airport at camera ACC buy PAST
 'When I went to Japan, I bought a camera at the Detroit Airport.'
- (49) Nihon-ni it-ta toki Akihabara-de kamera-o kat-ta
 Japan to go PAST when Akihabara in camera ACC buy PAST
 'When I went to Japan, I bought a camera in Akihabara.'

In both (48) and (49), the eventuality described in the matrix clause is presented to fall within the time frame introduced by the *toki*-clause. That is, the subject's act of purchasing a camera took place during the time of his/her occasion of visit to Japan. In (48), the subject's going to the Detroit Airport to take the airplane to Japan is also taken to be part of this occasion. That the eventuality described in the *toki*-clause in (49) is understood to be providing a temporal frame within which the matrix eventuality falls can be supported by the fact that it can be rephrased as in (50) below, where the past-tense form of a stative verb *iru* is used in place of a change-of-state verb *iku* without much difference in meaning.

- (50) Nihon-ni i-ta toki Akihabara-de kamera-o kat-ta
 Japan in be PAST when Akihabara in camera ACC buy PAST
 'When I was in Japan, I bought a camera in Akihabara.'

Since the subordinate eventuality in (50) is taken to be co-temporal with the matrix eventualities, it is possible to replace *-ta* in the subordinate clause in (50) with *-ru* without no change in the overall interpretation of the sentence. However, the grammaticality of a sentence like (50) in which *-ta* in the subordinate clause does not indicate the perfectivity of the eventuality with respect to the matrix eventuality is enough to refute the arguments of the proponent of the aspectual view of this temporal morpheme.

To find out what predicates allow the multiple ordering possibilities and what do not is not our concern for the purpose of the present study, as it is up to the pragmatics of the ontology of different kinds of eventualities. There are innumerable possibilities for

the combinations of eventualities which may or may not allow an inclusion relationship. One thing that may be of our interest is the effect of the aspect of the subordinate eventuality on availability of such inclusion relationship. For example, if the subordinate eventuality is durative as in (50), most likely we will have the inclusion relationship between the matrix and the subordinate eventualities due to the unbounded nature of the durative aspect. However, as we can see in (48) and (49) above, having a non-durative, or culminating eventuality, in the subordinate clause may not necessarily lead to a single ordering of the two eventualities. Thus, while aspect may contribute to the allowable interpretations, we must refer to the semantic content of both eventualities in the description and the pragmatic effect that arises from their interactions to determine the actual temporal ordering of the described eventualities. Certainly, this cannot be considered to be part of the system of temporal reference, and thus, we do not pursue the enumeration of the available possibilities in this study.

In the next section, we will extend our framework and the generalizations to the analysis of relative-clause tenses.

3.3 Tenses in relative clauses

3.3.1 Tenses in relative clauses in English

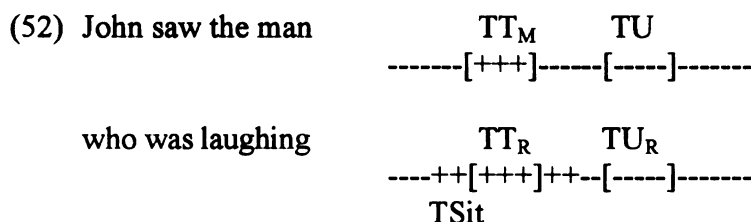
Tenses in relative clauses have been pointed out to behave differently from tenses in complement clauses in many previous studies on tense (e.g., Enç, 1987, Abusch 1988, Ogihara 1989, among others). This section examines whether those differences can be attributable to the system of temporal reference or to something else. The purpose of this section is not to provide an exhaustive listing of different types of relative clause

tense data and explain them in detail, but to provide a uniform account of the behavior of tense morphemes of English and Japanese, focusing on the crucial data that have been the object of much discussion by previous studies on tense.

The following example from Ogihara (1989: 96, (29)), which we discussed earlier in section 1.3.3 (example (11)), is repeated here to refresh our memory as to how tenses in relative clauses behave differently from tenses in complements.

(51) John saw [_{NP} the man [_{S'} who was laughing]].

As pointed out by Ogihara and by many others, a sentence like (51) can be interpreted to represent any temporal ordering between the matrix and the embedded eventualities. This is apparently very different from interpretation of tenses in the complement construction, in which the complement eventuality can not be interpreted to be posterior to the matrix eventuality. Since the lack of forward-shifted reading of the complement tense is attributable to the intensional context created by the matrix intensional verb in the complement construction as we discussed earlier in this chapter (cf. section 3.1.1), this difference between interpretation of tenses in complements and that of tenses in relative clauses can be explained as stemming from the absence of an intensional context in relative clause examples. Thus, the difference is not the result of the complications of the tense system. Let us see the temporal structure of (51), and see how our model explains this data.

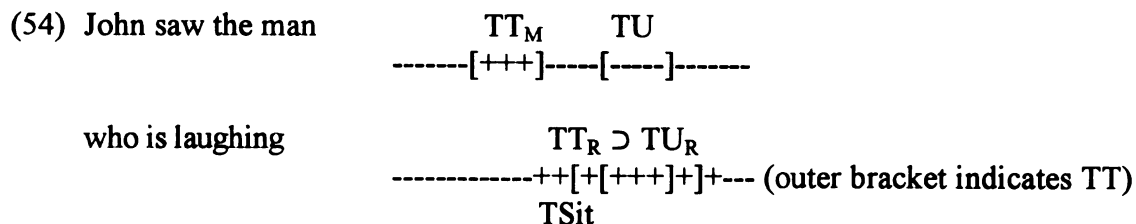


As there is no intensional context that makes the embedded eventuality temporally dependent on the time of the matrix eventuality, tenses in both clauses are independently evaluated with respect to the utterance time. Since the temporal ordering of TT_M and TT_R is not specified by the tense system (as there is no association line between the matrix and the relative clauses), TT_R can be either before, after, or simultaneous with TT_M . Thus, the model presented here correctly predicts all and only available readings for (51) without speculating anything that is special about tenses in relative clauses. Now, how about present-under-past in relative clauses?

(53) John saw [NP the man [S' who is laughing]].

The above sentence with present-tense marking on the predicate in relative clause can only be interpreted as John's past act of seeing the man who, at the time of speech, is laughing. Crucially, there is no DA (= double-access) reading available for (53), unlike present-under-past in complements such as *John said that Mary is pregnant*. The lack of a simultaneous reading for (53) can be ascribed to the absence of an intensional context in the relative clause as opposed to the complement structure just in the case of (51). However, the fact that we cannot interpret the time of man's laughing to be temporally overlapping with the time of John's seeing him is potentially problematic to our theory,

given the fact that the eventuality described in the relative clause in (53) is durative. Let us examine the temporal structure of (53) to see what the problem is.



While we do not assume any temporal dependency between the two eventualities represented in (54) that is required by the grammar, the model does not necessarily prohibit the temporal overlap of the TSit of the relative clause with the matrix TT_M . This is so because TSit of the relative clause has durative aspect. If we assume that the eventuality in the relative clause is unbounded on both ends, there is nothing in the model per se that prohibits the temporal overlap between the TSit of the relative clause and the topic time of the matrix eventuality. I would argue that unavailability of possible temporal overlap between *laughing* and *seeing* in (54) should be attributed, not to the system of tense, but to our pragmatic knowledge. In order to see this point, let us examine the interpretation of the following sentence which has the same relative clause structure and the tense marking but with different eventuality type in the relative clause.

(55) John saw the man who is allergic to fish.

In (55), despite the absence of intensionality, the man who is allergic to fish at the utterance time is understood to have had this symptom also at the time when John saw this man. This is simply because ‘being allergic to fish’ is a property of an individual, and our knowledge of the world tells us that such a property is not likely to change in a short

period of time. In contrast, ‘be laughing’ is only a stage-level property of an individual, and it is hard for us to imagine someone laughing continuously for an extended period of time. In case of (53), for example, we do not normally expect a man to be laughing from the last time John saw the man (which might have been a week ago or even a year ago) up to the speech moment. Thus, we have a piece of evidence to support that different interpretations for the same construction with the same tense marking can be ascribed to the aspect of the described eventuality and our pragmatic knowledge. This is a desired result, as the information encoded by the present and the past tenses may remain constant in this way.

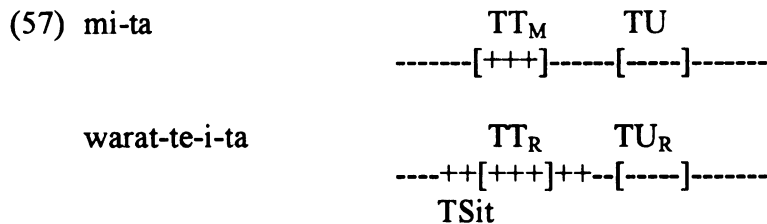
So far, the functions of present and past tenses in English presented as hypothesis in section 1.4.2 together with our theory of temporal reference have been able to account for the behavior of temporal morphemes in various different constructions. The next section examines the behavior of *-ru* and *-ta* in Japanese in relative clauses to see the semantics of past tense and the semantics of non-past or present tense can be maintained cross-linguistically.

3.3.2 Tenses in relative clauses in Japanese

The following sentence is a surface counterpart of English past-under-past in the relative clause construction in Japanese.

- (56) John-wa [NP warat-te i-ta otoko]-o mi-ta
 John-TOP laughing be-PAST man -ACC see-PAST
 ‘John saw a man who was laughing.’

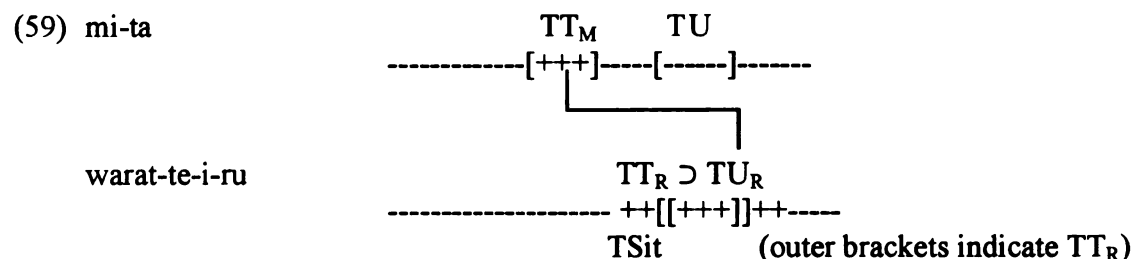
Unlike past-under-past in the complement construction, which does not allow a forward-shifted reading, (56) may allow any temporal ordering between the eventualities in the matrix and the relative clauses just like its English counterpart (51). This is indicative of the absence of an intensional context in the relative clause in the above example, which is expected if we assume that the theory of intensionality applies cross-linguistically. However, it also implies that there is no shift in the deictic center for the interpretation of the tense in the relative clause in (56), which is different from the cases of the complement structure in Japanese, where the shift of anchor for the complement tenses is obligatory. We assume that tenses in relative clauses in Japanese may receive an independent reading just like tenses in relative clauses in English, and (56) is assumed to have the following temporal structure.



Just as in the case of (52), temporal ordering of the matrix TT_M and the TT_R in the relative clause is indeterminant by the model, and we can correctly predict three different interpretations. However, we cannot assume that tenses in relative clauses in Japanese always receive an independent reading. Consider the surface equivalent of the English present-under-past in the relative-clause construction (as given in (53) above) in Japanese.

- (58) John-wa [NP warat-te i-ru otoko]-o mi-ta
 John-TOP laughing be-NON PAST man -ACC see-PAST
 a. John saw a man who was laughing. (simultaneous reading / *shifted reading)
 b. John saw a man who is laughing. (*DA-reading)

Unlike its English counterpart, which does not allow a simultaneous interpretation, (58) does allow a simultaneous interpretation by which the time of man's laughing is taken to be co-temporal with the time of John's seeing him as given in (58a). (58) may also have a reading in which the time of man's laughing has a present time relevance at the utterance time, but lacking a double-access reading. The second reading is the same as the interpretation of its English counterpart (53), but the first one cannot be accounted for with the temporal structure given in (54). How do we explain the existence of a simultaneous interpretation for (58)? To account for the similarities and differences between English and Japanese that we observe here, we assume that the system may optionally shift the deictic center for the interpretation of tenses in relative clauses in Japanese.



In (59), the association line between TU_R and TT_M indicates that the deictic center for the interpretation of the relative-clause tense is shifted from the utterance time to the matrix topic time, and the time of laughing is now co-temporal with the time of seeing. This will explain the simultaneous interpretation. Although the durative aspect of the eventuality in the relative clause should allow us to interpret the described situation to overlap the

utterance time, this is excluded on pragmatic grounds as we discussed in case of the English example in (53).

In fact, there is a reason to believe that past-under-past in relative clauses may also involve a shift of the deictic center for interpretation of tenses in embedded clauses.

- (60) Gakusee dat-ta hito to at-ta
student COP PAST person with meet PAST
'I met a person who had been a student.'

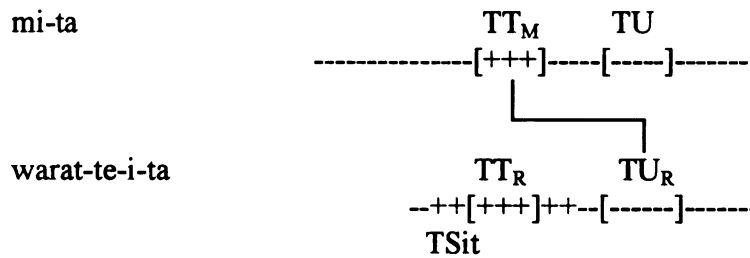
The eventuality in the relative clause in (60) is most naturally interpreted to obtain before the time of meeting as shown by the English translation. This indicates that there is a shift in the deictic center for the interpretation of the relative-clause tense. Yet, it is also possible to have a simultaneous interpretation with a relevant contextual information as illustrated by (60').

- (60') Sono toki gakusee dat-ta Tanaka-san to sono paatii de at-ta
that time student COP PAST Mr. Tanaka with that party at meet PAST
'I met Mr. Tanaka who was student then at the party.'

In (60'), the eventuality in the relative clause is understood to be co-temporal with the time of meeting, which suggests that there is no shift in the deictic center involved for the interpretation of the relative-clause tense.

Assuming that there are two possibilities for the temporal structure associated with past-under-past of the relative clause construction (i.e., the one that may allow the simultaneous interpretation, and the other for the shifted interpretation), the shifted reading of (56) is assumed to have the following structure rather than the temporal structure presented in (57).

(61) mi-ta



The temporal structure of (61) does not allow temporal overlap of TT_R with TT_M as the association line between TU_R and TT_M necessarily makes them temporally disjoint. This structure will give us a desired interpretation in which the speaker saw the man who was laughing at some earlier time (for example, when the speaker met this person at the party several months before the time of the matrix-clause eventuality).

The optional nature of the shift in deictic center in case of relative clauses can further be supported by the following data.

- (62) Asu koko-e ku-ru hito-ni kore-o watas-u
 tomorrow here-to come NON PAST person to this ACC hand NON PAST
 'I will give this to a person who will come here tomorrow.' a) kuru < watasu
 b) watasu < kuru

The ordering relation between the eventuality in the matrix clause and the one in the relative clause in (62) is underdetermined by the tense system, confirming our generalizations as to the availability of independent interpretation for the relative-clause tense. The interpretation in which the eventuality in the relative clause precedes the matrix eventuality can also be expressed by the following sentence in which the relative-clause event is marked with *-ta*.

- (63) Asu koko-e ki-ta hito-ni kore-o watas-u
 tomorrow here-to come PAST person to this ACC hand NON PAST
 'I will give this to a person who comes here tomorrow (when s/he arrives).'

In (63), the eventuality described with the past-tense morpheme *-ta* is evaluated with respect to the matrix TT rather than the utterance time, which unambiguously place the embedded TT anterior to the matrix TT. Now, let us see another example.

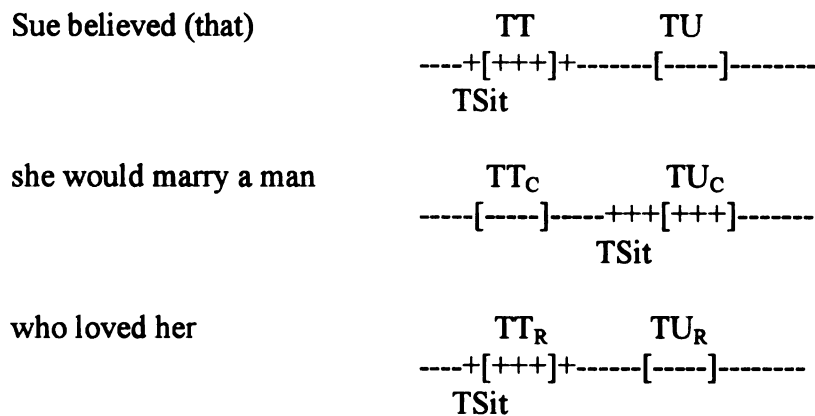
- (64) Paatii-ni ki-ta hito-ni at-ta
 party to come PAST person to meet PAST
 'I met a person who came to the party.'
- a) kita < atta
 b) atta < kita

Just as in the case of (62), the temporal ordering of the matrix and the relative-clause eventualities is underdetermined in (64), though the interpretation given in (64b) may be difficult to obtain without contextual information. The interpretation given in (64b) can unambiguously be obtained by replacing *-ta* in the relative clause with non-past tense marker *-ru*, which requires the shift of deictic center for its interpretation.

The above observation leads us to speculate that the use of relative tense for the relative-clause tense interpretation may be driven by pragmatic force to disambiguate the meanings underdetermined by the primary tenses. This is the same as in the case of the occurrence of relative tenses in adverbial-clauses.

Finally, the following data from Abusch (1997a), which we discussed in section 2.4.3 (example (32)), will receive more detailed explanation with the generalizations that we obtained so far.

(65) Sue believed that she would marry a man who loved her.



The problem with (65) in the previous theories is that we cannot explain why the time of loving (which is placed before the utterance time) is understood to be overlapping with the marrying time, which is understood to be posterior to the believing time. Also, an independent theory of tense combined with a theory of intensionality cannot explain why the relative-clause eventuality ‘loving her’, which is supposed to be outside the intensional context, cannot be placed after the time of marrying. In our theory of temporal reference, the temporal overlap between the marrying time and the loving time can straightforwardly be explained by unbounded nature of a stative verb *love*. An explanation of the second point lies in the fact that there is an alternative expression in English for that particular temporal relationship, namely, *Sue believed that she would marry a man who would love her*. Besides these advantages, our theory of temporal reference, by our revision of the nature of *will* in English as a marker for prospective aspect, allow us to capture the descriptive facts about the behavior of *would* (which is now understood to be a past-tense counterpart of *will*) in embedded clauses without introducing a highly ad-hoc morphology-copying rule (as proposed in the traditional

theory of SOT) or a tense-deletion rule (as proposed by Ogiwara (1989)), which is a desired result for the simplicity of the theory.

3.4 Division of labor and a note on the present tense

This chapter examined the behavior of temporal morphemes *-ru* and *-ta*, and the behavior of English present and past tense markers in different constructions: the complement structure, the adverbial *when*-clause structure, and the *relative*-clause structure. Tenses in all these constructions may receive an independent reading, unless there is additional operation required by other modules of the grammar (e.g., the theory of intensionality) or unless a particular ordering of the matrix and the subordinate eventualities are imposed due to the pragmatic force generated by the semantic content of the two eventualities in the description (section 3.2.2 and section 3.2.3). The temporal information of the sentence may also be affected by the aspect of the subordinate eventuality (section 3.1.2.3) and additional semantic contribution by another element in the sentence (e.g., *when* in *when*-clauses or *toki* in *toki*-clauses. The apparent complication in the interpretation of embedded tenses is, then, mostly attributable to the information outside the tense system, and the information encoded by the present and the past tenses in English and the non-past and the past tenses in Japanese can stay constant across different constructions. Furthermore, while most of the previous studies on tense in Japanese linguistics have highlighted the differences between the English tense system and the Japanese tense system, the preceding sections revealed that there are more commonalities than differences between them. The past tense encodes $TT < TU$ in both languages and the English present tense the Japanese non-past tense encodes $TT \supset TU$

(while the non-past tense in Japanese additionally encodes $TU < TT$). This makes perfect sense if we assume that tense is an abstract grammatical category that encodes certain TT-TU relation, and that temporal morphemes in a language, as tense markers, simply encode certain TT-TU relations as surface carriers of particular tenses. It is quite natural for us to believe that these tense morphemes may carry information other than that of tense (e.g., aspect and/or modal information), considering how tense systems in many languages of the world may have developed (cf. Dahl, 1985; Bybee et al., 1994). However, as carriers of tense, the past tense markers always encode the relation $TT < TU$ and the present tense markers encode the relation $TT > TU$ in any language that encodes these tenses in the grammar.

In terms of the mapping between the basic tenses and their carriers, as noted above, English and Japanese differ in that the non-past tense marker *-ru* in Japanese encodes both present and future tenses, while English PRESENT is assumed to encode the present tense only. However, now that we revised the nature of an auxiliary *will* in English from future tense marker to the marker of prospective aspect (cf. section 3.1.3), we wonder if the nature of the present tense marker in English should be redefined as well. It has been reported that there are many languages in the world in which the present tense is also used for the description of the future events (cf. Binnick, 1991). In English, this usage is also possible in some instances. Thus, the fact that the non-past tense marker *-ru* in Japanese may encode either present or future is not unusual in this respect.

Bybee et al (1994) points out that some of the functions of the PRESENT in English has been taken over by the PROGRESSIVE over several centuries, and that as the PROGRESSIVE takes over ‘part of an originally more general present’ of the PRESENT, the

English PRESENT was left to represent habitual and generic readings as a default reading (1994:150-151). Because of the fact that most occurrences of the English PRESENT actually have habitual or generic readings rather than the present, there are many linguists who claim that the PRESENT in English is not a present-tense marker. However, if we define the present tense as a grammatical device to encode the relation $TT \supset TU$, both habitual and generic readings of the English PRESENT will also fall under the occurrence of the present tense. Thus, under the system of tense employed in this study, there is nothing contradictory for calling the English present as the present tense marker, even though its occurrence with non-stative predicates normally have habitual or generic readings. The present progressive form of English also encodes the present tense in its composite part (i.e., the present tense form of the copula) for that matter. So one may safely conclude that the present-time denoting property of the present progressive form in English is simply inherited from its present-tense component which encode $TT \supset TU$.

In fact, what Jespersen (1924) and Reichenbach (1947) called as ‘tenses’ such as the present progressive and the present perfect forms in English encode both tense and aspect under the current framework. They were referred to by traditional grammarians as ‘complex tenses’ or ‘composite tenses’. If we apply Klein’s three temporal primitives that we employed for our analyses of the SOT data, both properties of these ‘complex tenses’ (i.e., their properties as tense markers and their properties as aspectual markers) can be captured nicely. The next section will examine the progressive form and the present perfect form in English in contrast to the *-te-iru* form in Japanese, and will demonstrate how we can unify the notions of tense and aspect in the current framework.

CHAPTER 4

Aspect and event representations

Introduction

The goal of this chapter is to show how Klein's theory of temporal reference can be fruitfully applied to aspect. It shows that Klein's system allows us to explain the cross-linguistic difference in the major aspectual forms in English and Japanese in a simple way with very few stipulations. Specifically, I will explain the puzzling behavior of the *Verb-te-iru* form in Japanese, which corresponds to both progressive and perfect readings in English.

In Klein's system of temporal representation, aspect is defined in terms of the relation between topic time (TT) and situation time (TSit). Assuming that the progressive in English encodes IMPERFECTIVE aspect, it provides the TT-TSit relation in which TT is included in TSit¹. On the other hand, the English perfect, which is realized as an auxiliary *have* followed by the past participial of a verb, encodes TSit < TT, representing the PERFECT aspect as defined by Klein. In Chapter 1 I presented a hypothesis that *-te-iru* in Japanese may represent either IMPERFECTIVE or PERFECT as defined by the respective TT-TSit relationship provided above. It is a puzzle, however, how it is possible for a single form to encode such different aspects as imperfective and perfect. In this chapter I present an analysis of *-te-iru* in which the source of the ambiguity of the sentences with this form can be explained in a very simple manner, applying Klein's theory of temporal reference in a novel way. I will argue that unlike the progressive and the perfect in English, each of which specifies a specific relationship

¹ This should not be taken to be synonymous as a claim that this is the only information encoded by the English progressive, since it is not the case.

between the topic time (TT) and the situation time (TSit), *-te* of *-te-iru* simply provides TT without information about its placement with respect to TSit. This temporal property of *-te* allows the sentences with *-te-iru* to have any placement of TT with respect to TSit that is allowable in the model.

The organization of this chapter is as follows: section 4.1 shows how Klein's system is applied to the analysis of the progressive and the present perfect in English and *-te-iru* in Japanese. The goal of this section is to show how the current system helps us capture the complex nature of these temporal morphemes and to clarify their similarities and differences for the later discussions. In section 4.2 I will propose a single semantic characterization of all the sentences of *-te-iru* so that the imperfective and the perfect aspects encoded by this form are unified. I will extend the analysis of *-te-iru* presented in Yamagata (1998) and argue that the form is a function on eventualities expressed by propositions to convert them into 'states' as defined as 'realizations of an individual' by Carlson (1977a). Section 4.3 presents a new account of the source of the ambiguity of the sentences with *-te-iru*, which ascribes this ambiguity to the *underspecification* of the temporal property of *-te* of *-te-iru* rather than to the lexical ambiguity of *-te*. Section 4.4 proposes aspectual distinctions in Japanese based on the characterization of *-te-iru* presented in this study, and Section 4.5 provides a summary of the chapter.

4.1 Application of Klein's model to aspectual forms

This section will show how aspectual classes encoded by major temporal morphemes of English and Japanese are represented in Klein's system. The model captures how the

progressives and the present perfect in English and *-te-iru* in Japanese encode both tense and aspect, and makes the common temporal properties of these morphemes transparent.

4.1.1 The progressive in English and eventive interpretations

Earlier in chapter 1 (section 1.4), we discussed how Klein's model allows us to account for the restriction on the distribution of the progressive form in English. In order to see a parallel between the English progressive and *-te-iru* in Japanese, this subsection reviews our earlier discussion on the restriction on the progressive form in English.

As I noted earlier in section 1.4, stative predicates in English may be divided into two classes: 1) property-denoting, individual-level predicates, and 2) predicates that denote stage-level properties of an individual. For example, the predicate *be a Canadian* in a sentence, *Alan is a Canadian*, is an example of the first class, and *be in her office* in a sentence, *Cristina is in her office*, is an example of the second class. Neither of these may co-occur with the progressive form, but the co-occurrences are prohibited on different grounds. Property-denoting stative predicates do not have the progressive form because eventualities described by these predicates may never be conceived to involve change over time. The second class of statives does not have the progressive form because eventualities described by these predicates are already stage-level, the property that is to be provided by the progressive form. The first class of stative predicates above corresponds to Klein's '0-state'. According to Klein, '0-state' predicates may not involve a TT-contrast. If an eventuality expressed by this type of predicate is linked to a particular TT, it is automatically linked to any other TT due to its lexical content. TSit for a 0-state predicate extends over the entire time of the existence of a described

individual, and TT is equated with TSit. See the temporal structure for a 0-state sentence in simple present.

(1) Alan is a Canadian.

$TT \supset TU$
 ++++++[++++]++++++
 TSit: *Alan be Canadian*

Recall that the temporal relation that is encoded by the progressive form of English as we defined in section 1.4.2 is $TT \subset TSit$. However, even if we place the TT of (1) within its TSit by the progressive form, it would not make any difference to the way we view an eventuality, since the TT is automatically hooked up to any other points on the time line. Thus, the model predicts that the above sentence will never co-occur with the progressive form, which is confirmed by the following example.

(2) *Alan is being a Canadian.

In contrast, a temporal structure of a sentence with a second class of stative predicate will look like (3) below.

(3) Cristina is in her office.

$TT \supset TU$ (outer brackets indicate TT)
 -----+++[[++++]]++++-----
 TSit: *Cristina be in her office*

In (3) the temporal relation that would be provided by the progressive (i.e., $TT \subset TSit$) already obtains in its absence. Thus, we would expect that this sentence may not have the progressive counterpart. The prediction is borne out as we see in (4).

- (4) *Cristina is being in her office.

A problem arises when we encounter a 0-state sentence, which turns out to have a progressive counterpart. See the following examples that illustrate this point.

- (5) a. David is a good boy.
b. David is being a good boy.

Since TSit, *David be a good boy*, in (5a) is a 0-state, we would expect that this sentence never co-occurs with the progressive. However, as evidenced by the grammaticality of (5b), it does occur in the progressive form. This is a puzzle if we assume that (5a) refers to a 0-state just like (1).

I would argue that an eventuality *David be a good boy* in (5b) is conceptualized as an event and that its realization is conceived as temporally delimited. Compare the temporal structures for (5a) and (5b), shown in (6) and (7) respectively.

- (6) David is a good boy.

TT $\supset$ TU
+++++++[[++++]]+++++++
TSit: *David be a good boy*

- (7) David is being a good boy.

TT $\supset$ TU (outer brackets indicate TT)
-----+++[[++++]]++++-----
TSit: *David be a good boy*

In contrast to TSit of (6), which extends unlimitedly on both ends, TSit of (7) is temporally bounded as an event. The progressive form places TT within this TSit, and we describe a stage of a process of David's intentional effort of behaving like a good boy.

The above discussion leads us to an important point: the availability of the progressive depends on how we conceive the eventuality described by the lexical content of a sentence. If an eventuality may receive an eventive interpretation that involves a change over time, the progressive is allowed. Thus, we can say, *David is lying on the couch*, but we say *The town lies on the mouth of the river*, and not *The town is lying on the mouth of the river*. Here, we do not want to make a hasty conclusion that agentivity of the subject is the key to the progressive, as such is not the case. We can say *The book is lying on the desk* without an agentive subject. The difference between this sentence and the ill-formed sentence, *The town is lying on the mouth of the river*, should be ascribed to our pragmatic knowledge of the world that change in the location of a book over time is likely while change in the location of a town is less likely. Then, a determinant factor for availability of an progressive interpretation is whether or not a given eventuality can be conceived as involving change over time (regardless of the agentivity of the subject). In this respect, our knowledge of the world interacts with the meaning of the progressive form to place a constraint on the co-occurrence of the form. I suspect that the reason agentivity of the subject often seems to determine progressive interpretation is that majority of our eventuality descriptions refer to human actions. Actions involve change over time. However, it is a property of 'change over time' that is requisite for the progressive interpretation, not agentivity. This is why we may have stative predicates in the progressive form with non-agentive subject as in the case of *The book is lying on the desk*.

To recapitulate the distinction between (5a) and (5b), the difference in interpretations stem not simply from the fact that the latter involves the progressive

form, but also from the fact that we take an eventuality expressed by *David be a good boy* in (5b) to be an event. The same generalization holds for the following pair of sentences with non-stative verbs.

- (8) a. John plays tennis.
b. John is playing tennis.

Non-stative verbs in English in their present-tense form may only allow an interpretation that denotes an individual-level property of a subject. Thus, when a verb *play* occurs in the present-tense form as in (8a), the proposition as a whole denotes 0-state, which involves no TT contrast. When it is changed into the progressive as in (8b), however, it receives an eventive interpretation. Since the contribution of the progressive form is just to provide a stage-level property to an eventuality, an eventive interpretation must be ascribed to the original eventuality *John play tennis*. In other words, we may conceptualize the lexical content of *John play tennis* in two different ways: one is a description of John's property, and another is a description of an event of John's playing tennis. It is only the latter interpretation of *John play tennis* that allows us to have a progressive counterpart, since only the latter involves change over time. Compare the temporal structures for these two interpretations.

- (9) John plays tennis.

TT $\supset$ TU
 ++++++[++++]++++++
 TSit: *John play tennis*

(10) John is playing tennis.

$TT \supset TU$ (outer brackets indicate TT)
 -----+++[+++++]++++-----
 TSit: *John play tennis*

We see the parallel between the contrast of (9) and (10) on the one hand and the contrast of (6) and (7) on the other: (6) and (9) show that TSit does not have a beginning nor an end, while (7) and (10) show that TSit is a discernible unit with left and right boundaries.

For both *David be a good boy* in (5) and *John play tennis* in (8), an eventive reading is available only with the progressive form. Even though the lexical contents of these sentences allow eventive readings, they may not be expressed with present tense. The only readings that are available for (5a) and (8a) are the ones in which the predicates refer to properties of an individual. Such an absence of eventive readings for eventualities expressed with present tense is a puzzle, since the definition of present tense as $TT \supset TU$ by itself does not provide such a constraint on its distribution. The next section briefly discusses this peculiar property of present tense in English.

4.1.2 Peculiarity of present tense in English

As we saw in the previous section, there is something peculiar about the present tense form in English. While the past tense form may co-occur with the description of either states or events, the occurrences of the present tense forms seem to be restricted to the description of states². We can talk about a generic property and a pattern of actions or events with the present tense form, as exemplified by sentences like *The whale is a mammal* and *John smokes a pipe*. We can also talk about states described by stative

predicates in the present-tense form as exemplified by a sentence like *Mary is in the garden*. However, we cannot express events with the plain present. A sentence, *John plays tennis*, can only refer to John's property, and never represents an event of John's current engagement in a game of tennis. On the other hand, *John played tennis* can either refer to the fact that John used to play tennis (i.e., a property that was ascribed to John in the past) or to an event of John's participation in a game of tennis in the past. Where does this asymmetry come from?

When we examine sentences described with the simple present, we realize that they all refer to a slice of a homogeneous eventuality that is captured at the speech moment. This is so regardless of whether we describe a property or a stage of an individual. The same state or the property that is captured at the speech moment also holds at some other point within the time interval conceived as TU. In contrast, when we have an event, since it involves dynamic change over time internally, the state that is captured at the speech moment is not identical to the event as a whole. For an eventuality to be an event, it must have both the beginning and the end in its scope. While events are discernible units from outside, they are indivisible inside. Thus, a fragment of an event that is captured at the speech moment with the present tense form cannot constitute an event itself. At least intuitively, this seems to be the reason why an event cannot be expressed with present tense.

On the other hand, the use of past tense confines the TT and the TSit of a sentence to the time frame before the speech moment regardless of an eventuality type of the sentence. All that is claimed by the past tense form as a tense marker is $TT < TU$,

² The term 'states' is used here in a broader sense that includes both description of a

regardless of the internal structure of an eventuality. Any eventuality becomes a discernible and indivisible unit placed before the speech moment and may become 'a countable occurrence' in Mourelatos's (1978) sense.

There is also a diachronic explanation for the peculiar distributional property of the present tense. Grammaticization theory advanced by Bybee et. al (1994) contends that it is common for the progressive to evolve into presents or imperfectives in the grammaticization path of the temporal grams, and that there are many languages in the world in which the progressive functions as a present marker with dynamic verbs. They go on to explain that in the case of English, as the use of the progressive extends to cover the description of the present of the dynamic verbs, the generic, gnomic, and habitual readings were left for the default readings for the present tense form. However, if the progressive was originally developed to express the agentive act as Bybee et. al explains, it is of no surprise that the occurrences of the present tense form with stative predicates were left intact, unless they represent agentive situation as a marked case.

4.1.3 Representation of the present perfect in English

In this section I apply Klein's three temporal primitives to the representation of the sentences with the present perfect form in English.

It has been noted elsewhere in the literature on aspect (e.g. Binnick, 1991; Comrie, 1976; Scheffer, 1975) that the present perfect represents 'current relevance' of a described situation. Klein accounts for this effect, stating that "the present perfect makes an assertion about a TT in the present" because "the tense component marks TU as being

property and description of a stage of an individual.

included in TT” (1994: 110). He also explains that the aspectual component relates the time of situations to the TT, from which the idea that the situation is relevant to the present comes about.

According to Klein, “the perfect form itself does not say anything about the distance between TT and TSit’, and hence ‘does not specify HOW FAR TSit is before TT” (1994: 104, emphasis his), which gives rise to various types of perfect such as the ‘resultant perfect’ or the ‘perfect of experience’. While Klein provides representations of different ways in which TT can be associated with TSit to show this effect, he does not provide the temporal representation of the perfect sentences with all three of his temporal primitives. In order to see the effect of the tense component of the perfect forms on the temporal structure of the sentence, I represent below the present and the past perfect sentences in English with Klein’s TT, TSit and TU in (11) and (12) respectively.

(11) Mary has left.

TT $\supset$ TU (outer brackets indicate TT)
 -----+++++-----[[-----]]-----
 TSit: *Mary leave*

(12) Mary had left.

TT < TU (brackets on the left indicate TT)
 -----+++++-----[-----]--[-----]----
 TSit: *Mary leave*

The aspectual component of the perfect forms place TT at the post time of the situation expressed by the lexical content of the sentence. Thus, in either of the above cases, the claim is made by TT about the post state of the situation of Mary’s leaving. Since the ‘post state’ cannot exist without reference to the situation, TT in the perfect is inherently

tied to the situation itself. In (11), this post state includes the utterance time, owing to its tense component (which encodes present tense). This gives rise to the effect of the current relevance of the situation by relating TSit and TU via TT. In (12), on the other hand, the tense component of the sentence confines the claim about the post state to some time frame before TU. Since past tense makes TT and TU necessarily disjoint, it deprives the sentence of the effect of the current relevance.

According to Smith (1997: 186), English perfect sentences “have a stative value, and they ascribe to the subject a property based on participation in the prior situation.” What Smith calls ‘stative value’ may be ascribed to the fact that TT is placed at the post state of the situation. Then, present perfect sentences may be characterized as expressing that a described individual is in the post state of having participated in the situation described by the predicate.

The assumption that the English present perfect encodes $TT \supset TU$ in its tense component also accounts for its co-occurrence restrictions with temporal adverbial expressions. Vlach (1993) points out that an adverbial expression like *since Thursday* is an *extended now*, or *XN* adverbial in that it specifies a time that extends up to the time of utterance. XN adverbials exhibit the following distributional property in contrast to past-time denoting adverbials (Vlach, 1993: 264).

- (13) a. I saw John Thursday.
b. *I saw John since Thursday.
- (14) a. *I have seen John Thursday.
b. I have seen John since Thursday.

As we see in (13) and (14) above, XN adverbials may co-occur with the present perfect, while it may not co-occur with the simple past. On the other hand, past-time denoting adverbials may not co-occur with the present perfect as shown in (13a). If we assume that XN adverbials introduce TT in the present and past-denoting adverbials introduce TT in the past, the above restriction may follow as the prohibition on the contradictory temporal information in a single sentence. Since the present perfect makes an assertion at the present time by its tense component (i.e., $TT \supset TU$), it may not co-occur with adverbials that confines its claim in the past time.

Both the present perfect and the present progressive in English encode $TT \supset TU$ in their tense components. Since their aspectual components relate the described situation (i.e., TSit) to TT, it is natural that the both the present progressive and the present perfect indicate the ‘current relevance’ of the situation at the utterance time. Since TT of the situations described by either form refers to a state, the fact that the present progressive and the present perfect have existential readings also follows naturally from their definitions provided in the current framework.

The next section examines how the current system can be applied to analyze the Japanese *-te-iru*.

4.1.4 Representation of sentences with *-te-iru* in Klein’s model

The *Verb-te-iru* form of Japanese consists of the continuative form of verbs and an auxiliary verb *iru*, which on its own means ‘be, exist’.

Sentences with *-te-iru* have been noted to express various different aspectual meanings such as progressive, resultative, habitual, iterative, and experiential (cf., Soga, 1983) as the following examples illustrate.

- (15) David-wa ima hon-o yon-de-i-ru.
David TOP now book ACC read-ASP-NON PAST (PROGRESSIVE)
'David is reading a book now.'
- (16) Mado-ga ai-te-iru.
window NOM open-ASP-NON PAST (RESULTATIVE)
'A window is open.'
- (17) Dave-wa mainiti kurasu-ni it-te-iru.
Dave TOP every day class to go-ASP-NON PAST (HABITUAL)
'Dave has been going to class every day.'
- (18) Sakki -kara nandomo to-o tatai-te-iru.
a while ago from many times door ACC knock ASP-NON PAST (ITERATIVE)
'S/he has been knocking at the door many times for a while'.
- (19) Yukari-wa izen Boulder-ni ki-te-iru.
Yukari TOP before Boulder to come-ASP-NON PAST (EXPERIENTIAL)
'Yukari has been in Boulder before.'

While interpretations of the sentences with *-te-iru* may vary, if we consider that progressive, habitual, iterative are the subtypes of imperfective aspect in that they place TT within the time of the situation or the patterns of events, and that resultative and experiential are the subtypes of perfect aspect in that they place TT in the post state of the situation, we may collapse five different interpretations as illustrated above into two distinct categories of aspect: IMPERFECTIVE and PERFECT.

Kunihiro (1982) points out that the sentences with *-te-iru* may correspond either to the progressive or the present perfect in English, and that the choice depends on the lexical aspect of the predicate. According to Kunihiro, if the verb is punctual (i.e., lacking

internal duration), a sentence with *-te-iru* indicates a perfect meaning, and if the verb has duration, a sentence with this form will indicate a progressive meaning. However, he also points out that a single verb can be associated with more than one lexical aspect, and therefore may represent either a progressive or a perfect meaning with *-te-iru* as in the case of *yomu* ‘to read’, for example.

- (20) a. Kare-wa ima shoosetu-o yon-de-i-ru
 he TOP now novel ACC read-ASP-NON PAST
 ‘He is reading a novel now.’
 b. Kare-wa Sooseki-no sakuhin-o subete yon-de-i-ru
 he TOP Sooseki GEN work ACC all read-ASP-NON PAST
 ‘He has read all the stories by Sooseki.’

As indicated by English translations, (20a) represents an action in progress, while (20b) represents an experiential state or a post state of ‘reading all the stories by Sooseki.’ The temporal structures of (20a) and (20b) are shown in (21a) and (21b) respectively.

- (21) a. Kare-wa ima shoosetu-o yon-de-i-ru.

$$\begin{array}{c} \text{TT} \supset \text{TU} \quad (\text{outer brackets indicate TT}) \\ \text{-----}++++[[++++]]++++\text{-----} \\ \text{TSit: } \textit{shoosetu-o yom} \end{array}$$

 b. Kare-wa Sooseki-no sakuhin-o subete yon-de-i-ru.

$$\begin{array}{c} \text{TT} \supset \text{TU} \quad (\text{outer brackets indicate TT}) \\ \text{-----}++++\text{-----}[[\text{-----}]]\text{-----} \\ \text{TSit: } \textit{Sooseki-no sakuhin-o subete yom} \end{array}$$

(21a) represents the temporal structure of IMPERFECTIVE aspect as we saw in case of English progressive (cf. section 4.1.1, (10)), and (21b) represents the temporal structure of PERFECT aspect as we saw in case of English present perfect (cf. section 4.1.3, (12)). If this difference stems from the lexical ambiguity of the verb *yomu* ‘to read’ as Kunihiro

argues, it amounts to saying that the distinctions encoded by grammatical categories of aspect in English are in the lexicon in Japanese. This certainly does not capture native speaker's intuition on the use of this form nor the verb *yomu* 'to read'. I will argue that the source of the ambiguity of the interpretations of the *-te-iru* sentences between imperfective and perfect lies in the form itself. We will investigate the source of this ambiguity in section 4.3.

For now, let us revert to our discussion of the temporal structures of the sentences with *-te-iru* to further clarify the common properties of *-te-iru* with the aspectual markers in English. The following examples show that *-te-iru* exhibits the same distributional properties as the progressive in English.

- (22) a. Kare-wa gakusee da.
 he TOP student COP
 'He is a student.'
- b. *Kare-wa gakusee de -i-ru.
 he TOP student COP-ASP-NON PAST
 '*He is being a student.'
- c. Kare-wa mada gakusee de -i-ru.
 he TOP still student COP-ASP-NON PAST
 '?He is still being a student.'

The property-denoting stative predicate *gakusee da* 'to be a student' in (22) may only indicate an individual-level property of a described individual in the plain present. Since the predicate is 0-state, we would expect that the progressive interpretation of this predicate with *-te-iru* would be ungrammatical, which is shown in (22b). However, if we interpret this eventuality to represent the subject's intentional efforts to remain a student with an appropriate context, the *-te-iru* version of this sentence would become

grammatical with a progressive interpretation as shown in (22c). Compare the temporal structures of (22a) and (22c) shown in (23a) and (23b) respectively.

- (23) a. Kare-wa gakusee da.

TT $\supset$ TU
 ++++++[++++]+++++
 TSit: *Kare-wa gakusee da*

- b. Kare-wa mada gakusee de-i-ru.

TT $\supset$ TU (outer brackets indicate TT)
 -----+[[++++]]+-----
 TSit: *Kare-wa mada gakusee de-iru*

In contrast to (23a), which does not involve any TT contrast, (24b), being an event, may have a TT contrast, and therefore allows a modification of its TT by *-te-iru*.

The relation between non-stative predicates and *-te-iru* also turns out to be the same as in the relation between non-stative predicates and the English progressive.

- (24) a. Kare-wa nihon-no shoosetu-o yom-u

he TOP Japan GEN novel ACC read-NON PAST
 'He reads Japanese novels.'

- b. Kare-wa ima nihon-no shoosetu-o yon-de-i-ru

he TOP now Japan GEN novel ACC read-ASP-NON PAST
 'He is reading a Japanese novel now.'

The eventuality in (24a) with the present tense interpretation of the non-past tense maker *-ru* may only be interpreted to denote an individual-level property of an individual, while the *-te-iru* version of this sentence in (24b) receives an eventive reading. The temporal structures of (24a,b) are shown in (25a,b).

(25) a. Kare-wa nihon-no shoosetu-o yom-u

TT $\supset$ TU
 ++++++[++++]++++
 TSit: *Kare-wa nihon-no shoosetu-o yom*

b. Kare-wa ima nihon-no shoosetu-o yon-de-i-ru

TT $\supset$ TU (outer brackets indicate TT)
 -----+[[++++]]+-----
 TSit: *Kare-wa nihon-no shoosetu-o yom*

Just as in the case of (23a), the eventuality of (25a), is 0-state in content, and therefore, does not involve any TT contrast. Linking its TT to the present time automatically indicates that the described property holds at any other time. Such an eventuality may not be expressed with *-te-iru*, since confining TT within TSit by this form does not make any difference to the overall temporal structure of this sentence. However, if we take the lexical content of *Kare-wa nihon-no shoosetu-o yom* ‘he read Japanese novels’ to represent an event, the TSit becomes a temporally delimited discernible unit. Such a temporally bounded occurrence may allow a TT contrast, and so *-te-iru* creates a new TT that is confined within the TSit of the original sentence, and renders an interpretation in which the described individual is in a particular stage of an event in which he participates.

4.1.5 Summary

In this section, we examined how Klein’s theory of temporality can be applied to the analysis of major aspectual categories of English and Japanese. The system clarifies the respective contributions of the tense and the aspectual components of the temporal morphemes under investigation, namely, the present progressive and the present perfect of English, and *-te-iru* of Japanese. All three of them encode TT $\supset$ TU in their tense

components, thereby constraining the assertion to the time span identified as present.

The aspectual component of the progressive, places its TT within the TSit of the described eventuality, while the aspectual component of the present perfect places its TT in the post state of the TSit. The former represent IMPERFECTIVE and the latter represents PERFECT aspect as defined by Klein. In the case of the latter, the post state of the situation may exist only by virtue of the situation itself. Thus, placing TT (which is co-temporal with TU owing to the tense component of the form) in the post time of the situation will link the situation with TU, giving rise to an effect of the current relevance of the described situation. If tense needs to agree with the co-occurring adverbials in the location of the TT they introduce, the fact that the present perfect in English may not co-occur with past denoting adverbials follows from the fact that the form encodes present tense in its tense component.

-Te-iru in Japanese turns out to encode both $TT \subset TSit$ like the English progressive and the $TSit < TT$ like the English perfect. As the imperfective aspect marker, the form subject to the same co-occurring restriction as the English progressive. Klein's system also proved to be a strong tool for us to capture the parallel between the relation between the English progressive and the availability of eventive readings on the one hand and the relation between *-te-iru* and the availability of eventive readings on the other. The system shows that 0-state eventualities, lacking TT contrast, may not co-occur with either progressive or *-te-iru*.

The fact that these temporal morphemes encode both tense and aspect in their respective composite parts is also captured nicely in Klein's system. However, while the system helps us see the common properties between English progressive and perfect on

the one hand and Japanese *-te-iru* on the other, it remains to be a puzzle why *-te-iru* may encode such distinct aspects as imperfective and perfect in a single form. The following sections further investigate the problems concerning the interpretations of *-te-iru*, and attempt to provide an explanation for this interesting fact.

4.2 The ambiguity of *-te-iru*

This section investigates the problem surrounding the interpretations of sentences with *-te-iru*. As mentioned in section 4.1.4, sentences with this form may receive various different interpretations. While research on this form is abundant and there have been several proposals on the unification of this form (cf. Teramura, 1984; Kudo, 1989; Jacobsen, 1992; Shinzato, 1993 among others), I will discuss two of the recent proposals on the semantic nature of this aspectual form.

4.2.1 The stage-level and individual-level distinction (Ogihara, 1999)

Based on Fujii's (1966, 1976) observation that both durative or instanteneous verbs may receive an experiential interpretation³ with *-te-iru*, Ogihara (1999) proposes that *-te-iru* exhibits an ambiguity between an experiential use on the one hand and progressive and resultative uses on the other. Ogihara points out that the experiential use of *-te-iru* is unique in that it is always found with an adverbial expression indicating 'past' with either an instanteneous verb or a durative verb. This is illustrated by the following examples from Ogihara (1999: 336, originally (16)).

³ The experiential use of *-te-iru* was introduced by Fujii (1966, 1976) in contrast to the resultative use of the form, and was characterized by the fact that they must always accompany the past-denoting adverbs such as *izen* "in the past" or *sono toki* "then".

- (26) a. Taroo-wa 1970-nen ni kekkonsi-te iru
 Taro- TOP 1970-year in marry-Te iru-PRES
 ‘Taro has the experience of having gotten married in 1970.’
- b. Taroo-wa kyonen itido hugu-o tabe-te iru
 Taro-TOP last year once globefish-ACC eat-Te iru- PRES
 ‘Taro has the experience of having eaten globefish once last year.’

According to Kindaichi (1976), a seminal work on verbal aspect in Japanese, durative verbs typically receive a progressive interpretation with *-te-iru*, while instantaneous verbs receive a resultative interpretation in the *-te-iru* form. However, in (26), both the instantaneous verb *kekkonsuru* ‘to get married’ and *taberu* ‘to eat’ receive an experiential interpretation with past-denoting adverbs.

Based on the above observation, Ogihara argues that the progressive and resultant state use of *-te-iru* should be grouped together as referring to stage-level properties of an individual while the experiential use refers to individual-level properties of an individual.

Although I agree with Ogihara in that the progressive and resultant state use of *-te-iru* both indicate stage-level properties of an individual, which was independently proposed in Yamagata (1998), I will present counter arguments to his analysis of *-te-iru* based on both the descriptive facts and the problems in theoretical details of his proposal.

Ogihara presents the following table for the classification of the interpretations associated with *-te iru* form (1999: 336, originally (17)).

(27)

<i>Verb class</i>	<i>“Current situation”</i>	<i>“Experiential”</i>
Durative verbs	Progressive	Experiential
Instantaneous verbs	(Concrete) result state	Experiential

Descriptively, this table misses an important fact about the occurrence of this aspectual form. As pointed out in section 4.1.4, there is no one-to-one correspondence between the verb type and the interpretation of *-te-iru*. For example, there are many durative verbs that may indicate either progressive or resultative meanings with *-te-iru*, which has been noticed by previous studies on this form (cf. Machida, 1989; Yamagata, 1994, 1997). See the following examples to illustrate this point.

- (28) a. Yukari-wa ima hirugohan-o tabe-te iru
 Yukari TOP now lunch ACC eat -ASP-NON PAST
 'Yukari is eating lunch now.'
- b. Yukari-wa moo hirugohan-o tabe-te iru kara issyoni ko-nai
 Yukari TOP now lunch ACC eat -ASP-NON PAST because together come-NEG
 'Yukari has already eaten lunch, and so she won't come with us.'
- (29) a. David-wa ima hon-o yon-de iru
 David TOP now book ACC read-ASP-NON PAST
 'David is reading a book now.'
- b. David-wa moo sono hon-o yon-de iru to omo-u
 David TOP already that book ACC read-ASP-NON PAST COMP think-NON PAST
 'I think that David has already read the book.'

Both *taberu* 'to eat' in (28) and *yomu* 'to read' in (29) are durative verbs. However, both (28b) and (29b) receive resultant state interpretations. Thus, whether a given verb is categorized as instantaneous or not is not a determinant factor for the resultative reading. Since Ogihara's account of *-te-iru* ascribes the progressive reading and the resultative reading of the sentences with *-te-iru* to the lexical difference between durative and instantaneous verbs, it cannot account for the resultative readings of durative verbs with *-te-iru*, unless he claims that *taberu* in (28b) and *yomu* in (29b) are instantaneous verbs.

Another point of my objections concerns the mechanism he proposes to obtain an experiential reading of *-te-iru*. Ogihara (in press) argues that the experiential reading

always co-occur with a past-denoting adverbial because *-te* of *-te iru* must receive the feature [+perfect] from the past-denoting adverbial, which locates the situation in the past for us to obtain an experiential reading. The proposed system is schematized in (30).

- (30) **Sono toki** moo gohan-o tabe-te iru
 that time already meal-ACC eat [+perfect] EXPERIENTIAL
 ‘He/she has already eaten the meal then.’

According to Ogihara, since (30) has a past-denoting adverbial *sono toki* ‘then’, *-te* of *-te-iru* receives a feature [+perfect], and the experiential reading is obtained. However, the readings in which the situation represented by *-te* is perceived as a closed event in the past is not limited to the experiential readings. The resultant state readings also place the situation represented by *Verb-te* in the past interval. Then, the difference between the experiential and the resultative interpretations must be explained in some way.

In the current framework, the resultative and the experiential interpretations are both considered to be PERFECT, and they are differentiated by the distance between TT and TSit. Thus, in the present analysis of *-te-iru*, the fact that both resultative and experiential readings provide a view of a situation as a closed event naturally follows from the TT-TSit relation encoded by PERFECT aspect.

Ogihara’s account that the experiential reading is provided by [+perfect] feature on the *-te* of *-te-iru* that is assigned by the past-denoting adverbial has another problem. It cannot explain the ambiguity of the following sentence.

- (31) Taroo-wa **sono toki** moo hugu-o tabe-**te-i-ta**
 Taroo TOP at that time already globefish ACC eat-ASP-PAST
 a. ‘Taroo was already eating globefish then.’ PROGRESSIVE
 b. ‘Taroo had already eaten globefish then.’ RESULTATIVE
 c. ‘Taroo already had an experience of eating globefish then.’ EXPERIENTIAL

(31) is ambiguous in three different readings. Given an appropriate context, any of the interpretations in (31) are possible readings for the sentence. If the past-denoting adverbials provide [+perfect] value to the *-te* of *-te-iru*, it is a puzzle why they do not do the same for the *-te* of *-te-ita* in (31), and uniquely assigns an experiential reading to the sentence. Furthermore, if an experiential reading can be freely assigned to any sentences with *-te-iru* regardless of the verb type as Ogihara suggests, we do not have an account of the ungrammaticality of (32) in contrast to the grammaticality of (33) below.

- (32) a. *Kare-wa san-nen mae ni koko-ni i-te iru
 he TOP three years ago in here LOC exist-ASP-NON PAST
 ‘He has an experience of being here three years ago.’
- (33) a. Kare-wa san-nen mae ni koko-ni ki-te iru
 he TOP three years ago in here LOC come- ASP-NON PAST
 ‘He has an experience of coming here three years ago.’

If *-te-iru* provides an experiential reading to any predicate with the help of past-denoting adverbials as Ogihara suggests, we would expect that *iru* ‘to be, to exist’ in (32) may also co-occur with *-te-iru*, just like *kuru* ‘to come’ in (33).

Yet another objection comes from his treatment of the transitive-intransitive asymmetry of *-te-iru*. Ogihara provides the following three principles shown in (34) (1999: 339, originally, (24)) to account for the transitive-intransitive asymmetry in the interpretations of the sentences with *-te-iru*, which is illustrated in (35) (originally (23)).

- (34) a. In general, a sentence in the *-te iru* form is used to assign a property to the entity denoted by the subject NP and to nothing else. (This is implicit in Okuda's remarks.)
 b. An agentive entity can be assigned a property of "engaging in" the action named by the predicate (i.e. VP), whereas a nonagentive entity cannot.
 c. An entity can be assigned a property of being in some state if its obtaining this state as soon as the event described by the sentence is part of the lexical meaning of the predicate.
- (35) a. Taroo-wa ki-o taosi-te iru
 Taro-TOP tree-ACC fell-Te iru-PRES
 'Taro is felling a tree.'
 b. Ki-ga taore-te-iru
 tree-NOM fall-down-Te iru-PRES
 'A tree is on the ground (as a result of having fallen).'

It is often discussed in the literature on *-te-iru* that the form with a transitive construction typically receives a progressive reading, while it receives a resultant state interpretation with an intransitive construction (e.g. Jacobsen, 1992). However, the following examples from Machida (1989: 47, originally (17a)), shown in (36), and Yamagata (1998: 254, originally (25b)), shown in (37), reveals that this generalization does not always hold.

- (36) Dareka-ga sai-hu-o otosi-te iru
 someone NOM wallet ACC drop-ASP-NON PAST
 'Someone has dropped a wallet (and the wallet is on the ground).'
- (37) Kaze-de mado-ga sukosizutu ai-te-iru
 wind by window NOM little by little open-ASP-NON PAST
 'The window is opening little by little due to the wind.'

The verb *otosu* 'to drop' in (36) is a transitive verb. However, (36) receives a resultant state interpretation rather than a progressive interpretation. This is because even though *otosu* 'to drop' is a transitive verb, we do not normally describe a process with this verb. Thus, the lack of the progressive interpretation should be ascribed to a temporal property

of ‘lack of process’ rather than whether or not the sentence has an agentive subject. (37)

with an intransitive verb *aku* ‘(for X) to open’, on the other hand, receives a progressive reading, suggesting that even a predicate with a non-agentive subject may receive a progressive interpretation. As long as there is a process or a change of state over time is conceived of an eventuality, *-te-iru* may render a progressive reading. Transitive constructions with an agentive subject tend to receive progressive interpretations simply because many transitive sentences refer to human actions that represent dynamic processes. Thus, the above principle provided by Ogihara does not have any contribution other than re-statement of often-discussed ‘tendency’ for the interpretations of the sentences with *-te-iru*; namely, the form tend to indicate action in progress with transitive verbs while it tends to express resultant state with intransitive verbs.

Another problem with Ogihara’s account lies in the observational adequacy of the data that he provides as evidence for his individual-level and stage-level distinction of *-te-iru* sentences. As introduced earlier, Ogihara argues that *-te-iru* sentences with progressive and resultative readings denote stage-level properties, while *-te-iru* sentences with experiential readings denote individual-level properties of an individual. In order to support this classification of *-te-iru* sentences, he points out that sentences with *-te-iru* that receive current-state interpretations (i.e., both progressive and resultative) contain a *ga*-marked NP that receives a ‘neutral description’, while *-te-iru* with experiential state interpretation must always contain a *ga*-marked NP that receives a focused interpretation. Based on Kuroda’s (1965a) observation that the *ga*-marked NP must receive a focused interpretation when a sentence contains a *ga*-marked NP and an individual-level predicate,

Ogihara uses this asymmetry to serve as evidence for his stage-level and individual-level distinction of *-te-iru*. He provides the following examples (1999: 338, originally (22)).

- (38) a. Taroo-ga ima ki-o taosi-te iru
 Taroo-NOM now tree-ACC fell-Te iru-PRES
 ‘Taro is now felling a tree.’
- b. Taroo-ga ima yooroppa-ni it-te iru
 Taroo-NOM now Europe-to go-Te iru-PRES
 ‘Taro is now in Europe (as a result of having gone there).’
- c. Taroo-ga imamade-ni hon-o zyussatu-mo kai-te iru
 Taroo-NOM till now-DAT book-ACC ten-as many as write-Te iru-PRES
 ‘Taro is the one who has the experience of having written as many as ten books.’
- d. Taroo-ga kyonen yooroppa-ni it-te iru
 Taroo- NOM last year Europe-to go-Te iru- PRES
 ‘Taro is the one who has the experience of having gone to Europe last year.’

Ogihara argues that the predicates in (38c) and (38d), which receive experiential interpretations, must denote individual-level property, in contrast to (38a) and (38b) with resultant state interpretations that are ‘neutral descriptive statements’. I would argue that this claim is not tenable. Although Ogihara only provides interpretations of (38a,b) which are neutral descriptive, it is possible for *ga*-marked NPs in (38a) and (38b) to receive a focused interpretation, given an appropriate context. (38a) will be a perfectly adequate reply to a question, *Dare-ga ki-o taosi-te-iru no* ‘Who is felling the tree?’, and in that case, it may receive an interpretation in which *Taroo-ga* is a focused NP. The same can be said with (38b).

One serious theoretical weakness of Ogihara’s analysis of *-te-iru* is that he has to assume that *-te* of *-te-iru* is lexically ambiguous. That is, there are *-te*₁ with a feature [-perfect] and *-te*₂ with a feature [+perfect] in the Japanese lexicon. Although such a

possibility is not entirely inconceivable, it is not desirable in consideration of the theoretical simplicity. It also goes against native speaker's intuition that there is a single *-te-iru* form, not two.

In sum, Ogihara's (1999) analysis of *-te-iru*, as interesting as it is in providing a new classification of the sentences with this form, cannot be supported either empirically or theoretically.

4.2.2 *-Te-iru* as a stage-level operator (Yamagata, 1998)

Unlike Ogihara (1999) who classifies the sentences with *-te-iru* into those denoting stage-level properties and those representing individual-level properties, Yamagata (1998) proposed that all the occurrences of *-te-iru* can be unified under the notion of stage-level property in Carlson's sense (1977a,b). In Yamagata (1998), I argued that evidence from the occurrences of *-te-iru* with some stative predicates and its occurrences with the type 4 verbs indicate that the form is a function on the eventualities described by the original sentences to convert them to the representations of stages of an individual. Below, I will go over the descriptive facts behind this proposal and see how it accounts for the data that cannot be explained by Ogihara's (1999) proposal.

4.2.2.1 The problem with the type 4 verbs

Dai-yon-shu-no doosi or the type 4 verbs were first identified by Kindaichi (1976), referring to a group of verbs that express that a described object '*bears* a certain state' in contrast to regular stative verbs which indicate that a described object '*is* in a certain state' (italic mine). A syntactic criterion provided to discern this class of verbs is that

they must always be used in the *-te-iru* form⁴. An example below with a type 4 verb, *sobieru* ‘to tower’, illustrates this point.

- (39) a. *Me-no mae-ni takai yama-ga sobie-ru
 eye-GEN front-LOC high mountain-NOM tower-NON PAST
 ‘A high mountain towers in front of us.’
- b. Me-no mae-ni takai yama-ga *sobie-te-iru*
 eye-GEN front-LOC high mountain-NOM tower-ASP-NON PAST
 ‘A high mountain is towering in front of us.’

As we see in (39a), a sentence with a type 4 verb, *sobieru*, turns out to be ungrammatical if the verb is used in the simple non-past form. The sentence with this type of verbs can be grammatical only if they co-occur with *-te-iru* as shown in (39b).

The fact that the type 4 verbs must always co-occur with *-te-iru* has long been ascribed to an idiosyncratic property of this class of verbs. However, what semantic or syntactic properties of this class of verbs are responsible for its peculiar distributional property was never fully accounted for. If we assume that the type 4 verbs are type of stative verbs, it is a puzzle why they occur exclusively in *-te-iru*, as stative verbs in general are assumed to be incompatible with *-te-iru* (cf. Kindaichi, 1976; Jacobsen, 1992).

However, as we discussed earlier in section 4.2.4 (example (22)), some stative predicates may actually co-occur with *-te-iru*. Additional examples are provided in (40) and (41).

⁴ Some type 4 verbs may occur in the simple non-past tense form in relative clauses as in *takaku sobieru yama* ‘a mountain that towers’.

- (40) a. Kono suponji-wa mizu-o takusan hukum-u.
 this sponge TOP water ACC much contain-NON PAST
 'This sponge absorbs a lot of water.'
- b. Kono suponji-wa mizu-o takusan *hukun-de-i-ru*.
 this sponge TOP water ACC much contain-ASP-NON PAST
 (Lit.) 'This sponge is containing a lot of water.'
- (41) a. Matt-wa kanji-ga kireeni kak-e-ru⁵.
 Matt TOP Chinese character NOM neatly write can-NON PAST
 'Matt can write Chinese characters neatly.'
- b. Matt-wa kono kanji-ga kireeni *kak-e-te-iru*.
 Matt TOP this Chinese character NOM neatly write can-ASP-NON PAST
 (Lit.) 'Matt is being able to write this Chinese character neatly.'

While (40a) and (41b) denote an individual-level property of the subject, (40b) and (41b) with *-te-iru* allow stage-level interpretations.

Based on this observation, in Yamagata (1998) I proposed that *-te-iru* is a function on an eventuality described by the original proposition to turn it into a description of a stage-level property of an individual in Carlson's (1977) sense. According to Carlson, individual-level predicates denote 'properties' of an individual and are directly in the property set of an individual, while stage-level predicates denote 'states' and are in the property set of one of the realizations of an individual (1977: 448-449). With this characterization of a 'stage-level' property, all the sentences with *-te-iru* presented earlier (i.e., examples (15) through (19), (21b), (39b), (40b), and (41b)) can be unified.

Under this proposal, the type 4 verbs in Japanese are characterized as follows: 1) the type 4 verbs lack temporal specification in the lexicon, 2) unlike non-stative verbs, they may not allow even a habitual or characteristic-disposition interpretations that holds

⁵ Following Kuno (1973), I assume that derivatives with potential morphemes *-(r)e/(r)are* are statives.

true at the speech moment in the simple non-past tense form, and hence, 3) they must co-occur with *-te-iru* to be anchored onto the time line.

The argument that the type 4 verbs without *-te-iru* denote generic properties without temporal specification while they indicate ‘states’ when they co-occur with *-te-iru* is supported by the following examples.

- (42) a. Ko-wa oya-ni ni-ru.
 child TOP parent(s) to resemble-NON PAST
 ‘A child/children will resemble his/their parent(s).’
- b. *Kono ko-wa oya-ni ni-ru.
 this child TOP parent(s) to resemble-NON PAST
 ‘This child resembles his/her parent(s).’
- c. Kono ko -wa oya-ni *ni-te-iru*.
 this child TOP parent(s) to resemble-ASP-NON PAST
 ‘This child resembles his/her parent(s).’

Niru ‘to resemble’ may exceptionally occur as a matrix predicate without *-te-iru*, but only as a proverb as in (42a)⁶. (42a) expresses a generalization that transcends a particular eventuality. The existential reading is unavailable in the simple non-past tense form of this predicate with the present time interpretation which is confirmed by (42b)⁷. The same verb in (42c) with *-te-iru*, on the other hand, is predicating of one of the realizations of an individual in specific time and space. Of (42a) and (42c), only (42c) denotes

⁶ (42a) can be rephrased as a gnomic statement, ‘Ko-wa oya-ni *niru mono da*’, which can be translated as ‘*It is usually the case that* a child resembles his parent(s).’

⁷ (42b) is fine with the future tense interpretation: ‘This child will come to resemble his parent(s).’

‘states’ as defined above, which is synonymous with what we call ‘stage-level’ property⁸.

The argument that propositions expressed with the type 4 verbs in *-te-iru* indicate stage-level properties is also supported by the fact that they never allow generic interpretation of the subject NPs. See the following example to illustrate this point.

- (43) *Kodomo-ga oya-ni *ni-te-iru*.
child NOM parent(s) to resemble-ASP-NON PAST
‘A child/children resemble(s) his/their parent(s).’

- (44) *Kenkyuu-ga *sugure-te-iru*.
research NOM excel-ASP-NON PAST
‘?A research is excellent.’

Following Carlson (1977), existential readings of the subject noun phrases are linked to stage-level readings of the predicates, and generic readings of the subject NPs are linked to individual-level readings of the predicates⁹. Both (43) and (44) with generic interpretation of the subject NPs are ungrammatical, which is amenable to the analysis that a predicate in *-te-iru* indicates a stage-level property of an individual.

⁸ One should be cautious about associating our knowledge of the world and what a proposition says. One's resembling his/her parents may be inherent property of that individual, and may be permanent within that individual's life span. However, this does not necessarily constitute evidence for the claim that the predicate *ni-te-iru* 'resemble' is directly in the property set of an individual.

⁹ See also Krifka et. al (1995) for the relation between the aspectual properties of predicates and generic interpretation of subject NPs.

4.2.2.2 *-Te-iru* with stative verbs

Assuming that *-te-iru* is a stage-level operator, we would expect that verbs that allow a stage-level interpretation on their own may not occur in this form. The following data will show that this prediction is borne out.

- (45) a. Yukari-wa niwa-ni i-ru.
 Yukari TOP garden LOC be NON PAST (S-LEVEL)
 ‘Yukari is in the garden.’
 b. *Yukari -wa niwa-ni i-te-iru.
 Yukari TOP garden LOC be
 ‘*Yukari is being in the garden.’

An existential verb like *iru* ‘be, exist’ in (45) only has a stage-level usage. Thus, it never co-occurs with *-te-iru* as shown in (45b). This is the same restriction that we observed in the case of the English progressive with stative verbs that indicate stage-level properties (cf. section 4.1.1, (3)). Consider the temporal structure of (45a) shown in (46).

- (46) *Yukari-wa niwa-ni i-ru*
- $$\begin{array}{c} \text{TT} \supset \text{TU} \quad (\text{outer brackets indicate TT}) \\ \text{-----}++++[[++++]]++++\text{-----} \\ \text{TSit: } \textit{Yukari-wa niwa-ni i} \end{array}$$

The temporal structure in (46) shows that the aspect that would be provided by *-te-iru*, namely, $\text{TT} \subset \text{TSit}$ is already present. Hence, the ungrammaticality of (45b) is explained.

In contrast, stative verbs that indicate individual-level properties of an individual may be converted to denote stage-level properties if the lexical content of the proposition allows an eventive interpretation, as we discussed in section 4.1.4.

- (47) a. Dave-wa ii ko da.
Dave TOP good boy COP (I-LEVEL)
'Dave is a good boy.'
- b. Dave-wa ii ko -de -i-ru.
Dave TOP good boy COP-ASP-NON PAST (S-LEVEL)
'Dave is being a good boy.'

While (47a) may only indicate an individual-level property, (47b) with *-te-iru* expresses an episodic property of a described individual. (47b) means that *Dave* is temporarily behaving like a good boy within the TT that is newly created by *-te-iru*. Thus, *-te-iru* converts ‘properties’ into ‘states’, which Carlson characterizes as ‘parts of a whole’.

4.2.2.3 -*Te-iru* with non-stative verbs

With the proposed function of *-te-iru*, we can also account for another distributional property of this temporal morpheme.

Non-stative verbs in Japanese in the non-past tense form may only indicate generic or habitual characterizing properties of an individual with a present tense interpretation, or otherwise they indicate the future event. They must always co-occur with *-te-iru* to receive present-time denoting eventive interpretations. This is because they may refer to the present time only as individual-level predicates. Recall the peculiar distributional property of the present tense form in English discussed earlier in section 4.1.2. The same generalization holds with the non-past tense form in Japanese, as discussed briefly in section 4.1.4. Non-stative verbs in the non-past tense form may not have present-time interpretations other than as individual-level predicates. Since part of an event that is captured at the utterance time cannot be identical to the entire event, they may never co-occur with the non-past tense form with eventive interpretations.

However, when *-te-iru* operates on the eventualities described with non-stative verbs, it would create a new TT frame that is placed within TSit. Then, an eventuality described with *-te-iru* expresses a temporally-bounded stage of realizations of an individual at the utterance time, which constitutes a part of the whole event denoted by a non-stative verb.

Having defined *-te-iru* as a stage-level operator, we wonder what property of *-te-iru* contributes to the stage-level or the existential interpretation of the sentences. The next section investigates this question under the current framework.

4.2.3 Existential readings and *-te-iru*

When we look into the composite parts of *-te-iru*, it seems natural to consider that the stage-level reading of *-te-iru* sentences stems from the semantics of the auxiliary *iru* of *-te-iru*. Recall that a verb *iru* ‘to be/to exist’ on its own always refer to a stage-level property of an individual and therefore, may never co-occur with *-te-iru* (section 4.2.2.2, (45)). Such restriction on the co-occurrence of a verb *iru* and *-te-iru* makes sense if we assume that the auxiliary *iru* of *-te-iru* inherits the stage-level property of a verb *iru*. This also explains why a verb *iru* ‘to be/to exist’ may not co-occur with *-te-iru* even with an experiential interpretation (section 4.2.1, (34)), which remained unexplained by Ogihara’s analysis of *-te-iru*.

However, if we assume that a stage-level property of an individual means that the described individual is in the midst of the situation represented by the lexical content of the original predicate before *-te-iru* is attached, neither resultant state interpretations nor experiential interpretations can be subsumed under the notion of ‘stage-level’. This is because both in resultant state and experiential interpretations, what is claimed by

sentences with *-te-iru* is that an individual is in the ‘post state’ of the situation described by the original verb. Thus, if the notion of stage-level is defined as above, it is not only problematic to the proposal by Yamagata (1998), but it is also problematic to Ogihara’s (1999) account of *-te-iru* which proposes that resultant state interpretations refer to stage-level properties of an individual on a par with progressive readings.

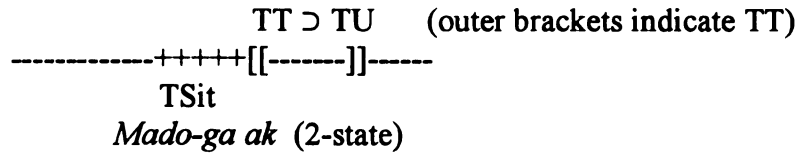
In order to solve this problem, I propose that the notion of ‘stage-level’ should be understood to represent one of the ‘realizations’ of an individual in terms of its participation in an eventuality. This is basically the same as a characterization of the English progressive proposed by Carlson (1977a). I propose that in the case of *-te-iru*, this participation in an eventuality must subsume an individual’s realization in a post state of the eventuality. This allows all the occurrences of *-te-iru* to be embraced under the notion of ‘stage-level’ properties of an individual, including the sentences with a resultant state interpretation and those with an experiential interpretation.

(48) Mado-ga ai-te iru
 window NOM open-ASP-NON PAST
 ‘The window is open.’

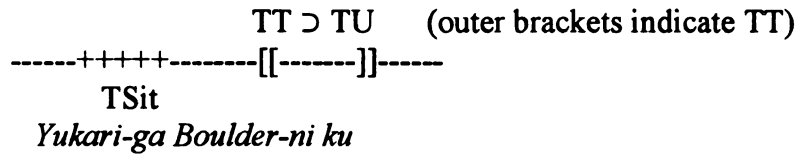
(49) Yukari-wa itido Boulder-ni ki-te iru
 Yukari TOP once Boulder to come- ASP-NON PAST
 ‘Yukari has come to Boulder once.’

Both (48) and (49) represent that the individual is in the post state of having participated in the situation described by the predicate. Consider the following temporal structures which illustrate the TT-TSit relations encoded by *-te-iru* in (48) and (49) respectively.

(50) Mado-ga ai-te iru



(51) Yukari-wa itido Boulder-ni ki-te iru



As is clear in the temporal structures provided in (50) and (51), both (48) and (49) refer to the PERFECT aspect as defined by Klein. That is, TT is placed in the post state of the situation represented by the lexical content of the original sentence. The difference between (48) and (49) is the distance between TT and TSit. In case of a resultant state interpretation, the TT is in the post state of TSit that is immediately after the situation as illustrated in (50), while such is not the case for an experiential interpretation, as shown by the temporal structure of (49) shown in (51). While TT in either case is disjoint from TSit, providing a view of the situation as a closed event, they both represents stage-level property of an individual in that they denote a stage of an individual as it is in the post state of the situation described by the original predicate.

While definition of ‘stage-level’ to include the post state of a situation allow us to unify all the occurrences of *-te-iru*, we still do not have an account of why this form allows us to represent such distinct aspects as IMPERFECTIVE and PERFECT. In the next section I will propose a new account of this dual nature of *-te-iru*, which does not ascribe this peculiarity of the form to the lexical ambiguity of *-te* but to its temporal property, by applying Klein’s theory in a novel way.

4.3 The semantics of *-te* in the aspectual composition

As we discussed in preceding sections, sentences with *-te-iru* may be ambiguous between imperfective and perfect aspect. In the following sections, I will show how Klein's system can be applied to explain this ambiguity in a simple manner with very few stipulations. I will also argue that a slight modification on Klein's system will allow us to capture the crosslinguistic variations of aspectual forms between English and Japanese economically.

4.3.1 The present and past participials in English and *-te* of *-te-iru*

So far we have assumed that the English progressive and perfect forms as a whole provide the respective TT-TSit relations. This is based on Klein's characterizations of these aspectual forms. In the same manner, the Japanese *-te-iru* was also assumed to encode $TT \subset TSit$ and $TSit < TT$ in its entirety. However, such characterizations of these aspectual forms do not allow us to identify the unique semantic contribution of the respective morphemes that constitute each aspectual form. Therefore, I will propose the following revision on Klein's representation of aspect so that the semantic contribution of each composite part of the morphologically complex aspectual forms would become clear:

1) Instead of ascribing the specific TT-TSit relation encoded by each aspect to the entire aspectual form, I propose that the information regarding the way TT is hooked up to TSit should be ascribed to the participials of the respective aspectual forms. 2) The stative properties of the progressive and the perfect should be attributed to the auxiliary verbs in the respective forms. The Table 5 shows that such modification on Klein's definition of

aspect makes the contribution of each component of the respective aspectual forms transparent, and help us explain an interesting crosslinguistic difference between English and Japanese.

English	PROG	<i>-ing</i>	TT $\subset$ TSit
		<i>be</i>	state
	PERF	<i>-en</i> ¹⁰	TSit $<$ TT
		<i>have</i>	state
Japanese	<i>-te-iru</i>	<i>-te</i>	[] _{TT}
		<i>i</i>	state

Table 5: Semantics of components of aspectual forms

Based on the proposed revision on Klein's system, it is the present participial that encodes TT $\subset$ TSit (i.e., the topic time is included in the time of the situation described by the predicate) rather than the progressive form in its entirety. The stative value of the progressive sentences is provided by the copula *be*. In case of the English perfect, the past participial encodes the TT-TSit relation in which TT is always placed after TSit. The stative value is provided by the auxiliary *have*.

Unlike English, which has separate morphemes to encode two distinctive TT-TSit relations, Japanese only has a single morpheme, namely, *-te*. However, the aspectual form *-te-iru* somehow must be able to encode both TT $\subset$ TSit and TSit $<$ TT. One possible solution is to propose that *-te* is lexically ambiguous. That is, there are *-te*₁ and *-te*₂ in the Japanese lexicon. This is tantamount to the claim that there are two *-te-iru* forms, each of which is associated with the distinctive *-te* for imperfective and perfect

aspect. Such a proposal is not only counterintuitive, but also undesirable in terms of theoretical simplicity. I argue that *-te* of *-te-iru* is *not ambiguous*, but simply *underspecified* in terms of the relation between TT that it provides and TSit of the original predicate. In other words, it creates a new TT for the complex predicate, but its placement with respect to TSit of the original predicate is simply unspecified. This is indicated by the bracket []_{TT} in the Table 5. Thus, while English has the present participial and the past participial, both of which define a specific way in which TT is related to TSit, Japanese has a single morpheme *-te* that simply provides a topic time without specifying how it relates to the situation time. This will allow *-te-iru* to occur in a description of as different aspects as perfect and imperfective¹¹.

4.3.2 Underspecification of *-te* and theoretical consequences

By assuming that *-te* of *-te-iru* is underspecified in terms of the placement of the TT that it creates with respect to TSit of the original predicate, we can maintain the uniform semantic characterization of *-te-iru*, namely, the form provides a stage-level interpretation of an individual. Thus, the analysis of *-te-iru* presented in Yamagata (1998) receives an additional support by Klein's system of temporal representation. The system allows us to unify all the occurrences of this form that we discussed in the preceding sections in a very simple way without any further stipulations.

¹⁰ I use *-en* to represent past participial here in order to distinguish this morpheme from the past-tense morpheme.

¹¹ There is one possibility of the placement of TT created by *-te* that must be excluded on independent grounds, namely, the placement of TT in the pre-state of TSit. This is because pre-state is synonymous with the non-occurrence of the situation, and non-occurrence of the situation, unlike the post state of the situation, cannot be identified as any situation that TT can be hooked up to.

By ascribing the reason for the ambiguity of *-te-iru* sentences between imperfective and perfect to the underspecification of the temporal nature of *-te*, we do not need to posit that *-te* is lexically ambiguous, nor do we need to assume that verbs such as *yomu* ‘to read’ and *taberu* ‘to eat’ that may have progressive, resultative, or experiential interpretations with *-te-iru* to have multiple entries in the lexicon. This is a desired result in terms of simplicity of the theory. It also captures native speaker’s intuition that there is a single *-te-iru* form, not *-te-iru*₁, *-te-iru*₂, and so on.

The current proposal on the nature of *-te* also allows us to explain the crosslinguistic difference between English and Japanese in a very simple way. The fact that English uses a copula *be* for the progressive and a possessive auxiliary *have* for the perfect and that Japanese uses an auxiliary *iru* for *-te-iru* is a historical accident in the development of each language. In the same way, it so happened that English has two separate morphemes: the present participial and the past participial for imperfective and perfect respectively. However, Japanese has developed only one morpheme, namely *-te*, for this aspectual distinction, and there had to be a way to express what would be expressed by two different morphemes in English. Positing that *-te* is underspecified for the distinctions encoded by the present and the past participials in English does not require any further stipulations nor the complication in the system of temporal reference that is defended in this study.

The current proposal on the nature of *-te* also has descriptive power and another intuitive appeal. Recall that for the sentences with *-te-iru* to receive a perfect interpretation, we do not necessarily need past denoting adverbials nor inherent telicity of described eventualities. Virtually any type of eventuality may co-occur with *-te-iru* to

render a perfect reading. All that we need for a perfect interpretation of a sentence with *-te-iru* is a contextual information that specifies the placement of TT created by *-te* in the post state of the situation. Likewise, all that we need for a progressive interpretation of a sentence with *-te-iru* is some contextual information which tells us that the eventuality is a process and that we are describing a particular stage of an individual in a development of an event over time. Such contextual information lets us place TT within TSit of the original eventuality, and provides a view of a situation from inside. In other words, what determines the placement of TT created by *-te* is the information provided by particular discourse in which the *-te-iru* form occurs. This explains native speaker's intuition about the context dependency of the interpretation of *-te-iru*. It also explains the fact that although some sentences with *-te-iru* may be ambiguous on the sentential level, they are not ambiguous on the discourse level because each occurrence of the form is associated with a unique placement of TT of *-te* provided by the discourse.

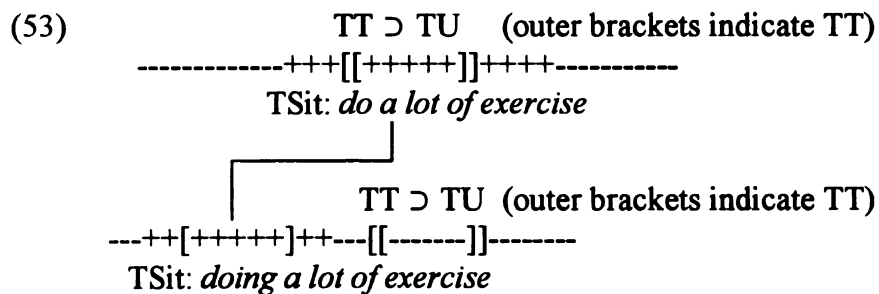
4.3.3 Remaining problems and further modification on Klein (1994)

As far as I am aware, there is one problem with Klein's model of temporal representation in its explanatory power in dealing with temporal expressions in English. Assuming that the English progressive encodes TTCTSit and that the English perfect encodes TSit < TT, we have no way to provide a temporal structure of a sentence in (52).

(52) Brian has been doing a lot of exercise lately.

This is a limitation of Klein's system of temporal reference as it is currently formulated. It does not consider the possibility that more than one aspectual form may apply to a single verb, as exemplified by (52).

In order to accommodate data like (52), we need further modification to Klein's system so that it can embrace the multiple application of aspectual forms. If we assume, for example, that a new TT that is created by the application of the progressive on the base verb serves as a new TSit for the application of the perfect in (52), it is possible to create another TT in the post state of this new TSit, which is exactly the state that (52) is referring to. This mechanism is schematized in (53).



The way aspectual forms apply recursively on top of another as described above will explain why the sentence in (54) is ungrammatical in contrast to (52).

(54) *Brian is having done a lot of exercise lately.

The ungrammaticality of (54) can be explained by the fact that the outcome of application of the perfect aspect is a state, which cannot be modified by the progressive form due to the co-occurring restrictions on the progressive with states that we discussed earlier in this chapter. Thus, the characterization of the progressive and the perfect provided in

this study also explains why applying the perfect aspect to the progressive as in (52) is perfectly legitimate, while the reverse is not permissible.

I believe that there is a way to modify Klein’s theory so that it can handle multi-layers of application of aspectual expressions as described above in a more sophisticated way than what is schematized in (53). Although it is interesting to see how Klein’s theory can be refined to incorporate the above data and possibly the data from other languages in which aspectual forms are even more complex than those in English and Japanese, I leave it as a subject of a future research.

4.4 Aspectual system of Japanese

This subsection proposes aspectual distinctions to be posited in Japanese based on our characterization of *-te-iru* presented in this chapter.

Dahl (1985) and Bybee and Dahl (1989) proposed that the typical system with a grammaticized perfective aspect is a tripartite system in which the perfective is restricted to the description of the past events and the imperfective is divided into present and past. The table in (55) illustrates this system (Bybee et. al, 1994, originally (49))

(55)

perfective	imperfective	
	present	past

Based on the above table, the fact that *-te-iru* has a past-tense counterpart *-te-i-ta* confirms our view that *-te-i* functions as an imperfective marker. In Chapter 1 (section 1.4.2), we hypothesized that the temporal morpheme that contrasts with the imperfective

use of *-te-i* is *-ta*, whose occurrence is limited to the description of the past. However, according to Bybee et. al (1994), if the gram in question co-occurs with imperfective, it should be considered a past. In this respect, the fact that *-ta* co-occurs with the imperfective *-te-i* as in *-te-i-ta* suggests that *-ta* also functions as a past-tense marker. Moreover, the fact that *-ta* co-occurs with stative predicates also indicates that *-ta* is more like a past tense marker than a perfective marker, since the notion of perfective is incompatible with states. And yet, it is none the less true that *-ta* provides a view of an eventuality as an indivisible whole, while an addition of *-te-i* as an imperfective marker provides a view of an eventuality from inside. Thus, as far as these facts are concerned, it seems to be fair to conclude that *-ta* in Japanese is a past-tense marker *and* a perfective aspect marker. As a past tense marker, *-ta* contrasts with a non-past tense marker *-ru*. However, as a perfective marker, what it contrasts with is not *-ru*, unlike what has been argued by many traditional grammarians, but the presence of *-te-i*, as an imperfective marker. The Table 5 provides an aspectual system in Japanese.

	perfective	imperfective	perfect
NON PAST	N/A	<i>-te-i-ru</i>	<i>-te-i-ru</i>
PAST	<i>-ta</i>	<i>-te-i-ta</i>	<i>-te-i-ta</i>

Table 6: Aspectual System in Japanese

The above table recapitulates the aspectual distinctions proposed in Chapter 1. Apart from the difference in terminologies, this basically supports Kudo's (1989) view on the discourse functions of *-te-iru* that is contrasted with the function of *-ta*. The contribution of the present study is that it unified the perfect and imperfective use of *-te-iru*, and

provides an explanation for why it may encode such different aspects in a single form without resorting to the idea of lexical ambiguity.

4.5 Summary

This chapter has shown that Klein's theory of temporal reference can be a strong tool for explaining behavior of aspectual forms as well. For languages like English and Japanese in which information on both tense and aspect resides in a single form, Klein's system turns out to be especially helpful in that it makes the contribution of the respective composite morphemes transparent. Furthermore, a slight modification on Klein's system allows us to explain an interesting crosslinguistic difference between English and Japanese. Unlike English, which has two separate morphemes for the perfect and imperfective, Japanese, having only one morpheme to represent the aspectual distinctions, have developed a way to endow this single morpheme a power to express these different aspects. By assuming that both the present participial and the past participial of English encode the specific way in which TT relates to TSit but *-te* of *-te-iru* is underspecified in how TT relates to TSit, the difference between English and Japanese can be explained in a very simple way and also in a way that it confirms to the assumptions of the theory without any extra stipulations. It also allows us to unify all the occurrences of *-te-iru* with a single semantic notion: stage-level of an individual, as proposed in Yamagata (1998).

CHAPTER 5

Conclusion

Tense and aspect are both grammatical devices to relate events and situations described by language to the time in our world. However, they are often treated as totally separate objects of research in contemporary formal linguistics. This study has shown that Klein's theory of temporal reference can unify these two distinct but related concepts.

I have shown that adopting Klein's theory of temporal reference in conjunction with Abusch's (1988, 1991, 1994, 1997a) theory of intensionality allows us to explain the complex behavior of tense morphemes within a language and the crosslinguistic differences between English and Japanese tense and aspect systems in a relatively simple manner. Both the English past and the Japanese *-ta* encode $TT < TU$ as a tense marker, while both the English present and the Japanese *-ru* encode $TT \supset TU$ for the present tense, though the Japanese *-ru* as a non-past tense marker additionally encodes $TU < TT$. I have shown that the complications with the interpretation of tenses in different types of embedded clauses can be ascribed to the presence or absence of an intensional context and the aspect of the predicates that appear with the tense morphemes. Then, the apparent complexity of the tense system within a language is attributed to the factors outside the tense systems. As a result, the semantics of tenses can be maintained consistent across different constructions and the tense system can be much simpler than those that have been proposed by some existing theories of tense.

I have also argued that the seeming differences between the English and the Japanese tense systems are attributed to the fact that Japanese employs the relative tense system on top of the absolute tense system. The position of the present study is different from the position of those who claim that Japanese tense morphemes are relative

tense markers (for example, Ogihara, 1999). I have pointed out that the use of *-ru* and *-ta* in embedded clauses may be either speech time oriented like English tenses or the matrix event time oriented. There is a piece of evidence which suggests that the Japanese verb-complement constructions involve an obligatory shift of the deictic center for tense interpretations. Thus, tenses in the verb-complements in Japanese are always evaluated with respect to the time of the matrix eventuality. However, such is not the case with the tenses in *toki*-clauses and relative clauses, since the shift of deictic center is not required in these constructions. Then, the presence of SOT in English and its absence in Japanese in the verb-complement constructions can be explained by the fact that the complement past in Japanese is always evaluated with respect to the matrix past, and therefore, may not be able to overlap with the matrix event. Thus, there is no need to posit a SOT rule for English nor different semantics for the past tenses in English and Japanese.

I have also shown that the reason some complex sentences only allow either *-ru* or *-ta* in embedded clauses in Japanese is not because of the idiosyncratic properties of the Japanese tenses, but because there is only one possibility for the ordering of two events in those cases (section 3.2.2). The ordering of the two events in a sentence partially depends on the aspect of the predicate and also on our knowledge of the world. I have shown that in both English and Japanese, aspect and pragmatics affect the ordering of events and the availability of tenses in the embedded clauses in an important way.

In this study I have defended the existence of the independent system of temporal reference in natural language grammar by separating the contribution of the theory of tense and aspect from that of other modules of the grammar. While the theory of intensionality, aspect of the original predicates, and our pragmatic knowledge interact

with the tense system in an interesting way, they should not be considered to be part of the system. Tense is a grammatical device to relate the topic time (TT) with the utterance time (TU), and past tense always encode $TT < TU$ and present tense always encode $TT \supset TU$ in any language that has present and past tenses in the grammar. Which morpheme encodes what TT-TU relation is part of the morpholexical information in a given language and is not part of the system of temporal reference. Such a view of the grammar of tense conforms to the assumptions of the current linguistic theory, namely, the Minimalist Program of linguistic theory (Chomsky, 1995), which argues that virtually all the linguistic information that is not innate is morpholexical information of the language that is being learned.

In Chapter 4 I have shown how Klein's system clarifies the respective contributions of tense and aspect components of the complex temporal morphemes like present progressive and present perfect forms in English, and the *-te-iru* form in Japanese. By applying Klein's theory in a novel way, separating the respective function of each component of these complex temporal morphemes, we can also account for the puzzling behavior of the Japanese *Verb-te iru* form, which can have both a perfect and a progressive interpretation. While the auxiliary verbs of the present progressive and the present perfect of English and the Japanese *-te-iru* form all provide the existential meaning and $TT \supset TU$ in their tense components, their aspectual components encode varied TT-TSit relations. I have argued that while the present participial and the past participial of English encode the specific way in which topic time is related to the situation time of the original predicate, *-te* of *-te-iru* simply introduces a topic time without specifying how it relates to the situation time. This will allow TT of the sentences with *-te-iru* to be placed

either within the situation time or in the post time of the situation, depending on the context, and accordingly, the sentence may receive either a progressive or a perfect interpretation.

Most of the arguments presented in this study are based on the simplicity of the theory. I have not proposed a brand-new theory of temporal reference, but have sorted out the contributions of the theory of temporal reference from the contributions of other modules of the grammar that interact with the temporal system. I hope that my study will serve as a stepping stone for an eclectic approach for the study of tense and aspect, which allow us to integrate any non-contradictory findings of syntax and semantics for fuller description and explanation of the natural language temporal expressions.

BIBLIOGRAPHY

- Abusch, Dorit. 1988. Sequence of Tense, Intensionality and Scope, *West Coast Conference on Formal Linguistics* 7, 1-14.
- Abusch, Dorit. 1991. The Present under Past as De Re Interpretation, *West Coast Conference on Formal Linguistics* 10, 1-12.
- Abusch, Dorit. 1994. Sequence of Tense Revisited: Two Semantic Accounts of Tense in Intensional Contexts, in H. Kamp (ed.), *Ellipsis, Tense and Questions*, Dyana-2 Esprit Basic Research Project 6852, Deliverable R2.2.B, 87-139.
- Abusch, Dorit. 1997a. Sequence of Tense and Temporal De Re, *Linguistics and Philosophy* 20 (1), 1-50.
- Abusch, Dorit. 1997b. Remarks on the State Formulation of De Re Present Tense. *Natural Language Semantics* 5, 303-313.
- Ando, Sadao. 1986. *Eigo no Ronri, Nihongo no Ronri: Taishoo- gengogakuteki Kenkyuu (The Logic of Japanese and the Logic of English: The Contrastive Study)*. Tokyo: Taishuukan.
- Bach, Emmon. 1986. The Algebra of Events. *Linguistics and Philosophy* 9, 5-16.
- Bennett, Michael. & Barbara Partee. 1978. Toward the Logic of Tense and Aspect in English, Indiana University Linguistics Club, Bloomington, Indiana.
- Bennett, Michael. 1981. Of Tense and Aspect: One Analysis. In Phillip J. Tedeschi and Annie Zaenen. (eds.) *Syntax and Semantics 14 : Tense and Aspect*. New York: Academic Press, 13-29.
- Binnick, Robert I. 1991. *Time and the verb: a guide to tense and aspect*. New York: Oxford University Press.
- Bybee, Joan. L., and Östen Dahl. 1989. The creation of tense and aspect systems in the languages of the world. *Studies in Language* 13: 51-103.
- Bybee, Joan. L., Revere Perkins, and William Pagliuca. 1994. *The Evolution of Grammar: Tense, Aspect, and Modality in The Languages of The World*. The University of Chicago Press.
- Carlson, Gregory N. 1977a. A Unified Analysis of the English Bare Plural. *Linguistics and Philosophy* 1, 413-457.

- Carlson, Gregory N. 1977b. Reference to Kinds in English. Ph.D. dissertation, University of Massachusetts, Amherst. Published 1980 by Garland Press, New York.
- Chomsky, Noam. 1981. *Lectures on Government and Binding*, Dordrecht Foris.
- Chomsky, Noam. 1986. *Knowledge of Language: its nature, origin and use*. New York: Praeger.
- Chomsky, Noam. 1995. *The Minimalist Program*, Cambridge, MA.: MIT Press.
- Comrie, Barnard. 1976. *Aspect*, Oxford: Cambridge University Press.
- Comrie, Barnard. 1985. *Tense*, Oxford: Cambridge University Press.
- Costa, Rachel. 1972. Sequence of Tenses in That-Clauses, in P. Peranteau, J. Levi, and G. Phares, eds., *Papers from the Eighth Regional Meeting*, Chicago Linguistic Society, University of Chicago, Chicago, Illinois.
- Cresswell, M.J. and A. von Stechow. 1982. De Re Belief Generalized, *Linguistics and Philosophy* 5 (4), 503-535.
- Dahl, Östen. 1985. *Tense and Aspect Systems*. New York: Basil Blackwell.
- Diesing, Molly. 1992. *Indefinites*. Cambridge, MA: MIT Press.
- Dowty, David. 1977. Toward a semantic analysis of verb aspect and the English 'imperfective' progressive. *Linguistics and Philosophy* 1, 45-77.
- Dowty, David. 1979. *Word Meaning and Montague Grammar. The Semantics of Verbs and Times in Generative Semantics and in Montague's PTQ*, Dordrecht: Kluwer Academic Publishers.
- Dowty, David. 1982a. Grammatical Relations and Montague Grammar. In P. Jacobson and G. Pullum (eds.), *On the Nature of Syntactic Representation*. Dordrecht: Reidel, 79-130.
- Dowty, David. 1982b. Tenses, Time, Adverbs and Compositional Semantic Theory. *Linguistics and Philosophy* 5, 23-55.
- Dowty, David. 1986. The Effect of Aspectual Classes on the Temporal Structure of Discourse: Semantics or Pragmatics? *Linguistics and Philosophy* 9, 37-61.
- Dowty, David. 1989. On the Semantic Content of the Notion "Thematic Role". In G. Chierchia, B.H. Partee and R. Turner (eds.), *Property, Theory, Type Theory and Natural Language, Vol.II: Semantic Issues*. Dordrecht: Reidel, 69-129.

- Enç, Mürvet. 1986. Toward a Referential Analysis of Temporal Expressions. *Linguistics and Philosophy* 9, 405-426.
- Enç, Mürvet. 1987. Anchoring Conditions for Tense. *Linguistic Inquiry* 18, 633-657.
- Garey, Howard B. 1957. Verbal aspect in French. *Language* 33, 91-110.
- Heim, Irene. 1994. Comments on Abusch's Theory of Tense. In: H. Kamp (ed.) *Ellipsis, Tense and Questions*, DYANA deliverable R2.2.B, University of Amsterdam, 143-170.
- Hinrichs, Erhard. 1986. Temporal Anaphora in Discourse of English, *Linguistics and Philosophy* 9, 63-82.
- Hopper, Paul. 1979. Aspect and Foregrounding in Discourse. T. Givon (ed.) *Syntax and Semantics Vol. 12: Discourse and Syntax*. New York: Academic Press.
- Hopper, Paul. and Thompson, Sandra. 1980. Transitivity in Grammar and Discourse. *Language* 56.
- Hornstein, Norbert. 1990. *As Time Goes By: Tense and Universal Grammar*. Cambridge, MA.: MIT Press.
- Jacobsen, Wesley. M. *The Transitive Structure of Events in Japanese*. Kurosio, Tokyo.
- Johnson, Marion. R. A Unified Temporal Theory of Tense and Aspect. In Phillip Tedeschi and A. Zaenen. (eds.) *Syntax and Semantics 14: Tense and Aspect*. New York: Academic Press, 145-175.
- Johnston, Michael. 1994a. When-clauses, Adverbs of Quantification, and Focus. In *Proceedings of WCCFL 13*. CSLI / Chicago University Press.
- Johnston, Micael. 1994b. *The Syntax and Semantics of Adverbial Adjuncts*. Doctoral Dissertation. University of California, Santa Cruz.
- Kamp, Hans. 1971. Formal Properties of 'now'. *Theoria* 37, 227-273.
- Kamp, Hans. 1984. A Theory of Truth and Semantic Representation. In J. Groenendijk, T.M.V. Janssen and M. Stokhof (eds.) *Truth, Interpretation and Information*. Dordrecht: Foris, 1-41.
- Kenny, Anthony. 1963. *Action, Emotion, and Will*. London and New York: Humanities Press.

- Kindaichi, Haruhiko. 1976. Kokugo Doosi no Iti-Bunrui (A Proposal for Japanese Verb Classification). Kindaichi (ed.), *Nihongo Doosi no Asupekuto (Aspect of Japanese Verbs)*, 7-26. Tokyo: Mugi Shoboo. Repr. -Originally published in *Gengo Kenkyuu* 15 (1950) 48-63.
- Klein, Wolfgang. 1994. *Time in Language*. London and New York: Routledge.
- Kratzer, Angelika. 1989. Stage-Level and Individual-Level Predicates. In *Papers on Quantification, NSF Grant Report*, Department of Linguistics, University of Massachusetts at Amherst.
- Krifka, Manfred, Francis Jeffry Pelletier, Gregory N. Carlson, Alice ter Meullen, Godehard Link, and Gennaro Chierchia. 1995. Genericity: An Introduction. In G. N. Carlson and F. Jeffry (eds.), *The Generic Book*, 1-124. University of Chicago Press.
- Kudo, Mayumi. 1989. Gendai Nihongo no Paafekuto o Megutte (On the Perfect of Modern Japanese). Gengogaku Kenkyuukai (ed.), *Kotoba no Kagaku*, 53-118. Tokyo: Mugi Shoboo.
- Kunihiro, Tetsuya. 1982. Nihongo, Eigo (Japanese and English) Teramura (ed.), *Kooza Nihongogaku: Gaikokugo to no Taishoo 11 (Lectures on Japanese Language: Contrastive Study)*, 2-8. Tokyo: Meiji Shoin.
- Kusumoto, Kiyomi. 1996. On the Syntax and Semantics of Tense. UMOP21.
- Landman, Fred. 1992. The Progressive. *Natural Language Semantics* 1, 1-32.
- Lewis, David. 1979. Attitudes De Dicto and De Se. *The Philosophical Review* 88, 513-543.
- McCawley, James. 1971. Tense and the Time Reference in English. In Charles J. Fillmore and D.T. Langendoen (ed.), *Studies in Linguistic Semantics*, 97-114. New York: Holt, Rinehart & Winston, Inc.
- McCawley, James. 1981. The Syntax and Semantics of English Relative Clauses. *Lingua* 53, 99-149.
- McClure, William. 1993. A Semantic Parameter: The Progressive in Japanese and English. In Soonja Choi (ed.), *Japanese/Korean Linguistics. Vol.3*, 254-270. CSLI, Stanford, CA
- McClure, William. 1995. *Syntactic Projections of the Semantics of Aspect*. Tokyo: Hituzi Syobo.

- McGloin, Naomi. H. 1989. *A Students' Guide to Japanese Grammar*. Tokyo: Taishuukan.
- Machida, Ken. 1989. *Nihongo no Jisei to Asupekuto*. Tokyo: Aruku.
- Miura, Akira. 1974. The *V-ru* Form vs. the *V-ta* Form. In *Papers in Japanese Linguistics* 3: 95-122. Los Angeles: University of Southern California.
- Nakamura, Akira. 1994a. Some Aspects of Temporal Interpretation in Japanese. In M. Koizumi & H. Ura (eds.), *Formal Approaches to Japanese Linguistics* 1, MIT Working Papers in Linguistics 24.
- Nakamura, Akira. 1994b. On the Tense System of Japanese. In N. Akatsuka (ed.), *Japanese Korean Linguistics* Vol. 4, CSLI, Stanford, 363-377.
- Nakau, Minoru. 1976. Tense, Aspect, and Modality. In M. Shibatani (ed.), *Syntax and Semantics 5: Japanese Generative Grammar*, New York: Academic Press. 421-482.
- Ogihara, Toshiyuki. 1989. *Temporal Reference in English and Japanese*. Ph.D. dissertation, University of Texas, Austin.
- Ogihara, Toshiyuki. 1995a. Double-Access Sentences and References to States. *Natural Language Semantics* 3, 177-210.
- Ogihara, Toshiyuki. 1995b. The Semantics of Tense in Embedded Clauses. *Linguistic Inquiry* 26, 663-679.
- Ogihara, Toshiyuki. 1999. Tense and Aspect, In Natsuko Tsujimura (ed.) *The Handbook of Japanese Linguistics*. Blackwell Publishers, 326-348.
- Ogihara, Toshiyuki. (in press.) The ambiguity of the -te iru form in Japanese. *Journal of East Asian Linguistics*.
- Okuda, Yasuo. 1978a. Asupekuto no kenkyuu o megutte, I (On the study of aspect, I). *Kokugo Kenkyuu* 53. 33-44.
- Okuda, Yasuo. 1978b. Asupekuto no kenkyuu o megutte, II (On the study of aspect, II). *Kokugo Kenkyuu* 54.14-27.
- Parsons, Terence. 1989. The progressive in English: events, states and processes. *Linguistics and Philosophy* 12, 213-241.
- Parsons, Terence. 1990. *Events in the Semantics of English*. Cambridge, Mass: MIT Press.

- Partee, Barbara. H. 1973. Some Structural Analogies between Tenses and Pronouns in English. *The Journal of Philosophy* Volume LXX. No. 18, 601-609.
- Partee, Barbara. H. 1984. Nominal and Temporal Anaphora. *Linguistics and Philosophy* 7, 243-286.
- Prior, Arthur N. 1967. *Past, Present, and Future*. Oxford: Oxford University Press.
- Pustejovsky, James. 1991. The syntax of event structure. *Cognition*, 41, 47-81.
- Quine, Willard, V.O. 1956. Quantifiers and Propositional Attitudes. *The Journal of Philosophy* 53: 183-194.
- Reichenbach, Hans. 1947. *Elements of Symbolic Logic*. New York: Macmillan.
- Scheffer, Johannes. 1975. *The progressive in English*. Amsterdam : North-Holland Pub. Co. ; New York : American Elsevier Publishing Company.
- Schmitt, Cristina. J. 1996. *Aspect and the Syntax of Noun Phrases*. Ph.D. dissertation, University of Maryland at College Park.
- Shinzato, Rumiko. 1993. A unified analysis of Japanese aspect marker *te iru*, *Language Quarterly* Vol 31: 1-2, 41-57.
- Shinzato, Rumiko. 1994. The modal and discoursal functions of temporal auxiliaries in Japanese - a cognitive account, *Journal of Asian Pacific Communication* Vol. 5: 1-2, 89-103.
- Smith, Carlota. S. 1978. The syntax and interpretation of temporal expressions in English, *Linguistics and Philosophy* 2, 43-100.
- Smith, Carlota. S. 1981. Semantic and Syntactic Constraints on Temporal Interpretation. In Philip J. Tedeschi and Annie Zaenen (eds.), *Syntax and Semantics Vol. 14: Tense and Aspect*, New York: Academic Press. 213-237.
- Smith, Carlota. S. 1997. *The Parameter of Aspect (Second Edition)*. Dordrecht: Kluwer Academic Press.
- Soga, Matsuo. 1983. *Tense and Aspect in Modern Colloquial Japanese*. Vancouver: University of British Columbia Press.
- Spejewski, Beverly and Greg N. Carlson. 1991. Reference Time Relations. In Proceedings of Eastern States Conference on Linguistics, The Ohio State University Press. 325-334.

- Spejewski, Beverly and Greg N. Carlson. 1993. Proceedings of the North East Linguistic Society 23, Vol.2. University of Ottawa. 479-493.
- Spejewski, Beverly. 1996. Temporal Subordination and the English Perfect. In Proceedings of Semantics and Linguistics Theory VI. Ithaca, NY: Cornell University. 261-278.
- Stowell, Tim. 1981. *Origins of Phrase Structure*, Doctoral dissertation, MIT, Cambridge, Mass.
- Stowell, Tim. 1995. What do the present and past tense mean? In Pier Marco Bertinetto et al (eds.), *Temporal Reference, Aspect and Actionality, Vol. 1: Semantic and Syntactic Perspectives*. Torino: Rosenberg & Sellier.
- Stowell, Tim. 1996. The Phrase Structure of Tense. In Johan Rooryck and Laurie Zaring (eds.) *Phrase Structure and the Lexicon*. Dordrecht: Kluwer Academic Publishers.
- Taylor, Barry. 1977. Tense and Continuity. *Linguistics & Philosophy*, 1. 199-220.
- Tedeschi, Phillip J. & Annie Zaenen (eds). 1981. *Syntax and Semantics 14: Tense and Aspect*. New York: Academic Press.
- Teramura, Hideo. 1984. *Nihongo no Sintakusu to Imi (Japanese Syntax and Meanings)*. vol. II. Tokyo: Kuroshio.
- ter Meulen, Alice. G. B. *Representing Time in Natural Language*. MIT Press: Cambridge, MA.
- Tsujimura, Natsuko. (ed.) 1999. *The Handbook of Japanese Linguistics*. Blackwell Publishers.
- Vendler, Zeno. 1967. Verbs and Times. *Linguistics and Philosophy*, 97-121. Ithaca, N.Y.: Cornell University Press. Reprint of 1957. Verbs and Times. *Philosophical Review* 66, 143-160.
- Verkuyl, Henk J. 1972. *On the Compositional Nature of Aspects*. Dordrecht: D. Reidel.
- Verkuyl, Henk J. 1989. Aspectual classes and aspectual composition. *Linguistics and Philosophy* 12, 39-95.
- Verkuyl, Henk J. 1993. *A theory of aspectuality: The interaction between temporal and atemporal structure*. Cambridge, N.Y.: Cambridge University Press.
- Vlach, Frank. 1973. 'Now' and 'Then', *A Formal Approach in the Logic of Tense Anaphora*, Ph.D. dissertation, UCLA.

- Vlach, Frank. 1981. The Semantics of the Progressive. In Philip J. Tedeschi and Annie Zaenen (eds.), *Syntax and Semantics* Vol. 14: *Tense and Aspect*, New York: Academic Press, 271-292.
- Vlach, Frank. 1993. Temporal Adverbials, Tenses and the Perfect. *Linguistics and Philosophy* 16, 231-283.
- Washio, Kenichi and K. Mihara. *Voisu to Asupekuto*. (Voice and Aspect.), Kenkyuusha, Tokyo.
- Yamagata, Ayako. 1994. An aspectual form *-te-iru* in Japanese. Unpublished MA thesis, Michigan State University.
- Yamagata, Ayako. 1997. Unified analysis of *-te-iru* as a stage-level operator. In *Proceedings of the 12th Annual Meeting of South Eastern Association of Teachers of Japanese (SEATJ)*. Geogia Institute of Technology, Atlanta, Georgia.
- Yamagata, Ayako. 1998. The stage-level/individual-level distinction: an analysis of *-te-iru*. In David J. Silva (ed.), *Japanese/Korean Linguistics* 8, CSLI, Stanford, 245-258.
- Zagona, Karen. 1989. On the non-identity of morphological tense and temporal interpretation. In *Studies in Romance linguistics*, Amsterdam; Philadelphia: John Benjamins Pub. Co.
- Zagona, Karen. 1991. Perfective Haber and the theory of tenses. In Hector Campos and Fernando Martínez-Gil. (eds.) *Current studies in Spanish linguistics*, Washington, D.C: Georgetown University Press. 379-403.
- Zagona, Karen. 1992. Tense binding and construal of present tense. *Theoretical analyses in Romance linguistics. Current Issues in Linguistic Theory* 74, 385-398. Amsterdam : John Benjamins. 385-398.
- Zagona, Karen. 1994. Perfectivity and Temporal Arguments. In Michael L. Mazzola (ed.) *Issues and Theory in Romance Linguistics*, Washington, D.C.: Georgetown University Press. 523-546.
- Zagona, Karen. 1997. Tenses and Anaphora: Is there a tense-specific theory of coreference? To appear in: Barssm Andrew and D. Terence Langendoen (eds.) *Current Topics in Anaphora* (preliminary title). Oxford, UK and Cambridge MA: Blackwell.

MICHIGAN STATE UNIVERSITY LIBRARIES



3 1293 02112 5830