

A Study on Movie-Based Context-Aware Learning

ハズリアニ

<https://hdl.handle.net/2324/1959141>

出版情報 : Kyushu University, 2018, 博士 (工学) , 課程博士
バージョン :
権利関係 :



A Study on Movie-Based Context-Aware Learning

Hazriani

July 2018

A Study on Movie-Based Context-Aware Learning

**A thesis submitted for the degree of
Doctor of Engineering
August 2018**

Hazriani

**Department of Advanced Information Technology
Graduate School of Information Science and Electrical Engineering
Kyushu University**

This thesis is dedicated to my parents and my husband
who have always encouraged me
to pursue my doctoral degree

Abstract

This thesis proposes the concept of movie-based context-aware learning (or MBCAL for short) as an alternative approach for context-aware ubiquitous learning (or u-learning for short) to facilitate contextual learning. As the movie contains extremely rich contexts, it can be a place in where quasi context-aware learning is practiced. In MBCAL, learners learn through situations (contexts) that occur in the movie scenes. MBCAL presents richer contextual learning contents with less effort than context-aware u-learning. This idea was raised after recognizing some limitations that may occur in managing a context-aware u-learning system.

Researches on context-aware u-learning systems have shown that learning from the real environment is attractive and effective for learners; however, it requires developers as well as operators of much effort to realize and utilize such a real environment. Developers have to develop a complicated learning application interacting with the real environment through various sensors and other smart devices. At the same time, operators have to prepare contextual learning materials that can provide learning tasks depending on dynamic situations. Furthermore, physical devices such as sensor network or other sensing devices deployed in the real environment can be out of order by long time operation under severe environment such as outdoors. In addition, availability of real time network connection is another important issue to be concerned. These complexities lead more development cost or limited functionality of the learning system. Therefore, concerning these issues, the author proposes the concept of MBCAL.

The learning process of MBCAL is conducted with context-awareness in the imaginary context situated in the movie scene. Since events in the movie are foreseeable, we can describe contexts of the movie scenes in advance. The developer does not need special devices to acquire the context as well as efforts to know unexpected circumstances in design of the context-aware learning system. Through this approach, various learning situations can be presented with less effort than a context-aware u-learning system does. These contribute to limit cost and effort to develop and operate the system. In short, MBCAL utilizes movies to facilitate learners proactive learning from the context situated in the scenes as ubiquitous learning systems try to do in the real environment. A necessary preparation to use the MBCAL system is to annotate the scenes in the movie with their contexts.

This thesis has the following contributions:

- (1) This thesis presents the concept of movie-based context-aware learning (MBCAL) concept, an alternative way of context-aware learning by facilitating learners proactive learning from the context situated in the movie scenes as ubiquitous learning systems try to do in the real environment.
- (2) This thesis introduces a novel approach for context description of the movie scenes, *i.e.* object-oriented context model (or OOCM for short). The OOCM incorporates with the case grammar concept which is popular in natural language processing. The OOCM functions as a reference model to describe the context of the movie scenes in a unified and formal manner to ensure uniformity of description.
- (3) This thesis defines the context description language (CDL), a textual language for context description which is subject to the OOCM. This CDL is designated to enable instant description of the movie scenes.

- (4) This thesis introduces the movie-based context-aware quiz generation, representing a movie-based context-aware language learning (MBCALL) system as one potential implementation domain of the MBCAL concept. The quiz generation system facilitates self-interactive language learning through movie. The system generates quizzes based on the context of the replaying scene, which is described under OOCM framework, trains and examines learners by the quizzes in an interactive manner.

Evaluation was conducted to assess performance of both the OOCM and the proposed MBCALL system. Evaluation of the OOCM shows that the current version of the OOCM can guide the instructor who composes the movie context to enrich information included in the movie context. It is indicated by a fact that more OOCM associations describing the movie contexts can be identified and described by the evaluation subjects by referencing the OOCM than without referencing the OOCM. Moreover, MBCALL contributed 3.4% to improve learners performance of learning than conventional learning. Furthermore, learner's subjective ratings shows that MBCALL can be positively adopted in terms of its activeness and authenticity.

Contents

List of Figures	viii
List of Tables	x
1 Introduction	1
1.1 Background	1
1.2 Objectives	3
1.3 Contributions	4
1.4 Structure of the Thesis	4
2 Context-Aware Learning	6
2.1 Fundamentals of the Context	6
2.1.1 Definition of the Context	6
2.1.2 Context Category	7
2.1.3 Context Framework for TEL	8
2.1.4 Context-Aware Application	9
2.2 Ubiquitous Learning (<i>u</i> -Learning)	10
2.2.1 Definition of Ubiquitous Learning (u-Learning)	10
2.2.2 Definition of Context-Aware u-Learning	13

2.2.3	Ubiquitous Language Learning	15
2.2.4	Challenges in Context-Aware u-Learning	15
2.3	x-Learning Paradigms	17
2.3.1	Emerging Learning Paradigms	17
2.3.2	Comparison of the x-Learning	18
3	Movie-Based Context-Aware Learning (MBCAL)	23
3.1	Movie for Learning	23
3.2	Movie-Based Context-Aware Learning (MBCAL)	24
3.3	Three Tiered Architecture of the Movie-Based Context-Aware Learning System	25
4	Object-Oriented Context Model	28
4.1	Context Description	28
4.2	Case Grammar Concept	29
4.3	Introduction of the Case Grammar Concept to the OOCM	33
4.3.1	OOCM Class Elements	33
4.3.2	OOCM Based on Verbert's Context Framework	35
4.4	Annotating Movies by Context Description	35
4.5	Context Description Language (CDL)	36
5	Movie-Based Context-Aware Quiz Generation	40
5.1	Movie-Based Context-Aware Quiz Generation	40
5.2	Commonality-Variability Analysis of Quizzes	41
5.3	Components of the System	42
5.4	Quiz Generation Template	44
5.5	Language Dependent Transformation	47
5.6	Prototype of the Context-Aware Quiz Generator	47

5.7	Comparison with Existing Quiz Generation Systems	48
6	Evaluation and Discussion	53
6.1	Evaluation of the OOCM	53
6.2	Evaluation of Learning	60
6.2.1	Overview of the Evaluation Process	60
6.2.2	Quiz Preparation	61
6.2.3	Evaluation of Learning Outcomes	62
6.2.4	Evaluation by the Learner's Subjective Ratings	66
6.3	Discussion	66
6.3.1	Learning Outcomes	66
6.3.2	Learner's Feedback	68
7	Conclusion and Future Work	70
	Acknowledgement	72
	Bibliography	73
A	OOCM Based on Verbert's Context Framework	79
B	OOCM Based on the Case Grammar Concept	84
C	Quizzes Used for Evaluation of Learning	88
D	Published Papers	93
D.1	Journals (peer reviewed)	93
D.2	International Conference (peer reviewed)	93

List of Figures

2.1	Context Dimension [8]	8
2.2	Verbert's Context Framework for TEL	9
2.3	Context-Aware u-Learning, u-Learning, Mobile Learning, and Ubiquitous Computing in Learning	14
2.4	Research on Computer Education [11]	17
3.1	Key Concepts of MBCAL	25
3.2	Three Tiered Architecture of the Movie-Based Context-Aware Learning System	26
4.1	Fillmore's Case Grammar First Base Rule	30
4.2	Object Oriented Context Model	34
4.3	Example of Annotation View Using ANVIL	36
4.4	Syntax of the CDL	37
4.5	Example Movie Scene	38
4.6	An Example OOCM	38
4.7	Example CDL	39
5.1	Feature Model of the Quiz Types	42
5.2	Components of an MBCALL Application and Their Interaction	43
5.3	An Example Movie Scene	44

5.4	Context Description in CDL	45
5.5	Example Quiz Template for English Learning	45
5.6	Example Quiz Template for Bahasa Learning	46
5.7	Generated Quiz and Expected Answer for English Learning	46
5.8	Generated Quiz and Expected Answer for Bahasa Learning	46
5.9	Be-Verb Transformation	47
5.10	Present Verb Transformation	48
5.11	Past Verb Transformation	48
5.12	Past Participle Verb Transformation	49
5.13	Auxiliary Verb Transformation	49
5.14	Quiz Generation and Validation of the Answers by a Prototyped Context-Aware Quiz Generator	51
5.15	Context Description	52
6.1	Comparison of Average Learning Outcomes	67
6.2	Comparison of Average Learning Outcomes for Japanese Vocabulary Learning	68
6.3	Comparison of Average Learning Outcomes for Bahasa Vocabulary Learning	68
6.4	Learner's Acceptance to MBCALL	69
A.1	OOCM Based on Verbert's Context Framework	83

List of Tables

2.1	Zimmermann's Categorization of the Context[8]	20
2.2	Ubiquitous Language Learning Applications	21
2.3	x-Learning Comparison	22
4.1	Survey on the Case Grammar[29]	32
6.1	Comparison Results	60
6.2	Time for Composing Movie Description	60
6.3	Criteria for Learner's Subjective Ratings	61
6.4	List of Quizzes to Learn Bahasa Vocabulary through MBCALL	62
6.5	List of Quizzes to Learn Japanese Vocabulary through MBCALL	63
6.6	List of Quizzes to Learn Bahasa Vocabulary through Conventional Learning .	64
6.7	List of Quizzes to Learn Japanese Vocabulary through Conventional Learning	64
6.8	Summary of Learner's Scores through MBCALL	65
6.9	Summary of Learner's Score through Conventional Learning	65
6.10	Average Learning Outcomes	66
6.11	Learning Outcomes for Japanese Vocabulary Learning	67
6.12	Learning Outcomes for Bahasa Vocabulary Learning	67

Chapter1

Introduction

1.1 Background

Context-aware ubiquitous learning has been considered as the most meaningful approach of technology enhanced learning (TEL) to provide contextual learning so far, since it enables dynamic learning with contexts in the real environment. However, it requires developers as well as operators of much effort to realize and utilize context-aware ubiquitous learning systems. Developers have to develop a complicated learning application interacting with the real environment through various sensors and smart devices. At the same time, operators have to prepare contextual learning materials that can provide learning tasks suitable for the current context of the real environment. Schmidt[1] and Verbert *et al.*[2] stated that context acquisition and representation of the real environment is one of the most significant challenges in development of the context-aware ubiquitous learning application.

Furthermore, physical devices such as various sensing devices and its network deployed in the real environment can be out of order by long time operation under severe environment such as outdoors. Puccinelli and Haenggi stated that the sensor network offers powerful combination of distributed sensing, computing, and communication for applications, yet at the same time it brings numerous challenges due to their failures[3]. In addition, availability of real

time network connection is another important issue to be concerned. These complexities lead more development cost or limited functionality of the learning system[4]. As also highlighted by Hwang *et al.*[5], expense of constructing context-aware ubiquitous learning systems is still very high.

Therefore, the author proposes the concept of movie-based context-aware language learning (or MBCAL for short) as an alternative approach for context-aware learning to overcome the above-mentioned difficulties on development and operation of the context-aware system working in the real environment. As the movie contains extremely rich contexts which can facilitate comprehension activities that are perceived as real [6], it can be a place in where quasi context-aware ubiquitous learning is conducted. Necessary preparation to use the movie-based context-aware learning system is to annotate the scenes in the movie with their contexts.

In MBCAL, learners learn through situations that occur in the movie scenes. In the other words, the learning process is given with context-awareness in the virtual context situated in the movie scenes. Since events in the movie are foreseeable, we can describe contexts of the movie scenes in advance. The developer does not need any special devices to acquire the context as well as efforts to know unexpected circumstances in design of the context-aware learning system[4]. Through this approach, various learning situations can be presented with less effort than the context-aware ubiquitous learning system. For instance, we can create learning tasks associated with the jungle context from the movie scenes of the jungle instead of managing learning tasks situated in the real jungle.

These contribute to limit cost and effort to develop and operate the system. In short, MBCAL utilizes movies to facilitate learner's proactive learning from the context situated in the movie scenes as ubiquitous learning systems try to do in the real environment. Considering that the language has been the most popular learning domain targeted by context-aware learning so far, the movie-based context-aware language learning (MBCALL) will be a suitable case study to realize the concept of MBCAL proposed in this study.

Two key concepts for MBCAL are the object-oriented context model (or OOCM for short) and context-aware quiz generation. The OOCM is a reference model to describe the context of the movie scenes in a unified and formal manner to ensure uniformity of description. The movie is annotated by context description of its scenes under this OOCM. Context-aware quiz generation is a function of the learning system that generates quizzes based on context description of the currently replaying scene, which described in the above-mentioned OOCM, and trains and examines learners by the quizzes in an interactive manner. It enables self-learning and also interactive training. Response of the learner to the generated quizzes is the only means to check if he/she followed learning contents in the movie and understood them.

1.2 Objectives

Corresponding to the ultimate goal of this study, that is to facilitate learners to enjoy interactive learning with contexts through movie scenes, the author defines the following research questions:

- (1) Does the proposed OOCM has sufficient descriptive power in describing contexts of the movie scenes?
- (2) Does the proposed context-aware quiz generation system generate quizzes sourced from context description of the movie scenes appropriately?
- (3) Does MBCALL contribute to improve learner's learning outcome?
- (4) Do the learners feel satisfied to the proposed MBCALL system?

To settle these questions, the author defines different versions of OOCM and examines their descriptive power by describing context description. The author also constructs a prototype system of MBCALL that performs context-aware quiz generation from context description

of the movie scenes. Through experiments with the prototype system, it is evaluated how MBCALL can improve learner's learning outcome and how learners feel during learning.

1.3 Contributions

Contributions of this study are summarized as follows:

- (1) The study defines the concept of movie-based context-aware learning as an alternative way of context-aware learning ubiquitous learning.
- (2) The study establishes the object-oriented context model (or OOCM), which is an object-oriented model to describe contexts of the movie scenes. The OOCM incorporates the idea of the case grammar concept.
- (3) The study defines a context description language (or CDL) to describe an instance of the OOCM in a textual form.
- (4) The study realizes the movie-based context-aware quiz generation for language learning. The system applied the context of the replaying scene to templates to generate quizzes and their expected answers.

1.4 Structure of the Thesis

The thesis is organized by seven chapters. The contents of the remaining chapters are as follows.

Chapter 2 presents concepts of the context-aware learning including comparison of recent learning paradigms, relevance of MBCAL to them, and also the standing point of MBCAL.

Chapter 3 introduces the concept of MBCAL as well as existing works on movie-based learning promoting suitability of the MBCAL approach.

Chapter 4 proposes the OOCM as a framework for context description of the movie scenes.

Chapter 5 describes the movie-based context-aware quiz generation for MBCALL as one of possible implementation of the MBCAL concept.

Chapter 6 conducts evaluation of MBCALL and shows the results.

Chapter 7 concludes this thesis as well as points future direction of MBCAL.

Chapter2

Context-Aware Learning

2.1 Fundamentals of the Context

2.1.1 Definition of the Context

Dey and Abowd defined the *context* in context-aware computing as *any information that can be used to characterize the situation of an entity*[7]. The entity can be a person, place, or object that is considered relevant to interaction between a user and the application, including the user and application themselves. This definition is adopted extensively in various application domains. In case of technology enhance learning, the context can be defined as *any information characterizing the situation of entities involved in the learning system*.

Furthermore, Dey and Abowd stated that a system is considered as *context-aware* if it uses the context to provide relevant information and/or services to the user. The information and services to be provided depend on the user's task. In the field of context-aware ubiquitous learning, Hwang *et al.* state that learning is *context-aware* if learner's situation and/or the real environment surrounding the learner can be sensed, implying that the learning system is able to conduct learning activities in the real world. Based on the definitions, it can be concluded that context-awareness is ability of a system to provide relevant information or services to the

user depending on user's situation and/or other related situations comprised in the system.

2.1.2 Context Category

Categorization of the context type helps developers of context-aware applications define which contexts are likely useful for the applications. The context-aware application observes *whos*, *wheres*, *whens*, and *whats* of entities and use them to determine reasons why the situation is occurring. Dey and Abowd introduced four primary context types to characterize the situation of a particular entity including *identity*, *location*, *time*, and *activity*[7].

Furthermore, they stated that these context types not only answer the questions of *who*, *where*, *when*, and *what* but also act as indices to other sources of contextual information. For example, given a person's identity, we can acquire many pieces of related information such as his/her phone number, address, email address, birthday, list of friends, relationship to other people in the environment *etc.*. Given an entity's location, we can determine what other objects or people are near the entity and what activity is occurring near the entity.

Afterwards, Zimmermann *et al.* extended Dey's context types and promoted five context dimensions illustrated in Figure 2.1[8]. Zimmermann *et al.* added the *relation* context and expanded Dey's *identity* context as *individuality* context. Table 2.1 presents more comprehensive information on Zimmerman's context dimension.

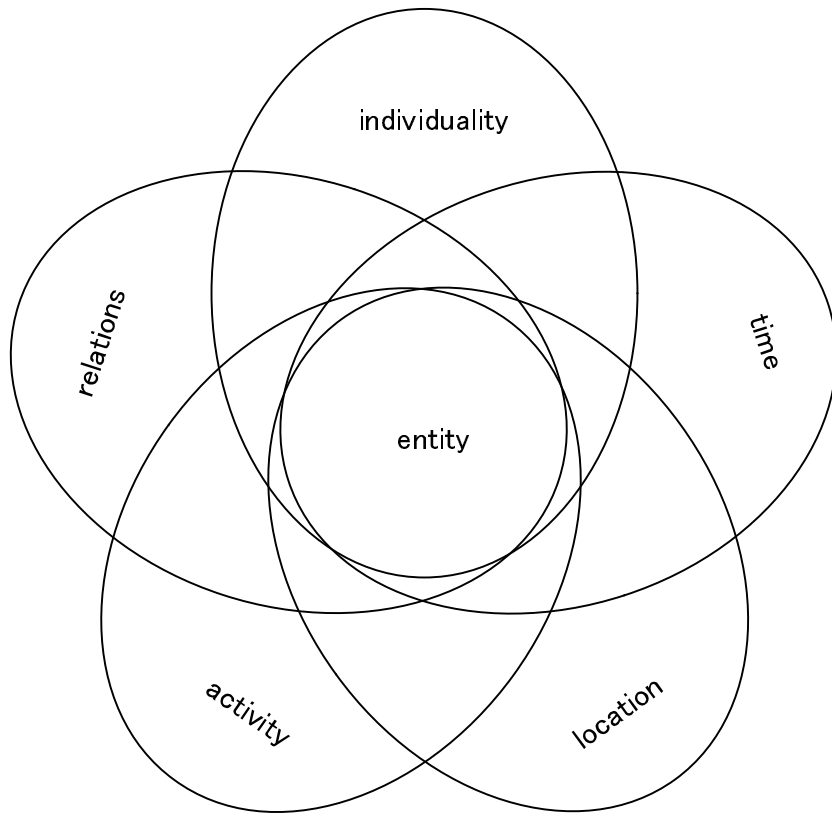


Figure 2.1: Context Dimension [8]

2.1.3 Context Framework for TEL

Since the above-mentioned definition of the context is still too abstract or general to be applied in implementation of the system, some researchers in different domains have attempted to re-define the context or its framework that can be applied to systems of their domains. A candidate of the comprehensive context framework in the domain of TEL will be one presented by Verbert *et al.*[2]. Figure 2.2 shows Verbert's context framework. The framework has seven context dimensions that relevant to context-aware applications in TEL, namely *computing*, *activity*, *time*, *resource*, *location*, *physical condition*, *user* and *social relations*. In addition, the framework defines the context elements of each context dimension as well as identifies relevant standards and specifications which can be used to represent data elements within those

categories.

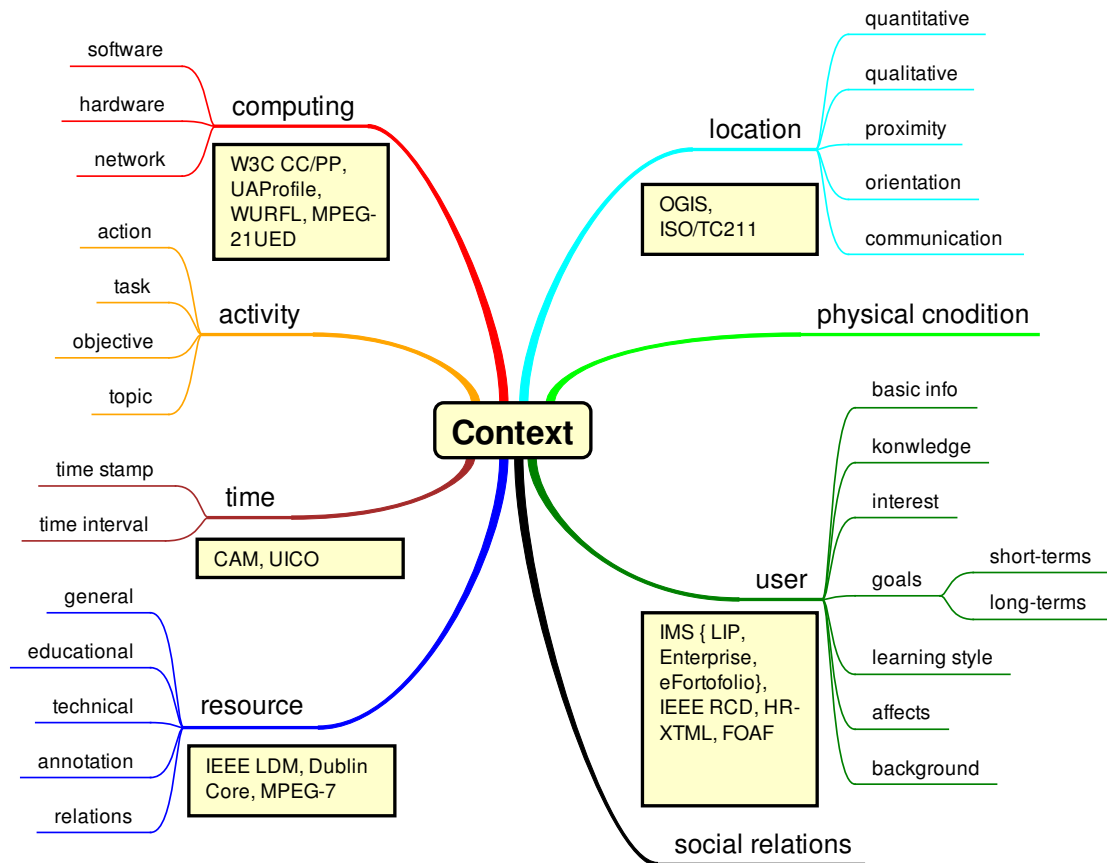


Figure 2.2: Verbert's Context Framework for TEL

2.1.4 Context-Aware Application

Based on survey on context-aware computing, Dey and Abowd introduced three categories of context-aware functions that the context-aware application may implement[9].

- Presenting information and services; in more concrete, presenting context information to the user or suggesting alternatives actions appropriate for the context to the user. For example, showing nearby gasoline stations, recommending a list of learning materials

suitable for learner's interest, showing a list of nearby learners who can be collaborative peers, *etc.*.

- Executing services automatically appropriate for the context; in more concrete, triggering a command or reconfiguring the system based on the user's context automatically or providing adaptive services to the user. For example, recomputing a path to the destination when the user misses the corner in car navigation system, delivering English speaking practices based on the learner's location, *etc.*.
- Attaching context information for later retrieval; in more concrete, tagging captured data with relevant context information and delivering suitable services later. For example, tagging notes taken by the user with the location and time of the observation, providing an interface to access informal meeting notes based on the attendants, time, and place of the meeting, *etc.*.

The contexts utilized by each context-aware application may vary depending on its learning purposes and can be derived from Verbert's framework in different manners. For example, context-aware learning systems in [10, 11] consider the contexts of location and activity, while systems in [5, 12] consider the contexts of user and location.

2.2 Ubiquitous Learning (*u*-Learning)

2.2.1 Definition of Ubiquitous Learning (*u*-Learning)

Basically, the definitions of ubiquitous learning (or *u*-learning for short) have been discussed in many literatures. However, many researchers have different views in definition of ubiquitous learning. Yang [13] defined the term *ubiquitous learning* by relying on ubiquitous computing technologies. However, Hwang *et al.* argued that Yang's definition of *u*-learning is

more appropriate for the definition of mobile learning[5]. He also proposed twelve context-aware u-learning models which can be used to conduct learning activities and to access learning performance based on their real-world and online behaviors as presented below[5].

- Learning in the real world with online guidance

Learners learn in the real world and are guided by the system based on the personal profile, portfolio and real world data collected by the sensors. For example, for the learner who takes an experimental chemistry course, the learning system provides hints automatically based on his/her behaviors in the real world during the experiment procedure.

- Learning in the real world with online support

Learners learn in the real world and are supported automatically by the system based on the personal profile, portfolio, and real world data collected by the sensors. For example, for the student who learns to identify the types of plants on campus, the learning system provides relevant information concerning the features of each type of plant automatically based on his/her location and the plants around him/her.

- Online test based on real-world object observations

Learners asked to answer questions presented on the screen of the mobile device by observing the real world objects. For example, the learning system asks the learner to answer the type of the tree located in front of him/her.

- Real object observation

Learners are asked to find the object in the real world based on the question presented in the mobile device. For example, the learning system asks the learner to observe the plants around him/her and find the plant that is most similar to the one shown on the screen.

- Collect data in the real world via observations

Learners are asked to collect data by observing objects in the real world and transfer

the data to the server via wireless communications. For example, the learning system asks the learner to observe the plants around him/her and transfer the photos taken or description of the feature of each plant written by him/her to the server.

- Collect data in the real world via sensors

Learners are asked to collect data by sensing objects in the real world and report what they have found. For example, the learning system asks the learner to find three different samples of water and report any contaminant found by using the sensors.

- Identification of a real world object

Learners are asked to answer the questions concerning the identification of the real-world objects. For example, the learning system asks the learner to answer the name of the insect shown by the instructor.

- Observation of the learning environment

Learners are asked to answer the questions concerning observation of the learning environment around them. For example, the learning system asks the learner to observe the school garden and upload the names of all of the insects he/she find.

- Problem-solving via experiments

Learners solve problems by designing experiments in the real world and finding hints on the internet. For example, the learning system asks the learner to consider the balloon given by the instructor and design an experiment to find the relationship between the payload mass and the altitude of the balloon.

- Real world observation with online data searching

Learners are asked to observe the real world objects and find solutions by accessing the network. For example, the learning system asks the learner to observe the building in front of him/her and find detailed data about it online.

- Cooperative data collecting

A group of learners is asked to cooperate to collect data in the real world and discuss their findings with others via mobile devices. For example, the learning system asks learners to draw a map of the school by measuring each area and integrate the collected data with other learners.

- Cooperative problem solving

Learners are asked to cooperate to solve problems in the real world by discussing through mobile devices. For example, the learning system asks learners to search each corner of the school and find the evidence that can be used to determine the degree of air pollution.

On the other hand, Zhan and Yuan defined ubiquitous learning as a new learning paradigm that learners can learn anything, anytime and at anyplace by utilizing ubiquitous computing technologies[14]. Consistent to their definition, Yahya *et al.* defined ubiquitous learning as a learning paradigm such that learning is provided in a ubiquitous computing environment that enables learning the right thing at the right place and time in the right way[15].

2.2.2 Definition of Context-Aware u-Learning

The term *context-aware ubiquitous learning* (or *context-aware u-learning* for short) was introduced by Hwang *et al.*[5]. He defined context-aware u-learning as *learning that employs mobile devices, wireless communication and sensor technologies in learning activities, which is able to support learning from real world, adaptive to learners behavior and context, actively provide personalized supports (services) and enable to adapt the subject content to meet the function of various devices*. Furthermore, to distinguish context-aware u-learning with u-learning, mobile learning, and ubiquitous computing in learning, Hwang *et al.* depicted the relationship among them as shown in Figure 2.3.

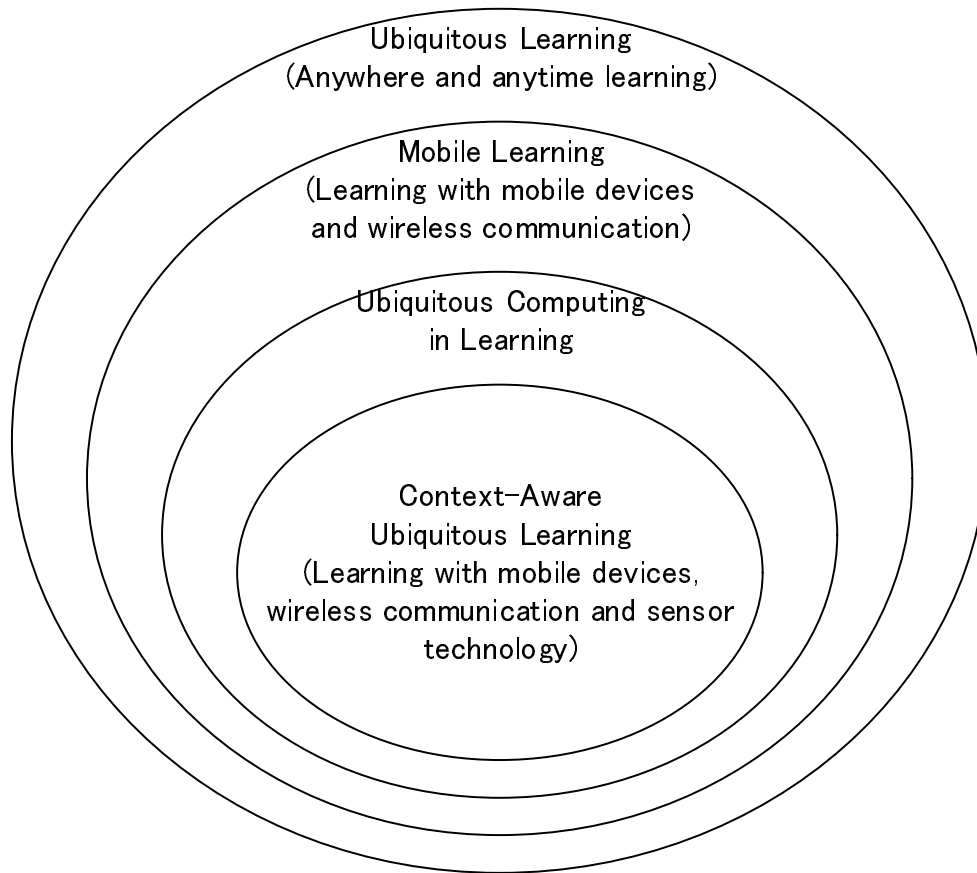


Figure 2.3: Context-Aware u-Learning, u-Learning, Mobile Learning, and Ubiquitous Computing in Learning

Accordingly, Hwang *et al.* summarized five potential criteria of a context-aware u-learning environment as follows:

- A context-aware u-learning environment is context-aware; that is, the learners situation or the situation of the real environment in which the learner is located can be sensed, implying that the system is able to conduct the learning activities in the real world.
- A context-aware u-learning environment is able to offer more adaptive supports to the learners by taking into account their learning behaviors and contexts in both the cyber world and the real world.
- A context-aware u-learning environment can actively provide personalized supports or

hints to the learners in the right way, in the right place, and at the right time, based on the personal and environmental contexts in the real world, as well as the profile and learning portfolio of the learner.

- A context-aware ubiquitous learning environment enables seamless learning from place to place within the predefined area.
- A context-aware ubiquitous learning environment is able to adapt the subject content to meet the functions of various mobile devices.

2.2.3 Ubiquitous Language Learning

A lot of research work on ubiquitous language learning have been conducted so far. Various sophisticated u-learning systems are produced. Some of them are context-aware systems. Ogata *et al.* introduced TANGO[10], LORAMS and SCROLL [11]. TANGO is an RFID applied system to learn words of physical objects. LORAMS enables learner's experience to be recorded on video with automatically linking to real objects in the scene by scanning the RFID tags attached to the objects. SCROLL provides an interface that helps the learners record, share, and reuse ubiquitous learning logs (ULLOs) with Android mobile devices. Chen *et al.* developed a personalized mobile English vocabulary learning system based on the item response theory and the learning memory cycle [16]. Li *et al.* developed an adaptive Chinese character learning system using mobile phones[12]. Underwood *et al.* developed miLexicon, a mobile collaborative language learning application[17]. Table 2.2 summarizes work on ubiquitous language learning.

2.2.4 Challenges in Context-Aware u-Learning

Verbert *et al.* figured out seven challenges in present context-aware recommender systems: context acquisition, context representation, evaluation, data set sharing, privacy, interaction,

and global data infrastructures[2].

On the other hand, Hwang *et al.* stated that the expense of constructing context-aware u-learning systems is still very high[5]. Therefore, they suggested that one might need to take the following considerations into account in implementation of the context-aware u-learning system.

- Do the learners need support from the system?
- Do they need personalized instructions?
- Do the instructions or support need to be given actively?
- Do the learners need to move from place to place during the learning process?
- Do the learners need to learn in the real world?
- Is u-learning in the real world better than simulation, e-learning or simple lectures?
- Does the context (*e.g.* location or environmental temperature) of the learner affect the learning process?
- Can the system be used in the future for other learners?

If the answers to most of those questions are *Yes*, the instructor can consider implementation. In addition, they suggested that subjects such as science (*e.g.* the observation of plants or animals), oral language practice, physical education, or e-training in some factories may have more potential for the application of u-learning.

2.3 x-Learning Paradigms

2.3.1 Emerging Learning Paradigms

Development of information and communication technology (ICT) has derived a series of emerging learning paradigms. Ogata and Uosaki presented a research map on use of computers in education that identifies six emerging learning paradigms including CAI/ITS, GBL, CSCL, web based learning (or e-learning), mobile learning (or m-learning), and ubiquitous learning (or u-learning), along with involved ICTs and their contribution to user behaviors as presented in Figure 2.4[11].

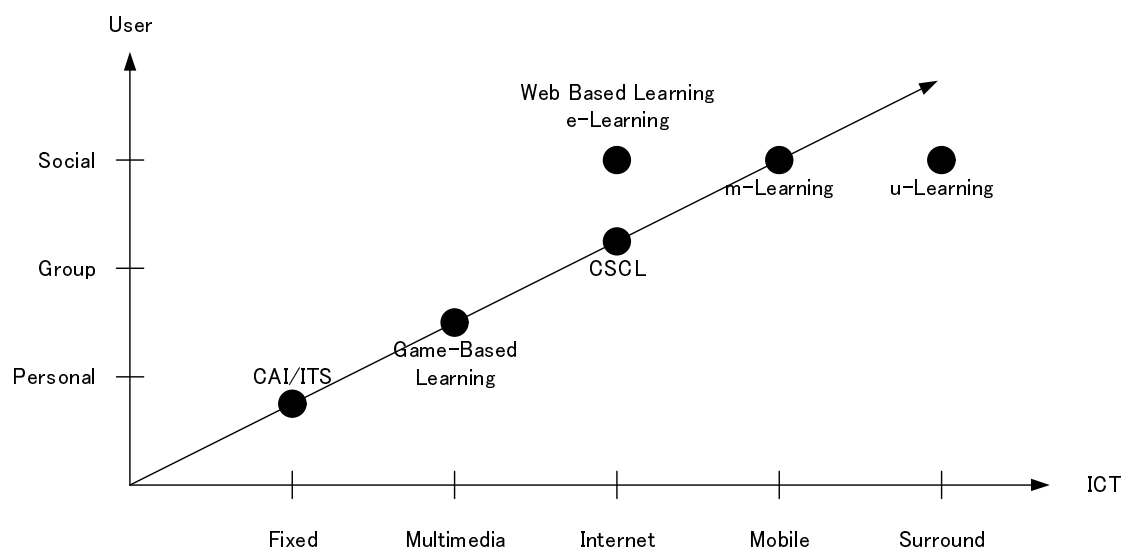


Figure 2.4: Research on Computer Education [11]

CAI (Computer Assisted Instruction) is a form of computer supported learning and teaching using fixed PC, which was introduced in 1960s. Later, researchers improve CAI by incorporating artificial intelligence to improve learning and problem solving and CAI evolved to intelligent CAI (ICAI) or ITS (Intelligent Tutoring System)[18]. GBL (Game-based learning) is a type of game play with a defined learning outcomes, which was introduced in 1980s along

with the emergence of multimedia technology. GBL incorporates the subject matter with the game play and facilitates the player (learner) retain and apply what they have learned to the real world. CSCL (Computer Supported Collaborative Learning) facilitates interaction and collaboration among group of learners connected each other, evoked by the invention of World Wide Web (WWW) in 1990.

E-learning facilitates learning on the web based services in the internet. M-learning facilitates learning in mobility, which is realized by mobile and wireless network and technologies. U-learning facilitates learning through learner's surroundings captured by mobile, wireless, and ubiquitous network and technologies such as sensors and RFIDs. E-, m-, and u-learning have similar potential in enhancing social interaction and cooperative learning among learners through network.

Moreover, u-learning enables learning from experience through interaction between cyber and physical spaces. Still in u-learning scope, Hwang *et al.*[5] introduced specific form of u-learning system named *context-aware u-learning system*, which incorporates real world context-awareness in its services. In his further work, Hwang promotes *smart learning environment* as enhancement of the context-aware learning system[19]. Various research and commercial applications of these learning paradigms are available in various learning domains.

2.3.2 Comparison of the x-Learning

To compare the movie-based context-aware learning (MBCAL) with e-learning, m-learning, and u-learning, the author surveyed existing work in the field of mobile and ubiquitous learning and adopted six characteristics including permanency, accessibility, immediacy, interactivity, context-awareness, and content of learning portfolio. The first five are the characteristics adopted by many researchers [5, 10, 15, 16, 20], while the last characteristic is solely mention by Hwang *et al.*[5].

However, those characteristics are too abstract to show more obvious comparison among

learning paradigms; thus, the author defines some properties with explanation for each characteristic and compare the x-learning paradigms in more comprehensive manner as presented in Table 2.3.

As shown in Table 2.3, the four learning paradigms can be accessed through wireless network, show context-awareness by accessing learning portfolio, and record on-line behaviours of the learner. Both e-learning and MBCAL can be accessed also via wired network, while m-learning and u-learning must be accessed via wireless network for its inherent characteristic. Both u-learning and MBCAL are able to keep learner's work permanently and provide dynamic adaptive services. In term of interactivity, u-learning shows effective interactivity with peers, teachers and experts. While both m-learning and MBCAL may have active interactivity, e-learning can have limited interactivity.

Furthermore, the most recognized advantage of *u*-learning is context-awareness by both learners physical context and real environment context. On the other hand, MBCAL shows context-awareness by utilizing virtual environment; *i.e.* movie scenes but with some technical limitations in performing real/physical context in u-learning. Lastly, u-learning records all aspects as learning portfolio; namely, on-line behaviours of the learners, real world behaviours, and correspondent environment behaviours. The others record learner's on-line behaviours only.

In short, MBCAL has closest functionalities to u-learning. Although it has six properties common to u-learning, MBCAL has one unique property; that is, context-awareness in virtual environment in a limited form from u-learning. Note that properties listed in Table 2.3 summarizes all possible properties of each learning paradigm may possess; yet, in practice, an application of each learning paradigm may not possess all those properties. Applications of a certain learning paradigm may have slightly different properties depending on their purposes.

Table 2.1: Zimmermann's Categorization of the Context[8]

Context Dimension	Description	Subdivision or Example
Individuality	Information comprises anything that can be observed about an entity, typically about the entity's state	<i>natural entity context</i> , <i>human entity context</i> , <i>artificial entity context</i> , and <i>group entity context</i>
Time	This category classifies time information like the time zone of the client, the current time or any virtual time	<i>categorical scales</i> (e.g.working hours or weekends), <i>process-oriented view of the time concept</i> (e.g.work flows), <i>intervals of time</i> , and <i>express recurring events</i> (e.g.always on Sundays)
Location Context	This category describes location models that classify physical or virtual (e.g. the IP address as a position within a computer network) residence of an entity, as well as other related spatial information like speed and orientation. In addition, a location may be described as an <i>absolute location</i> means the exact location of something; or as a <i>relative location</i> means the location of something relative to something else.	Physical locations can be split into <i>quantitative (geometric) location</i> (e.g. geographic coordinate); and <i>qualitative (symbolic) location</i> (e.g. buildings, rooms, streets, countries, etc.)
Activity Context	The activity context covers the activities the entity is currently and in future involved in. It answers the question <i>What does the entity want to achieve and how?</i>	Explicitly can be described as <i>goals</i> ; <i>tasks</i> ; and <i>actions</i> ;
Relation Context	A relation expresses a semantic dependency between two entities that emerges from certain circumstances these two entities are involved in. Each entity plays a specific role in a relation. Potentially, an entity can establish any number of different relations to the same entity. Relations are not necessarily static, it may emerge and disappear dynamically.	Divided into <i>social relation context</i> ; <i>functional relation context</i> ; and <i>compositional relation context</i> .

Table 2.2: Ubiquitous Language Learning Applications

Application	Description
Vocab Tutor (Stockwell, 2007)	A prototype system of mobile-based intelligent vocabulary learning
TANGO (Ogata and Yano, 2008)	A system for vocabulary learning using the physical objects with RFID tags
Chung (2008)	A personalized mobile English vocabulary learning system based on Item Response Theory and the learning memory cycle
HELLO (T. Y. Liu, 2009)	A Context-Aware Ubiquitous Learning Environment for Language Listening and Speaking
Li <i>et al.</i> (2010)	An adaptive kanji learning System using mobile phones
M-iLexicon (Underwood, 2010)	A mobile self-initiated vocabulary learning application
LORAMS (Ogata and Uosaki, 2012)	A system recording learners experience in video, which is automatically linked to the real object in the scene by scanning their RFID tags
SCROLL (Ogata and Uosaki, 2012)	A system providing an interface that supports the learners to record, share and reuse ubiquitous learning logs (ULLOs) with android mobile devices

Table 2.3: x-Learning Comparison

Dimension	Aspects	e-learning	m-learning	u-learning	MBCAL
Permanency	Learners can lose their work	X	-	-	-
	Learners may lose their work	-	X	-	-
	Learners can never lose their work	-	-	X	X
Authentic	Access via computer network	X	-	-	X
	Access via wireless network	X	X	X	X
	Communicate with ubiquitous technology	-	-	X	-
Adaptive Services	Learners cannot get services immediately	X	-	-	-
	Learners get information immediately in fixed environment	-	X	-	X
	Learners get information immediately in dynamic environment (active services)	-	-	X	-
Interactivity	Limited interaction	X	-	-	-
	Active interaction with peers, teacher and experts	-	X	-	X
	Effective interaction with peers, teacher and experts	-	-	X	-
Context-Awareness	CA by accessing learning portfolio	X	X	X	X
	CA by accessing learners personal context (physical context)	-	-	X	-
	CA by accessing real environment	-	-	X	-
	CA by utilizing virtual environment	-	-	-	X
Content of Learning Portfolio	Record online behaviors of the learners	X	X	X	X
	Record real world behavior	-	-	X	-
	Record corresponding environment behavior	-	-	X	-

Chapter3

Movie-Based Context-Aware Learning (MBCAL)

3.1 Movie for Learning

The movie has been practically used for learning so far. Movie-based learning enables learners to learn from virtual experience through the movie.

Kerber *et al.* encouraged learners to learn the medical theory from the movie about it[21]. Lumlertgul *et al.* conducted a cinemeducation project, which aimed to help learners learn medical professionalism using movies[22]. Jungraithmayr and Weder improved the process learning about techniques of orthotopic mouse lung transplantation by full-length movie visualization[23].

Furthermore, the movie is a powerful means that facilitates learning context dependent aspects of the language. Chan and Herrero use movies to teach various languages in the classroom in their project and stated that plentiful contexts in the movie can be utilized for language learning as same as the real environment, as it is known that the movie can facilitate comprehension activities perceived as real[6]. Ismaili showed effectiveness of using movies in the EFL (English as a Foreign Language) classroom[24].

It is demonstrated by these researchers that movies improve learner's understanding significantly and also attract learner's attention successfully. Additionally, in the field of language learning, the movie presents languages naturally as well as shows quite rich linguistic and

environmental contexts.

The above-mentioned things motivated the author to incorporate the movie into a learning platform and develop a system for movie-based context-aware language learning (MBCALL).

3.2 Movie-Based Context-Aware Learning (MBCAL)

The movie can be a place in where quasi context-aware ubiquitous learning is practiced. *Movie-based context-aware learning (MBCAL* for short) is a concept proposed as an alternative approach for context-aware ubiquitous learning to overcome difficulties on development and operation of the context-aware ubiquitous learning system in the real environment[4]. MBCAL utilizes movies to facilitate proactive learning from the context situated in movie scenes as ubiquitous learning systems try to do in the real environment. Moreover, it does not require any physical devices to be deployed in the real environment and any specific learning applications interacting with them, since every learning activity is performed in the movie player equipping learning functions. Instead, instructors have to prepare context description of the movie scenes which is used as inputs to the learning functions of the movie player.

In movie-based context-aware learning, learners learn from imaginary environment or through imaginary experience in the movie; that is, learning is performed differently depending on the virtual context situated in the movie. The movie contains extremely rich contexts which can be utilized for learning as same as the real environment, as it is known that the movie can facilitate comprehension activities that are perceived as real as well as show natural expressions and flows of speech to learners[6]. Moreover, since events in the movie are foreseeable, we can know contexts of the movie scenes in advance. We do not need any special devices to acquire the context and any effort to know unexpected circumstances in design of the context-aware learning system. These contribute to limit cost and effort required to develop and operate the system. Instead, instructors using the movie-based context-aware learning system are required

to annotate the movie scenes with their contexts.

Two key concepts underlying the movie-based context-aware learning are *object-oriented context modeling* and *movie-based context-aware quiz generation*. The object-oriented context model (or OOCM for short) is a framework for movie context description to help instructors compose suitable context description of movie scenes. Based on context description with the OOCM, the movie-based context-aware learning system generates the context-aware quizzes in a template-based manner. Figure 3.1 illustrates key concepts of MBCAL concepts.

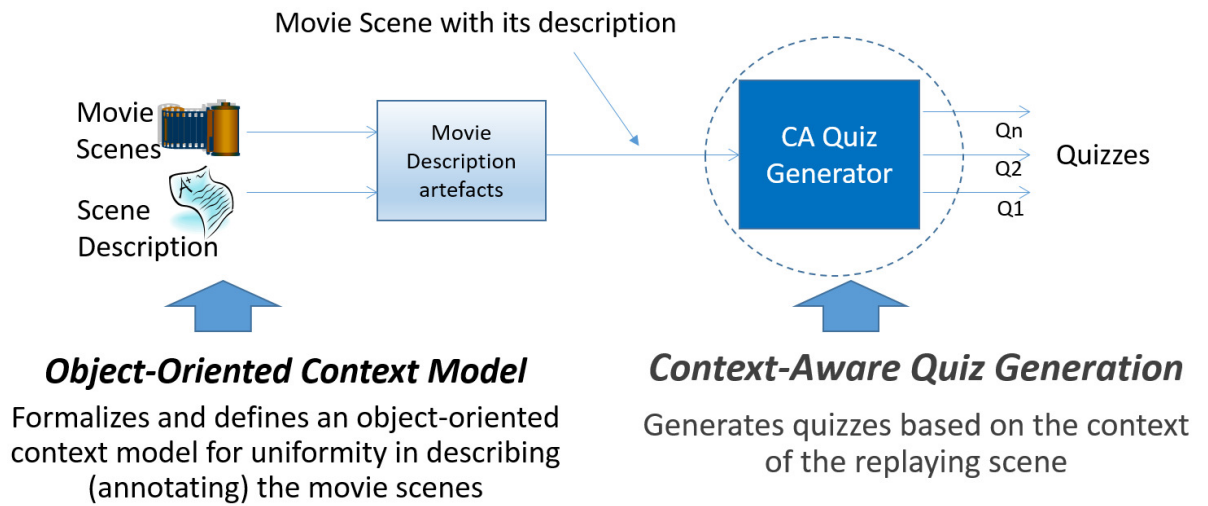


Figure 3.1: Key Concepts of MBCAL

3.3 Three Tiered Architecture of the Movie-Based Context-Aware Learning System

The author proposes a three tiered architecture to implement movie-based context-aware learning applications. The architecture consists of three layers: physical layer, context representation layer, and context-aware learning application layer as shown in Figure 3.2.

Each layer has specific responsibility.

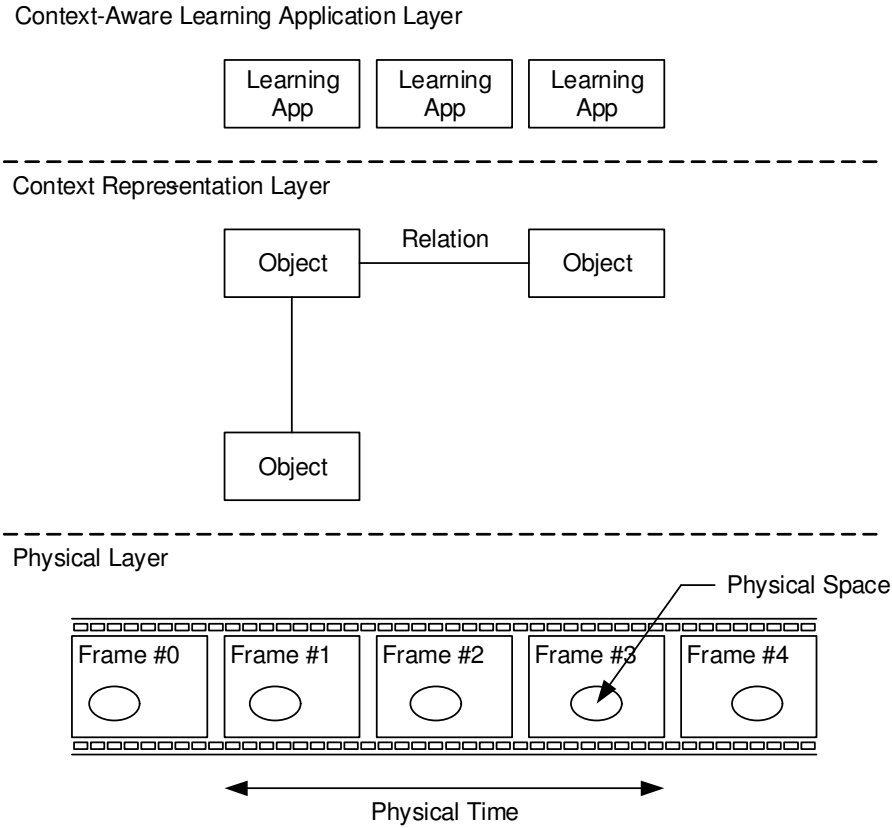


Figure 3.2: Three Tiered Architecture of the Movie-Based Context-Aware Learning System

The physical layer is responsible for replaying the movie and extracting context description annotating the movie as ordered by context-aware learning applications. Physical temporal and spacial information of the movie is handled in this layer.

The context representation layer is responsible for handling of context representation shared by context-aware learning applications. As the movie is replayed, the story goes and the movie context changes. Therefore, each item in the context description has physical life time and space in the movie. In other words, there exist temporal and spatial links from the context representation layer to the physical layer.

The context-aware learning application layer provides applications utilizing context-aware

representation and giving various learning tasks.

Chapter4

Object-Oriented Context Model

4.1 Context Description

In movie-based context-aware learning, learners learn along contexts in the movie. While physical context-aware learning requires sensing devices and other data sources to acquire the context of the physical environment, movie-based context-aware learning requires none of them. Instead, context description of the movie scenes must be prepared. Context description can be described manually by instructors and, in ultimate and probably limited cases, it can be generated automatically with image recognition. Note that context description can include not only visible or audible things but also something the viewer can know based on the story or implication of the story, since it can be described by human beings manually. Later, context-aware learning contents, *e.g.* quizzes and their expected answers, are generated from the context description annotating the movie scene replayed currently.

Context description needs a certain amount of effort for its preparation; therefore, reusability of the description is an essential property. To guarantee reusability of context description for different languages, context description should be uniform, universal, and natural language independent. In addition, considering state-of-the-art image recognition technology, full automatic context capturing from the movie is impossible. Therefore, to achieve the goal of

movie-based context-aware learning, the author defines an object-oriented context model (or OOCM for short) by introducing the object-oriented paradigm for context description. The author also defines a context description language (or CDL for short) which is subject to the OOCM to describe the movie context easily and instantly in a textual form by human work.

In object-oriented modeling of the context in the real environment, it is important how to find classes and associations among them that can describe the context in the real environment. Moreover, the object-oriented model must be in a form that can generate quizzes and their expected answers. The author adopted Verbert's context model[2] and defined the object-oriented model based on it as shown in Appendix A; however, there were some difficult cases in context description of movie scenes. The authors proposes an object-oriented context model incorporating the case grammar concept, which is a popular concept to represent semantic roles of words and phrases in natural language processing.

4.2 Case Grammar Concept

Fillmore proposed the case grammar concept[25], which got popular as semantic roles later. Fillmore introduced the basic structure of the sentence which is divided into a couple of constituents: proposition and modality as shown in Figure 4.1. The proposition consists of a tenseless set of relationships involving verbs and cases. The verb represents an event or activity, while cases are nouns having semantical relationships to the verb. Furthermore, Fillmore defines the following case lists, *i.e.* semantic role lists:

- Agentive: Who causes the event or activity?
- Instrument: What is used for the event or activity?
- Dative: Who experienced the event or activity?
- Factive: What is the result of the event or activity?

- Locative: Where the event or activity happens?
- Objective: What is the target of the event or activity?

The modality includes such modalities on the sentence as-a-whole as negation, tense, mood, and aspect. For example, “Alice slices the bread with the knife.” can be broken into *Alice* as agentive, *cut* as verb, *bread* as objective, and *knife* as instrument.

$$\begin{aligned} \text{Sentence} &\rightarrow \text{Modality} + \text{Proposition} \\ \text{Proposition} &\rightarrow \text{Verb} + \text{Case}_1 + \text{Case}_2 + \cdots + \text{Case}_n \end{aligned}$$

Figure 4.1: Fillmore’s Case Grammar First Base Rule

Many linguists have refined Fillmore’s case grammar concept and introduced more cases. The author surveyed some work and summarized them as shown in Table 4.1. Note that there are some cases which have the same or similar meaning but are named differently by different linguists; for example, Fillmore’s *agentive* is renamed as *agent* by Jurafsky and Martin[26] and as *cAgent* by Parunak[27]. The cases in the table have the following meanings:

- An *agent* means a perceived instigator of the action of the verb.
- An *instrument* means force or an inanimate object casually involved in the state or action of the verb.
- An *experiencer* means an entity whose mental or emotional state is affected by the action of the verb.
- A *factitive* means a thing produced by the action of the verb.
- A *location* identifies the place or spatial orientation of the state or action.

- An *object* means a thing affected by the action or state of the verb.
- A *goal* means a place to where something moves or thing toward which an action is directed.
- A *measure* means quantification of the event.
- A *path* describes the pathway of a motion.
- A *source* means the place of the origin, entity from which a physical sensation emanates, or original owner in a transfer.
- A *commutative* refers to somebody else who performs an action with the agent.
- A *beneficiary* means an entity that possesses an object or participates with an agent in transfer of an object.
- A *causer* means the referent which instigates an event rather than doing it actually.
- A *manner* means how the action, experiences, or process of an event is carried out.
- A *time* indexes the action of the verb in time.

Since there has been no standardized case grammars so far, it is unclear which and how many semantic roles (or cases) are needed. That is considered as one of the major problem of the case grammar theory[28]. However, the objective of this study is not to investigate the case grammar theory deeply. Note that this study surveys the case grammar theory to adopt its concept and find suitable cases for construction of the OOCM that can describe contexts of the movie scenes appropriately.

Table 4.1: Survey on the Case Grammar[29]

Case Lists	Fil.	Long.	Lar.	Par.	Berk	Jur.
Agentive (A) / Agent	X	X	X	X	X	X
Instrumental (I) / Instrument	X	X	X	X	X	X
Dative (D) / Experiencer (E)	X	X	-	X	X	X
Factitive (F) / Result / Range	X	X	X	-	-	-
Locative (L) / Location	X	X	-	X	X	X
Objective (O) / Object / Patient / Theme	X	X	X	X	X	X
Goal (G)	-	X	X	-	-	-
Measure	-	X	X	-	-	-
Path	-	X	-	-	-	-
Source	-	X	-	-	-	-
Accompaniment / Comitative	-	-	X	X	-	-
Beneficiary / Benefactor / Recipient / Possessor	-	-	X	X	X	X
Causer / Force	-	-	X	-	-	X
Manner	-	-	X	-	-	-
Time	-	-	X	X	-	-
Content	-	-	-	-	-	X

Fil.: Fillmore[25], Long.: Longacre[30], Lar.: Larson[30], Par.: Parunak[27], Berk: Berk[30],
Jur.: Jurafsky[26]

4.3 Introduction of the Case Grammar Concept to the OOCM

4.3.1 OOCM Class Elements

Since the cases represent what the human being can recognize and verbalize, the case grammar concept can be used to describe movie scenes in a semantically formalized manner. The author defines the OOCM integrating the case grammar concept for this purpose[29]. The classes are defined by selecting cases from the summary shown in Section 4.2 and mapping them on the OOCM as roles of the class. Figure 4.2 shows the OOCM in the class diagram.

The OOCM has eleven classes. Classes *Thing*, *Action*, and *Mode* are the most fundamental ones.

- Class *Thing* means a thing appearing in the movie.
- Class *Action* means an action performed in the movie.
- Class *Mode* means a mode of the thing or the action.

Thing has the following descendant classes:

- Class *Animate*, a sub class of *Thing*, has its own intention.
- Class *Inanimate*, a sub class of *Thing*, does not have any intention.
- Class *Object*, a sub class of *Inanimate*, is something passively involved in the action.
- Class *Space*, a sub class of *Inanimate*, represents a location.
- Class *Time*, a sub class of *Inanimate*, represents a time instant of a duration.

Action has the following sub classes:

- Class *Transitive Action* is an action that finishes immediately.
- Class *Continuous Action* is an action that continues for a certain duration.



4.3.2 OOCM Based on Verbert's Context Framework

As mentioned in Section 4.1, the author adopted Verbert's context framework initially to define the OOCM[2]. However, it was revealed that the OOCM based on Verbert's context framework had limitation in its descriptive power by the result of preliminary evaluation. The evaluation was conducted by writing context description of a short movie with the OOCM. Some of the context description are presented in Appendix A. Limitations of the OOCM were summarized into two aspects, uniqueness and comprehensiveness.

Uniqueness: There could be multiple context descriptions to describe a single movie scene in the OOCM based on Verbert's context framework. We could describe appearance of the object with a certain space, such as "*The sky looks clear.*", as *Object-seen-Appearance* and also as *Place-seen-Physical Condition*.

Comprehensiveness: There were some cases that we could not describe contexts of the movie scene. The *actor* was centric in the OOCM based on Verbert's context framework. Therefore, we felt inconvenience to explain the *activity* in the movie scene. For example, there were no ways to describe how and where the activity was done.

Considering the above-mentioned observation, the author adopted the case grammar concept, instead of Verbert's context framework, for definition of the OOCM. The OOCM based on the case grammar concept is easier to explain activities since the case grammar concept is verb-centric than the OOCM based on Verbert's context framework.

4.4 Annotating Movies by Context Description

In MBCAL, the movie used for context-aware learning should be annotated by the OOCM in advance. The author surveyed several existing movie annotation tools including ANVIL (Annotation of Video and Language)[31], AAV (Annotating Academic Video Tool)[32], and VATIC (video annotation tool for computer vision research)[33], and adopted ANVIL as the

most suitable annotation tools for MBCAL.

ANVIL enables multimodal annotation; that is, we can annotate the movie by various types of information. ANVIL can annotate the movie by auditory information including words, prosody, syntax, dialogue acts, utterance, rhetorics *etc.* and visual information such as gesture, posture, graphics *etc.* Figure 4.3 shows a snapshot of ANVIL.

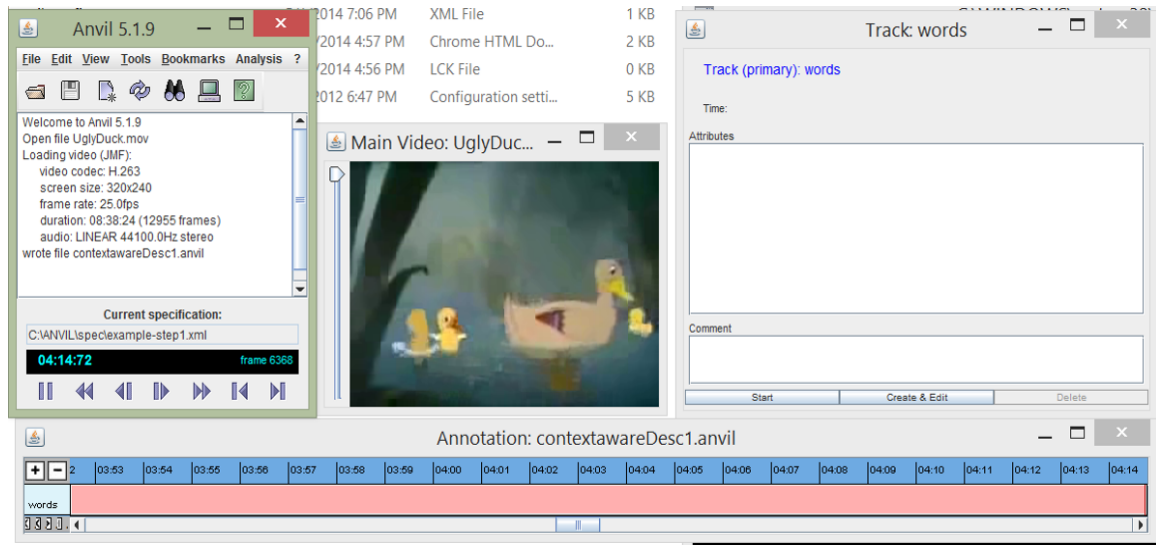


Figure 4.3: Example of Annotation View Using ANVIL

4.5 Context Description Language (CDL)

The movie context can be represented by the object diagram instantiated from the OOCM. Since it is time-consuming to edit the object diagram, we define the context description language (or CDL for short) to compose the movie context easily and quickly in a textual form. Figure 4.4 shows the syntax of the CDL in the extended BNF in accordance with the OOCM.

The start symbol of the grammar is `ContextScript`. Symbol `ClassName` is the name of a class defined in the OOCM. Symbol `AssocName` is the name of an association between classes

defined in the OOCM. `ClassName` and `AssocName` must be prefixed by the backslash (\). `Symbol LanguageCode` is a language code specified by ISO 639-1[34] and 639-3[35] such as `en` for English, `fr` for French, `zh` for Chinese, *etc.* `Symbol DQ` represents a double quotation mark (").

```
ContextScript  = (InstanceDesc | ContextDesc)* .
InstanceDesc  = ClassName "{" InstanceName "}" "{" WordDescList "}" .
WordDescList  = WordDesc | WordDescList "," WordDesc .
WordDesc      = LanguageCode "=" DQ Word DQ .
ContextDesc   = "\context" "{" AssocDesc* "}" .
AssocDesc     = AssocName "{" SourceInstance "}" "{" SinkInstance "}" .
SourceInstance = InstanceName .
SinkInstance  = InstanceName .
InstanceName  = [A-Za-z0-9]+ .
```

Figure 4.4: Syntax of the CDL

For example, the movie scene shown in Figure 4.5 can be modeled by the OOCM as shown in Figure 4.6 and described in CDL as shown in Figure 4.7.

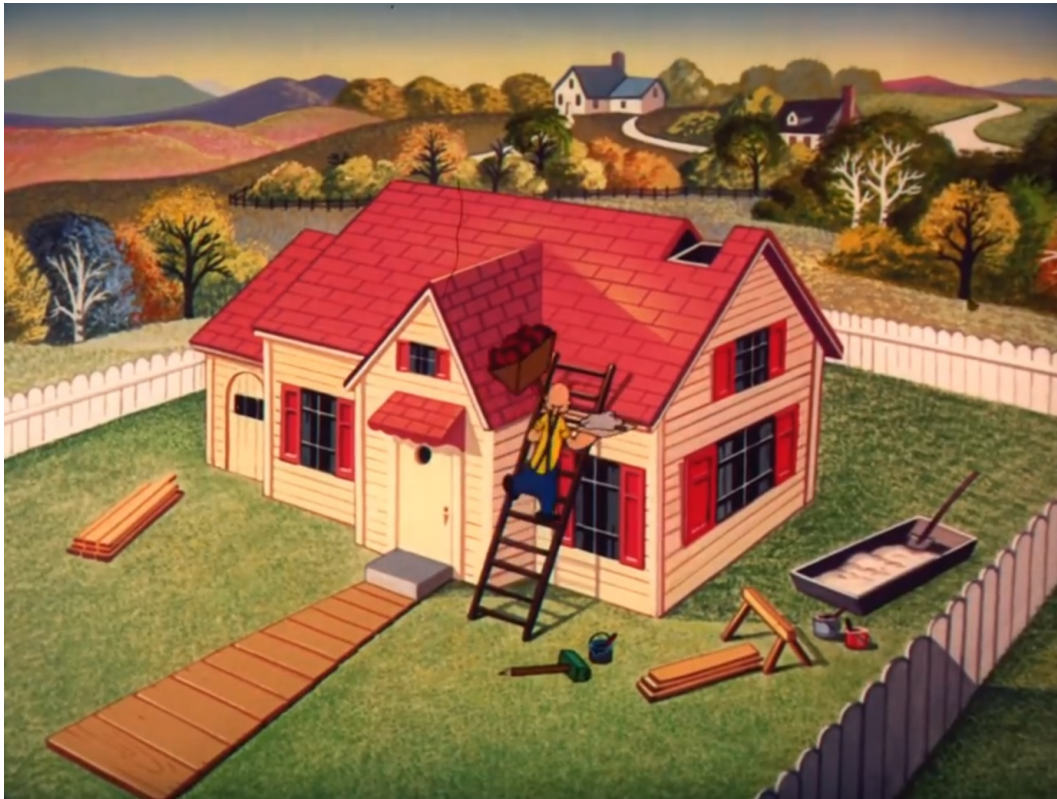


Figure 4.5: Example Movie Scene

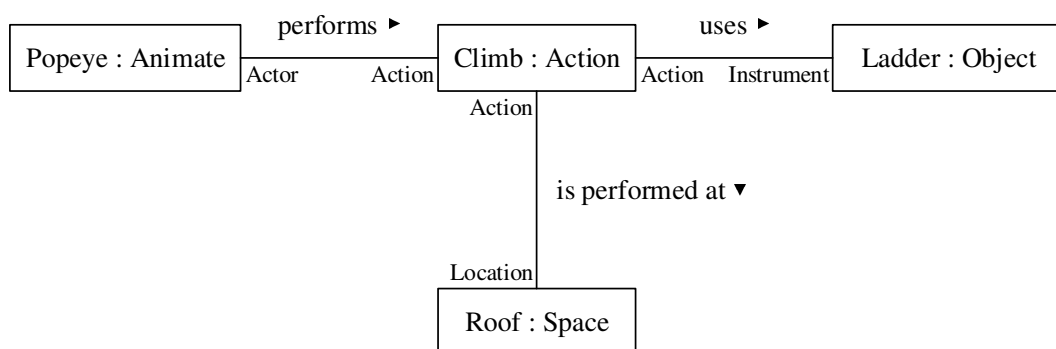


Figure 4.6: An Example OOCM

```
\context{
  \Actor{Popeye}{en="Popeye", id="Popeye"}
  \Action{Climb}{en="climb", id="menaiki"}
  \Space{Roof}{en="roof", id="atap"}
  \Object{Ladder}{en="ladder", id="tangga"}
  \performs{Popeye}{Climb}
  \isPerformedToward{Climb}{Roof}
  \uses{Climb}{Ladder}
}
```

Figure 4.7: Example CDL

Chapter5

Movie-Based Context-Aware Quiz Generation

5.1 Movie-Based Context-Aware Quiz Generation

As an implementation of the movie-based context-aware learning (MBCAL), the author developed a prototype system performing movie-based context-aware quiz generation. The prototype is an application of movie-based context-aware language learning (MBCALL) which was mentioned before as one of the promising learning domains for MBCAL.

In MBCALL, the learner watches the movie and the movie player with learning functions generates quizzes automatically based on both *movie's context* and *learner's context*.

The movie's context means the current scene of the replayed movie and any information derived from it, which is associated with the scene's time stamp. It is described and embedded in the movie by the instructor in advance using language independent context description. The movie context is likened to the environmental context in context-aware learning performed in the real environment.

On the other hand, the learner's context consists of information relating to the learner such as his/her learning level, completed learning tasks, watched/unwatched movies, and preferred movies. It is entered initially by the learner him/herself and updated by the system along with his/her learning activity. The MBCAL system, namely the movie player with learning func-

tions, generates quizzes and their expected answers from the context description with replaying the movie. The system can generate different quizzes to different learners with different contexts depending on these contexts even for the same movie.

Note that quiz generation is template-based. The quiz sentence and its expected answer are generated from the template prepared in advance. The generator parses the OOCM which is active in the movie, references properties of the objects in the OOCM, applies them to the template, and generates quizzes and their expected answers.

5.2 Commonality-Variability Analysis of Quizzes

To construct templates for quiz generation, the author surveyed some language learning textbooks[36, 37, 38] and analyzed commonality and variability of the quizzes appeared in them to construct templates for quiz generation. The result of the analysis is organized in a feature model[39] as shown in Figure 5.1.

This feature model shows variability of the quiz objectives, question types, and ways to answer. The author identified seven quiz objectives which represent desired skills assessed by the quiz including vocabulary, grammar, translation, composition, reading, speaking, and listening skills. The question types can be classified into five general forms: asking facts, filling blanks, combining sentences, translating, and composing. Ways to answer the quiz are also specified in association with the question types. There are five general ways to answer the quiz: composing sentence, filling the blanks, multiple choices, true-false, and translating. In addition, answering language has alternatives: learning language and instructing language. This feature model reveals possible types of quiz application to be developed.

The quiz template is designed for each variation represented by this feature model.

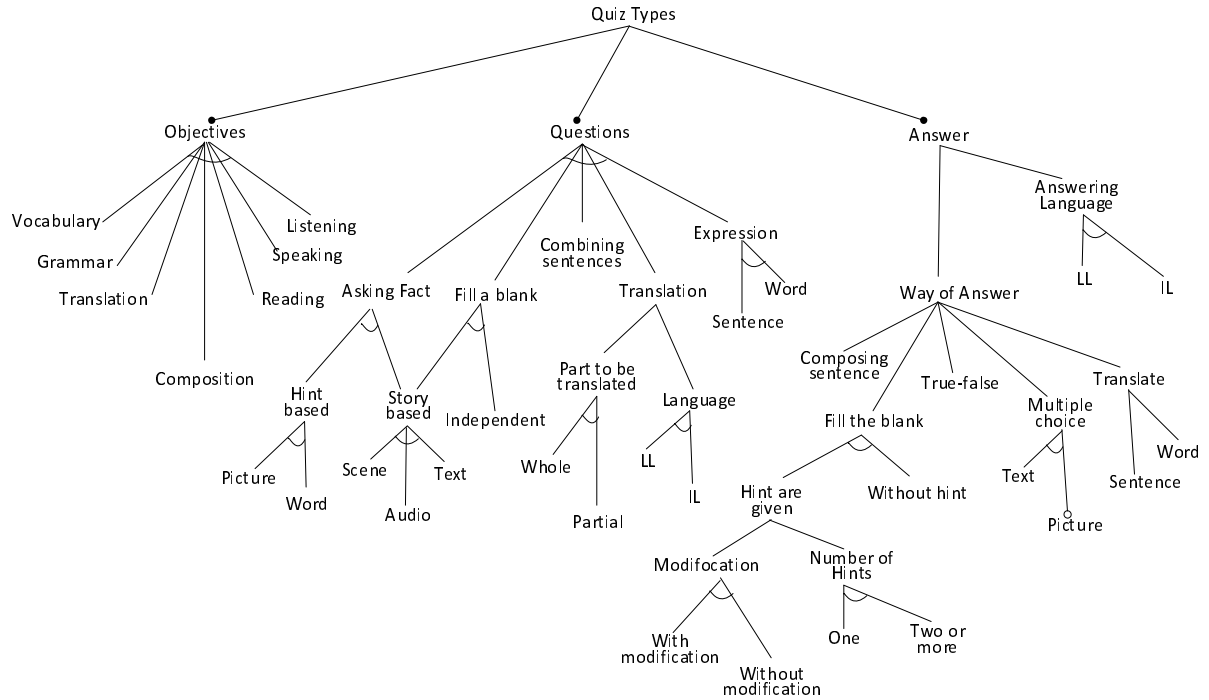


Figure 5.1: Feature Model of the Quiz Types

5.3 Components of the System

Components of the MBCALL application are defined based on the MBCAL architecture presented in Section 3.3. The MBCALL application consists of four principle components: contents manager, context manager, quiz generator, and movie & quiz player. Interaction among the components are illustrated in Figure 5.2.

- The *contents manager* gets the learning contents consisting of the movie and its context description and provides them to the context manager and the movie and quiz player, respectively.
- The *context manager* maintains the context model based on the context description transferred from the contents manager. The context description is traceable from the physical

time of the movie. The context manager constructs or destructs objects and relations among objects of the context model when they get active or inactive.

- The *quiz generator* generates quizzes and their expected answers based on the current context model which the context manager maintains.
- The *movie & quiz player* replays the movie transferred from the contents manger and shows the quizzes generated by and transferred from the quiz generator. The player prompts the leaner to answer the quiz with replaying the movie and sends the answer back to the quiz generator. Then the quiz generator checks if the answer is the expected one.

Context-Aware Learning Application Layer

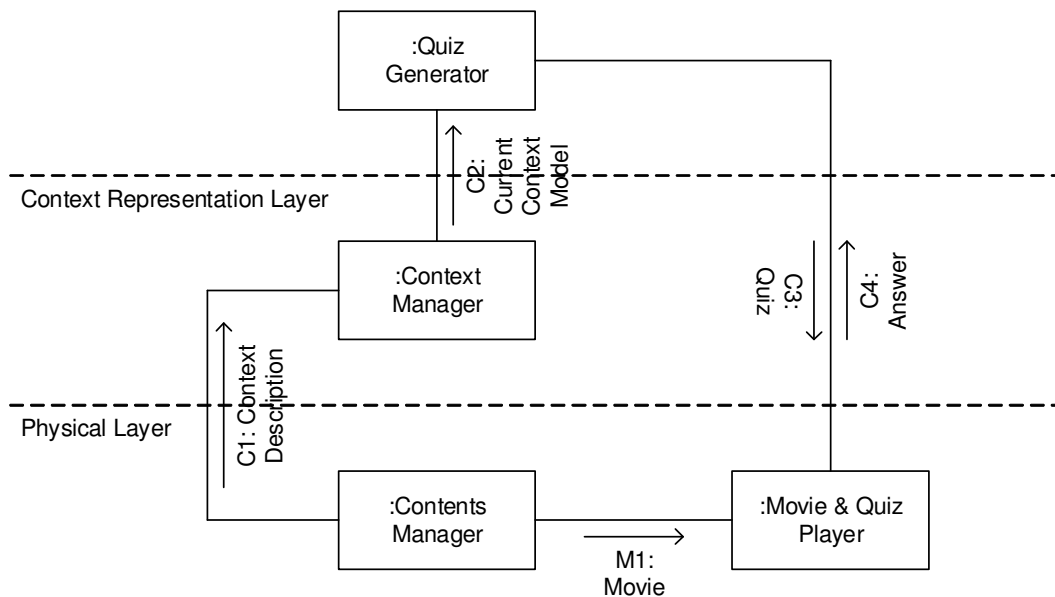


Figure 5.2: Components of an MBCALL Application and Their Interaction



Figure 5.3: An Example Movie Scene

5.4 Quiz Generation Template

Automatic quiz generation from the OOCM is the keystone of the movie-based context-aware language learning system. It should be allowed to construct multiple quiz generators which generates different quizzes for different learning objectives and run them concurrently in the learning system. The quizzes are generated by a template based manner. The quiz sentence and its expected answer are generated from prepared templates. The generator parses the current context model, references properties of the objects in the context model, applies them to the template, and generates quizzes and their expected answers.

For the movie scene shown in Figure 5.3, we can describe its context description as shown in Figure 5.4. Figures 5.5 and 5.6 show example templates for English and Bahasa learning, respectively. These templates generate quizzes asking the learner what he/she can see at the specified place in English and Bahasa.

The template is described in XML. The element `FillTheBlank` represents one pattern

```

\context{
  \Animate{Duck}{en="duck", id="bebek"}
  \Space{Pond}{en="pond", id="kolam"}
  \ContinuousAction{Play}{en="play", id="bermain"}
  \performs{Duck}{Play}
  \isPerformedAt{Play}{Pond}
}

```

Figure 5.4: Context Description in CDL

```

<FillTheBlank>
  <Question>What is in {en_prepend_definite_article(<Space.en>)}?</QuestionSentence>
  <ExpectedAnswer>{en_capitalize(en_prepend_indefinite_article(<Animate.en>))} is in {BLANK(en_prepend_definite_article(<Space.en>))}.</ExpectedAnswer>
</FillTheBlank>

```

Figure 5.5: Example Quiz Template for English Learning

of the fill-a-blank type quiz generation and includes necessary information for quiz generation. The element `Question` contains a template sentence and the element `ExpectedAnswer` contains the expected answer to the question.

Basically, the contents of these elements are used for the quiz sentence and expected answer as they are. However, generation or transformation of texts is performed for the part surrounded by braces `{}`. An term surrounded by angle brackets `<>` in the part surrounded by braces represents a reference to the object of the OOCM. The term should be written in the form of `<Class.Property>`. The quiz generator replaces it with the property value of the object specified by the class to show the the quiz sentence to the learner and check learner's answer. For example, `<Animate.id>` in Figure 5.6 is replaced with the value of `id` property of an `Animate` object, namely "bebek." Pre-defined transformation to the referenced property value is also allowed. A term `tra(a)` in the part surrounded by braces

```

<FillTheBlank>
  <Question>Apa yang ada di {<Space.id>}?</Question>
  <ExpectedAnswer>{BLANK(id_capitalize(<Animate.id>))} ada di
  {<Space.id>}.</ExpectedAnswer>
</FillTheBlank>

```

Figure 5.6: Example Quiz Template for Bahasa Learning

represents that a pre-defined transformation `tra` is applied to the text `a` passed as its parameter. The parameter of the pre-defined transformation can be multiple. For example, `en_prepend_indefinite_article(<Animate.en>)` in Figure 5.5 means that a pre-defined transformation `en_prepend_indefinite_article(enword)` to prepend an appropriate English indefinite article to an English word `enword` is applied to the value of `en` property of an `Animate` object; that is, “duck” is transformed to “a duck.” Transformation `BLANK(a)` in the `ExpectedAnswer` element included in the `FillTheBlank` element makes the word a blank to be filled by learners. We can get quizzes and expected answers for English and Bahasa learning shown in Figures 5.7 and 5.8 if we apply the templates shown in Figures 5.5 and 5.6 to the context description shown in Figure 5.4.

Question:
 What is in the pond?
 Expected Answer:
 [A duck] is in the pond.

Figure 5.7: Generated Quiz and Expected Answer for English Learning

Question:
 Apa yang berada di kolam?
 Expected Answer:
 [Bebek] ada di kolam.

Figure 5.8: Generated Quiz and Expected Answer for Bahasa Learning

5.5 Language Dependent Transformation

In template-based quiz generation, we can reference properties of the OOCM object. An important role of language dependent pre-defined transformation is to absorb difference among languages to keep the context description language independent. Concerning that each language has specific grammatical rules, the author prototyped Python based libraries of language dependent pre-defined transformation to perform various grammatical modification. The pre-defined transformation defined in the library can be invoked from the template in quiz generation.

Figures 5.13, 5.9, 5.10, 5.11, and 5.12 are example code fragments of pre-defined transformation for English: transformation to generate the be verb suitable for the person of the actor, to modify the present verb suitable for the person of the actor, to modify the past verb suitable for the person of the actor, to modify the past participle suitable for the person of the actor, and to modify the present auxiliary verb suitable for the person of the actor.

```
def en_tobe (actor):  
    if actor == "I":  
        return "am"  
    elif actor == "you" or actor == "we":  
        return "are"  
    else:  
        return "is"
```

Figure 5.9: Be-Verb Transformation

5.6 Prototype of the Context-Aware Quiz Generator

The author implemented a prototype of the context-aware quiz generator. The prototype is implemented by using Python. Figure 5.14 shows quiz generation and validation of the answer by the prototype in case we use the context description shown in Figure 5.15.

```

def en_modified_present_verb (verb, actor):
    # The verb does not change if the subject is "I", "we", or "you"
    if actor == "I" or actor == "we" or actor == "you":
        return verb
    # Otherwise, it changes depending on the ending sound of the verb
    # Regular change
    if re.match (r'.*[^aeiou]y$', verb):
        return verb[:-1] + "ies"
    for pattern in [r'.*ch$', r'.*s$', r'.*sh$', r'.*x$', r'.*z$', r'.*o$']:
        if re.match (pattern, verb):
            return verb + "es"
    return verb + "s"

```

Figure 5.10: Present Verb Transformation

```

def en_modified_past_verb (verb, actor):
    past = irregular_past_verb_dict.get (verb)
    if past != None:
        return past
    if re.match (r'.*e$', verb):
        return verb + "d"
    if re.match (r'.*[^aeiou]y$', verb):
        return verb[:-1] + "ied"
    return verb + "ed"

```

Figure 5.11: Past Verb Transformation

5.7 Comparison with Existing Quiz Generation Systems

The author surveyed work on quiz generation systems and compared them to the MBCALL system.

Nanjiani *et al.* proposed an automatic quiz generation system of a live broadcast or recorded events[40]. In this system, the instructor writes questions based on the prepared information. Later, the questions will be delivered to the learner and the system will assess their answer automatically.

```
def en_modified_past_participle (verb):
    past_participle = irregular_past_participle_dict.get (verb)
    if past_participle != None:
        return past_participle
    if re.match (r'.*e$', verb):
        return verb + "d"
    if re.match (r'.*[^aeiou]y$', verb):
        return verb[:-1] + "ied"
    return verb + "ed"
```

Figure 5.12: Past Participle Verb Transformation

```
third_party_present_auxiliary_verb = {
    "do" : "does",
    "have to" : "has to"
}

def en_modified_present_auxiliary_verb (auxiliary_verb, actor):
    if actor == "I" or actor == "you" or actor == "we":
        return auxiliary_verb
    third_party_form = third_party_present_auxiliary_verb.get (auxiliary_verb)
    if third_party_form != None:
        return third_party_form
    return auxiliary_verb
```

Figure 5.13: Auxiliary Verb Transformation

Liu and Lin created *Sherlock*, a semi-automatic quiz generation system using linked data[41]. In this system, the quiz and its answer are defined by the instructor based on suggestion from the system. The system suggests question as well as distractor to the instructor. Selected quizzes and distractors then stored into quiz database for later quiz generation. This system assesses learners answer automatically.

The advantage of our approach compare to these two quiz generation systems is that the questions are composed automatically by the system without human intervention. The quizzes and expected answers are sourced from context description of the movie scenes to the provided

quiz generation templates. This system also display quizzes to the learner, assess learners answer, and give feedback automatically.

Movie-Based Context-Aware Quiz Generation

What does the dolphine do?

The dolphine [].

Your answer: swims

Correct!

Where does the dorphine swim?

The dorphine swims in the [].

Your answer: river

Correct!

What does the cat do?

The cat [].

Your answer: runs

Correct!

Where does the cat run from?

The cat runs from the [].

Your answer: house

Incorrect!

The correct answer is: park

How does Popeye look?

He looks [].

Your answer: in panic

Correct!

What does Popeye feel?

Your answer: sad

Incorrect!

The correct answer is: satisfied

of Questions: 6

of Correct Answers: 4

of Incorrect Answers: 2

Total Score: 65

Figure 5.14: Quiz Generation and Validation of the Answers by a Prototyped Context-Aware Quiz Generator


```
\context{
  \Animate{Dorphine}{en="dorphine", id="lumba-lumba"}
  \Action{Swim}{en="swim", id="berenang"}
  \Space{River}{en="river", id="sungai"}
  \performs{Dorphine}{Swim}
  \isPerformedToward{Swim}{River}
}
\context{
  \Animate{Cat}{en="cat", id="kucing"}
  \Action{Run}{en="run", id="berlari"}
  \Space{Park}{en="park", id="taman"}
  \performs{Cat}{Run}
  \isPerformedFrom{Run}{Park}
}
\context{
  \Animate{Popeye}{en="popeye", id="popeye"}
  \Action{Eat}{en="eat", id="makan"}
  \Thing{Sushi}{en="sushi", id="sushi"}
  \Space{Restaurant}{en="restaurant", id="restoran"}
  \performs{Popeye}{Eat}
  \affects{Eat}{Sushi}
  \isPerformedAt{Eat}{Restaurant}
}
\context{
  \Animate{Popeye}{en="popeye", id="popeye"}
  \State{Shocked}{en="shocked", id="kaget"}
  \looks{Popeye}{Shocked}
}
\context{
  \Animate{Popeye}{en="popeye", id="popeye"}
  \State{Satisfied}{en="satisfied", id="lega"}
  \looks{Popeye}{Satisfied}
}
```

Figure 5.15: Context Description

Chapter6

Evaluation and Discussion

6.1 Evaluation of the OOCM

In order to evaluate the descriptive power of the OOCM, the author conducted an experiment by asking three volunteer subjects to describe the context of a public domain movie from *Popeye the Sailor Series* which is entitled as *Insect to Injury*[42]. The length of the movie clip is six minutes seven seconds. The process of the experiment consists of two steps. First, the author asked the subjects to describe the movie scenes in the free text without any more direction. Then, the author gave a short lecture on the OOCM and asked them to do the same thing again but with referencing the OOCM. All the subjects are Indonesian and they describe the movie scenes either in English or in Bahasa, which is the common language in Republic of Indonesia.

The descriptions written in the first step by the subjects were transformed into the OOCM for comparison. Movie contexts included in a single sentence can vary significantly depending on forms of writing. A complex sentence with more modifiers can have more meanings than a simple sentence. It enables fair comparison in a normalized form to transform a sentence into the OOCM and compare them based on the associations included in the OOCM. In comparison, the author checked if the same type of association exists in the OOCMs describing the

same movie scene by the subjects with ignoring trivial difference of words. Rather, we focused on difference of meanings.

In this evaluation, forty-two (42) equivalent sentences from different subjects were collected. The forty-two equivalent sentences derives eighty-five (85) OOCM associations. The author checked if these OOCM associations appear in each subject's description one by one. The sentences and associations are listed in the the enumeration below. The item in the enumeration shows an equivalent sentence in the first line and the OOCM associations derived from the sentence in the following line(s).

- (1) *The surroundings of the house looks shady and cool.*

Object(*the surroundings of the house*) **looks** Appearance(*shady*)

Object(*the surroundings of the house*) **looks** Appearance(*cool*)

- (2) *Popeye climbs to the roof by a ladder.*

Actor(*Popeye*) **performs** Action(*climb*)

Action(*climb*) **isPerformedToward** Goal(*the roof*)

Action(*climb*) **uses** Instrument(*a ladder*)

- (3) *Popeye lays bricks very quickly.*

Actor(*Popeye*) **performs** Action(*lay*)

Action(*lay*) **affects** Object(*bricks*)

Action(*lay*) **byTheWayOf** Manner(*quickly*)

- (4) *Popeye feels satisfied.*

Animate(*Popeye*) **feels** Feeling(*satisfied*)

- (5) *Popeye slides down the stair.*

Actor(*Popeye*) **performs** Action(*slide down*)

Action(*slide down*) **isPerformedAt** Location(*the stair*)

(6) *Popeye puts in the mailbox in front of the fence.*

Actor(*Popeye*) **performs** Action(*put in*)

Action(*put in*) **affects** Object(*the mailbox*)

Action(*put in*) **isPerformedAt** Location(*in front of the fence*)

(7) *Popeye wears a yellow T-Shirt.*

Actor(*Popeye*) **performs** Action(*wear*)

Action(*wear*) **affects** Object(*a yellow T-Shirt*)

(8) *Popeye writes his name on the mailbox.*

Actor(*Popeye*) **performs** Action(*write*)

Action(*write*) **affects** Object(*his name*)

Action(*write*) **isPerformedAt** Location(*on the mailbox*)

(9) *The insects eat up the mailbox.*

Actor(*the insects*) **performs** Action(*eat up*)

Action(*eat up*) **affects** Object(*the mailbox*)

(10) *Popeye closes the gate.*

Actor(*Popeye*) **performs** Action(*close*)

Action(*close*) **affects** Object(*the gate*)

(11) *Popeye is in a panic.*

Animate(*Popeye*) **looks** Appearance(*in a panic*)

(12) *Popeye blocks the insects.*

Actor(*Popeye*) **performs** Action(*block*)

Action(*block*) **affects** Object(*the insects*)

(13) *Popeye catches the insects into the trash can.*

Actor(*Popeye*) **performs** Action(*catch*)

Action(*catch*) **affects** Object(*the insects*)

Action(*catch*) **isPerformedToward** Goal(*the trash can*)

(14) *The insects escape from the trash can.*

Actor(*the insects*) **performs** Action(*escape*)

Action(*escape*) **isPerformedFrom** Source(*the trash can*)

(15) *Popeye repairs the fence by the hammer.*

Actor(*Popeye*) **performs** Action(*repair*)

Action(*repair*) **affects** Object(*the fence*)

Action(*repair*) **uses** Instrument(*the hammer*)

(16) *Popeye looks stressful.*

Animate(*Popeye*) **looks** Appearance(*stressful*)

(17) *Popeye casts out the insects with wheels.*

Actor(*Popeye*) **performs** Action(*cast out*)

Action(*cast out*) **affects** Object(*the insects*)

Action(*cast out*) **uses** Instrument(*wheels*)

(18) *The insects run towards the river.*

Actor(*the insects*) **performs** Action(*run*)

Action(*run*) **isPerformedToward** Goal(*the river*)

(19) *The insects ruin the bridge.*

Causer(*the insects*) **causes** Action(*ruin*)

Action(*Action*) **affects** Object(*the bridge*)

(20) *Popeye falls into the river.*

Experiencer(*Popeye*) **experiences** Action(*fall*)

Action(*fall*) **isPerformedToward** Goal(*the river*)

(21) *The insects rush for the house.*

Actor(*the insects*) **performs** Action(*rush*)

Action(*rush*) **isPerformedToward** Goal(*the house*)

(22) *Popeye digs a trench by the shovel.*

Actor(*Popeye*) **performs** Action(*dig*)

Action(*dig*) **affects** Object(*a trench*)

Action(*dig*) **uses** Instrument(*the shovel*)

(23) *Popeye fills the trench with water.*

Actor(*Popeye*) **performs** Action(*fill*)

Action(*fill*) **affects** Object(*the trench*)

Action(*fill*) **uses** Instrument(*the water*)

(24) *The insects look angry.*

Animate(*the insects*) **looks** Appearance(*angry*)

(25) *Popeye feels secure.*

Animate(*Popeye*) **feels** Feeling(*secure*)

(26) *Popeye enters to the house.*

Actor(*Popeye*) **performs** Action(*enter*)

Action(*enter*) **isPerformedToward** Goal(*the house*)

(27) *Popeye locks the door.*

Actor(*Popeye*) **performs** Action(*lock*)

Action(*lock*) **affects** Object(*the door*)

(28) *Popeye reads a book.*

Actor(*Popeye*) **performs** Action(*read*)

Action(*read*) **affects** Object(*a book*)

- (29) *The insects cross the trench with a can.*
Actor(*the insects*) **performs** Action(*cross*)
Action(*cross*) **isPerformedAt** Location(*the trench*)
Action(*cross*) **uses** Instrument(*a can*)
- (30) *The insects enter to popeye's house.*
Actor(*the insects*) **performs** Action(*enter*)
Action(*enter*) **isPerformedToward** Goal(*Popeye's house*)
- (31) *Popeye looks shocked.*
Animate(*Popeye*) **looks** Appearance(*shocked*)
- (32) *The insects destroy the house.*
Causer(*the insects*) **causes** Action(*destroy*)
Action(*destroy*) **affects** Object(*the house*)
- (33) *Popeye falls into the bathtub.*
Actor(*Popeye*) **performs** Action(*fall*)
Action(*fall*) **isPerformedToward** Goal(*the bathtub*)
- (34) *The insects look satiate.*
Animate(*the insects*) **looks** Appearance(*satiate*)
- (35) *Popeye eats spinach.*
Actor(*Popeye*) **performs** Action(*eat*)
Action(*eat*) **affects** Object(*spinach*)
- (36) *Popeye looks strong.*
Animate(*Popeye*) **looks** Appearance(*strong*)

- (37) *Popeye rebuilds his house with the metal.*
 Actor(*Popeye*) **performs** Action(*rebuild*)
 Action(*rebuild*) **affects** Object(*his house*)
 Action(*rebuild*) **uses** Instrument(*the metal*)
- (38) *Popeye rebuilds the house in a moment.*
 Actor(*Popeye*) **performs** Action(*rebuild*)
 Action(*rebuild*) **affects** Object(*his house*)
 Action(*rebuild*) **byTheWayOf** Manner(*in a moment*)
- (39) *The insects look tired.*
 Animate(*the insects*) **looks** Appearance(*tired*)
- (40) *Popeye triumphs over the insects.*
 Actor(*Popeye*) **performs** Action(*triumph over*)
 Action(*triumph over*) **affects** Object(*the insects*)
- (41) *The insects eat a cigar.*
 Actor(*the insects*) **performs** Action(*Action*)
 Action(*eat*) **affects** Object(*a cigar*)
- (42) *Popeye feels relieved.*
 Animate(*insects*) **feels** Feeling(*relieved*)

Table 6.1 shows the results. The subjects could identify more associations of the OOCM describing the movie scenes with referencing the OOCM than without referencing the OOCM. That means the OOCM can guide the composer to enrich information included in the movie context; therefore, we can generate more types of quizzes based on the movie context.

Furthermore, the author also measured time needed by the subjects to describe the contexts of the movie scenes subject to the OOCM. The times needed by the volunteers are presented

Table 6.1: Comparison Results

Approaches	Subjects		
	#1	#2	#3
Without Referencing OOCM	46	47	18
With Referencing OOCM	63	55	53

Table 6.2: Time for Composing Movie Description

Subjects	Description Elements	Time (minutes)
#1	63	48
#2	55	45
#3	53	55

in Table 6.2.

Based on observation during the movie description process, the time needed for context description was affected by the writing speed of the subject as well as his/her comprehension level on the OOCM. The volunteers were not experienced in context description of the movie scenes, since the OOCM was introduced to them just a few moment before the experiment. It is expected that we can reduce time to compose context description of movie scenes and also to create CDL significantly, if experienced descriptors do them with suitable tools.

6.2 Evaluation of Learning

6.2.1 Overview of the Evaluation Process

Evaluation was conducted under cooperation of twelve volunteers (hereafter referred as *learners*). The evaluation focused on learner's learning outcomes achieved by MBCALL as well as learner's subjective ratings to MBCALL. Learning outcomes were evaluated by comparing test scores of the learner before and after training in MBCALL. On the other hand, learner's subjective ratings were obtained from their feedback based on Chiu's evaluation scheme[20]. Chiu presented three dimensions for evaluation of meaningful ubiquitous learn-

ing: *activeness*, *authenticness*, and *constructiveness*. The author refined the dimensions into five criteria for MBCAL as presented in Table 6.3.

Table 6.3: Criteria for Learner's Subjective Ratings

Dimension	Criteria
Activeness	Did the learner feel involved?
	Did the learner make self-learning interactively?
Authenticness	Were the quizzes appropriate for the replaying movie scene?
	Was the learner engaged in various realistic events?
Constructiveness	Could the learner increase his/her knowledge?

To demonstrate that MBCALL presented in this study can work independently from languages to be learned, in evaluation, the author adopted two different languages to be learned: Bahasa, which is a common language in Republic of Indonesia, and Japanese. Bahasa was assigned to two learners whose native languages are not Bahasa, while Japanese was assigned to ten learners whose native languages are not Japanese.

6.2.2 Quiz Preparation

In evaluation, the author used a short copyright-free movie entitled as *Insect to Injury*. The length of the movie was 367 seconds. To examine language independency of MBCALL, or to prove MBCALL is applicable to learn different languages in other words, the author prepared context description and quiz templates to learn Bahasa and Japanese. The author described twenty-eight context descriptions of the movie scenes in various associations in the OOCM as listed in Appendix B. The amount of context description represents potential of various contextual language learning. Context description of the movie scenes can be used as learning contents that we can process to provide various types of learning tasks and modify for different learners with different learning levels. The modification of context description is based on templates and rules for quiz generation. The author constructed the 5W1H (namely, *what*, *who*, *where*, *why*, *whose*, and *how*) quiz templates, applied them to context description of

the movie scenes, and generated thirty-four quizzes for vocabulary learning with their expected answers. The generated quizzes are listed in Section C. Indeed, the type of quizzes as well as the number of the quizzes can thrive as needed. Context description elements are activated or deactivated synchronously with the replaying scene of the movie. The quizzes relating to the current movie scene are displayed, since they are generated from currently active context description elements.

As mentioned before, the author assigned six quizzes for each learner both in MBCALL and conventional learning. The list of quizzes for each learning are presented in Tables 6.4, 6.5, 6.6, and 6.7.

Table 6.4: List of Quizzes to Learn Bahasa Vocabulary through MBCALL

Quiz	Answer	Quiz (English)	Answer (English)	Score
Apa yang Popeye gunakan memanjat ke atap?	tangga	What does Popeye use to climb to the roof?	ladder	15
Apa yang mengikis kotak pos?	serangga	Who breaks the mailbox?	insect	20
Bagaimana kondisi Popeye?	stress	How does Popeye look?	stressful	15
Popeye jatuh ke mana?	sungai	Where does Popeye fall?	river	20
Bagaimana kondisi serangga?	kekenyangan	How do the insects look?	satiated	20
Bagaimana perasaan Popeye?	lega	How does Popeye feel?	relieved	10

6.2.3 Evaluation of Learning Outcomes

Evaluation of learning outcomes was conducted by asking the learners to perform vocabulary learning tasks. Each learner was asked to make conventional learning as well as MBCALL

Table 6.5: List of Quizzes to Learn Japanese Vocabulary through MBCALL

Quiz	Answer	Quiz (English)	Answer (English)	Score
Yane ni noboru tameni popeye wa nani wo tsukaimasuka?	hashigo	What does Popeye use to climb to the roof?	ladder	15
Nani ga yubin bako wo kowashite imasuka?	mushi	Who breaks the mail-box?	insect	15
Popeye wa donoyouni mie masuka?	iraira shite	How does Popeye look?	stressful	20
Popeye wa doko ni ochi masuka?	kawa	Where does Popeye fall?	river	10
Mushi wa donoyouni mie masuka?	manzoku ni	How do the insects look?	satiated	20
Popeye wa donoyouni kanjite imasuka?	anshin ni	How does Popeye feel?	relieved	20

but with different quizzes. So, learners learned different words in conventional learning and MBCALL.

Six quizzes were assigned for each learner in conventional learning and MBCALL. Each learner was asked to have a pre-test and post-test before and after each learning. The pre-test is aimed to know their initial knowledge on the vocabulary to be learned, while the post-test is aimed to know learner's learning achievement after performing each learning.

MBCALL was performed by watching a short movie with answering related quizzes to be explained in Section 6.2.2. The quizzes relating to the currently replaying scene were identified based on the physical time of the movie and displayed. The expected answer was displayed after the learner wrote his/her answer in order to verify the answer and also to help the learner to memorize the word easily in case the answer was incorrect. The learning process, which was ended with a post-test, was repeated twice to give more opportunities that the learner memorizes the assigned words. The learner's scores achieved by MBCALL are summarised

Table 6.6: List of Quizzes to Learn Bahasa Vocabulary through Conventional Learning

Bahasa Word	English Word	Score
sangat cepat	very quickly	20
nama-nya	his name	15
panik	panic	10
sekop	shovel	15
tempat sampah	trash can	20
lelah	exhausted	20

Table 6.7: List of Quizzes to Learn Japanese Vocabulary through Conventional Learning

Japanese Word	English Word	Score
totemo hayaku	very quickly	10
jibun no namae	his name	15
awatete	panic	20
shoberu	shovel	15
gomibako	trash can	20
tsukarete	exhausted	20

in Table 6.8. The learners #1 to #10 learned Japanese words, while the learners #11 and #12 learned Bahasa words.

Conversely, conventional learning was performed by assigning a list of words that the learner should memorize. The learner is allowed to memorize the words in his/her own way. The learner's scores achieved by conventional learning are summarised in Table 6.9.

To ensure fair evaluation, the author performed testing by equalizing degree of difficulty, scoring weight, and time assignment for both conventional learning and MBCALL. Time assigned for each test was 30 minutes.

Table 6.8: Summary of Learner's Scores through MBCALL

Learner ID	Language	Pre-Test	L1	L2	Post-Test	Outcomes
#1	Japanese	0	0	45	80	80
#2	Japanese	0	0	60	95	95
#3	Japanese	10	10	75	100	90
#4	Japanese	20	30	100	100	80
#5	Japanese	0	0	60	90	90
#6	Japanese	0	30	100	100	100
#7	Japanese	0	0	80	100	100
#8	Japanese	0	0	60	100	100
#9	Japanese	0	0	60	90	90
#10	Japanese	0	0	50	100	100
#11	Bahasa	0	0	45	90	90
#12	Bahasa	0	0	80	100	100
Average	-	2.50	5.83	67.92	95.42	92.92

Table 6.9: Summary of Learner's Score through Conventional Learning

Learner ID	Language	Pre-Test	Post-Test	Outcomes
#1	Japanese	0	100	100
#2	Japanese	0	100	100
#3	Japanese	10	95	85
#4	Japanese	20	100	80
#5	Japanese	0	55	55
#6	Japanese	10	100	90
#7	Japanese	0	100	95
#8	Japanese	0	100	100
#9	Japanese	0	80	80
#10	Japanese	0	100	100
#11	Bahasa	0	93	93
#12	Bahasa	0	95	95
Average	-	3.33	93.17	89.83

6.2.4 Evaluation by the Learner's Subjective Ratings

After conducting both learning, the author collected learner's feedback regarding their learning experience in MBCALL. The learner's feedback was in the form of his/her subjective ratings based on the evaluation criteria derived from Chiu's work as explained in section 6.2.1. The results of the evaluation are presented in section 6.3.

6.3 Discussion

6.3.1 Learning Outcomes

As shown in Table 6.10 and Figure 6.1, it is demonstrated that learners recorded significant post-test scores both in MBCALL and in conventional learning. The average learning outcomes achieved by MBCALL and conventional learning are 92.92 and 89.83, respectively. The score achieved by MBCALL is 3.4% higher than the score achieved by conventional learning.

Table 6.10: Average Learning Outcomes

Learning	Pre-Test	Post-Test	Outcomes
MBCALL	2.5	95.42	92.92
Conventional Learning	3.33	93.17	89.83

In addition, interestingly, one learner was able to obtain 90 points in MBCALL, although she recorded only 55 points for conventional learning; moreover, she recorded 0 points in pre-tests for both learnings. Her learning outcome in MBCALL was 35 points or 38.89% higher than that of conventional learning, while other learners only 10 points at maximum. It indicates that for a certain type of learners, MBCALL makes significant contribution in enhancing learning outcomes. Furthermore, learners aged between 20 to 30 years tended to be more enjoyable in using the MBCALL. That is reflected on higher MBCALL scores than learners aged above 30 years.

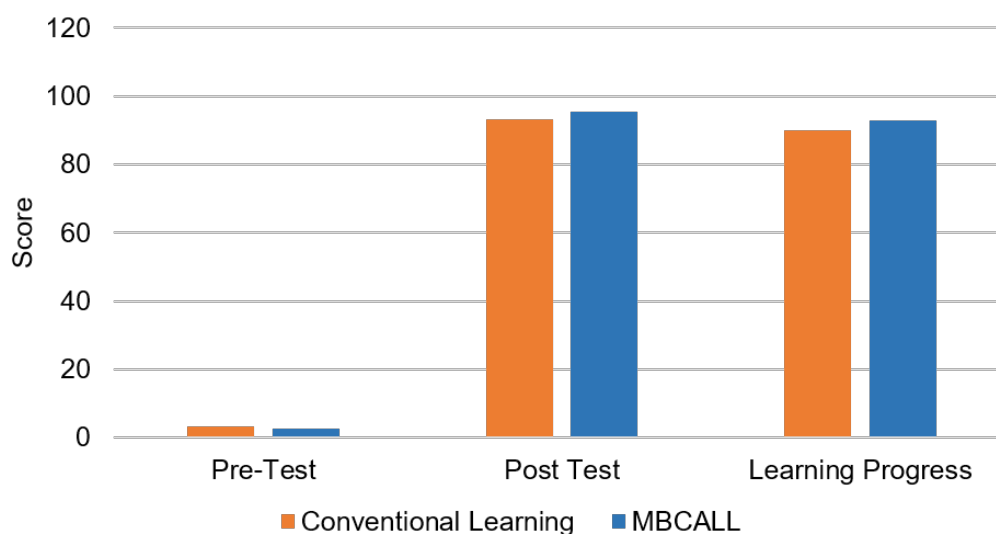


Figure 6.1: Comparison of Average Learning Outcomes

In order to make clear comparison between two languages to be learned (LL), the author shows separated evaluation results of Japanese vocabulary learning and Bahasa vocabulary learning in Figures 6.2 and 6.3, respectively. The results show that there is no significant difference of learner's scores in learning Japanese and Bahasa by MBCALL. Similar with the average learning outcomes, MBCALL contributes to enhance achievement in learning both Japanese and Bahasa vocabulary than conventional learning.

Table 6.11: Learning Outcomes for Japanese Vocabulary Learning

Learning	Pre-Test	Post-Test	Outcomes
MBCALL	3.00	95.50	92.50
Conventional Learning	4.00	93.00	89.00

Table 6.12: Learning Outcomes for Bahasa Vocabulary Learning

Learning	Pre-Test	Post-Test	Outcomes
MBCALL	0.00	95.00	95.00
Conventional Learning	0.00	94.00	94.00

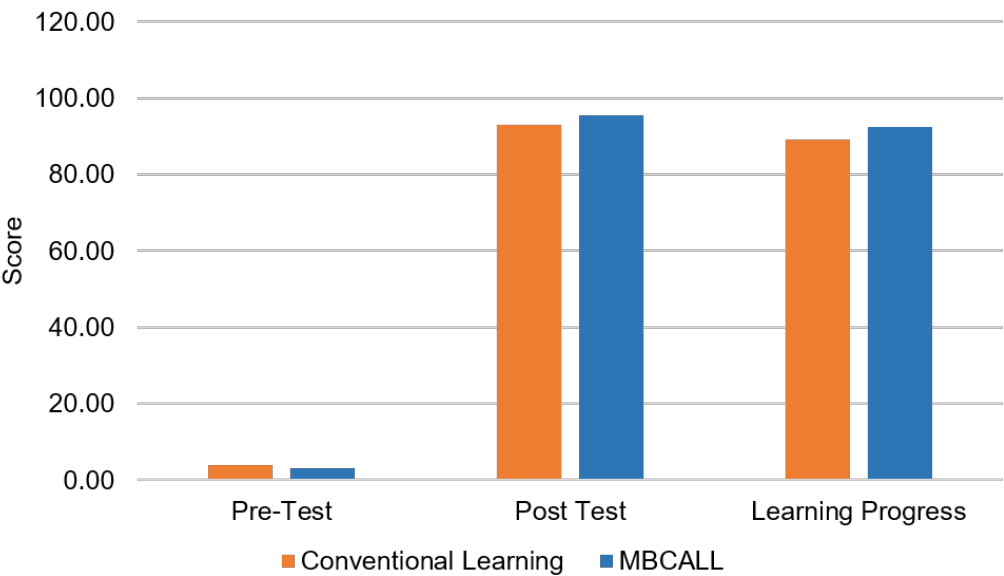


Figure 6.2: Comparison of Average Learning Outcomes for Japanese Vocabulary Learning

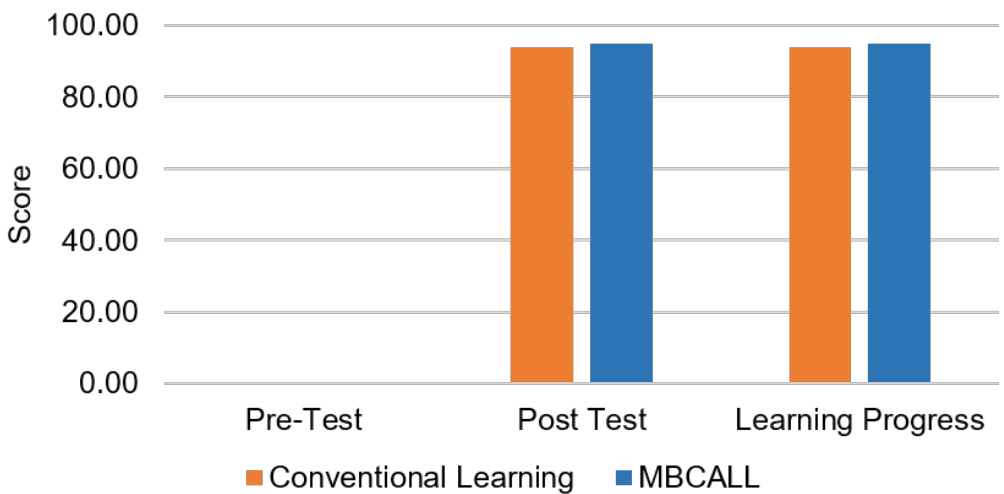


Figure 6.3: Comparison of Average Learning Outcomes for Bahasa Vocabulary Learning

6.3.2 Learner's Feedback

Learner's feedback is also important to know if MBCALL can be adopted learners. As mentioned above, learner's subjective ratings on MBCALL was collected in a form of five

criteria as presented in Table 6.3.

Figure 6.4 shows the averaged subjective rating of each criteria. It seems that the learners react in a positive way to MBCALL. The ratings in four of the five criteria were around 4.0.

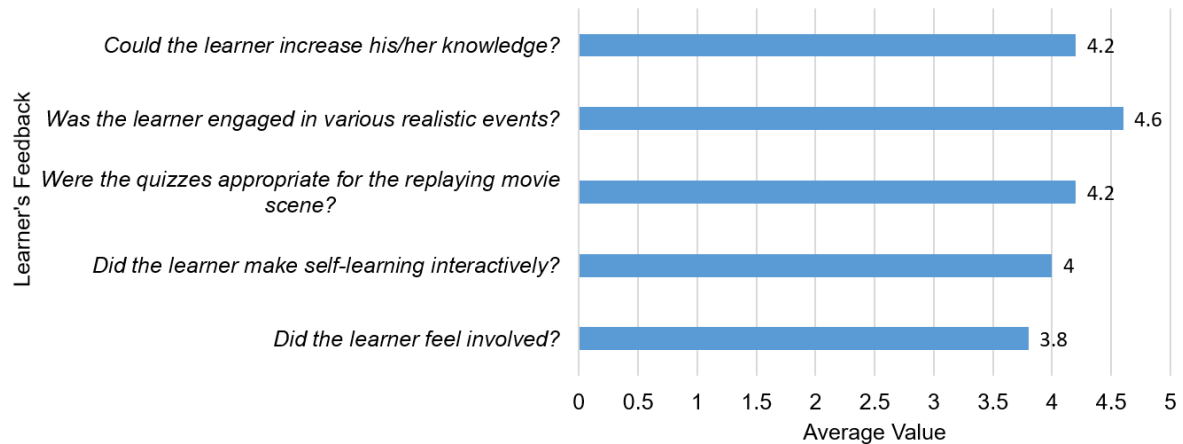


Figure 6.4: Learner's Acceptance to MBCALL

It can be safely stated that the learners could feel involved and engage in MBCALL. They also agree that MBCALL facilitates interactive self-learning as well as contextual learning. In addition, the learners could increase their knowledge through MBCALL. For instance, one learner said that it was the first time for her to see a shovel. Another learner said that she just knows that water can block insects.

Chapter7

Conclusion and Future Work

The author presented a concept of movie-based context-aware learning (MBCAL) that facilitates learner's proactive learning from the context situated in the movie scenes as an alternative approach to context-aware ubiquitous learning that performs learning in the real environment. The key ideas of MBCAL presented in this thesis are the object-oriented context model (OOCM) and context-aware quiz generation.

The OOCM is a framework for context description of the movie scenes that combines the object-oriented approach, which has been adopted in software engineering, with the case grammar concept, which is a semantic framework used in natural language processing. In OOCM, the context of the movie scene is represented by a set of objects and associations among them. The object is an instance of the class such as such as thing, action, mode *etc.* and its role involved in the association represents a case of the case grammar. The author also presented a context description language (CDL), a textual context description language subject to the OOCM, to help the instructor write context description easily and instantly, since it is time-consuming to edit object diagrams of the OOCM for context description.

Context-aware quiz generation is conducted by generating quizzes and their expected answers from the context description by a template-based approach. Context description items of the present movie scene is activated and context description items of the past movie scene

is deactivated. The MBCAL system generates quizzes from active context description items with playing the movie.

Evaluation of MBCAL was conducted with a prototype system of movie-based context-aware language learning (MBCALL), an instance of MBCAL. Its results show that the OOCM can guide the instructor to enrich context description of the movie scenes. The subjects could identify and describe more context description by referencing the OOCM. Concerning learning outcome of the learner, MBCALL could contribute 3.4% improvement in learning tasks memorizing Japanese and Bahasa words than conventional learning.

Learner's feedback regarding their learning experience in MBCALL was also surveyed. The feedback was collected as learner's subjective ratings between 1 to 5 with five evaluation criteria. The average score of each of the criteria was around 4, implying that learners adopted MBCALL positively.

Future work is to apply MBCAL to more learning domains other than language learning and construct a practical system of MBCAL. MBCAL is quasi context-aware ubiquitous learning that learners can learn something in the virtual environment as if they made context-aware ubiquitous learning in the real environment. Originally, a considerable amount of effort to develop and operate context-aware ubiquitous learning motivated realization of MBCAL. Therefore, it is the prime concern to minimize effort needed for context description of the movie scenes and maximize reusability of the context description written once. Peripheral tools for context description based on the OOCM for this purpose are also needed.

Acknowledgement

First of all, I would like to acknowledge my supervisor, Prof. Akira Fukuda, for his valuable guidance, suggestion and support during my stay in his laboratory for the doctoral program. I am extremely grateful to my vice-supervisor, Prof. Tsuneo Nakanishi, for his consistent support, insights and patience in advising and reviewing my work. I convey special thanks to my reviewers, Prof. Tsunenori Mine and Prof. Kenji Hisazumi, for their valuable feedback and suggestion regarding my work.

I express my sincere gratitude to my sponsors: Ministry of Research, Technology and Higher Education of Republic of Indonesia for its scholarship and STMIK Handayani Makassar for its approval and support to my study. I appreciate Kyushu University for offering me the infrastructural and logistics facilities during my study as well as enjoyable and convenient atmosphere of learning.

I thank to Ms. Yoko Otsuru for her support and assistance that helped me to concentrate on my research and also Prof. Shigemi Ishida for his help and guidance in dealing with some technical issues on my work.

I am thankful to the support and encouragement of my lab members, especially to Mr. Baso Habibi and Mr. Asifur Rahman, for their constant support and encouragement. Last but not least, I am deeply thanks to my parents, my husband, and other family members for their continuous support, understanding and encouragement to fulfil my endeavour.

Bibliography

- [1] A. Schmidt, "Potentials and Challenges of Context-Awareness for Learning Solutions," *Proc. IEEE 13th Annual Workshop of SIG Adaptivity & User Modeling in Interactive System (ABIS)*, pp. 63–68, Oct. 2005.
- [2] K. Verbert, N. Manouselis, X. Ochoa, M. Wolpers, H. Drachsler, I. Bosnic, and E. Duval, "Context-Aware Recommender System for Learning: A Survey and Future Challenges," *IEEE Trans. on Learning Technologies*, Vol. 5, No.4, pp. 318–335, Apr. 2012.
- [3] D. Puccinelli, and M. Haenggi, "Wireless Sensor Networks: Applications and Challenges of Ubiquitous Sensing," *IEEE Circuit and System Magazine*, Vol. 5, No. 3, pp. 19–29, Sep. 2005.
- [4] Hazriani, T. Nakanishi, K. Hisazumi, and A. Fukuda, "A Movie Based Context-Aware Learning System: Its Concept and System," *Proc. IEEE 6th Int. Conf. on Technology for Education (T4E)*, pp. 164–167, Dec. 2014.
- [5] G. J. Hwang, C. C. Tsai, and S. J. H. Yang, "Criteria, Strategies and Research Issues of Context-Aware Ubiquitous Learning," *Int. J. of Educational Technology and Society*, Vol. 11, No. 2 pp. 81–91, Apr. 2008.
- [6] D. Chan, and C. Herrero, *Using Film to Teach Languages*, Cornerhouse, 2010.

- [7] G. D. Abowd, A. K. Dey, P. J. Brown, N. Davies, M. Smith, and P. Steggles, "Toward a Better Understanding of Context and Context-Awareness," *Proc. 1st Int. Symp. on Hand-held and Ubiquitous Computing (HUC)*, pp. 304–307, Sep. 1999.
- [8] A. Zimmermann, A. Lorenz, and R. Oppermann, "An Operational Definition of Context," *Proc. 6th International and Interdisciplinary Conference (CONTEXT)*, pp. 558–571, Aug. 2007.
- [9] A. K. Dey, G. D. Abowd, and D. Salber, "A Conceptual Framework and a Toolkit for Supporting the Rapid Prototyping of Context-Aware Applications," *Int. J. of Human-Computer Interaction*, Vol. 16, No. 2, pp. 97–166, Dec. 2001.
- [10] H. Ogata, R. Akamatsu, and Y. Yano, "Computer Supported Ubiquitous Learning Environment for Vocabulary Learning using RFID Tags," *Proc. IFIP TC3 Technology Enhanced Learning Workshop (TeL' 04)*, pp. 121–130, Aug. 2004.
- [11] H. Ogata, and N. Uosaki, "A New Trend of Mobile and Ubiquitous Learning Research: Toward Enhancing Ubiquitous Learning Experiences," *Int. J. of Mobile Learning and organization*, Vol. 6, No. 1, pp. 64–78, May 2012.
- [12] M. Li, H. Ogata, B. Hou, S. Hashimoto, N. Uosaki, Y. Liu, and Y. Yano, "Development of Adaptive Vocabulary Learning via Mobile Phone E-mail," *Proc. 6th IEEE Conf. on Wireless, Mobile and Ubiquitous Technologies in Education (WMUTE)*, pp. 34–41, Apr. 2010.
- [13] S. J. H. Yang, "Context Aware Ubiquitous Learning Environments for Peer-to-Peer Collaborative Learning," *Int. J. of Educational Technology and Society*, Vol. 9, No. 1, pp. 188–201, Jan. 2006.
- [14] Q. Zhan, and M. Yuan, "The Design of a Ubiquitous Learning Environment from the

- Holistic View,” *Proc. Int. Conf. on Networking and Digital Society*, pp. 53–56. May 2009.
- [15] S. Yahya, E. A. Ahmad, and K. A. Jalil, “The Definition and Characteristics of Ubiquitous Learning: A Discussion,” *Int. J. of Education and Development Using ICT (IJEDICT)*, Vol. 6, No. 1 pp. 117–127, Jan. 2010.
- [16] C. M. Chen, and C. J. Chung, “Personalized Mobile English Vocabulary Learning System Based on Item Response Theory and Learning Memory Cycle,” *Int. J. of Computers and Education*, Vol. 51, No. 2 pp. 624–645, Sep. 2008.
- [17] J. Underwood, R. Luckin, and N. Winters, “Milexicon: Harnessing Resources for Personal and Collaborative Language Inquiry,” *Proc. 1st Int. Conf. on Interdisciplinary Research on Technology, Education and Communication (ITEC)*, pp. 87–98, May 2010.
- [18] A. Molnar, “Computers in Education: A Brief History,” *Int. J. of Transforming Education Through Technology (THE)*, Vol. 24, No. 11, pp. 63–68, June 1997.
- [19] G. J. Hwang, “Definition, Framework and Research Issues of Smart Learning Environments: A Context-Aware Ubiquitous Learning Perspective,” *Int. J. of Smart Learning Environment*, Vol. 1, No. 4 pp. 1–14, Nov. 2014.
- [20] P. S. Chiu, Y. H. Kuo, Y. M. Huang, and T. S. Chen, “The Ubiquitous Learning Evaluation Method Based on Meaningful Learning,” *Proc. 3rd Int. Conf. on Computer Education (ICCSE)*, pp.257–264, July 2008.
- [21] C. H. Kerber, D. Clemens, and W. Medina, “Seing is Believing: Learning about Mental Illness as Potrayed in Movie Clips,” *Int. J. of Nurse Education (JNE)*, Vol. 43, No. 10, pp. 479, Oct. 2004.

- [22] N. Lumlertgul, N. Kijpaisalratana, N. Pityaratstian, and D. Wangsaturaka, “Cinemeducation: A Pilot Student Project Using Movies to Help Student Learn Medical Professionalism,” *Int. J. of Medical Teacher*, Vol. 31, No. 7, pp. e327–e332, July 2009.
- [23] W. Jungraithmayr and W. Weder, “The Technique of Orthotopic Mouse Lung Transplantation as a Movie: Improved Learning by Visualisation,” *American J. of Transplantation*, Vol. 12, No. 6, pp. 1624–1626, June 2012.
- [24] M. Ismaili, “The Effectiveness of Using Movies in the EFL Classroom: A study Conducted at South East European University,” *Academic J. of Interdisciplinary Studies*, Vol. 2, No. 4, pp. 121–132, May 2013.
- [25] J. Fillmore, *The Case for Case*, *Proc. Texas Symp. on Language Universals*, pp. 1–134, Apr. 1967.
- [26] D. Jurafsky and J. H. Martin, “Speech and Language Processing: Semantic Role Labeling,” 3rd Edition draft, pp. 377–395. June 2017.
- [27] V. Parunak, “Case Grammar: A Linguistic Tool for Engineering Agent-Based System,” *A Technical Report*, 1995.
- [28] ELiES, “Semantic Role Lists,” <http://elies.rediris.es/elies11/cap51111.htm>, Last Accessed on May 24, 2018.
- [29] Hazriani, T. Nakanishi, and A. Fukuda, “Introducing the Case Grammar Concept to Object Oriented Movie Description,” *Proc. of IEEE Int. Conf. on Teaching, Assessment, and Learning for Engineering (TALE)*, pp. 261–265, Dec. 2016.
- [30] SIL International, “Semantic Role,” <https://glossary.sil.org/term/semantic-role/>, Last Accessed on April 1st, 2018.

- [31] M. Kipp, “Multimedia Annotation, Querying and Analysis in ANVIL,” in M. Maybury (ed.), *Multimedia Information Extraction*, Chapter 19, Wiley, 2012.
- [32] Annotating Academic Video (AAV), <https://github.com/entwinemedia/annotations>, Last Accessed on July 26th, 2018.
- [33] Video Annotation Tools from Irvine, California (VATIC), <https://www.cs.columbia.edu/~vondrick/vatic/>, Last Accessed on July 26th, 2018.
- [34] ISO 639-1:2002, Codes for the Representation of Names of Languages – Part 1: Alpha-2 Code, 2002.
- [35] ISO 639-3:2007, Codes for the Representation of Names of Languages – Part 3: Alpha-3 Code for Comprehensive Coverage of Languages, 2007.
- [36] Tokyo YMCA Inst. Japanese Studies, *Express Japanese*, Hakusuisha, 1987. (in Japanese)
- [37] E. Banno, Y. Ikeda, Y. Ohno, C. Shinagawa, and K. Tokashiki, *GENKI: An Integrated Course in Elementary Japanese*, Vol. 1, The Japan Times, 2011.
- [38] E. Banno, Y. Ikeda, Y. Ohno, C. Shinagawa, and K. Tokashiki, *GENKI: An Integrated Course in Elementary Japanese Workbook*, Vol. 1, The Japan Times, 2011.
- [39] K. C. Kang, J. Lee, and P. Donohoe, “Feature-Oriented Product Line Engineering,” *IEEE Software*, Vol. 9, No. 4, pp. 58–65, Aug. 2002.
- [40] N. A. Nanjiani, M. R. Kuhlke, M. Mishra, and B. Ramakrishnanda, “Automated Quiz Generation System,” US Patent 9028260, May 2015.
- [41] D. Liu, and C. Lin. “Sherlock: A Semi-Automatic Quiz Generation System using Linked Data,” *Proc. 13th Int. Conf. on Semantic Web*, 4 pages, Oct. 2014.

- [42] D. Tendlar (direction), M. Reden, T. Moore (animation), I. Klein (story), A. Loeb (scenics), and W. Sharples (music), *Insect to Injury, Popeye the Sailor Series*, 1956.

AppendixA

OOCM Based on Verbert's Context Framework

As mentioned in Section 4.1, the author adopted Verbert's context framework initially to define the OOCM[2]. Figure A.1 shows the OOCM based on Verbert's context framework. The OOCM has the following nine classes:

- *Actor*: animate beings appeared in the movie, such as *person*, *animal*, other animated things *etc.*.
- *Object*: inanimate beings appeared in the movie such as *car*, *flower*, *etc.*.
- *Place*: place at where an activity performed or an object exists.
- *Activity*: activity performed in the movie.
- *Feeling*: actor's feeling such as *happy*, *angry*, *etc.*.
- *Appearance*: actor's or object's looks such as *beautiful*, *strong* *etc.*.
- *Physical_Condition*: situation of how a place looks such as *cloudy*, *dark*, *etc.*.
- *Time Instant*: time instant at when an actor performs an activity or has a feeling such as *at 3 am*, *at morning*, *on monday*, *etc.*.

- *Time Interval*: time interval for how long an actor performs an activity or has a feeling such as *three hours, two days, etc.*.

Eighty-nine associations of the OOCM, namely context description of the movie scenes, were obtained from the PR movie clip of Kyushu University. The length of the movie was twelve minutes and fifty seconds. Some of these eighty-nine associations are listed below.

(1) *The sky looks clear.*

Object(*the sky*) **seems/looks like** Appearance(*clear*)

(2) *The Buildings are towering.*

Object(*the buildings*) **seems/looks like** Appearance(*towering*)

(3) *The students are doing exercise.*

Actor(*the students*) **performs** Activity(*exercise*)

(4) *A student is doing research.*

Actor(*a student*) **performs** Activity(*research*)

(5) *The students are doing exercise.*

Actor(*the students*) **performs** Activity(*exercise*)

(6) *A lecturer is giving a lesson.*

Actor(*a lecturer*) **performs** Activity(*giving a lesson*)

(7) *A girl is studying.*

Actor(*a girl*) **performs** Activity(*studying*)

(8) *Students perform sport competition.*

Actor(*students*) **performs** Activity(*sport competition*)

(9) *Students are walking in the campus courtyard.*

Actor(*students*) **performs** Activity(*walking*) Actor(*students*) **be at** Place(*the campus courtyard*)

(10) *Frans and Michiko are greeting each other.*

Actor(*Frans and Michiko*) **performs** Activity(*greeting*)

(11) *Students attends a lecture.*

Actor(*students*) **performs** Activity(*attend a lecture*)

(12) *A blue car is in the street.*

Object(*a blue car*) **be at** Place(*street*)

(13) *A green campus environment.*

Place(*campus environment*) **seems/looks** PhysicalCondition(*green*)

(14) *Kyushu University's main entrance.*

Object(*Kyushu University's main entrance*)

(15) *Cluster of stones in front of the building*

Object(*cluster of stones*) **mapped against** Object(*the building*)

(16) *Cherry trees lined in the yard.*

Object(*cherry trees*) **seems/looks like** Appearance(*lined*) **be/ at** Place(*yard*)

(17) *The computers are on the table.*

Object(*the computers*) **mapped against** Object(*table*)

(18) *Beautiful campus view.*

Object(*campus view*) **seems/ looks like** Appearance(*beautiful*)

(19) *Two female students are walking.*

Actor(*two female students*) **performs** Activity(*walk*)

(20) *It takes 24 minutes from Tenjin to Ito Campus by JR.*

Activity(*Tenjin to Ito Campus by JR*) (occurs for) Duration(*24 minutes*)

(21) *Beautiful sea side.*

Object(*sea side*) **seems/looks like** Appearance(*beautiful*)

(22) *The sea surrounded by hills and mountain.*

Place(*the sea*) **seems/looks** PhysicalCondition(*surrounded by hills and mountain*)

(23) *The wave waves smoothly.*

Object(*the wave*) **seems/looks like** Appearance(*waves smoothly*)

(24) *The girls are enjoying the seascape.*

Actor(*the girls*) **performs** Activity(*enjoying the seascape*)

(25) *The men is surfing.*

Actor(*the man*) **performs** Activity(*surfing*)

(26) *People are watching performances.*

Actor(*people*) **performs** Activity(*watching performances*)

(27) *The boy is playing.*

Actor(*the boy*) **performs** Activity(*playing*)

(28) *The sky is cloudy.*

Object(*the sky*) **seems/looks like** Appearance(*cloudy*)

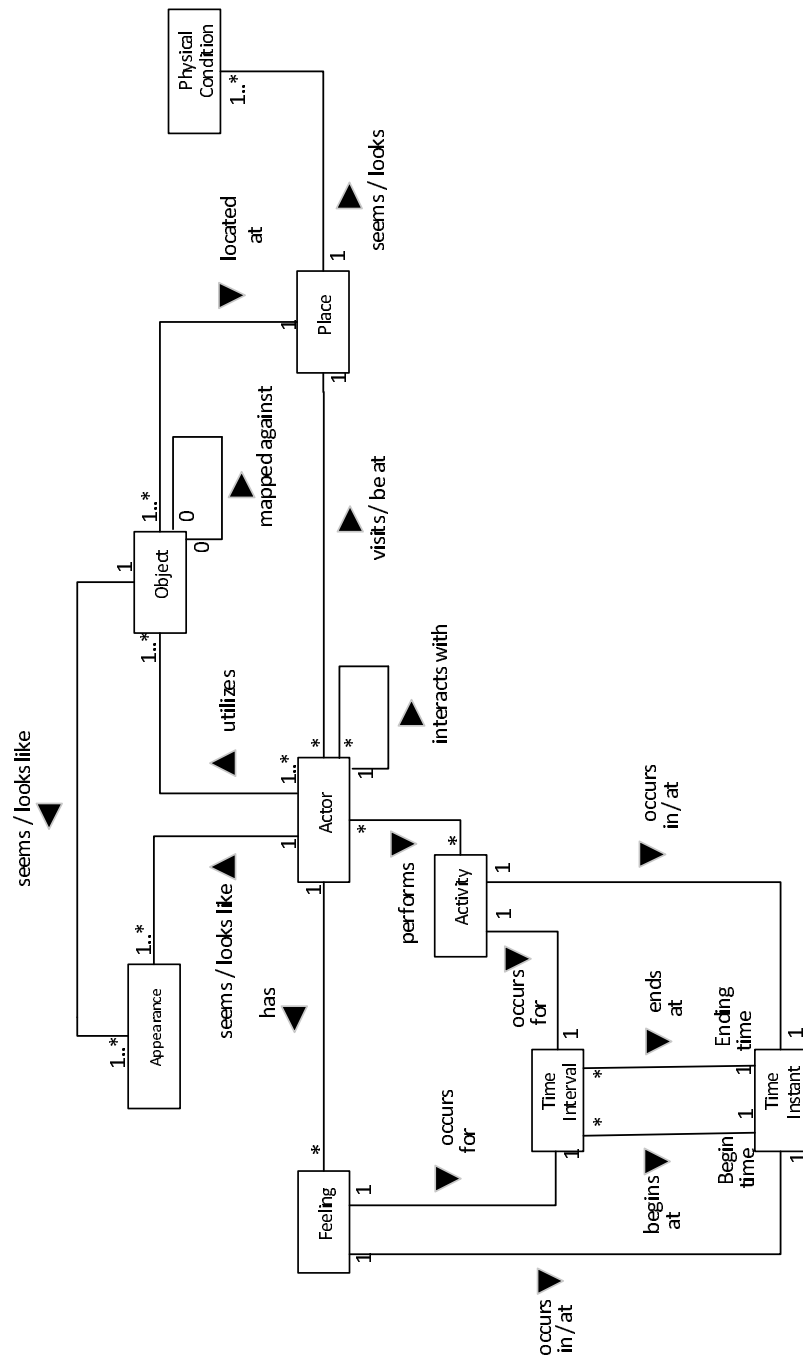


Figure A.1: OOCM Based on Verbert's Context Framework

AppendixB

OOCM Based on the Case Grammar Concept

In Section 6.2, twenty-eight context descriptions under the OOCM based on the case grammar concept were described from a short copyright free movie, entitled as *Insect to Injury*. The length of the movie was 367 seconds. All the context description are listed below.

- (1) The house looks beautiful.

Object(*house*) **looks** Appearance(*beautiful*)

- (2) Popeye climbs to the roof with a ladder.

Actor(*Popeye*) **performs** Action(*climb*)

Action(*climb*) **isPerformedToward** Goal(*roof*)

Action(*climb*) **uses** Instrument(*ladder*)

- (3) Popeye installs the bricks very quickly.

Actor(*Popeye*) **performs** Action(*install*)

Action(*install*) **byTheWayOf** Manner(*very quickly*)

- (4) Popeye writes his name on the mailbox.

Actor(*Popeye*) **performs** Action(*write*)

Action(*write*) **affects** Object(*his name*)

Action(*write*) **isPerformedAt** Location(*the mailbox*)

- (5) The insects break the mailbox.

Actor(*insect*) **performs** Action(*break*)

Action(*break*) **affects** Object(*the mailbox*)

- (6) Popeye looks panic.

Actor(*Popeye*) **looks** Appearance(*panic*)

- (7) Popeye blocks the insects.

Actor(*Popeye*) **performs** Action(*block*)

Action(*block*) **affects** Object(*insects*)

- (8) Popeye puts the insect into the trash can.

Actor(*Popeye*) **performs** Action(*put*)

Action(*put*) **affects** Object(*the insect*)

Action(*put*) **isPerformedToward** Goal(*the trash can*)

- (9) The insects escape from the trash can.

Actor(*insects*) **performs** Action(*escape*)

Action(*escape*) **isPerformedFrom** Source(*trash can*)

- (10) Popeye repairs the fence with a hummer.

Actor(*Popeye*) **performs** Action(*repair*)

Action(*repair*) **affects** Object(*fence*)

Action(*repair*) **uses** Instrument(*hummer*)

- (11) Popeye looks stressful.

Actor(*Popeye*) **looks** Appearance(*stressful*)

- (12) Popeye casts out the insects with a wheel.

Actor(*Popeye*) **performs** Action(*cast out*)

Action(*cast out*) **affects** Object(*insects*)

Action(*cast out*) **uses** Instrument(*wheel*)

- (13) The insects break the bridge.

Actor(*insects*) **performs** Action(*break*) Action (*break*) **affects** Object(*bridge*)

- (14) Popeye falls into the river.

Experiencer(*Popeye*) **experiences** Action(*fall*)

Action(*fall*) **isPerformedToward** Goal(*river*)

- (15) The insects rush into the house.

Actor(*insects*) **performs** Action(*rush into*)

Action(*rush into*) **isPerformedToward** Goal(*house*)

- (16) Popeye digs a trench with a shovel.

Actor(*Popeye*) **performs** Action(*dig*)

Action(*dig*) **affects** Object(*trench*)

Action(*dig*) **uses** Instrument(*shovel*)

- (17) Popeye fills the trench with water.

Actor(*Popeye*) **performs** Action(*fill*)

Action(*fill*) **affects** Object(*trench*)

Action(*fill*) **uses** Instrument(*water*)

- (18) The insects look angry.

Actor(*insects*) **looks** Appearance(*angry*)

- (19) Popeye feels secure.

Actor(*Popeye*) **feels** Feeling(*secure*)

- (20) The insects pass the trench with a can.
Actor(*insects*) **performs** Action(*pass*)
Action(*pass*) **affects** Object(*trench*)
Action(*pass*) **uses** Instrument(*can*)
- (21) Popeye looks shocked.
Actor(*Popeye*) **looks** Appearance(*shocked*)
- (22) The insects destroy Popeye's house.
Actor(*insects*) **performs** Action(*destroy*)
Action(*destroy*) **affects** Object(*Popeye's house*)
- (23) The insects look satiate.
Actor(*insects*) **looks** Appearance(*satiate*)
- (24) Popeye looks strong.
Actor(*Popeye*) **looks** Appearance(*strong*)
- (25) Popeye rebuilds his house in a jiffy.
Actor(*Popeye*) **performs** Action(*rebuild*)
Action(*rebuild*) **byTheWayOf** Manner(*in a jiffy*)
- (26) The insects feel exhausted.
Actor(*insects*) **feels** Feeling(*exhausted*)
- (27) Popeye ridicules the insects.
Actor(*Popeye*) **performs** Action(*ridicule*)
Action(*ridicule*) **affects** Object(*insects*)
- (28) Popeye feels relieved.
Actor(*Popeye*) **feels** Feeling(*relieved*)

AppendixC

Quizzes Used for Evaluation of Learning

In Section 6.2, thirty-four quizzes for vocabulary learning were generated by applying the 5W1H (*what, who, where, why, whose, and how*) quiz templates to context description of the movie scenes in Appendix B. The author assigned six questions for each learner out of this quizzes while they watched the movie. The generated quizzes are listed below. The expected answer is written next to each question in the parentheses.

- (1) (en) How does the house look? (beautiful)
(id) Bagaimana kondisi rumah itu? (indah)
(ja) Ie wa donoyouni mie masuka? (utsukushiku)
- (2) (en) What does Popeye use to climb to the roof? (ladder)
(id) Apa yang Popeye gunakan memanjat ke atap? (tangga)
(ja) Yane ni noboru tameni popeye wa nani wo tsukatte imasuka? (hashigo)
- (3) (en) How does Popeye install the bricks? (very quickly)
(id) Bagaimana Popeye memasang batu bata? (sangat cepat)
(ja) Popeye wa donoyouni renga wo tsunde imasuka? totemo hayaku)
- (4) (en) What does Popeye install? (bricks)

- (id) Apa yang Popeye pasang? (batu bata)
(ja) Popeye wa nani wo tsunde imasuka? (renga)
- (5) (en) How does Popeye feel? (satisfied)
(id) Bagaimana perasaan Popeye? (senang/puas)
(ja) Popeye wa donoyouni kanjite imasuka? (manzoku ni)
- (6) (en) What does Popeye write on the mailbox? (his name)
(id) Apa yang Popeye tulis pada kotak pos? (nama-nya)
(ja) Popeye wa yubin bako ni nani wo kaite imasuka? (jibun no namae)
- (7) (en) Where does Popeye write his name? (mailbox)
(id) Di mana Popeye menuliskan nama-nya? (pada kotak pos)
(ja) Popeye wa doko ni jibun no namae wo kaite imasuka? (yuhbin bako)
- (8) (en) What do the insect break? (mailbox)
(id) Apa yang dihancurkan oleh serangga? (kotak pos)
(ja) Mushi wa nani wo kowasite imasuka? (yuhbin Bako)
- (9) (en) Who breaks the mailbox? (insects)
(id) Apa yang mengikis kotak pos? (serangga)
(ja) Nani ga yubin bako wo kowashite imasuka? (mushi)
- (10) (en) How does Popeye look? (panic)
(id) Bagaimana kondisi Popeye? (panik)
(ja) Popeye wa donoyouni mie masuka? (awatete)
- (11) (en) What does Popeye do? (blocks the insects)
(id) Apa yang Popeye lakukan? (menghalangi pergerakan serangga)
(ja) Popeye wa nani wo shite imasuka? mushi wo tomete iru

- (12) (en) What does Popeye put into the trash can? (insects)
(id) Apa yang Popeye pindahkan ke tempat sampah? (serangga)
(ja) Popeye wa nani wo gomibako ni irete imasuka? (mushi)
- (13) (en) Where does Popeye put the insects? (trash can)
(id) Ke mana Popeye meindahkan seranga? (tempat sampah)
(ja) Popeye wa mushi wo doko ni irete imasuka? (gomibako)
- (14) (en) Who escapes from the trash can? (insects)
(id) Siapa yang kabur dari tempat sampah? (serangga)
(ja) Nani ga gomibako kara nigete imasuka? (mushi)
- (15) (en) Where do the insects escape from? (trash can)
(id) Dari mana serangga kabur? (tempat sampah)
(ja) Doko kara mushi ga nigete imasuka? (gomibako)
- (16) (en) What does Popeye repair? (fence)
(id) Apa yang diperbaiki oleh Popeye? (pagar)
(ja) Popeye ha nani wo naoshite imasuka (hei)
- (17) (en) What does Popeye use to repair the fence? (hammer)
(id) Apa yang popeye gunakan memperbaiki pagar? (palu)
(ja) Popeye wa hei wo naosu tameni nani wo tsukatte masuka? (hanma)
- (18) (en) How does Popeye look? (stressful)
(id) Bagaimana kondisi Popeye? (stress)
(ja) Popeye wa donoyouni mie masuka? (iraira shite)
- (19) (en) What do the insects break? (bridge)
(id) Apa yang serangga rusak? (jembatan)
(ja) Mushi wa nani wo kowashite imasuka? (hashi)

- (20) (en) Where does the Popeye fall? (river)
(id) Popeye jatuh ke mana? (sungai)
(ja) Popeye wa doko ni ochi masuka? (kawa)
- (21) (en) Where do the insects rush in? (popeyes house)
(id) Ke mana serangga bergegas? (ke rumah Popeye)
(ja) Mushi wa doko ni kakekonde imasuka? (Popeye no ie)
- (22) (en) What does Popeye use to dig a trench? (shovel)
(id) Apa yang digunakan Popeye menggali parit? (sekop)
(ja) Popeye wa hori wo horu tameni nani wo tsukatte imasuka? (shoberu)
- (23) (en) What does Popeye fill into the trench? (water)
(id) Apa yang Popeye isi kedalam parit? (air)
(ja) Popeye wa hori ni nani wo irete imasuka? (mizu)
- (24) (en) How does the insects look? (angry)
(id) Bagaimana keadaan serangga? (marah)
(ja) Mushi wa donoyouni mie masuka? (okotte)
- (25) (en) How does Popeye feel? (feel secure)
(id) Bagaimana perasaan Popeye? (merasa aman)
(ja) Popeye wa donoyouni kanjite imasuka? (anshin ni)
- (26) (en) What do the insects use to cross the trench? (can)
(id) Apa yang digunakan oleh serangga menyeberangi parit? (kaleng)
(ja) Mushi wa hori wo koeru tameni nani wo tsukatte imasuka? (kan)
- (27) (en) How does Popeye look? (shocked)
(id) Bagaimana kondisi Popeye? (kaget/terkejut)
(ja) Popeye wa donoyouni mie masuka? (odoroite)

- (28) (en) Who destroys Popeyes house? (insects)
(id) Siapa yang merusak rumah Popeye? (serangga)
(ja) Dare ga popeye no ie wo kowashite imasuka? (mushi)
- (29) (en) How do the insects look? (satiated)
(id) Bagaimana kondisi serangga? (kekenyangan)
(ja) Mushi wa donoyouni mie masuka? (manzoku ni)
- (30) (en) How does Popeye look? (strong)
(id) Bagaimana kondisi Popeye? (kuat)
(ja) Popeye wa donoyouni mie masuka? (tsuyoku)
- (31) (en) How long did Popeye rebuild his house? (in a jiffy)
(id) Berapa lama Popeye membangun rumahnya? (dalam sekejap)
(ja) Donokurai nagaku Popeye wa ie wo tatenaoshi mashitaka? (attoiuma)
- (32) (en) How do the insects feel? (exhausted)
(id) Bagaimana perasaan serangga? (lelah)
(ja) Mushi wa donoyouni kanjite imasuka? (tsukarete)
- (33) (en) What does Popeye do to the insects? (ridicule)
(id) Apa yang Popeye lakukan terhadap serangga? (menertawakan)
(ja) Popeye wa mushi ni nani wo shite imasuka? (waratte)
- (34) (en) How does Popeye feel? (relieved)
(id) Bagaimana perasaan Popeye? (lega)
(ja) Popeye wa donoyouni kanjite imasuka? (anshin ni)

AppendixD

Published Papers

D.1 Journals (peer reviewed)

- (1) Hazriani, Tsuneo Nakanishi, Kenji Hisazumi, and Akira Fukuda, *Object-Oriented Context Description for Movie Based Context-Aware Language Learning*, International Journal of Advanced Computer Science and Applications (IJACSA), Volume 9, Issue 4, pp.350-357, April 2018.

D.2 International Conference (peer reviewed)

- (1) Hazriani, Tsuneo Nakanishi, and Akira Fukuda, *Introducing the Case Grammar Concept to Object Oriented Movie Description*, Proceeding of IEEE International Conference on Teaching, Assessment, and Learning for Engineering (TALE 2016), pp. 261-265, Bangkok, Thailand, December 7-9 2016.
- (2) Hazriani, Tsuneo Nakanishi, and Akira Fukuda, *Architecture, Textual Context Description, and Quiz Generation Scheme for the Movie Based Context-Aware Learning System*,

Proceeding of IEEE Region 10 InternationalConference (TENCON 2016), pp. 2412-2415, Singapore, November 22-25 2016.

- (3) Hazriani, Tsuneo Nakanishi, Kenji Hisazumi, and Akira Fukuda, *A Movie Based Context-Aware Learning System: Its Concept and System*, Proceeding of the 6th IEEE International Conference on Technology for Education (T4E 2014), pp.164167, Kerala, India, December 18-21 2014.