



外国語学習における形成的フィードバックシステム 構築についての検討

邵, 帥
大月, 一弘
清光, 英成
康, 敏

(Citation)

日本教育工学会研究報告集, 19(1):9-16

(Issue Date)

2019

(Resource Type)

journal article

(Version)

Version of Record

(URL)

<https://hdl.handle.net/20.500.14094/90005738>



外国語学習における形成的フィードバックシステム構築 についての検討

A Study on Formative Feedback System for Foreign Language Learning

邵 帥
Shuai Shao

大月 一弘
Kazuhiro Ohtsuki

清光 英成
Hidenari Kiyomitsu

康 敏
Min Kang

神戸大学大学院国際文化学研究科
Graduate School of Intercultural Studies, Kobe University

＜あらまし＞近年、ライティング学習における形成的フィードバックが強調されている。フィードバックを時系列に学習者に提供し、学習効果の向上が期待されている。形成的フィードバックを提供するシステムも提案されているが、具体的な学習内容を取り込まず、枠組みについて提案しているシステムがほとんどである。本研究では具体的な学習内容を踏まえて、日本人学習者向けの和文中訳について課題分析を行い、特定の文法項目に注目して、検出と追跡フィードバックを提示できるアルゴリズムを提案し、形成的フィードバックを行うシステムの構築について検討を行った。

＜キーワード＞ 外国語学習支援 誤用分析 形成的フィードバック 依存構造

1. はじめに

フィードバックは外国語学習のプロセスに重要な役割を果たしている（木村ほか 2010）。近年、ライティング指導におけるフィードバックの有効性が数多くの研究に検証され（Duppenthaler 2004, Oi *et al.* 2000, Yousef 2016）、教育現場でも教師が学習者にフィードバックを提供しつつある。フィードバックの種類やタイミングなどの要素が有効性に大きな影響を与えている（Attali and van der Klei 2017, van der Kleij *et al.* 2012）。最近では、形成的フィードバックが学習に不可欠であると重要視されている（Shute 2008, Goldin *et al.* 2017）。Shute（2008）により、形成的フィードバックは学習プロセスを支援し、学習効果を向上させるために学習者に与える情報として定義されている。数多くの種類が存在し、学習プロセスの異なるタイミングで学習者に提供している。学習者のライティングへのフィードバックは不可欠だが、教育現場では主に教師によるフィードバックに頼る。大人数の授業で、教師が学習者の文章を一つずつチェックし、修正やコメントなどを提供することが大きな負担となり、即時にフィードバックを与えることも困難である。そのため、自動フィードバックシステムが注目を浴びつつある。

近年、AI 技術を知的学习支援システムに導入

し、形成的フィードバックを提供するためのシステム開発に関する研究は盛んであるが、多くの研究は汎用的な知的システムの枠組みの提案に焦点を合わせている（Goldin *et al.* 2017, Easterday *et al.* 2017）。McNamara ら（2013）は、英語を母語とする大学一年生の英語ライティングを支援するために、自然言語処理に基づく知的学習支援システムである Writing Pal を提案した。しかし、外国語ライティングを支援する形成的フィードバックシステムに関する研究まだまだである（Golonka *et al.* 2014）。

本研究では、中国語学習者を対象とした授業において、具体的な学習内容を取り込んだ形成的フィードバックシステムの構築を目指している。日本人学習者向けの和文中訳について課題分析を行い、特定の文法項目に注目して、検出と追跡フィードバックを提示できるアルゴリズムを提案し、システムの構築について検討を行った。

2. システムの枠組み

具体的な外国語学習内容に合わせた形成的フィードバックシステムを構築するには、初級レベルにおいて、学習する文法、学習者の課題への解答を分析して得る文法への理解状況を常に学習するプロセス（順番）に従って処理することが重要なポイントとなる。授業では、カリキュラム上

新しい内容を次々学習していくため、これまでに学習したものをその後の学習を経て忘れがちになる。ある文法項目を学習し、課題をこなして、指摘や正答などのフィードバックも受けたにも関わらず、その後の課題では、以前と同じ間違いをしてしまうことは教育現場でしばしば遭遇することである。このような状況を改善するためには、形成的フィードバックを行うことが有効と考えられる。

図1はシステムの枠組みを示すものである。形成的フィードバックを提供できるように、各文法項目を学習する順番でシステムに取り込み、学習者の解答を分析した上、フィードバックは、1回の課題だけでなく、以前の同じ文法項目についての解答と照合し、提示する必要がある。すなわち、すべての課題への解答に対して、具体的な文法項目を追跡する必要がある。

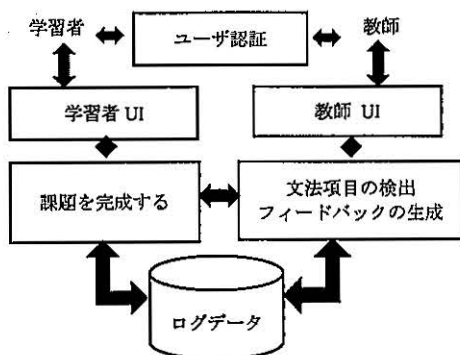


図1 システムの枠組み

システムによって、具体的な文法項目について追跡を行うために、本研究では、実際の授業で収集した課題を分析することにより、具体的な文法項目を選出している。

3. ライティング過程に着目した課題分析

具体的な文法項目を定めるに当たって、本学で学部1、2年生向けの初級中国語と中級中国語の授業を対象とした。中級中国語はSVO型の単文による和文中訳など基本的な文法学習がすでに終わっており、接続詞を使用した複文による作文が目標とした授業である。ここでの複文は日本語の複文と異なる意味を持つ。日本語のセンテンスは単文、重文、複文のように分けられるが、中国語では単文と複文の分類しかない(李 1998)。広辞苑(第六版)により、日本語では、単文は「主語と述語の関係を一組だけ含む文」、重文は「主語・述語の関係が成り立つ部分が、対等の資格で

結ばれている文」、複文は「主節と従属節から成る文。主節の一部に従属節が含まれている文。」のように分類され、中国語では、後方の重文と複文を合わせて複文と呼ぶ。

3.1. VO型フレーズについて

本研究はまず「動詞+目的語」のVO型フレーズに主眼をおき分析を行った。多くの言語で、動詞が全語彙において大きな比率をしめている、中国語も例外ではなく、動詞句が基本的な文型の1つとされている。さらに、「動詞+目的語」のVO型フレーズは動詞句における比率が高く、多くの中国語文法書に単独の文法ポイントとして扱われて、中国語学習者にとってライティングする際に避けられない基本的な文法項目である。

VO型フレーズを含む課題への解答は2年生34名と34名、計2つのクラスからのものであった。それらの学習者に対して、8週間に渡り、1回の授業で1個のセンテンスに対する和文中訳課題を計3回(S1-S3)課した。S1とS2の間隔は1週間であり、S2とS3の間隔は6週間である。3個のセンテンスいずれの内、「花見」という特定の単語が含まれる。なお、第1週目の時、Class1の学習者に「花見に行く」というフレーズの中国語訳「看櫻花」をフィードバック(ヒント)として提示したが、Class2の学習者はヒントなしの状況で課題を完成させた。2週目に、「花見」について説明せず、2つのクラスともヒントなしで再び「花見」を含む課題を完成させた。3週目に、教師側が「花見」の異なる回答について解説し、最初にClass1に提示したヒントが最も適切な表現であることを2つのクラスに口頭で説明した。そして、5週間後、「花見」が含まれる課題を再度学習者に課した。以下はS1-S3の日本語原文と中国語の参考訳である。

● 日本語原文:

- S1.「もし明日雨が降らなければ、私たちは花見に行くつもりです。」
- S2.「もし花見に行くなら、京都が一番いい。」
- S3.「来年3月末に私は神戸に来る予定だが、花見に来るのではなく、出張に来るのだ。」

● 中国語参考訳:

- S1.「如果明天不下雨，我们就去看樱花。」
- S2.「如果去看樱花，京都是最好的。」
- S3.「明年3月底我打算来神戸，但不是为了来看樱花，而是来出差。」

これらの課題の目的は「もし...ならば、...」と「...ではなく、...だ」に相当する複文のライティング学習であった。

分析は「花見」の訳を対象にして、2つのクラスにおける訳の変化を調べることによって行った。「花見」の中国語訳は「看櫻花」（「桜の花を見る」）という「動詞+目的語」のVO型フレーズになる。学習者の翻訳文において、「花見」の正しい中国語訳は「看櫻花」、「看花」、「賞花」と「观赏櫻花」があり、いずれも「動詞+名詞」のVO型フレーズであるが、「看櫻花」が最も適切である。

表1 学習者の「花見」の正解率

	1週目	2週目	8週目
Class 1	100%	94.1%	94.1%
Class 2	73.5%	88.2%	100%

学習者の3つの課題における「花見」の正解率の値は表1に示す。表からは以下のことが読み取れる。まず、1週目に、Class 1と2はそれぞれヒントありとヒントなしの状況で課題を完成したので、Class 1の学習者は100%の正解率に達し、Class 2より上回ることになっている。Class 1の学習者全員最も適切な訳「看櫻花」を使用したことに対して、Class 2の学習者の内「花見」を「看櫻花」に翻訳した人が2人だけであり、日本語の「花見」そのまま使いあるいは欠落のような誤りを犯した人が26.5%を占めた。しかし、2週目に、Class 1の正解率が落ちたことに対して、Class 2の正解率が向上した。3週目の時、教師側が「花見」について、異なる翻訳の差異を説明し、ヒントであった「看櫻花」の適切さを強調した。さらに、8週目の第3回の課題における正解率により、最初にヒントなしのClass 2がClass 1より上回ることが明らかになった。以上の結果から、それぞれの課題に関して正答やヒントを示すだけのフィードバックは効果的と言えない。

表2 8週間における「花見」についての変更率

	1-2週目	2-8週目	1-8週目	「看櫻花」への 変更
Class 1	85.3%	64.7%	73.5%	23.5%
Class 2	52.9%	64.7%	73.5%	23.5%

8週間における「花見」についての訳の変更率は表2に示す。ここの変更率は「花見」について前回と異なる表現を変えた割合を指す。上で述べたように、Class 1の学習者は1週目に全員最も適切な訳文を産出したが、85.3%の人は次の週に異なった表現に変更した。これに対して、Class 2の2週目の翻訳文と1週目を比較し得た変更率は52.9%であった。そして、8週目の第3回の課題の回答における変更率は2クラスとも同様であり、第3回で「看櫻花」の訳に変更した学習者の数は2クラスとも7人であった。教師側の説明がない状況で行った2週目の練習の翻訳文において、Class 1の34名の内2名だけが「看櫻花」を使い続け、Class 2の34名の内1人だけが「看櫻花」に変更した。8週目に2つのクラスとも7人が最も適切な翻訳文を産出したのは3週目に教師側の説明の効果だと考えられる。このような分析結果により、具体的な文法項目を追跡して学習者の習得状況に合わせて説明や指摘などのフィードバックは有効と考えられる。

3.2. 特定の文法項目について

表3 課題における高頻度の誤り

	誤り
S01	動詞フレーズが連体修飾語になる語順
S02	誤字
S03	「やはり」;「けれど」
S04	過去形;「多くの」の訳し方
S05	「授業」の訳し方;「試験」の訳し方
S06	誤字;「多くの」の訳し方
S07	連用修飾語;量詞
S08	「ならば」後半の省略
S09	場所の位置;「どちらも」省略
S10	「より」前後の語順
S11	「始めた」の省略
S12	「ばかり」の位置;「花見」の訳し方
S13	目的を表す「ため」;「できる」の省略
S14	「の」の訳し方;後半の「ある」の省略
S15	「是」の省略
S16	逆接の「が」の省略;「感じた」の訳し方
S17	「皆」の訳し方;「のように」の訳し方
S18	-
S19	「すでに」の省略
S20	「喋る」の訳し方
S21	過去形;「卒業」と「できない」のつながり
S22	場所の位置
S23	「すればするほど」の訳し方
S24	「聞くところによる」の訳し方;目的の「ため」

VO 型フレーズ以外に、注目したのは、使用頻度が高く、誤用頻度も高い文法項目であった。初級中国語受講者 1 クラスの 1 学期に渡る 24 個の課題を分析して、最も頻繁に (10 回以上) 出現した誤用が以下の表 3 で示す。S18 の難易度が低いため、10 回以上出現した誤用が現れなかった。間違いが最も出現したのは、「多くの」に関する表現のセンテンスであった。「多くの」という特定の文法項目が S04, S06 および S13, 3 つの課題に 4 回現れ、日常的使用頻度の高い文法でありながら、学習者の間違える可能性が高いことがわかった。

4. システムの試作

上記の分析により、VO 型フレーズと「多くの」という文法項目に着目して、形成的フィードバックを行うことができるシステムの構築を試みた。

4.1. VO 型フレーズについて

VO 型フレーズを追跡して学習者の習得状況を把握するには、VO 型フレーズを構築する動詞と目的語の間の依存関係の抽出が必要である。この抽出は dependency parser を利用することにした。さらに、抽出を繰り返し、比較することで、学習者の変更履歴を記録し表示することが可能である。以下の手法は 1 つの VO 型構文を含む単文に対応している。2 段階に分けられる：

1) VO 型フレーズの用意と抽出：

- a) 教師が学習効果を確認したい VO 型フレーズ (intended phrase, IP と呼ぶ) を選択する。
- b) IP を含むはずの翻訳文を中国語解析器にかけ、中に含む VO 型フレーズ (learner's phrase, LP と呼ぶ) が dependency parser の結果に基づき抽出する。

2) 形成的フィードバックの提供：

- a) 抽出した LP とともに LP を含む訳文の提出時間も抽出することにより、LP が時系列のデータになり、前後の LP が比較することが可能になる。
- b) LP の抽出と比較は VO 型フレーズについてのフィードバックだけではなく、学習者が翻訳文を変更したか否かの情報も学習者と教師両方に提供する。

システムは主に以下の 4 つのステップに分けて実装した。

- i. 単文を抽出する前処理
- ii. 抽出された単文の分かち書きと形態素解析
- iii. 依存構造解析による VO 型フレーズの抽出

iv. 抽出された VO 型フレーズ間の比較

このシステムでは、Python NLTK (Natural Language Toolkit) インターフェイスを通して Stanford Parser (Manning *et al.* 2014) を使用し、学習者の和文中訳データを分析した。システムのインターフェイスは PHP により開発された。以下は実装したステップに沿ってシステムの説明を行う。入力データは学習者が練習問題で提出した考察したい IP が含まれるはずの翻訳文中の単文である。「看櫻花」を VO 型フレーズの具体例とした。

i) 単文を抽出する前処理

自然言語処理では文構造を自動的に分析できることが重要である。複数の言語に対応する構文解析器は、複文を単文に分けることがすでにできるようになった。しかし、誤用が含まれている文に対して、単文でも自動解析することは困難な課題である。現在多くのアプローチは英語のみを扱い、中国語の構造に対応するものは少ない。さらに、アプローチのほとんどは学習者コーパスに基づき、中国語の学習者コーパスは極めて少ない。したがって、教師の経験的観察を通して、学習者のセンテンスを分割するために、句読点と特定の文法項目、例えば、接続詞、副詞などに基づく方法を提案する。学習者の複文センテンスに対し、まず句読点の位置に基づく分割と抽出を行い、次に文法項目に基づく抽出を行う。

学習者のデータにより、中級レベルの学習者は、教室で学んだばかりの文法を非常によく把握しているが、数週間または数ヶ月前に学んだ文法を間違える可能性が高い。学習者が授業で勉強した主な文法項目は接続詞と特定の文構造であるため、これらに関連する誤りはほとんど現れていないが、以前に学んだ前置詞や VO 型フレーズなどについての誤りが出現している。これにより、課題における接続詞に関する文法がより上手く把握され、単文抽出の手懸かりになる。

2 つの句を含む学習者のセンテンスには、句が句読点で区切られた場合が一般的である。そして、文中の句読点の位置が正しい可能性が非常に高い。この場合、句読点の位置により、IP を含む単文を抽出するのは簡単である。これを句読点位置ステップと呼ぶ。表 4 は、「中級中国語」学習者から収集した翻訳文の中における正しい句読点位置を持つ訳文のパーセンテージを示している。

句読点位置ステップでは、参考訳を使用して IP を含む単文の位置を判定する。たとえば、S1 の参考訳は「如果明天不下雨，我们就去看櫻花。」であり、IP は「看櫻花」であるため、抽出したいのはセンテンス後半の「我们就去看櫻花」にな

ることがわかる。したがって、学習者のセンテンスにおけるコンマを見つけ、次に文の後半部分を検索することにより、学習者の翻訳文から単文を抽出できる。

表4 句読点、文法項目と主語の正答率

	句読点		文法項目		主語	
	Class 1	Class 2	Class 1	Class 2	Class 1	Class 2
S1	100%	91.2%	14.7%	5.9%	88.2%	76.4%
S2	100%	100%	100%	100%	91.2%	88.2%
S3	82.4%	70.1%	85.3%	79.4%	100%	97.1%

句読点のない、または複数の句読点を含む2つ以上の句で構成される学習者の翻訳文に対して、文法項目ステップを導入し、抽出を行う。このステップでは、参考訳におけるIPの位置に完全に依頼するのではなく、上述したように、学習者が学んだばかりの特定の文章構造で使用される接続詞、副詞などの文法項目に基づく条件を追加する。このような授業で学んだばかりの文法項目は学習者に正しく翻訳される可能性が高い。したがって、文法項目は文を分割するときに有用な要素として考慮されるべきである。さらに、英語と同様に、中国語のセンテンスには通常、主語が含まれている。したがって、主語も、センテンスを分割する際に役立つ重要な要素として考えられる。表4は、正しい文法項目を含む翻訳文の割合と、主語が含まれている翻訳文の割合を示している。

文法項目ステップでは、教師が参考訳から文法とIPを指定し、指定された文法とIPにより、このステップで必要となる文法項目と主語が決定される。ここでの文法は、2つ以上の句を含むセンテンスを構成するために必要な文法と指し、例えば、接続詞または副詞である。Stanford Segmentor が参考訳を分かち書き、教師は上記の文法に関する文法要素を選択できるようにする。参考訳内のIPを検索することにより、IPを含む単文とその単文の位置を取得することができ、含まれる文法項目と其中的の主語を判別できる。例えば、S1の参考訳「如果明天不下雨，我们就去看樱花」である場合、IPが「看樱花」であり、文法項目はそれぞれ「如果」と「就」となる。「就」はIPを含む単文の中にあるべきため、学習者の翻訳文の中から目的となる単文を抽出する鍵となる。さらに、Stanford Dependency Parserも参考訳におけるIPを含む単文を解析するために使用され、「nsubj」（名詞である主語）が存在する場合は、対応する主語がキーワードとして抽出され、目的の単文を抽出するためのもう

一つの手懸かりとなる。この例文の場合、主語「我们」も検出されている。文法項目と主語に関する情報、そして参考訳におけるIPを含む単文の位置情報を取得した後、学習者の翻訳文からIPを含むべき単文を抽出できるようになる。

抽出アルゴリズムは、まず句読点に基づいて学習者の翻訳文を部分に分割する。複数の句読点を含む翻訳文である場合、参考訳から抽出された文法項目または主語が存在するかどうかを調べる。片方または両方が存在する場合、たとえば、「如果明天，不下雨，就去看樱花。」「就」という文法項目を含む部分が抽出される。文法項目や主語が翻訳文に存在しない場合、例えば、「如果明天，不下雨，观赏樱花。」なら、位置情報（2番目の部分）に基づき、「不下雨」が抽出される。翻訳文には句読点が存在しなく、参考訳から抽出された主語を含む場合、例えば、「如果不下雨我们打算去观赏樱花。」その翻訳文から、抽出された主語「我们」を検索する。そして主語は通常単文の始めに現れるので、主語から後半の部分を抽出する。それ以外の場合、抽出を行わず、次のステップで翻訳文全体を利用する。

したがって、学習者の複数の句が含まれる翻訳文に対して、句読点の存在に問わず、指定された文法項目または主語を含む場合、IPを含むべき単文を抽出される。含まない場合は、センテンス全体が次の段階で使用される。

ii) 抽出された単文の分かち書きと形態素解析

抽出された単文を対象に、Stanford Parserを用いて、分かち書きと形態素解析を行う。

iii) 依存構造解析によるVO型フレーズの抽出

Stanford Parserの依存関係パーサーは、分かち書きされた単文の文構造情報を提供し、「dobj」（direct object）タグがある場合はその中のLPが抽出される。「dobj」タグが存在する場合、その内容とその形態素情報が抽出されるが、存在しない場合、出力は「*」になる。

iv) 抽出されたVO型フレーズ間の比較

抽出されたLPは時系列に従うため、システムは後のLPを前のものと比較し、学習者が翻訳を変更したか否かを検出する。変更があった場合、赤字で表示する。

4.2. 「多くの」について

特定の文法表現「多くの」については構文解析器による抽出は煩雑になるため、正規表現による抽出を利用した。

表3に示すS04、S06とS13という3つの課題に「多くの」が出現した。Stanford Parserを利用して、学習者の翻訳文に対する解析結果によ

り、S04の翻訳文だけに11個のパターンが存在することが明らかになった。

表5で示しているように、「多くの」は一見簡単な文法項目であるが、S04だけでも「很多」「许多」「很多的」「许多的」「多」「多的」「大批」「大批的」という8種類の翻訳が現れた。また、「多くの」が修飾している名詞「報告」も「報告」「報告」という2種類の翻訳がある。さらに、同じ「多くの」+修飾している単語の文節に対して、前後のコロケーションにより、依存構造の結果が異なる場合がある、例えば、パターン3と10である。そのため、組み合わせたパターンが11個に至ることになる。

表5 S04における「多くの」翻訳パターン

パターン	文節
1	按照很多报告
2	按照很多的报告
3	按照许多报告
4	按照许多报告
5	按照许多的报告
6	按照多报告
7	按照多的报告
8	按照大批报告
9	按照大批的报告
10	据许多报告
11	多报告说

正規表現による抽出手法は以下になる。

参考訳からStanford Parserにより、「多」を含む単語が修飾する単語をキーワードとして抽出する。

学生の翻訳文に参考訳から抽出されたキーワードの存在を確認し、分類する：

a. キーワードが翻訳文の中に存在すれば：

翻訳文のはじめからキーワード前までの文字列を取り出す。

取り出した文字列から正規表現：「[许|很]?多|大批」に符合している単語を検索し、単語の位置から文字列の最後まで部分+キーワードを抽出する。

b. キーワードが翻訳文の中に存在しなければ：

正規表現：「[许|很]?多|大批」に符合している単語を検索し、翻訳文に単語の位置から単語の位置+キーワードの長さの位置までの部分を抽出する。

5. システムの出力について

上記の手法に従い、システムを試作し、課題分析で利用した学習者のデータを入力データとして、出力の一部が図2で示している。抽出されたVO型フレーズと参考訳が異なるなら、赤字で表示され、以前の翻訳文と異なるなら、赤字で「変

更」欄に表示される。

単文を抽出するための前処理アルゴリズムは、全体で99.5%の正しい抽出率を達成し、204の翻訳文から203文が正しく抽出された。唯一の誤りは、次の文から抽出されたものである。「明年3月末我将来神户，不过不是来，观赏樱花，来出差。」この翻訳文の中に、学習者は「不是来」と「观赏樱花」の間にカンマを追加した。システムはまず文法項目「不是」を検索し、参考訳とほとんどの学習者の翻訳文において、句読点なしでVO型フレーズが続くため、単文は正しく抽出された。しかし、この学習者は単文を分割するというまれな誤りを犯し、間違った抽出をもたらした。

単文からVO型フレーズを抽出する確率は表6で示している。抽出されなかった最も大きな原因は、「看花」と「赏花」のような2文字のVO型フレーズがStanford ParserにVO型フレーズではなく、名詞として扱われた。

表6 LPの抽出率

	Class 1	Class 2
Week 1	100%	91.2%
Week 2	100%	97.1%
Week 8	76.5%	82.4%

ID	課題1	VO In S1	課題2	VO In S2	変更(S1>S2)
1	如果明天不下雨，我们就去赏樱花。	赏樱花	如果明天不下雨，我们就去赏樱花。	赏樱花	赏樱花=>去赏樱
2	如果明天不下雨，我们就去赏樱花。	赏樱花	如果明天不下雨，我们就去赏樱花。	赏樱花	No
3	如果明天不下雨，我们就去赏樱花。	赏樱花	如果明天不下雨，我们就去赏樱花。	赏樱花	赏樱花=>去赏樱
4	明天不雨，我们还在赏樱花。	赏樱花	如果明天不下雨，我们就去赏樱花。	赏樱花	No
5	如果明天不下雨，我们就去赏樱花。	赏樱花	如果明天不下雨，我们就去赏樱花。	赏樱花	赏樱花=>去赏樱
6	如果明天不下雨，我们就去赏樱花。	赏樱花	如果明天不下雨，我们就去赏樱花。	赏樱花	赏樱花=>去赏樱

図2 VO型フレーズに関する出力

さらに、図2で示したように、課題間における変更の追跡フィードバックも提供している。全204翻訳文における変更の正しい検出率は94.6%に達し、この高い精度はシステム利用の可能性を証明した。

図3は特定の文法表現「多くの」について抽出・追跡した例である。正規表現に基づく抽出を利用することにより、121個の翻訳文から100%の抽出率を達成している。

S04	抽出	S06	抽出
1 按照很多报告，中国在GDP上超过日本，成为世界第二的经济大国。	很多	与很多人不同，我是在喝很多大报告。	很多人
2 按照许多报告，在国内生产总值上，中国超过日本，成为世界第二的经济大国。	许多	与许多人不同，我读了很多大学。	许多人
3 按照多的报告，中国的国内生产总值超过了日本的，成为世界第二的经济大国。	多的	与许多人不同，我在调查许多大学。	许多

図3 「多くの」文節抽出の出力例

6. まとめ

語学教育現場で技術の利用の普及と教師がフ

フィードバックを提供するとその効果を検証する難しさを背景に、自動的にフィードバックを提供するシステムの利用について検討した。まず、対面式授業で異なるフィードバックを受けた学習者のライティングを収集し、その中におけるVO型フレーズを追跡して、習得状況を分析した。分析の結果をより迅速にフィードバックとして学習者と教師に提供するため、VO型フレーズに関する学習者の翻訳の変化を追跡するシステムを試作した。また、「多くの」という特定の文法項目が支援する必要性が高くと分析からわかった。「多くの」とそれが修飾する名詞の文節について検出の手法を提案した。本研究は学習者のライティング過程に形成的フィードバックを提供するに重点を置いた。

対面式授業で収集した学習者の翻訳文におけるVOフレーズへの追跡することにより、文法項目について説明せず、訳だけをヒントとして学習者に提示することには支援効果が期待できないことが示唆された。このようなフィードバックが学習者の記憶に残らず、推敲を促すことができなく、長期的な支援効果を期待できないと考えられる。さらに、ヒントや自律学習だけがある場合、学習者が一定の正解率に達することが可能であるが、教師の説明がない限り、VO型フレーズの異なる翻訳文の中から最も適切な文を選択することが難しいことがわかった。このような追跡の結果を分析することにより、教師が教室内提供した説明や教授法の効果を確認でき、学習者が学習プロセスの中における自分の不足の部分を把握する効果が期待できる。

以上の人手の分析結果をより迅速学習者と教師に提供するため、自動的に追跡フィードバックを提示するシステムを提案し、システムの構築を試みた。対面式授業で収集した学習者データをテストデータとして実験を行い、VO型フレーズへの追跡と人手評価を比較することにより、システムが人手評価と一致する結果が得られた。変更の判定を正確に行ったが、VO型フレーズを抽出する際、人手評価と異なる結果があることが明らかになった。

さらに、学習者の1学期の課題の解答を分析することにより、「多くの」という特定の文法項目についての誤用の出現頻度が高いことが明らかになった。したがって、「多くの」が存在する文節を抽出できる手法を提案し、システムの機能として開発した。

今後の課題として、2文字のフレーズを判別する機能を開発し、現在のシステムのVO型フレーズの抽出方法を改善する必要がある。また、「多くの」以外の高頻度の誤用に対応するアルゴリズムの追加する予定である。さらに、対面式授業で

行った教師側のフィードバックをシステムに取り入れ、提案手法の有効性について教育現場での実験と評価を行い、学習者のフィードバックにより、システムの設計やフィードバックの提示方を改善する必要がある。

謝辞

本研究の一部は学術研究基金助成金基盤研究(C)(課題番号:17K01081)の助成を受けている。

参考文献

- Attali, Y., & van der Kleij, F. (2017). Effects of feedback elaboration and feedback timing during computer-based practice in mathematics problem solving. *Computers & Education*, 110, 154-169.
- Duppenthaler, P. (2004). The effect of three types of feedback on the journal writing of EFL Japanese students. *JACET Bulletin*, (38), 1-17.
- Easterday, M. W., Lewis, D. R., & Gerber, E. M. (2017). Designing crowdcritique systems for formative feedback. *International Journal of Artificial Intelligence in Education*, 27(3), 623-663.
- Goldin, I., Narciss, S., Foltz, P., & Bauer, M. (2017). New directions in formative feedback in interactive learning environments. *International Journal of Artificial Intelligence in Education*, 27(3), 385-392.
- Golonka, E. M., Bowles, A. R., Frank, V. M., Richardson, D. L., & Freynik, S. (2014). Technologies for foreign language learning: a review of technology types and their effectiveness. *Computer assisted language learning*, 27(1), 70-105.
- 木村博是, 木村友保, 氏木道人. (2010). 『リーディングとライティングの理論と実践』, 大修館書店
- 李汉威. (1998). 简论划分汉语单句复句的标准. *华中师范大学学报: 人文社会科学版*, (4), 111-116.
- Manning, C. D., Surdeanu, M., Bauer, J., Finkel, J. R., Bethard, S., & McClosky, D. (2014, June). The stanford corenlp natural language processing toolkit. In *ACL (System Demonstrations)* (pp. 55-60).
- McNamara, D. S., Crossley, S. A., & Roscoe, R. (2013). Natural language processing in an intelligent writing strategy tutoring system. *Behavior research methods*, 45(2), 499-515.

- Oi, K.kamimura, T. Kumamoto, T. & Matsumoto, K. (2000). A Search for the Feedback That Works for Japanese EFL Students: Content-based or Grammar-based. *JACET Bulletin*, (32), 91-108.
- 新村出. (2008). 広辞苑第六版. 新村出編, 岩波書店.
- Shute, V. J. (2008). Focus on formative feedback. *Review of educational research*, 78(1), 153-189.
- van der Kleij, F. M., Eggen, T. J., Timmers, C. F., & Veldkamp, B. P. (2012). Effects of feedback in a computer-based assessment for learning. *Computers & Education*, 58(1), 263-272.
- Yousef, A. M. F., Wahid, U., Chatti, M. A., Schroeder, U., & Wosnitza, M. (2015, May). The impact of rubric-based peer assessment on feedback quality in blended MOOCs. In *International Conference on Computer Supported Education* (pp. 462-485). Springer, Cham.