

文章生成における推敲機能の実現について

乾健太郎, 徳永健伸, 田中穂積
東京工業大学 工学部

{inui,take,tanaka}@cs.titech.ac.jp

文章生成には語彙選択や語順などさまざまな要素に関する選択が必要である。従来のシステムでは、規則を用いて各選択点における選択肢の優先度を局所的に計算するのが一般的だった。しかしながら、選択肢間の優先度の根拠となる評価基準のなかには、評価に将来の選択の予測や大域的な情報を必要とするものもある。それらのすべてを局所的な規則で扱うアーキテクチャは、規則が複雑になりすぎるため、実現や保守の点で望ましくない。本稿では、この問題に対して推敲機能の導入が有効であることについて議論し、推敲機能を備えた生成モデルとその実装方法を提案する。また、その手法を実装した日本語文章生成システムについて述べ、その性能を評価する。

Implementation of Revision in Text Generation

INUI Kentaro, TOKUNAGA Takenobu and TANAKA Hozumi
Department of Computer Science, Tokyo Institute of Technology
(2-12-1 Ōokayama Meguro-Ku Tokyo 152 Japan)

To generate good text, many kinds of decisions should be made. Many researchers have been engaged in searching for the heuristics that would make an appropriate decision locally at each decision point. However, even if such heuristics were found, there are still certain kinds of problems that are difficult to consider during the generation process. Those problems can be more easily solved by introducing a revision process after generation. In this paper, we argue the importance of text revision with respect to natural language generation, and propose a computational model of text revision. We also discuss its implementation issues and introduce dependency directed backtracking in order to realize efficient revisions. Finally, we evaluate our method on an experimental Japanese text generation system.

1 はじめに

文章生成システムは一般に書き手の情報伝達ゴールを入力として受け取り、文章を出力する。文章生成に必要な処理は、どのような内容をどのような構成で述べるかに関する選択(what-to-sayの選択)とhow-to-sayの選択(どのように言語表現するかに関する選択(how-to-sayの選択))に大きく分けることができる[10]。多くのシステムでは、まずwhat-to-sayを選択し、その結果を修辞構造[6, 10]と呼ばれる構造化された命題の集合で表現する。つぎに、how-to-sayを選択し、文字列を出力する。本稿では、とくにhow-to-sayの選択に関する従来の手法の問題点を取り上げ、この問題を解決するために、推敲機能を備えた生成モデルを提案する。

how-to-sayの選択には、複数の命題を一文に統合するなどの文章レベルの統語的選択、文型、語順などの文内の統語的選択、および語選択が含まれる。これらの選択は図1のような探索空間の各アークに対応する。how-to-sayの生成過程では、システムは修辞構造を入力として受け取り、探索空間を根ノードから順にたどる。中間ノードはそれまでの選択を蓄積した状態を表し、これは次の選択に対する制約となる。各ノードから出るアークは、その状態で必要な統語的・語彙的制約を満たす選択肢を表す。システムは、最終的に葉ノードに到達し、文章を1つ出力する。図1には葉ノードが複数あるが、これは同じ入力に対応する言語表現が複数存在することを意味する。システムの目標は、統語的・語彙的制約を満たす複数のパスの中から、質の高い文章を生成するパスを選択することである。

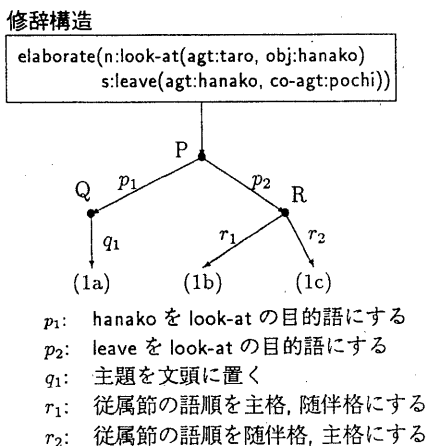


図1 文章生成の探索空間

多くのシステムでは、個々の選択点で選択肢の局所的な優先度を求め、それにしたがって選択を行なう。各選択点での局所的な優先度は一般にヒューリスティクスを用いて与える。本稿ではこのような規則を選好規則と呼ぶ。選好規則は文章の質を決めるさまざまな評価

基準に基づいて作られる。文章生成の研究の一つ目的はどのような選好規則が適当かを究明することである。たとえば、照応表現や態の選択などに関して、すでにさまざまな規則が提案されている[?, 10]。しかしながら、文章の質を局所的な優先度だけで実現する方法には限界がある。

係り受けのあいまい性を文章の評価基準の一つとして考えてみよう。入力として図1のような修辞構造を仮定する。いま、システムが図の選択肢 p_1, q_1 を選択して、文(1a)を生成したとしよう。

(1a) 太郎はボチと去っていく花子を見ていた。

(1b) 太郎は花子がボチと去っていくのを見ていた。

(1c) 太郎はボチと花子が去っていくのを見ていた。

この文は「ボチと」の係り先があいまいなため、複数の解釈が可能であり、望ましい文とは言えない。一方、 p_2, r_1 を選択すると、あいまい性のない文(1b)を生成することができる。この例では、 p_1 を選択するとあいまい性を持つ文を生成するので、あいまい性を避けるためには選択点Pで p_2 を選択すべきだったことがわかる。

しかしながら、出力文があいまい性を持つか持たないかはP以降の選択にも依存する。たとえば、 p_2 を選択しても、Rで r_2 を選択すれば、文(1c)のようなあいまい性が生じてしまう。つまり、選択点Pでは、 p_1 よりも p_2 の方が望ましいとは必ずしも言えない。選択点Pでは、まだ語順や語などが決まっていないため、出力文のあいまい性を評価するために必要な情報が部分的にしか得られないからである。このように、評価基準の中には、選択肢の優先度を局所的に評価するのが困難なものがある。このような評価基準の例に、係り受けのあいまい性、文や節の長さなどの統語的複雑さなどがある。

この問題は選択の順序をうまく制御すれば解消するように見えるが、最適な選択の順序を静的に決めるのは一般に難しい。また、選択の順序を動的に決定する手法もいくつか提案されているが[4]、処理が複雑になるため、その実現や保守の点で問題がある。

そこで、本稿では、一度生成した文章を推敲することができる生成モデルを提案し、その実現について述べる。推敲機能を導入することによって、係り受けのあいまい性や文の複雑さのような評価基準については、文章を生成した後で評価することができる。また、一度決めた選択を変更することができるため、選択の順序が必ずしも最適である必要はない。

文章生成における推敲機能の役割とその重要性についてはすでにいくつかの議論がある。Meteer[7]、Yazdani[9]は、一度生成した文章に対する評価・修正を繰り返すモデルをそれぞれ提案し、それについて次のような利点をあげている。

1. 推敲過程には心理学的妥当性がある。

2. 構造的に類似した2つの文を1つに統合するなど、表層化したほうが考慮しやすい問題について、表層から途中の選択へのフィードバックが可能になる。
3. 比較的単純な修正を繰り返すことによって、全体として複雑な処理を実現することができる。

本稿では、2と3の観点から推敲機能を導入し、より具体的なモデル化を行なう。また、従来ほとんど議論されてこなかった推敲機能の実装について一つの手法を提案する。

以下、2節で文章生成モデルについて述べ、3、4節で効果的かつ効率的な文章の修正を可能にする実装方法を提案する。さらに、5節で日本語の文章を生成する実験システムの定量的評価について述べる。

2 生成モデル

図2に推敲機能を備えた生成モデルを示す。本稿では how-to-say の選択に限って議論するので、入力として what-to-say すなわち修辭構造を仮定する。文章生成過程は初期生成過程とそれにつづく推敲過程からなる。

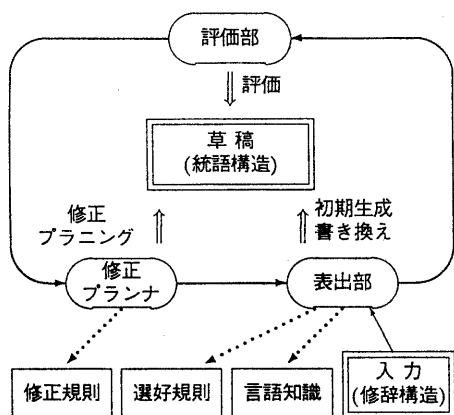


図2 文章生成モデル

初期生成過程では、表出部が修辭構造を受け取り、統語的選択と語選択を行ない、文章を生成する。我々のモデルではこの文章を書き換える可能性があるため、これを草稿と呼ぶ。統語的選択と語選択は言語知識と選好規則を参照して行なう。言語知識は統語的・語彙的制約を与える。一方、選好規則は参照表現や主題化などの語用論的・文脈的優先度を与える。1節で述べたように、ある選択点に統語的・語彙的制約を満たす選択肢が複数ある場合、表出部は選好規則を用いて優先度の高い選択肢を選ぶ。草稿は、すべての統語的選択と語選択を行なった結果であり、統語構造として表現される。草稿は中間表現でなく、文字列であることに注意してほしい。

推敲過程は修正サイクルの繰り返しからなる。各修正サイクルでは、評価、修正プランニング、書き換えをおこなう。まず、評価部が草稿を評価し、問題点を検出する。つぎに、修正プランナが検出された問題点の中から1つを選択し、修正規則を参照して、その問題を解消するために有効な修正方法を決定する。最後に、表出部が修正プランナの決定にしたがって実際に草稿を書き換える。これで一回の修正サイクルが終わる。修正サイクルを繰り返すことによって、草稿の質が徐々に向上することが期待できる。

モデルの特徴は、文章の質を判断するさまざまな評価基準を選好規則と修正規則という2つの資源を用いて記述する点にある。選好規則については、多くの研究者がさまざまなヒューリスティクスを提案し、その有効性を確認している。我々のモデルではそれらの成果を自然に利用できる。しかしながら、1節で述べたように、評価基準のなかには、生成過程の途中では評価に必要な情報が部分的にしか得られないものがある。そのような評価基準を生成過程の途中で評価するには将来の選択に関する慎重な予測が必要になる。このため、それらのすべてを選好規則の中に押し込めると、選好規則が複雑になり、システムの実現や保守の観点から望ましくない。我々のモデルでは、選好規則を単純化するという立場から、将来の選択に関する予測を用いず、局所的な情報だけで選択肢の優先度を計算するように選好規則の能力を限定する。一方、大域的な情報を必要とする評価基準に関する優先度については、選好規則だけでは十分に与えることができないので、選好規則が与えた優先度を評価部の評価に基づく修正規則の適用によって修正する。これにより、個々の規則を単純化だけでなく、選好規則だけでは十分に扱えない評価基準についても適切な優先度を与えることができる。本稿では、評価部で扱う評価基準の例として、係り受けのあいまい性と文の構造の複雑さ(4.1項)を取り上げるが、どのような評価基準を評価部で扱うべきかは今後の研究課題である。

3 言語知識と初期生成

3.1 機能単一化文法

1節で述べたように、文章を生成する過程は探索空間の1本の探索パスを選択することに対応する。一方、草稿の書き換えは、一度選択したパス上のある選択点について選択を変更する処理とみなせる。もとの草稿と書き換え後の草稿はどちらも同じ修辭構造から生成したものであるため、草稿の意味は書き換えた後も保存される。

問題を解消するために変更すべき選択は、統語的な選択や語選択など、さまざまな言語知識に関係する。したがって、言語知識を統一的な表現形式で扱うことが

望ましい。このような理由から、我々のシステムでは言語知識の表現形式として機能単一化文法 (FUG)[5] を採用している。FUG では、意味構造や統語構造を表現する作業機能記述 (WFD) と言語知識を表現する文法機能記述 (GFD) とを単一化することによって生成または解析を行なう。

入力となる WFD の例を図 3 に、GFD の一部を図 4, 5, 7 に示す。修辞関係には Mann の RST[6] を用いる。n は主部 (nucleus) を、s は従属部 (satellite) を表す。各素性は、対象とする領域中の個体 (document#1 など)、統語的・語彙的素性値 (sentence など)、機能記述 (FD) のいずれかを値にとる。素性 path:(...) はパス(...) で指された FD と素性を共有することを表す。↑は親の FD に 1 段さかのぼることを示す。↑が 1 つの場合は自分自身を指し、無い場合は絶対パスを示す。@で始まる素性名・素性値は型つき変数を表す。

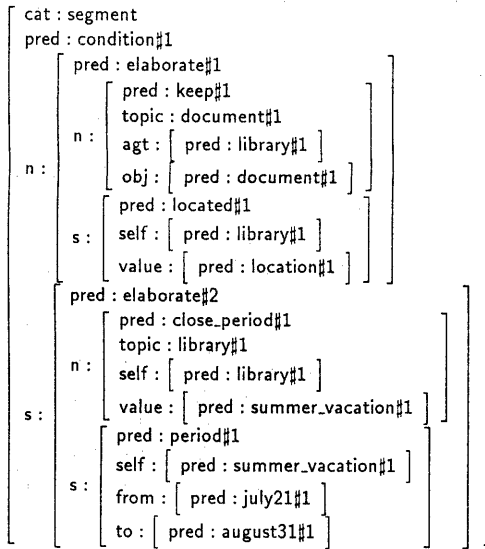


図 3 WFD の初期状態

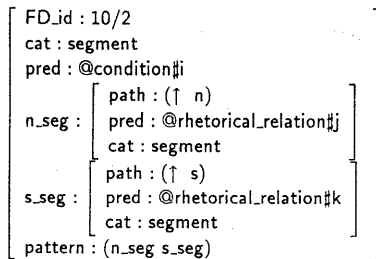


図 4 GFD の一部

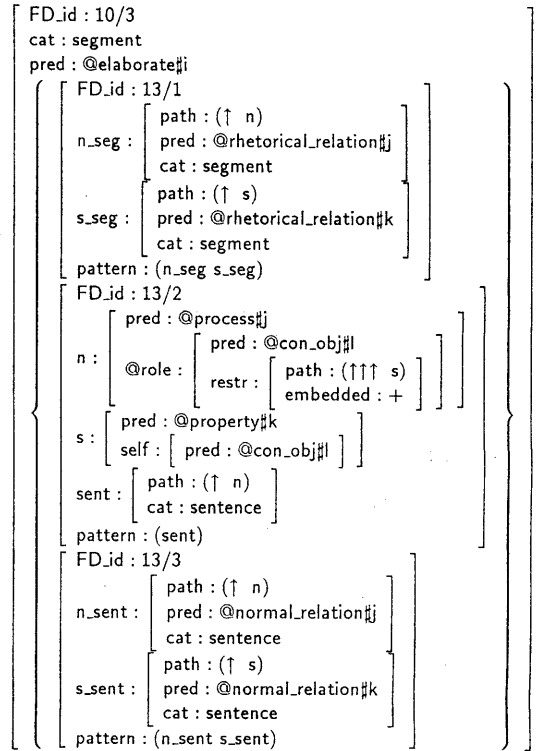


図 5 命題の埋め込みに関する GFD

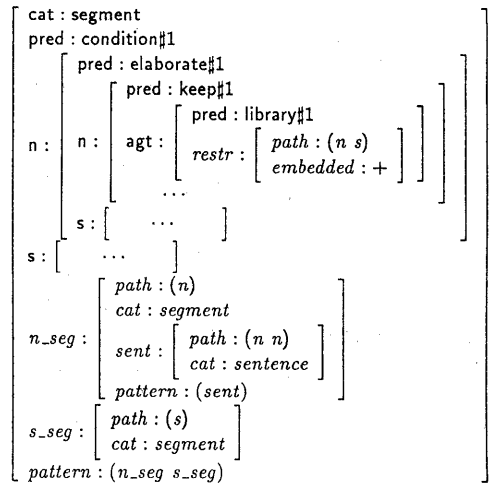


図 6 単一化後の WFD

生成過程では、まず WFD のパス () の下にある FD と図 4 の GFD との単一化が成功し、つぎにパス (n_seg) の下にある FD と図 5 の GFD との単一化が成功する。図 5 のように GFD が選言を含む場合は、一番上の選択

肢から順に単一化を試み、最初に成功するものを1つだけ選択する。ここではFD(13/2)との単一化に成功し、図6のWFDを得る。斜字体で示した素性が単一化によって付加された素性である。このように、単一化はWFD上をトップダウンに進み、表層の文字列の生成に必要な統語的・語彙的素性がすべて決定した時点で終了する。この例では草稿(2a)を生成する。

- (2a) あなたが探している資料は正門にもっとも近い建物の4階にある資料室に保管されています。ただし、資料室は7月21日から8月31日までの夏休みの間は利用できません。

生成過程には参照表現の選択も含まれる。たとえば、document#1が「あなたが探している資料」と表現されているが、これは、読み手が「個体 document#1」を同定するために必要な情報を表出部が対象領域のモデルから取り出し、表層化した結果である。

3.2 言語知識と生成過程

図6は、(n_seg)との単一化でFD(13/2)を選択したことにより、(n n agt)にrestrが加わった状態である。これにより命題located#1が命題keep#1に埋め込まれることになる。修辞構造から表層の文章構造への写像規則に関しては、Scottらの研究[8]をはじめ、いくつかのヒューリスティクスが提案されている。それらの多くは図5のようなGFDとして記述することができ、文内レベルの統語的選択・語選択と同様、単一化を用いて処理できる。このように、本システムでは文を越えたレベルの選択と文内レベルの選択を統一的に扱う。

GFDが選言を含む場合、3.1項で述べたように、表出部は上の選択肢から順に単一化を試みる。成功するものが複数存在する場合でも、最初に成功した選択肢だけを選んで先に進むため、選択肢の間には優先度が存在することになる。この優先度は1節で述べた選好規則に相当する。

選言中の選択肢の位置は選択肢の間に静的な優先度を与えるが、照応表現の選択など、文脈に依存して動的に優先度が変化する場合もある。たとえば、草稿(2a)の第2文の主題library#1は「資料室は」という後置詞句として表層に現れているが、仮に直前の文ですでに「資料室は」という提題が起こっていたとすると、省略する方が自然であろう(4.1項の草稿(2b))[?]。この知識は図7のような選言として記述できる。FD(25/1)中のexternalは、任意の外部手続きを呼び出し、その結果を値として返す特別な関数で、Elhadadらが生成システムFUF[3]に導入したものと同じである。たとえば、FD(25/1)のexternal(current_theme)は直前の文で主題になっていた個体を返す役割をはたし、この情報を用いることによって主題の省略を決定できる。

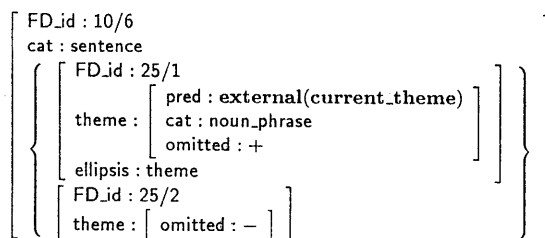


図7 主題の省略に関するGFD

AppeltのTELEGRAM[1]は参照表現を選択するための外部手続きをGFDから呼び出す機構を備えているが、externalはTELEGRAMの機構の一般化と見なせる。externalは、読み手のモデルやそれまでの選択の結果など、文法をとりまく環境とのインタフェースであり、これを利用すれば動的に変化する制約や優先度を扱うことができる。

4 草稿の修正

4.1 修正サイクル

初期生成が終了すると、システムは最初の修正サイクルに入る。3.1項で示した例では、まず、評価部が草稿(2a)を評価し、問題を検出する。2節で述べたように、本稿では評価部の評価基準として係り受けのあいまい性と文の複雑さを取り上げる。日本語の文の複雑さとして表1に示す基準を考える。

表1 文の複雑さの評価基準

評価項目	上限	下限
文の長さ(語数)	40	5
係り受けの深さ	3	0
中心埋め込みの深さ	3	0

表中の係り受けの深さは、“[[[[太郎の] 友達の] 家]]”のように係り受け関係が入れ子になるときの構造の深さを意味する。この基準に照らすと、第1文の「資料室」にかかる名詞修飾節「正門に...4階にある」の係り受け関係は深さが4なので深すぎる。つぎに、修正プランナが、この問題を解消するためには(n_seg)におけるFD(13/2)の選択を変更すればよいことを判断する。さいごに、表出部がFD(13/2)の代わりにFD(13/3)を選択し、草稿(2b)を生成する。

- (2b) あなたが探している資料は資料室に保管されています。資料室は正門にもっとも近い建物の4階にあります。ただし、[ゆハ] 7月21日から8月31日までの夏休みの間は利用できません。

これで一つの修正サイクルが終了する。システムは以下同様の処理を繰り返す。

修正の際に選択の変更を行なう選択点、すなわちバックトラックの戻り先を CCP (culprit choice point) とよぶ。上の例では、(n_seg) における選択点 FD(13) (図 5) が CCP に当たる。上で述べたように、修正プランニングは CCP を特定する処理に相当し、草稿の書き換えは単一化の処理を CCP までもとし、新しいパスをたどることに相当する。本システムでは、真偽値維持システム (TMS) [2] で用いられている依存指向性バックトラック (DDB) を応用することによって、これら 2 つの処理を効率的に行なう。DDB を用いる利点は次の 2 つに要約できる。第 1 に、バックトラックの戻り先をヒューリスティクス (修正規則) によって特定することができるため、単純な時系列バックトラックに比べ、CCP の同定に要するバックトラックの数を大幅に抑制できる。第 2 に、草稿の書き換えの際、変更した選択に依存する選択だけを再計算するため、時系列にそって選択を取り消す手法に比べ、CCP から新たな草稿を生成する際の計算量が少ない。

4.2 依存関係ネットワーク

システムは、TMS と同様、依存関係ネットワークを用いて DDB を実現する。WFD 中の各素性に対し固有のノード (素性ノードと呼ぶ) を割り当て、素性ノードの集合として WFD を表現する。素性ノードは次のようにパスと素性の組 (path, feature) で表現される。

((n n agt), pred:library#1)

初期生成を開始する時点での依存関係ネットワークは、入力 WFD が持つ素性ノードの集合である。生成過程では、素性ノードの集合として表現した WFD を GFD と単一化する。このとき、GFD 中の各選択点と WFD 中の素性との間の依存関係を、対応する FD ノードと素性ノードの間に正当化のアーキを張ることによって記録する。FD ノードはパスと選択した FD の識別子の組 (path, FD_id) で表現する。

図 8 は、依存関係ネットワーク中の 3.1 項で述べた FD(13/2) の選択に対応する部分である。FD ノード ((n_seg), FD(13/2)) に向かうアーキを持つ素性ノード (((n n agt), pred:library#1) など) は、FD(13/2) に含まれる素性のうち単一化以前から WFD に存在した素性に相当し、これらの連言が FD(13/2) の正当化 (必要条件) に当たる。他方、FD ノードから外に向かうアーキの先の素性ノード (((n n agt restr), embedded:+) など) は単一化によって新たに WFD に付加された素性に相当する。システムは、これら FD ノードと素性ノードの間の依存関係の他に、FD ノード間の依存関係もネットワーク上に記録する。((n_seg), FD(10/3)) から ((n_seg), FD(13/2)) に張られたアーキがその例である。

表出部は FD を単一化する度にネットワークを更新

する。単一化では WFD 上に生成された素性はそれ以後つねに制約として働くので、ネットワークは必ず非循環有向グラフになる。また各ノードは高々 1 つの正当化しか持たない。FD ノードは in, out, unknown のいずれかの状態を、素性ノードは in, out のいずれかの状態をとる。各ノードは、最初に生成されたとき in をとり、親ノードの少なくとも 1 つが out のとき out になる。アーキで表現されるノードの正当化は必要条件にすぎないため、一般に out のノードを状態の伝播によって in に変えることはできない。したがって、out のノードはネットワークから取り除く。unknown は external によって生じる依存関係を扱うための状態である。external の値は文脈に依存して変化する可能性があるため、システムは external を含む FD が選択の変更後も単一化できるかどうか調べる必要がある。しかしながら、external を含む FD の選言は一般に語用論的・文脈的優先度に関する選択点に相当するので、選択の変更後も再利用できる可能性が高い。そこで、選択の変更後、external を含む FD に対応する FD ノードに状態 unknown を与える。unknown は、out と異なり、子ノードに伝播しない。

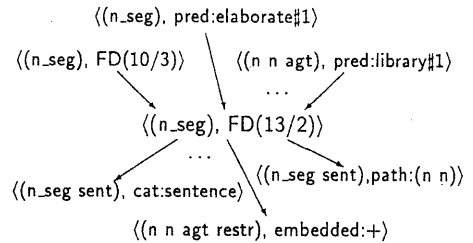


図 8 依存関係ネットワーク

4.3 修正プランニング

修正プランニングでは修正規則を参照して CCP を特定する。4.1 項の例ではつぎの修正規則を適用する。Path, Path1 は変数を表す。

```
If too_deep(Path) ∧
state_in((Path, cat:rel_clause)) ∧
constituent_of(Path1, Path) ∧
state_in((Path1, embedded:+))
then remove((Path1, embedded:+)).
```

このように、修正規則には、現在の草稿からどの素性を取り除けばよいかをプロダクション規則の形式で記述する。修正プランニングでは、適用可能な修正規則の中からもっとも優先度の高い規則を選択し、取り除くべき素性を特定する。あとは、依存関係ネットワークを参照すれば、容易に CCP が特定できる。この例では、((n n agt restr), embedded:+) を取り除くべき素性として特定し、図 8 に示したネットワークを参照して、FD ノード ((n_seg), FD(13/2)) を CCP として特定する。

4.4 草稿の書き換え

草稿の書き換えアルゴリズムを示す。ノードの状態は各処理ごとに伝播するものとする。

1. 修正プランニングで CCP として特定された FD ノードの状態を *out* にする。状態伝播後の *in* の素性ノードの集合が現在の WFD である。
2. *external* を含む FD に対する FD ノードが存在し、そのパスが統語構造中で CCP のパスの下または右にある場合、その FD ノードを *unknown* にする。
3. 未決定の選択点を統語構造についてトップダウン左優先に探す。すべての選択が決定していれば書き換え成功で、処理を終える。未決定の選択点は、選択の変更で新たに生じた選択点の他に、FD ノードが *out* または *unknown* であるものを含む。
4. 3 で見つけた選択点の FD ノードが:
 - a. CCP であるとき 4d へ。
 - b. *unknown* のとき、親の *external* ノードについて単一化を試み、成功すれば FD ノードを *in* にして 3 へ、失敗すれば *out* にして 4d へ。
 - c. 存在しないとき、通常の単一化を試みる。成功すれば 3 へ、失敗すれば 5 へ。
 - d. 残りの選択肢について通常の単一化を試みる。成功すれば 3 へ、失敗すれば 5 へ。

5. 修正プランナに失敗メッセージを出して終了。

主な処理は、1 で CCP に依存する選択だけを取り消し、3 と 4 でそれらの選択点について再計算することである。CCP に依存する選択だけをやりなおすため、CCP の後で行なった選択のすべてを取り消す時系列バックトラックより効率的な書き換えが実現できる。5 は書き換え失敗の場合に相当し、このときシステムは別の修正規則を適用するか、別の問題点の解消を試みる。

つぎに、4.1 項の修正例を使って処理 2, 4b の意味を説明する。草稿 (2b) では第 2 文の主題「資料室は」が省略されている。これは、選択 $\langle(n_seg), FD(13/2)\rangle$ を変更した結果、(2a) の最初の文が 2 つに分割され、 $\langle(s_seg\ sent), FD(25/1)\rangle$ (図 7) 中の *external* の評価値が *document*#1 から *library*#1 に変わったためである。この場合、システムは FD(25/1) から FD(25/2) に選択を変更しなければならない。このように、選択間には *external* を介しての依存関係が存在するが、これらはネットワーク中に明示することができない。2 と 4b は *external* を介した依存関係を管理するための処理である。

5 評価

一般に、TMS で用いられている DDB は、むだな再計算を防ぐことができる半面、ネットワーク管理のためのコストが無視できない。そこで、本節で述べたアルゴリズムの有効性を実験システムを用いて定量的に評価した。

システムはすべて Prolog 上に実装し、ネットワークの更新には *assert*, *retract* を用いた。比較対象として時系列バックトラックと Elhadad の手法を用いた。Elhadad の生成システム FUF では、FUG の単一化に失敗したとき効率よく CCP まで戻れるように、特定の素性が単一化に失敗したときの戻り先の候補を *bk-class* という特殊なクラスとして文法中に記述することができる [3]。この手法を我々の例題に応用するためには、「係り受けの深さが深すぎる」という問題が生じたときの戻り先の候補として図 5 の選択点 FD(13) を *bk-class* に指定すればよい。

実験では、図 3 に示した WFD を入力し、草稿 (2a) の生成を経て草稿 (2b) を生成するまでの計算量を比較した。草稿 (2b) の生成に到達するまでに生成した草稿の数とそれに要した素性の単一化の回数を表 2 に示す。

表 2 生成した草稿の数と必要な単一化の数の比較

	草稿の数	単一化の数
時系列バックトラック	17	10855
<i>bk-class</i>	3	2401
本手法	2	1177

本節で述べたように我々の手法では 2 度目の生成で草稿 (2b) が得られる。しかしながら、*bk-class* を用いると、草稿 (2b) を生成する前に (2b') のような unnecessary 草稿を生成してしまう。

- (2b') あなたが探している資料は正門にもっとも近い建物の 4 階にある資料室に保管されています。ただし、資料室は夏休みの間は利用できません。夏休みは 7 月 21 日から 8 月 31 日までです。

これは、*bk-class* に指定された選択点 FD(13) が第 2 文にも存在するためである。また、名詞句に関する選択点のように頻繁に生じる選択点が *bk-class* に指定された場合、(2b') のようなむだな草稿の生成はさらに増えると予測できる。草稿の評価に要するコストを考慮に入れた場合、この問題は深刻である。これに対し、我々の手法では WFD におけるパスと選択肢の組を用いて CCP を特定するため、1 度で適切な CCP (この例では第 1 文の選択点 $\langle(n_seg), FD(13/2)\rangle$) に戻ることができる。また、*bk-class* を用いる場合、時系列にそって選択を取り消すため、CCP から新たな草稿を生成する過程でもむだな再計算が多くなる。

つぎに、CPU 時間の比較を表 3 に示す。表は草稿 (2b) を生成するまでの実行時間を示しており、ネットワーク

管理のコストを支払っても、本手法の方が bk-class を使うより効率的であることがわかる。

表3 CPU 時間の比較

	初稿	第2稿	第3稿	合計
時系列バックトラック	32.1	11.0	35.4	440.3
bk-class	32.0	30.5	31.5	94.0
本手法	34.5	23.5	—	58.0

(SONY NEWS 3860; [sec])

6 おわりに

本稿では、文章生成における推敲機能の重要性について議論し、how-to-say に関する推敲機能を備えた生成モデルとその実装方法を提案した。

文章生成を探索問題と見なした場合、各選択点で局所的に優先された選択の集合がシステムの出力に対応する。しかしながら、係り受けのあいまい性や文の複雑さなど、いくつかの評価基準については、評価に必要な情報が生成の途中では部分的にしか得られないため、局所的な選好規則としては実現しにくい。本稿では、選好規則が与える不十分な優先度を修正規則によって修正するモデルを提案した。このモデルでは、選好規則が局所的な情報しか用いないため、個々の規則を単純化することができ、システムの実装と保守に有利である。どのような評価基準を評価部で評価し、修正規則として実現するかについては、今後具体的な応用システムの構築を通じて考察する必要がある。また、修正規則は、選好規則と同様、経験から得られるヒューリスティクスであり、その有効性についても検証が必要である。

推敲機能を備えたモデルの実装についてはこれまでほとんど議論されていない。これについては文章の評価の難しさが一つの要因にあげられる。評価の実現については本稿でもほとんど議論しなかったが、我々のモデルではある特定の評価基準についてのみ草稿を評価するため、あらゆるレベルの評価を仮定した他のモデルに比べ、評価部の実装は比較的容易であると考えられる。ただし、係り受けのあいまい性の評価については高度な解析処理技術の開発を待たねばならず、現在のシステムでは人手で行なっている。

草稿の書き換えは探索空間における選択の変更として実現することができる。本稿では、選択を効率的に変更するために、DDB のメカニズムを FUG の単一化の制御に導入する手法を提案した。DDB の実現は選択の順序を動的に入れ換えることに相当する。この点で選択の順序を規定しない単一文法は DDB に適している。書き換えの際には照応表現など、選択の変更の文脈的な影響も考慮する必要があるが、これについては external を介した選択間の依存関係を用いて対処できることを示した。DDB を導入した生成システムの例として Elhadad の FUF があげられる。5節では、FUF と我々の手法をおもに時間的効率の面から比較し、我々

の手法が効率の向上に有効であることを示した。また、FUF では、バックトラックの戻り先という制御に関する情報を bk-class として文法中に記述する必要があるが、これは文法の開発や保守の点で望ましくない。本稿で示した手法はこの問題も解決している。

参考文献

- [1] D. E. Appelt. TELEGRAM: A grammar formalism for language planning. In *Proceedings of the International Joint Conference on Artificial Intelligence*, pp. 595-599, 1983.
- [2] J. Doyle. A truth maintenance system. *Artificial Intelligence*, pp. 231-272, 1979.
- [3] M. Elhadad and J. Robin. Controlling content realization. In R. Dale, E. Hovy, D. Rösner, and O. Stock, editors, *Aspects of Automated Natural Language Generation*, pp. 89-105. Springer-Verlag, 1992. Lecture Notes in Artificial Intelligence Vol. 587.
- [4] E. H. Hovy. *Generating Natural Language under Pragmatic Constraints*. Lawrence Erlbaum Associates, 1988.
- [5] M. Kay. Functional Unification Grammar: A formalism for machine translation. In *Proceedings of the International Conference on Computational Linguistics*, pp. 75-78, 1984.
- [6] W. C. Mann and S. A. Thompson. Rhetorical Structure Theory: Description and construction of text structures. In G. Kempen, editor, *Natural Language Generation*, chapter 7, pp. 85-96. Martinus Nijhoff, 1987.
- [7] M. M. Meteer (Vaughan) and D. D. McDonald. A model of revision in natural language generation. In *Proceedings of the Annual Meeting of the Association for Computational Linguistics*, pp. 90-96, 1986.
- [8] D. R. Scott and C. S. de Souza. Getting the message across in RST-based text generation. In R. Dale, C. Mellish, and M. Zock, editors, *Current Research in Natural Language Generation*, chapter 3, pp. 47-74. Academic Press, 1990.
- [9] M. Yazdani. Reviewing as a component of the text generation process. In G. Kempen, editor, *Natural Language Generation*, chapter 13, pp. 183-190. Martinus Nijhoff, 1987.
- [10] 桃内佳雄, 高橋晃. 生成文章の結束性の良さを考慮した文章生成に関する研究. 言語情報処理の高度化研究報告 7, 言語情報処理の高度化の諸問題, pp. 359-378, 1989.
- [11] 徳永健伸, 乾健太郎. 1980 年代の自然言語生成 1-3. 人工知能学会会誌, Vol. 6, No. 3-5, 1991.