

読み情報をもつ日本語エディタの作成と評価

畠山 勉 角田 博保
電気通信大学 情報工学科

かな漢字変換型ワープロの利用者は、漢字を入力する際、漢字かな混じりの表記を読みで指定する。この読みを捨てずに保持し利用すれば、漸増探索指令や動的展開を、英文に対する時と同じ感覚で実行でき、かな漢字変換のやり直し(再変換)の効率も増す。本研究ではこの機能をもつエディタ(JEM)を漢字かな混じりテキストを保持するバッファ(漢字層)と読み情報を捨てずに保持するバッファ(読み層)の多層構造で実現することを考え、GNU Emacs 系のエディタである Nepoch 上に上位互換に近い形で実装した。機能評価実験を行なったところ、JEM は漸増探索指令において従来の方法より有利とわかった。

The design and evaluation of a Japanese Editor with pronunciation-information

Tsutomu Hatakeyama Hiroyasu Kakuda

Department of Computer Science, The University of Electro-Communications

When one uses a word-processor of Kana-to-Kanji conversion type, he/she specifies a sentence with its pronunciation. If an editor preserves and uses these pronunciation-information, a user can use I-search(incremental search) or dynamic abbreviate expansion in Japanese sentences as same as in English ones, and can perform Kana-to-Kanji re-conversion efficiently. So the authors have designed and developed an editor, called "JEM", with multi-layered-structure of Kanji-layer (contains normal text) and Yomi-layer (contains pronunciation-information). JEM is implemented on Nepoch that is one of the GNU Emacs series. The authors have evaluated its effectiveness through experiments on I-search.

1 はじめに

日本語の文章を作成するのに かな漢字変換が広く利用されている。かな漢字変換では、利用者は 漢字かな混じりの表記を直接指定することはできないので、読みを入力して、適当な文節に区切ったり、複数の変換候補から一つを選ぶ作業を経て、目的の漢字かな混じり文を得る。そして、基となった読みは捨ててしまう。

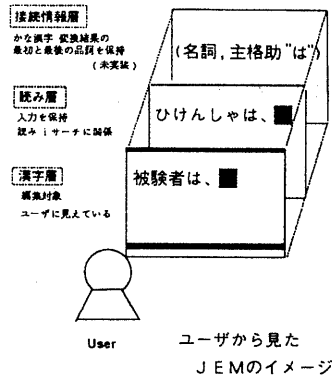
一方で、利用者はテキスト内での移動や探索のためにしばしば逐次探索指令 (incremental search, 以下、漸増探索と略す) を使う。しかし日本語を含む文章で日本語の部分に対して、それを行なうのは、変換が成功しそうな、ある程度の文字数まで入力して漢字に変換する作業が必要である。そのため、ある程度以上離れた位置への探索ならばともかく、ちょっとした移動のために英文のときと同じ頻度で使われることはない。また、プログラムなど、同じ名前が頻繁に登場するテキストに対して特に有効な動的展開による入力についても、やはり、変換が成功しそうな、ある程度の文字数まで入力して漢字に変換する作業が必要なので、英文のときと同じ頻度で使われることはない。

しかし漢字かな混じりの表記のレベルで照合させようとするのではなく、読みのレベルで照合させるならば、変換が成功しそうなところまで入力する必要はなくなる。

それには漢字の入力の時に、いつも捨てられていた読みと文節区切り (合わせて読み情報と呼ぶ) を保持する機構が必要になる。本研究では漢字かな混じりテキストを保持するバッファ (漢字層) と読み情報を捨てずに保持するバッファ (読み層) の多層構造を持ったエディタを、GNU Emacs 系のエディタである Nepoch 上に実装した。これを JEM (Japanese Editor with Multi-layered Structure) と名付ける。

2 JEM V2 の設計

本節では JEM の外部仕様について述べる。



品詞：文字列と接続情報が同じになるものを同士でグループ化したグループ名

図 1: ユーザからみた JEM のイメージ

2.1 多層構造

2.1.1 ユーザから見た多層構造

ユーザから見ると、JEM は次のように振舞わなければいけない (図 1)。

- (1) 各文節の入力時の読みを理解している。
- (2) 各層のカーソルが同期して動いている。
- (3) 漢字層の変更は各層に反映されている。

2.2 漢字層と読み層

漢字層は利用者が普段 編集対象としているバッファのことである。そして、漢字層の読み情報を利用できる形で保持するバッファを読み層とよぶ。読み層に関連づけられたファイル (読み層ファイル) は漢字層に関連づけられたファイルの名前に拡張子 “.yomi” を付加した名前がつけられる。もし読み層ファイルが存在しない場合は、Wnn の逆変換を用いて漢字層ファイルから自動生成される。

2.3 削除・挿入の扱い

一方の層のオフセットに両層の対応する文字列の長さの比率を掛けて四捨五入した値を他方の層のオフセットとすることで、両層の対応づけを決定する (図 2)。例で、漢字層の “第” を削除すると、読み層の “だい” が同期して削除される。漢字層の “1” と “回” の間に文字列を

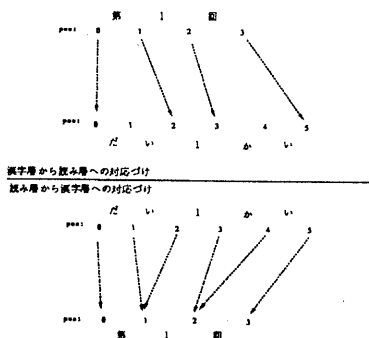


図 2: 両層の対応づけ

挿入すると、それは読み層の“1”と“かい”の間に同期して挿入される。¹

2.4 読みによる漸増探索

英字版と同じ使用感を得ることを目標としている。

2.4.1 読みによる漸増探索の外部仕様

“C-s”に続けて読みをローマ字で打鍵すると順方向の探索が始まる。ローマ字はかな漢字変換され、探索文字列の平仮名が確定するたびに探索が実行されることが漸増探索の最大の特徴である。

逆方向に探索したければ“C-r”で起動するが、順方向探索を起動した後、いつでも逆方向探索に切り換えることができる。全角半角は区別しない。また、探索文字列は通常文節の先頭にあたる文字から入力されるのが普通であるので、文節先頭からのみ照合するオプションを設けている。複数行に渡る探索文字列については、文字と文字の間に0または1個の改行コードが入っても照合するので、改行を意識しないで探索できる。入力モードをトグルするキーで切換えて入力することで、探索文字列にアルファベットを加えることが可能である。探索文字列に英字を含む場合、大文字小文字の区

¹ 図の例は、たまたま両層の対応づけがうまくいっている場合だが、「第1位」と「だい1い」の時は、「第」、「1」、「位」に対し、「だ」、「い1」、「い」が対応づけられるので、漢字層の“1”を削除すると読み層の“い1”が削除されてしまう。

別することも可能である。正規表現は指定することはできない。

2.5 再変換

仮名漢字変換の利用者はしばしば思考の流れを中断したくないなどの理由で大して確認もせず確定し、どんどん先に進む。そのようなとき、後で直す作業が必要になる。最も単純な直し方は誤った部分を削除して再度入力することである。しかし、変換結果は望んだものでなくても、その読みは望んだものと同じ読みなのであるから、読みだけ再利用すれば、再び打ち込む必要がない。読みを再利用して変換することを再変換と呼ぶことにする。

再変換は読み情報を使わなくても可能であるが、その場合、精度が高くないのが難点である。しかし JEM では入力時に保存した読み情報を利用して再変換することができる。

2.6 読みによる動的展開

動的展開 (dynamic abbreviation expansion) は、現在入力中の単語 (文節) と同じ表記で始まる単語を現在編集中のテキストなどから探して、補完する入力方法である。

動的展開はプログラムなど、同じ単語を頻繁に使うテキストの編集で特に有効である。

日本語では漢字を含む語を動的展開しようという場合、入力した文字列 (ローマ字ないし平仮名) を漢字に変換しなければ対応づけることができない。しかしながら JEM は読みを保持しているので、そのまま読み層において略称と展開結果の対応づけを行ない、展開結果に相当する位置の漢字層文字列を展開結果とすることで動的展開をごく自然に実現している。

JEM の読みによる動的展開は日本語入力モードで起動する。日本語入力モードで文字列の先頭の方を打ち込んでから“M-/”を打鍵すると読みによる動的展開が起こる。

展開の次候補を得るには「スペースバー」または「メタキーを押しながら /」を打鍵する。現在ポイントから逆方向に順に探される。バッファの先頭まで行ったら、現在ポイントから順方向に順に探される。なんらかの、次にやりた

いコマンドをやれば確定されてそのコマンドが実行される。

また、新たな展開候補となるための条件は、「これまでの全ての展開候補と漢字部分が異なること」としている。なお、入力はローマ字仮名変換の途中までで良い。(例“き sh”は“きしゃ”または“きしゅ”または“きしょ”)として展開される。)

3 JEM V2 の実装

JEM の試作版 [2] では文字単位の編集ができない(文節単位のみ)、漸増探索が高速でないという欠点があった。それらの反省を踏まえて実装した。

3.1 実装の基本方針

複数のバッファによる多層構造を用いて、読み情報・文節区切り情報を保持する。これは、編集対象のバッファ(漢字層バッファ)を読み情報を保持しているバッファ(読み層バッファ)と関連づけることで行なう。

また、計算機資源との兼ね合いとして、プログラムサイズが大きくなったり、メモリを多く消費することがあっても、原則として高速化と気のきいたユーザインタフェースのほうを優先する。

3.2 JEM V2 のシステム構成

JEM は X11 R5 の下で Nepoch 1.1² に手を加えたものとして実現されている。日本語環境としては Wnn V4 を、日本語フロントエンドプロセッサには EGG (たかな) を用いている。

JEM の部分のソースコードは Emacs Lisp で現在約 1 万行余りである。そのうち、読みによる漸増探索の部分は約 1650 行、動的展開の部分は約 1200 行である。図 3 に JEM V2 のシステム構成を示す。

²JEM の母体となるエディタは、Nepoch の第 1.1 版を使用した。それは、GNU Emacs の X11 環境下の使用を前提にして、X11 のサービスを利用できるように GNU Emacs に対しいくつかの拡張を行なった epoch の第 3.2 版の日本語版である。epoch は GNU Emacs と互換性を持ち、従来の Emacs Lisp コードはそのまま動作する。

特徴としては、X11 環境のマルチウィンドウのサポート、X11 の描画属性を使ったテキストリージョンのマーキング(color, underline, stipple, font)などが挙げられる。

Nepoch は広く使われており、ソースコードの入手・改変・再配布が自由である。今回 JEM を Nepoch 上に実装することにしたのは、文節区切りが表示できるからである。文節単位で編集する場合、文節区切りが表示されない、次の予想がつかず使いにくい [2]。epoch の表示属性機能を使わないのであれば、JEM は X Window の下でなくても実現可能である。

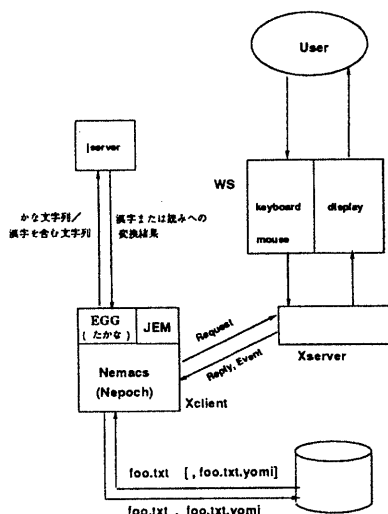


図 3: JEM のシステム構成

3.3 読み情報などの持たせ方

読み層の第 i 行を漢字層の第 i 行に対応させる。

- (1) (読み層バッファ) ::= (読み情報の行)*
 - (2) (読み情報の行) ::= {(文節レコード)}*(改行文字 | バッファの終り)
 - (3) (文節レコード) ::= (レコード先頭を表す 2 バイト文字)(漢字層での長さを表す 2 バイト文字 2 個)(読み方情報)(レコード末尾を表す 2 バイト文字)
 - (4) (読み方情報) ::= (ひらがな | 記号)*
- ここで (読み方情報) は日本語フロントエンドプロセッサを介していない時の自己挿入文字のときは、そのものの 1 文字以上の接続であり、そうでないときはローマ字仮名変換の結果 (ひらがなや句読点などの 1 文字以上の接続) である。

3.3.1 文節レコード

文節レコードと読み方文字列について。

1 文節の、JEM の読み層における表現を文節レコードとよぶことにする。その文節の読み情報を持っている。文節レコードは次の構成要素からなる。

- (1) 漢字層での漢字表記 (文字列)
- (2) 読み方文字列 (文字列)
- (3) 漢字層での漢字表記の長さ (バイト単位の非負整数)

漢字層が次の 3 行からなるテキストを考える。

```
彼は新しい  
gnus.el を  
インストールしました。
```

このとき、読み情報として

4	かれは	6	あたらしい
5	gnus	5	el を
22	いんすとーるしました。		

を持たせる。

3.4 エディタ (Nepoch) 本体と JEM モジュールとのインタフェース

3.4.1 読み層の同期の取り方

epoch のコンパイル時に “HAVE_LEIF” オプションを指定して得られる、before-change-function と after-change-function の 2 つの hook に、漢字層の変更を読み層に反映させる hook (jem-synchronous-hook) を登録し、それが呼ばれることで実現する。

jem-synchronous-hook は漢字層の変更が「削除」ならば読み層の対応リージョンを計算して削除を行ない、「挿入」ならば逆変換を行なって読み層の対応ポイントに挿入する。

4 JEM の評価実験

読みによる漸増探索について、評価実験を行った。

4.1 実験 1 読みによる漸増探索と 変換型探索の比較実験

(実験 1h) 変換型探索, (実験 1y) 読みによる漸増探索。

キーストローク数からは読みによる漸増探索のほうが速いはずだが、実際に確かめることが目的。

4.1.1 実験に使うコマンドの外部仕様

読み 漸増探索の外部仕様は別項の通りである。また、変換型 漸増探索は Wnn+EGG の利用者に一般に良く使われているものを多少手直ししたものである。それは C-s で変換型 漸増探索に入り、C-k に続いて探索文字列として漢字を入力することができる。(ミニバッファを改めて日本語入力状態にするキーは不要にしてある。漢字入力には EGG を用いている。) ミニバッファの変換を確定した後、リターンキーで探索が起こる。

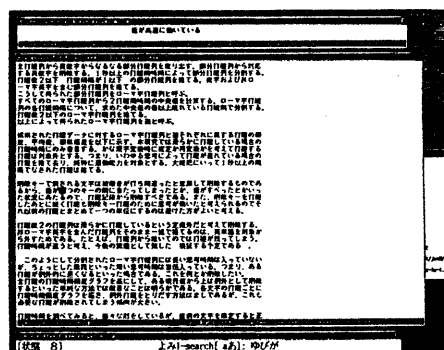


図 4: 「指が高速に動いている」を探している途中

4.1.2 方法

被験者 被験者は、筆者らが所属するの研究室の 9 名 (男性 6 名・女性 3 名) で、年齢は全員 22 歳前後である。全員とも普段から Emacs を使用しているが、他のエディタの方を使い慣れている被験者がいた。また、通常は漸増探索を使うことは滅多にないという被験者もいた。

原稿 英字と、句読点以外の記号を削除し、さらに改行コード以外はすべて全角文字に変換した 750 行分、各行 70 カラム (35 字) の日本文のテキストファイルを用意した。題材は ヒューマンインタフェース関係の論文 (4 編) からとった。問題文の読み層ファイルは逆変換で生成したものを入力時の読みとなるように、手動で補正した。

ターゲットの文字列

必ず文節の先頭から始まる。平均長 10.76 文字 (標準偏差 2.56 文字) 文字列の開始位置は行頭・行末の 5 文字分を外して乱数で決めた。

ワークステーション Sun Sparc Station SLC と ELC を併用した。パフォーマンスの差が問題になるようなタスクではない。

ディスプレイ 17 inch, グレースケールディスプレイ。フォントは本文の部分と問題表次窓の部分には a18 と kanji18 を、ミニバッファの部分には 12x24 と k24 を使用した。

本文用のウィンドウのサイズは 112x36 に統一し (バッファには 36 行、114 カラム表示される)、ミニバッファはその下方においた。本文用のウィンドウの上に問題文用のウィンドウを描いておいた (図 4)。

収集されるデータ キーとその押された時刻 (精度は 30 msec 程度)、キーシーケンスから決定されたコマンド名、そのコマンドが実行される前の画面の状態に関するデータ (画面の開始ポイント値、現在ポイント値、画面内での行番号、現在ポイントのカラム値、現在行のカラム数、ポイントの前/後の文字)、バッファ名

手続き

(1) **練習** 被験者は十分な練習をしてから実験に入った。特に変換型の探索をよく練習してもらった。

(2) **課題** 原稿修正タスクを模擬した次のタスクを課す。画面上にミニバッファとは別に横に細長い窓 (問題窓) が上部に置かれ、そこに問題として「行くべき場所」を指示される。その場所まで、(変換型漸増探索・読みによる漸増探索) のいずれか 1 つを主に使用して、必要ならば基本コマンドを補助的に使いながら、指示された場所へ移動する。そこで削除キー C-d (コントロールキーを押しながら d) または C-h を押して正しい答ならば正しい答のリージョンに下線がつき (「正解」の意)、次の問題に進むことができる。

(3) **正解基準** 問題の文字列上とその前後 1 文字の範囲とした。

(4) **問題数** 問題は 1 セッションに 50 問であり、被験者は全問を順に解かなければいけない。実際の所要時間は 10 分程度であった。

4.1.3 実験 1 の結果

C-s を押してから問題を獲得するか獲得してしまってから C-s を押すかを特定できないので、こうした。第 n 問の正解部分に対する C-d

(または C-h.) から次の問題の正解部分に対する C-d (または C-h.) までを「第 n 問にかかった時間」に着目する。データは C-s C-k の C-k の分を補正していないままの値である。得られたデータとそのサマリを巻末の表 1 と 2 に示す。ただし下記の除去基準が施されている。

次のデータは捨てた

- (1) 1 セッション (50 問) の初めの 5 問
- (2) サーチ中に C-g または DEL を打鍵したもの (ミス直し) とコマンド “abort-recursive-edit” のキーが打鍵された問題。
- (3) ターゲットが順方向にない問題 (約半数) と C-r を打鍵したもの。
- (4) 差が 20 秒 (20000 msec) よりも大きかった問題 (4 件)。これは被験者が実験者にやりかたについての質問をしていたときなどがこれに当るからである。

4.1.4 実験 1 の考察

有効な問題数は 7 から 19 しか残らなかったが、1 問当たり 1.7 秒から 6.5 秒の範囲で読み漸増探索の方が速い (分散は 3.0 から 49.9)。変換型は変換のため以外に不要な打鍵が 3 打鍵ほどあるが、これを 1 問当たり 0.6 秒として補正してもやはり読み漸増探索の方が速い。変換のための打鍵が必要か不要かの違いがそのまま現れた結果といえる。

4.2 実験 2 読み漸増探索による移動とマウスによる移動の比較実験

(実験 2m) マウス型 (実験 2y) 読み型。

ターゲットが画面内にある場合、打鍵レベル模倣型 [3] からはマウスによる移動方法が有利と推測されるが、それが読み漸増探索と比べてどのくらい有利なのかを調べる。

4.2.1 実験に使うコマンドの外部仕様

読み 漸増探索の外部仕様は別項の通りである。マウスによる移動はマウスの左ボタンを押すとマウスカーソル (鉛筆の形をしていて大きさ約 1cm) の位置に、カーソルが移動する。

4.2.2 方法

被験者 実験1に同じ。

原稿 原稿は実験1に同じであるが、実験1では問題を画面内に問題窓として提示したのに対し、実験2ではハードコピーにオレンジ色の蛍光ペンを施したものとして提示する。

ターゲットの文字列 実験1と同様の基準で新たに選び直した。平均長 11.32 文字 (標準偏差 2.96 文字) である。

ワークステーション・ディスプレイ・収集されるデータ 実験1に同じ

手続き

- (1) 練習 被験者は十分な練習をしてから実験に入った。特に変換型の探索をよく練習してもらった。
- (2) 課題 原稿修正タスクを模擬した次のタスクを課す。ハードコピーに示された問題 (50 問で行番号順に整列されている) を「行くべき場所」を指示される。その場所まで、(マウスによる移動・読み漸増探索による移動) のいずれか1つを主に使用して、必要ならば基本コマンドを補助的に使いながら、指示された場所へ移動する。そこで削除キー C-d (コントロールキーを押しながら d) または C-h を押して正しい答えならば正しい答のリージョンに下線がつき (「正解」の意)、次の問題に進むことができる。

- (3) 正解基準・問題数 実験1に同じ

4.2.3 実験2の結果

実験1と同様に第 n 問の正解部分に対する C-d (または C-h.) から次の問題の正解部分に対する C-d (または C-h.) までを「第 n 問にかかった時間」に着目する。得られたデータを巻末の表3に示す。ただし下記の除去基準が施されている。

次のデータは捨てたサーチ中に C-g または DEL を打鍵したもの (ミス直し) のキーが打鍵された問題。

4.2.4 実験2の考察

画面内にターゲットがあるならば、打鍵レベル模型の示す通り、ほとんどの場合マウスのほうが速かった。

しかし、個人差はあるものの、距離 (次の問題のターゲットまでの行数) が10を超えると、ターゲットが画面内に入っているにもかかわらず読めないことが読みとれる。

5 今後の課題

今後の課題としてはまず、実験データの精密な解析と再実験が必要だと思う。また、システムの実装面では、品詞情報の保持と、標準にはありながら、JEM では削られている narrowing 機能を実現すべきである。

さらに主要なメジャーモードの JEM 化として、Jem $\text{\LaTeX}$ Mode ($\text{\LaTeX}$ Mode と JEM Mode を共存させたもの)、Jem 日本語 C Mode を考えている。後者はプログラムの識別子として、動的展開が多用されることが考えられるので、有力である。

読みによる動的展開は現在の外部仕様がベストであるとは思えない。今後、十分な使用を経て次の点について、より妥当な仕様を探すべきである (おくりがな や 助詞も展開結果に付随するのが良いか悪いのか、展開結果の文節の長さ (現在は1)、展開結果があまりに短い時は展開しないのが良いか悪いのか、前の候補に戻ることはできないのは良いか悪いのか、など)

参考文献

- [1] 角田博保: 多層テキスト構造を持つ構造エディタ, 第27回プログラミングシンポジウム (1986.1)
- [2] 北村譲治: 多層構造の日本語エディタの試作, 平成3年度電気通信大学 情報工学科 卒業研究
- [3] Card, S.K., Moran, T.P. and Newell, A. The Psychology of Human Computer Interaction, Lawrence Erlbaum Associates, 1983

