

(12) 特許協力条約に基づいて公開された国際出願

(19) 世界知的所有権機関  
国際事務局

(43) 国際公開日  
2018年5月31日(31.05.2018)



(10) 国際公開番号

WO 2018/097022 A1

(51) 国際特許分類:  
G06F 17/28 (2006.01)

東京都小金井市貫井北町4-2-1 Tokyo (JP).

(21) 国際出願番号: PCT/JP2017/041249

(22) 国際出願日: 2017年11月16日(16.11.2017)

(25) 国際出願の言語: 日本語

(26) 国際公開の言語: 日本語

(30) 優先権データ:  
特願 2016-227583 2016年11月24日(24.11.2016) JP

(71) 出願人: 国立研究開発法人情報通信研究機構 (NATIONAL INSTITUTE OF INFORMATION AND COMMUNICATIONS TECHNOLOGY) [JP/JP]; 〒1848795

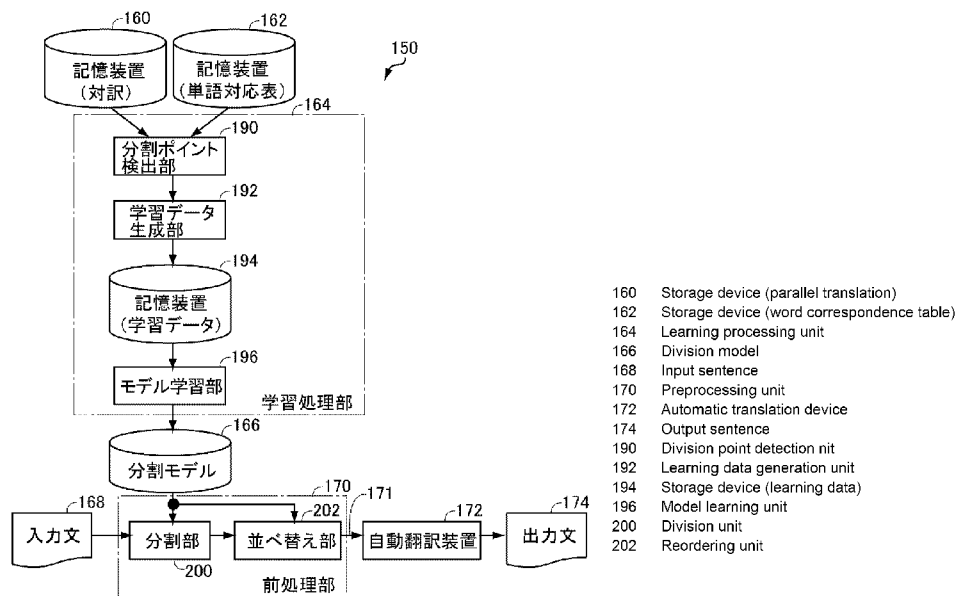
(72) 発明者: 富士 秀(FUJI, Masaru); 〒1848795 東京都小金井市貫井北町4-2-1 国立研究開発法人情報通信研究機構内 Tokyo (JP). 内山 将夫(UCHIYAMA, Masao); 〒1848795 東京都小金井市貫井北町4-2-1 国立研究開発法人情報通信研究機構内 Tokyo (JP).

(74) 代理人: 清水 敏(SHIMIZU, Satoshi); 〒5300002 大阪府大阪市北区曽根崎新地二丁目5番3号 堂島TSSビル4階 Osaka (JP).

(81) 指定国(表示のない限り、全ての種類の国内保護が可能): AE, AG, AL, AM, AO, AT, AU, AZ, BA, BB, BG, BH, BN, BR, BW, BY, BZ, CA, CH,

(54) Title: AUTOMATIC TRANSLATION PATTERN LEARNING DEVICE, AUTOMATIC TRANSLATION PRE-PROCESSING DEVICE, AND COMPUTER PROGRAM

(54) 発明の名称: 自動翻訳パターン学習装置、自動翻訳の前処理装置、及びコンピュータプログラム



(57) Abstract: [Problem] Provided is an automatic translation pattern learning device for learning patterns for, prior to automatic translation, transposing constituent component strings of an original text into an order suitable for automatic translation. [Solution] A learning processing unit 164 includes: a dividing point detection unit 190 for detecting in a parallel translation, as the division position of two sentences, the position at which the corresponding relationship between



CL, CN, CO, CR, CU, CZ, DE, DJ, DK, DM, DO,  
DZ, EC, EE, EG, ES, FI, GB, GD, GE, GH, GM, GT,  
HN, HR, HU, ID, IL, IN, IR, IS, JO, KE, KG, KH,  
KN, KP, KR, KW, KZ, LA, LC, LK, LR, LS, LU, LY,  
MA, MD, ME, MG, MK, MN, MW, MX, MY, MZ,  
NA, NG, NI, NO, NZ, OM, PA, PE, PG, PH, PL, PT,  
QA, RO, RS, RU, RW, SA, SC, SD, SE, SG, SK, SL,  
SM, ST, SV, SY, TH, TJ, TM, TN, TR, TT, TZ, UA,  
UG, US, UZ, VC, VN, ZA, ZM, ZW.

- (84) 指定国(表示のない限り、全ての種類の広域保護が可能): ARIPO (BW, GH, GM, KE, LR, LS, MW, MZ, NA, RW, SD, SL, ST, SZ, TZ, UG, ZM, ZW), ユーラシア (AM, AZ, BY, KG, KZ, RU, TJ, TM), ヨーロッパ (AL, AT, BE, BG, CH, CY, CZ, DE, DK, EE, ES, FI, FR, GB, GR, HR, HU, IE, IS, IT, LT, LU, LV, MC, MK, MT, NL, NO, PL, PT, RO, RS, SE, SI, SK, SM, TR), OAPI (BF, BJ, CF, CG, CI, CM, GA, GN, GQ, GW, KM, ML, MR, NE, SN, TD, TG).

添付公開書類:

- 国際調査報告(条約第21条(3))

the order of occurrence of constituent components of a sentence in the original language and the order of occurrence of constituent components of the corresponding translated sentence changes; a learning data generation unit 192 for generating, as learning data for mechanical learning, data comprising a constituent component string of a continuous prescribed number of items within the sentence in the original language, and correct answer information indicating whether the division position is within the constituent component string; and a model learning unit 196 for, when an input constituent component string of a prescribed number of units in a sentence in the original language is provided by machine learning, learning a division model 166 for determining whether the division position is at the prescribed position in the constituent component string, using the generated learning data.

(57) 要約: 【課題】自動翻訳に先立って、原文の構成部品列を自動翻訳に適した順序に入れ替えるパターンを学習する自動翻訳パターン学習装置を提供する。【解決手段】学習処理部164は、対訳のうち、原言語の文の構造部品の出現順序と、対応する訳文の構造部品の出現順序との対応関係が変化する位置を両者の文の分割位置として検出する分割ポイント検出部190と、原言語の文内の、連続する所定個数の構成部品列と、その中に分割位置があるか否かを示す正解情報とからなるデータを機械学習の学習データとして生成する学習データ生成部192と、生成された学習データを用いて、機械学習により、原言語の文の所定個数の入力構成部品列が与えられたときに、その中の所定位置に分割位置があるか否かを判定する分割モデル166の学習を行うモデル学習部196を含む。

## 明 細 書

発明の名称：

自動翻訳パターン学習装置、自動翻訳の前処理装置、及びコンピュータプログラム

### 技術分野

[0001] この発明は自動翻訳技術に関し、特に、定型性の高い文書に対する自動翻訳の精度を向上させる技術に関する。

### 背景技術

[0002] 自動翻訳では、定型的な文書を翻訳することが多い。典型的には特許出願書類のうちの権利範囲を定める部分の翻訳である。例えば、米国における先行技術調査を日本の技術者が行う場合、問題となるのはクレームの部分である。英語の文で記載された米国特許のクレームを技術者、弁護士、及び弁理士等がそのまま理解できれば理想的であるが、限られた時間に大量の英文クレームを読む負担は大きい。そのため、英文クレームを日本語に精度高く翻訳する技術が望まれている。

[0003] 特許文献１において、そうした技術が提案されている。特許文献１に記載された技術では、英文クレームと、その英文クレームに対応する日本語クレームが共通した定型性を持っていることを利用する。

[0004] より具体的には、日英翻訳では、原文及び訳文の間で各構造部品の出現順序の対応関係が入替わることがしばしば起こる。例えば図１（Ａ）に示される対訳３０を考える。図１（Ｂ）に示すように、この対訳３０の構造部品列の対応関係を考えると、日本語の構造部品列４０、４２及び４４が、英語の構造部品列５４、５２及び５０にそれぞれ翻訳されており、構造部品列４０及び構造部品列４４の出現順序が、英語では構造部品列５０及び５４となって互いに入れ替わっている。図２（Ａ）に示される対訳７０についても同様である。日本語では構造部品列８０、８２及び８４がこの順序で出現しているのに対し、英語では構造部品列８０及び８４に対応する部分が入れ替わり

、構造部品列 9 0， 9 2 及び 9 4 の順序で出現している。

[0005] 特許文献 1 に開示された技術は、英文クレームと日本語クレームとのこのような関係を利用することで、自動翻訳の精度の向上を図っている。すなわち、予め英文クレームの構造部品と日本語クレームとの双方について、構造部品の並びのパターンを特定し、構造部品の対応関係（並べ替え）のパターンを特定する。このパターンにより、対応する英文クレームと日本語クレームとについて、構造部品の対応関係が特定できる。翻訳時に英文クレームが入力されると、英文クレームのパターンを特定し、そのパターンにしたがって英文クレームを構成要素に分割する。特定されたパターンに対応する日本語クレームのパターンを特定し、英文クレームの構成要素を日本語クレームのパターンに応じて並べ替える。このように構成要素を並べ替えた英文クレームを日本語に自動翻訳することにより、日本語への翻訳の精度が向上すると報告されている。

## 先行技術文献

## 特許文献

[0006] 特許文献1：特開 2 0 1 6 - 1 6 4 7 0 7 号公報

## 発明の概要

## 発明が解決しようとする課題

[0007] 特許文献 1 に記載された技術により、英文クレームを日本語クレームに精度高く翻訳することが可能になる。しかし、特許文献 1 に記載された技術では、英文クレームと日本語クレームとのパターンの作成及びそれらの構成要素の対応関係の特定はいずれも人手で行う必要があり、そのために自動翻訳の精度を向上させるためのコストが高いという問題がある。英文クレームを日本語クレームに翻訳する際のような、定型的な文章の翻訳をより精度高くかつ低コストで行うためには、上記したパターンの作成及び構成要素の対応関係をより低コストで行う必要がある。

[0008] それゆえに本発明は、定型的な文章を高い精度で自動翻訳するために、自

動翻訳に先立って、文の構成部品列を自動翻訳に適した順序に並べ替えるパターンを学習する自動翻訳パターン学習装置、そのパターンを用いて自動翻訳の前処理を行う前処理装置、及びそのためのコンピュータプログラムを提供することである。

### 課題を解決するための手段

[0009] 第1の局面に係る自動翻訳パターン学習装置は、第1の言語の文を第2の言語の文に自動翻訳により翻訳するために、第1の言語の文を複数の構造部品に分解して順序を並べ替えるパターンを学習する。この自動翻訳パターン学習装置は、第1の言語及び第2の言語の対訳を記憶するための対訳記憶手段と、対訳記憶手段に記憶された対訳の各々について、第1の言語の文の構造部品の出現順序と、対応する第2の言語の文の構造部品の出現順序との対応関係が変化する位置を第1の言語の文及び第2の言語の文の各々の分割位置として検出するための検出手段と、第1の言語の文の、連続する所定個数の単語列と、当該単語列内の所定位置に分割位置があるか否かを示す正解情報を付したデータを機械学習の学習データとして生成するための学習データ生成手段と、学習データ生成手段により生成された学習データを用いて、機械学習により、第1の言語の文の所定個数の入力単語列が与えられたときに、当該入力単語列内の所定位置に分割位置があるか否かを判定するための判定モデルの学習を行うための学習手段とを含む。

[0010] 好ましくは、学習データ生成手段は、第1の言語の文の、連続する所定の偶数個数の単語列と、当該単語列内の中央に分割位置があるか否かを示す正解情報を付したデータを機械学習の学習データとして生成するための手段を含む。

[0011] より好ましくは、学習データ生成手段は、第1の言語の文の、連続する所定個数の単語列と、当該単語列内の所定位置に分割位置があるか否かを示す情報、及び分割位置の前後の単語列の各々について、第2の言語の文における当該単語が属する構成部品の並べ替えパターンを示す情報とからなる正解情報を付したデータを機械学習の学習データとして生成するための手段を含

む。

[0012] 第2の局面に係る前処理装置は、第1の言語の文を第2の言語の文に自動翻訳により翻訳する前処理として、第1の言語の文を複数の構成部品列に分割し、その順序を並べ替える前処理装置であって、第1の言語の文の所定個数の入力単語列が与えられたときに、当該入力単語列内の所定位置に、第1の言語の文の分割位置があるか否かを判定するように、予め学習済の判定手段と、第1の言語の入力文が与えられたことに応答して、当該入力文を構成する、連続する所定個数の単語列を生成するための単語列生成手段と、単語列生成手段により生成された所定個数の単語列の各々を判定手段に与えることにより、入力文を構成する構成部品の分割位置を特定する分割位置特定手段と、分割位置特定手段により特定された分割位置で入力文を複数の構成部品に分割し、判定手段に対応して定められる並べ替えパターンにしたがって、当該複数の構成部品を並べ替える並べ替え手段とを含む。

[0013] 好ましくは、判定手段は、第1の言語の文の所定個数の入力単語列が与えられたときに、当該入力単語列内の所定位置に、第1の言語の文の分割位置があるか否かを判定し、当該入力単語列内の、分割位置前後の単語が属する構成部品が第2の言語の文のどの位置に配置されるかを示す分割位置情報を出力するように予め学習済の手段を含む。並べ替え手段は、分割位置特定手段により判定された分割位置で入力文を複数の構成部品に分割し、判定手段により出力された分割位置情報にしたがって、当該複数の構成部品を並べ替えるための手段を含む。

[0014] 第3の局面に係るコンピュータプログラムは、コンピュータを、上記したいずれかの装置の各手段として機能させる。

### 図面の簡単な説明

[0015] [図1]日英翻訳における構成部品列の並べ替えパターンの例を説明するための図である。

[図2]日英翻訳における構成部品列の並べ替えパターンの他の例を説明するための図である。

[図3]本発明の第1の実施の形態における分割ポイントの検出手法を説明するための模式図である。

[図4]本発明の第1の実施の形態における学習データの構成を説明する模式図である。

[図5]本発明の第1の実施の形態に係る自動翻訳システムの構成を示すブロック図である。

[図6]本発明の第1の実施の形態において、文の分割ポイントを検出するためのプログラムの擬似コードを示す図である。

[図7]図6に示す擬似コードにしたがうプログラムの動作を説明するための模式図である。

[図8]図5に示す自動翻訳システムにおいて、入力された文の分割を行うためのプログラムの制御構造を示すフローチャートである。

[図9]本発明の第2の実施の形態における学習データの構成を説明する模式図である。

[図10]本発明の各実施の形態に係る自動翻訳システム並びにそれを構成する学習装置及び前処理装置を実現するコンピュータの外観を示す図である。

[図11]図10に示すコンピュータのハードウェア構成を示すブロック図である。

### 発明を実施するための形態

[0016] 以下の説明及び図面では、同一の部品には同一の参照番号を付してある。したがって、それらについての詳細な説明は繰返さない。

[0017] [基本的考え方]

以下に本願発明の実施の形態を説明する前に、本願発明の基礎となっている基本的考え方を説明する。図3を参照して、「信号を発信できる通信装置を提供する」という日本語の文112と、「To provide a communication apparatus capable of sending signals」というその対訳文110との間の単語の対応関係を示す。日本語と英語との間では、一般的に語順が逆転するが、図3に示すように、両者の単語の対応関係は3個のブロック114, 11

6 及び 118 に分けられる。第 1 のブロック 114 及び第 3 のブロック 118 では、英語の語順と日本語の語順とが互いに逆の関係になっている。これに対して第 2 のブロック 116 では、英語の語順と日本語の語順とが部分的ではあるが同一の関係になっている。図 1 及び図 2 と、図 3 とを比較すると、第 1 のブロック 114、第 2 のブロック 116、及び第 3 のブロック 118 を構成する日本語の単語列及び英語の単語列が、それぞれ一団の構成部品列となっていて、翻訳時にはこれらがそれぞれまとめて互いに並べ替えられる事がわかる。そして、その境界の 1 つは、第 1 のブロック 114 と第 2 のブロック 116 の境界（これを分割ポイント B という。）であり、他の境界の 1 つは、第 2 のブロック 116 と第 3 のブロック 118 の境界（これを分割ポイント A という。）であることが分かる。

[0018] 以下の実施の形態では、日本語と英語との間の対訳の、このような性質を利用し、入力される文の分割ポイントを推定し、その分割ポイントで分けられた構成部品列からなるブロックの順番を、翻訳のターゲット言語の順番に並べ替え、その後に自動翻訳を行う。こうすることで、ターゲット言語での翻訳が精度よく行えるようになる。

[0019] なお、以下の説明は日本語から英語への翻訳を例にしているが、言語の組み合わせはこれには限定されない。特に、語順が変わる事が多い、日本語と欧州言語との間の翻訳に本発明を適用できる。

[0020] [第 1 の実施の形態]

<構成>

第 1 の実施の形態では、分割ポイントを推定するために統計的モデルを用いる。本実施の形態では統計的モデルとしてサポートベクターマシン（SVM）を用いる。

[0021] SVM を用いるにあたって、その学習データが問題となる。本実施の形態では、自動翻訳装置を用いて対訳の間の単語対応表を求め、その情報を利用して図 3 に示すような単語の順序の間の関係を調べ、分割ポイントの位置を特定する。対訳の間の単語対応表は、高い信頼性で求められることが分かっ



ている。

[0022] 図4に、学習データの構成を示す。図4を参照して、日本語から英語に翻訳する際に、図3のような関係を利用して、入力文112の内の2個の分割ポイント120及び122を検出する。入力文内で連続する4個の単語からなる単語列を全て生成し、その中央に分割ポイントが存在するか否かを調べ、その情報を各単語列に付す。想定される分割ポイント位置に左側（前側）から隣接する単語列を左隣接単語群、右側（後側）から隣接する単語列を右隣接単語群という。図4では、これら各単語群に、その中央位置に分割ポイントがあるか否かを示す情報（タグ）を正解データとして付す。こうしてできたデータ130の各行が学習データとなる。

[0023] 図5を参照して、第1の実施の形態に係る自動翻訳システム150について、日本語を英語に翻訳する事例を説明する。自動翻訳システム150は、学習データの元となる対訳を記憶する記憶装置160と、記憶装置160に記憶された対訳の各々について、自動翻訳装置を用いて求めた単語対応表を各対訳と関連付けて記憶する記憶装置162と、記憶装置160に記憶された対訳及び記憶装置162に記憶された単語対応表を用いて、SVMからなる分割モデル166の学習を行うための学習処理部164とを含む。

[0024] 自動翻訳システム150はさらに、日本語の入力文168を受けて、分割モデル166を参照して、入力文168に対して図3に示すような並べ替えが可能であれば、入力文168を分割ポイントで3個のブロックに分割し、図3に示す対応関係にしたがって各ブロックの順番を英語の順番に並べ替える前処理を行う前処理部170と、前処理部170により前処理が行われた入力文171に対して、日英の自動翻訳を行って英語の出力文174を生成するための自動翻訳装置172とを含む。自動翻訳装置172としてはどのようなものを用いても良いが、本実施の形態では、後述するように、単語の事前並べ替えを行うものを採用すると特に高い精度で自動翻訳が行えることが分かった。

[0025] 学習処理部164は、記憶装置160に記憶された対訳の各々の日本語に

ついて、記憶装置 162 に記憶された対応する単語対応表を用いて分割ポイントを検出するための分割ポイント検出部 190 と、分割ポイント検出部 190 により分割ポイントが検出された各対訳の日本語の文から、図 4 に示す構成にしたがった学習データを生成するための学習データ生成部 192 と、学習データ生成部 192 により生成された学習データを記憶するための記憶装置 194 と、記憶装置 194 に記憶された学習データを用いて、SVM からなる分割モデル 166 の学習を行うためのモデル学習部 196 とを含む。

[0026] 前処理部 170 は、入力文 168 を受けて、入力文 168 から学習データと同じ単語数の、連続する単語からなる単語列を全て生成し、これら単語列の全てについて分割モデル 166 を参照することによりそれら単語列の中央に分割ポイントがあるか否かを判定し、入力文 168 に分割ポイントがあればその分割ポイントで分割する分割部 200 と、分割部 200 により分割されたブロックを、図 1 及び図 2 に示すような入替パターンにしたがって並べ替える処理を行って前処理後の入力文 171 を出力する並べ替え部 202 とを含む。

[0027] 図 6 は、図 5 に示す分割ポイント検出部 190 を実現するためのプログラムの制御構造を表す擬似コードであり、図 7 はその動作を説明するための図である。図 6 の擬似コードは、日本語の入力文の構造が図 7 に示すブロック 118、ブロック 116、及びブロック 114 の順番で、対応する英語の出力文の構造がブロック 114、ブロック 116、及びブロック 118 である場合に、日本語の入力文の 2 個の分割ポイントを検出する。入力文の構造が上記したものではない場合には、このプログラムは何も返さない。したがって図 5 に示す分割ポイント検出部 190 は、このプログラムから 2 個の分割ポイントが得られたときにはその分割位置を出力するが、何も得られない場合には分割ポイントは出力しない。この結果、分割ポイントが得られたときには処理対象の対訳から学習データが作成されるが、得られないときには処理対象の対訳は無視され、対応する学習データは作成されない。

[0028] 図 6 及び図 7 において、「align(i)」は、日本語の文の i 番目の単語に対

応付けられた英訳文の単語位置を示す。このプログラムは、以下のような条件を満たす  $i$ 、 $j$  の組を求めるものである。

[0029] (1) 日本語の文の 1 番目の単語が英訳文の末尾の単語に対応し、それ以後、日本語の単語を  $i$  番目まで順番に進めると、対応する英語の単語は文の先頭に向かって順番に進み、

(2)  $i + 1$  番目の日本語の単語に到達すると、対応する英語の単語が、先頭に向かっていくつかの単語を飛び越して進み、以後日本語の単語を  $j$  番目まで順番に進めると、対応する英語の単語は末尾に向かって逆順に進み、

(3) 最後に、 $j + 1$  番目の日本語の単語に到達すると、対応する英語の単語が、(2) でジャンプした単語よりさらに文の先頭に近い部分にジャンプし、以後、日本語の単語を文の末尾の  $X$  番目の単語まで順番に進めると、英語の単語は文の先頭に向かって順番に進む。

[0030] 以上から分かるように、このプログラムは、図 1 及び図 2 に示すような並べ替えパターンに適合する対訳について、その分割ポイントを検出するためのものである。図 1 及び図 2 に示した並べ替えパターンと異なる並べ替えパターンに適合する学習データを作成する場合には、図 6 に示したプログラムと同様の考え方をを用いて、そのためのプログラムを作成する必要がある。

[0031] より具体的に、図 7 を参照しながら図 6 の擬似コードを説明する。日本語文の単語の総数を  $X$  とする。図 7 において日本語の単語の位置を  $x$  座標、英語の単語の位置を  $y$  座標で表す。いずれも文頭の単語位置が 1 である。図 7 では、原点は左上にあり、 $x$  座標は左から右に、 $y$  座標は上から下に伸びる。ブロック 118 とブロック 116 との境界の左側の日本語の単語位置を  $m$  とし、ブロック 116 とブロック 114 の境界の左側の単語位置を  $n$  とする。ここでの処理は、日本語及び英語の対応関係から、図 7 に示すように英文クレームと日本語クレームとのパターンに応じて、入れ替えの単位と成る構成要素（ブロック）の境界となる条件を満たす  $m$  及び  $n$  を求めることである。

[0032] 第 3 行目で計算される  $\min 1$  の値は次のように変化する。なお、以下では、

説明を簡明にするため、 $i$  番目の日本語単語が各ブロック（例えばブロック 1 1 6）内に属するものであることを、「 $i$  がブロック 1 1 6 内にある」と記載する。

[0033] (A 1)  $i$  がブロック 1 1 8 内にあるときには  $\text{align}(i)$  で表される英単語の  $y$  座標

(A 2)  $i$  がブロック 1 1 6 内にあるときには  $\text{align}(m)$ （図 7 では "capable"）で表される英単語の  $y$  座標

(A 3)  $i$  がブロック 1 1 4 内にあるときには  $\text{align}(i)$  で表される英単語の  $y$  座標

第 4 行で計算される  $\text{min}2$  の値は次のように変化する。

[0034] (B 1)  $i$  及び  $j$  がいずれもブロック 1 1 8 内にあるときには  $i$  の値に関係なく  $\text{align}(j)$  で表される英単語の  $y$  座標

(B 2)  $j$  がブロック 1 1 6 内にあるときには  $i$  の値に関係なく  $\text{align}(m+1)$  で表される英単語（図 7 では "a"）の  $y$  座標

(B 3)  $j$  がブロック 1 1 4 内にあるときには  $i$  の値に関係なく  $\text{align}(j)$  で表される英単語の  $y$  座標

図 6 の 5 行目は、 $i + 1$  番目の単語から  $j$  番目の単語までに対応する英単語の  $y$  座標がいずれも 1 番目の日本語単語から  $i$  番目の日本語単語に対応する英単語の  $y$  座標より小さいことを要求している。この条件が満足されるのは、 $i$  がブロック 1 1 8 内にある場合と  $i$  がブロック 1 1 6 内にあり、かつ  $j$  がブロック 1 1 4 内にある場合だけである。

[0035] 図 6 の 6 行目は、 $j + 1$  番目の単語から最後の単語までに対応する英単語の  $y$  座標が、いずれも  $i + 1$  番目の単語から  $j$  番目の単語の  $y$  座標より小さいことを要求している。この条件が満足されるのは、 $i$  及び  $j$  がいずれもブロック 1 1 8 内にあるか、 $i$  がブロック 1 1 8 にあってかつ  $j$  が  $n$  以上であるか、 $i$  及び  $j$  がいずれもブロック 1 1 4 内にある場合に限られる。

[0036] 図 6 の 5 行目の条件と 6 行目の条件の論理積をとると、 $i$  及び  $j$  がいずれもブロック 1 1 8 内にあるか、 $i$  がブロック 1 1 8 にあってかつ  $j$  が  $n$  であ

るか、のいずれかになる。

[0037] さらに、図6の7行目は、 $i + 1$ 番目の単語に対応する英単語の $y$ 座標が、 $j$ 番目の単語に対応する英単語の $y$ 座標より小さいことを要求している。この条件が満足されるのは、 $i < j$ であることを考慮すると、 $i + 1$ 及び $j$ がいずれもブロック116内にある場合に限られる。この条件と先の条件との論理積をとると、結局、 $i$ がブロック118内にあるかつ $i + 1$ は $m$ より大きく、かつ $j$ が $n$ である、ということになる。 $i$ がブロック118内にあるかつ $i + 1$ が $m$ より大きくなるのは $i = m$ のときに限られる。したがって、こうした条件を満足する $i$ 、 $j$ を求めれば、 $i$ は入れ替えの対象となるブロック118の右端の日本語単語の位置 $m$ 、 $j$ は同じく入れ替えの対象となるブロック116の右端の日本語単語の位置 $n$ を表すことになる。

[0038] 図6の第5行から第7行は、上記した条件を吟味するための処理である。このようにして $m$ と $n$ が見つければ、それ以上処理を続ける必要はないので、この2個の値をそれぞれboundary1及びboundary2という戻り値として（第8行～第9行）このルーチンを終了し、このプログラムを呼び出した親ルーチンに制御を戻せば良い（第10行）。

[0039] なお、図6に示した擬似コードはあくまで一例であって、これ以外にも上記した条件を満たす $m$ 及び $n$ を見つけ出すアルゴリズムであればどのようなものを採用しても良い。

[0040] 図8は、図5に示す分割部200の機能を実現するプログラムの制御構造を示すフローチャートである。図8を参照して、このプログラムは、図5に示す日本語の入力文168の形態素解析を行うステップ250と、ステップ250の処理により得られた形態素列を用いて、図5に示す分割モデル166に入力する単語列を作成するステップ252とを含む。ステップ252で作成される単語列は、入力文168の内の、連続する4個の単語からなる全ての単語列である。

[0041] このプログラムはさらに、ステップ252で得られた全ての単語列を分割モデル166に与えることにより、各単語列の中央位置に分割ポイントがあ

るか否かの判定結果を得るステップ254と、ステップ254で2個の分割ポイントが得られたか否かを判定するステップ256と、ステップ256の判定で2個の分割ポイントが得られたことに応答して、その位置で入力文168を構造部品に分割し、所定の分割がされたことを示す信号とともに出力して処理を終わるステップ258と、ステップ256の判定で2個の分割ポイントが得られなかったことに応答して、さらにステップ256の判定で分割ポイントが3個あったか否かを判定するステップ260とを含む。

[0042] このプログラムはさらに、ステップ260の判定が肯定であるときに、翻訳前の言語がいわゆる主辞後置型言語（日本語、韓国語等）か否かを判定するステップ262と、ステップ262の判定が肯定であることに応答して、ステップ260の判定で得られた3個の分割ポイントのうち、翻訳前の文の文末側の2個の分割ポイントで入力文を分割し、3個ではなく2個の分割ポイントが得られたことを示す信号とともに出力して処理を終わるステップ264と、ステップ262の判定が否定であること、すなわち翻訳前の言語がいわゆる主辞前置型言語（英語、中国語等）であることに応答して、翻訳前の文の文頭側の2個の分割ポイントで入力文を分割し、3個ではなく2個の分割ポイントが得られたことを示す信号とともに出力して処理を終わるステップ266と、ステップ260の判定が否定であること、すなわち分割ポイントが1個以下、又は4個以上であることに応答して、2個の分割ポイントが得られなかったことを示す信号とともに入力文168をそのまま出力して処理を終わるステップ268とを含む。

[0043] <動作>

以上に説明した自動翻訳システム150は以下のように動作する。図5を参照して、自動翻訳システム150の動作には、分割モデル166の学習と、分割モデル166を用いた入力文168の自動翻訳という2つのフェーズがある。これらを順番に説明する。

[0044] <分割モデル166の学習フェーズ>

これに先立ち、記憶装置160には多数の対訳が記憶されているものとす

る。それら対訳の各々について、自動翻訳装置による対訳間の単語対応付がされ、単語対応表が作成された各対訳と関係付けて記憶装置 162 に記憶される。

[0045] 分割ポイント検出部 190 は、記憶装置 160 から順番に対訳を読み出し、さらに記憶装置 162 から対応する単語対応表を読み出し、図 6 に示す疑似コードで示されるような制御構造を持つプログラムによって、各対訳のうちの日本語の文の分割ポイントを特定する。なお、これら対訳は日本語と英語のように、いずれもブロック 114、116、118 が翻訳によりブロック 118、116、114 という順番に並べ替えられるような対訳であるものとする。分割ポイント検出部 190 は、各対訳データの日本語の文に、検出された 2 個の分割ポイントを特定する情報を付与して学習データ生成部 192 に出力する。

[0046] 学習データ生成部 192 は、分割ポイント検出部 190 から与えられた対訳の各々の日本語の文について、連続する 4 個の単語からなる単語列を全て生成し、その中央に分割ポイントがあるか否かを示す情報を各単語列に付与し記憶装置 194 に学習データとして蓄積する。

[0047] モデル学習部 196 は、記憶装置 194 に蓄積された学習データを用いて分割モデル 166 のトレーニングを行う。このトレーニングの結果、分割モデル 166 は、4 個の単語からなる単語列が与えられると、その中央位置に分割ポイントがあるか否かを判定し、判定結果をその信頼度を示すスコアとともに出力するようにそのパラメータの値が調整される。

[0048] 全ての学習データによる分割モデル 166 のトレーニングが完了すると、そのトレーニング結果を用いた前処理部 170 及び自動翻訳装置 172 による自動翻訳を行うことができる。

[0049] 〈自動翻訳フェーズ〉

図 5 を参照して、自動翻訳システム 150 による自動翻訳は以下のようにして行われる。日本語の入力文 168 が与えられると、分割部 200 は入力文 168 を形態素解析し（図 8 のステップ 250）、その結果を用いて、入

力文 1 6 8 の中で連続する 4 個の単語からなる単語列を全て生成する（ステップ 2 5 2）。次いで分割部 2 0 0 は、生成された各単語列を分割モデル 1 6 6 への入力とし、当該単語列内の分割ポイントの有無に関する分割モデル 1 6 6 の判定結果を、上述した信頼度を示すスコアとともに受け取る（ステップ 2 5 4）。ステップ 2 5 4 の処理により、全ての単語列の内に 2 個のみ分割ポイントがあるという判定結果が得られると、分割部 2 0 0 は、その分割ポイントで入力文 1 6 8 を分割し、2 個の分割ポイントが得られたという信号とともに並べ替え部 2 0 2 に出力する（ステップ 2 5 8）。もしも分割ポイントが 1 個以下、又は 4 個以上の場合には、分割部 2 0 0 は入力文 1 6 8 をそのまま出力する（ステップ 2 6 8）。分割部 2 0 0 はこのとき、2 個の分割ポイントが得られなかったことを示す信号を並べ替え部 2 0 2 に出力する。さらに、分割ポイントが 3 個あったときには、翻訳前の言語が主辞後置型言語であれば翻訳前の文の文末側の 2 個の分割ポイントで文を分割して出力し（ステップ 2 6 4）、翻訳前の言語が主辞前置型言語であれば文頭側の 2 個の分割ポイントで文を分割して出力する（ステップ 2 6 6）。このとき、ステップ 2 6 4 とステップ 2 6 6 のいずれにおいても、2 個の分割ポイントが得られたという信号が並べ替え部 2 0 2 に出力される。

[0050] 並べ替え部 2 0 2 は、分割部 2 0 0 から 2 個の分割ポイントが得られたという信号を受け取ったときには、分割モデル 1 6 6 により予め定められた並べ替えパターンにしたがって、分割部 2 0 0 から受信した 3 個のブロックを並べ替え、入力文 1 7 1 として自動翻訳装置 1 7 2 に与える。もしも 2 個の分割ポイントが得られなかったという信号を受け取ったときには、並べ替え部 2 0 2 は分割部 2 0 0 の出力をそのまま入力文 1 7 1 として自動翻訳装置 1 7 2 に与える。

[0051] 自動翻訳装置 1 7 2 は、入力文 1 7 1 を入力として、日英の翻訳を行い、英語の出力文 1 7 4 を出力する。分割部 2 0 0 による分割と並べ替え部 2 0 2 による並べ替えとが行われた場合、自動翻訳装置 1 7 2 による翻訳では、入力文が出力である英語の語順に近い形となる。その結果、翻訳の精度が高



まることが知られている。以下に説明するように、実験でもそれが確認できた。

[0052] <実験結果>

実験では、日本の公開特許公報の英文抄録（P A J）と、もとになった日本語抄録とについて、自動的に単語対応付けを行った対訳コーパスを用いた。統計的機械翻訳の学習用にこれらのうち1, 000, 000文を使用し、開発・テスト用にそれぞれ1, 000文を使用した。この対訳コーパスのうち100, 000文を対象に、上記した並べ替えの必要な38, 194文を抽出し、これらから学習データを作成して分割モデル166としてSVMのトレーニングを行った。なお、以下の表では、上記実施の形態でのブロックごとの並べ替えを「グローバルな並べ替え」と呼び、以前から統計的機械翻訳で用いられている、単語レベルでの並べ替えを「単語並べ替え」と呼んでいる。

[0053] 自動翻訳は以下の4種類の設定で行い、翻訳結果はRIBES値により評価した。RIBES値は単語の並べ替えの情報を重視する評価尺度として知られている。チェックマークはその列の見出しに記載された手法を採用したことを示す。

[0054] [表1]

|    | グローバルな並べ替え | 単語並べ替え | RIBES値 |
|----|------------|--------|--------|
| T1 |            |        | 44.9   |
| T2 | ✓          |        | 61.0   |
| T3 |            | ✓      | 64.9   |
| T4 | ✓          | ✓      | 72.1   |

この表から分かるとおり、グローバルな並べ替えを用いると（T2）、並べ替えの手法を何も採用しない場合（T1）と比較してRIBES値が16.1ポイント上昇した。これは、単語並べ替えを単独で用いた場合（T3）のRIBES値の20ポイント上昇という数値という値には及ばないが、それでも精度の改善値としては大きな値である。さらに、グローバルな並べ替えを単語並べ替

えと併用する（Ｔ４）と、単語並べ替えを単独で採用した場合（Ｔ３）と比較してさらに７．２ポイントと、精度を大幅に改善できた。したがって、この手法を単語並べ替えと併用することにより、日英翻訳での精度を大きく改善できることが分かった。翻訳の方向としては逆でも同様の結果が得られる。また、日英の組み合わせ以外にも、任意の言語の組み合わせに適用可能であり、特に語順が互いに大きく異なる言語の組み合わせに適用すると特に効果的である。

[0055] なお、この実施の形態に係る自動翻訳システム１５０で使用する分割モデル１６６は、図１及び図２に示されるような並べ替えパターンに対してのみ有効である。しかし、これ以外にも並べ替えパターンはあり得る。そうした並べ替えパターンについても同様の手法を適用可能であり、学習データを変えればよいことは容易に理解可能である。さらに、そのように異なる分割モデルを併用し、同じ入力文に対してそれら分割モデルを別々に適用することで、種々の並べ替えを行って翻訳することが可能になる。この場合には、複数の分割モデルの出力する分割ポイントのスコアを比較し、最も高いスコアを与えた分割モデルの判定を用いることが有効と考えられる。分割モデルによる分割数が２個以下及び４個以上の場合についても同様である。

[0056] [第２の実施の形態]

上記第１の実施の形態では、１個の分割モデルによりブロックの並べ替え方法は一通りに定められる。しかし本発明はそのような実施の形態には限定されない。例えば、１個の分割モデルにより分割個数が同じ複数通りの並べ替えパターンを推定することも可能である。そのためには、学習データとして、分割個数が同じで並べ替えパターンが互いに異なる対訳を多数集めることが必要である。これら対訳から学習データを作成し、分割モデルの学習を行う点では第１の実施の形態と同様である。問題は、学習データの形式である。

[0057] 第１の実施の形態では、各単語列の中央位置に分割ポイントがあるか否かのみを判定していた。しかし、分割モデルはそのようなものには限定されな

い。例えば、図9に示すような構成の学習データを準備することにより、1個の分割モデルで同じ分割数で複数通りの並べ替えパターンを特定できる。

[0058] 図9を参照して、文112は分割ポイント120及び122により3個のブロック300、302、及び304に分割される。英語に翻訳すると、これらはブロック304、302、及び300の順番に並べ替えられる。このような条件で作成されるのが学習データ330である。学習データ330の構成が図4に示すものと異なるのは、単純に単語列の中央に分割ポイントがあるか否かを示す情報ではなく、左隣接単語群と右隣接単語群がそれぞれ並べ替え後にどのブロックに属するかを示す情報であるブロック番号ペア欄340を持つことである。このブロック番号ペアが同じ数値なら（例えば「3，3」）、左隣接単語群と右隣接単語群とは、並べ替え後も同じブロックに属する、すなわちこの単語列は分割ポイントを持たない事がわかる。また、ブロック番号ペアが異なる数値、例えば「3，2」であれば、左隣接単語群と右隣接単語群とが、並べ替え後にそれぞれ3番目と2番目のブロックに分けられることが分かる。すなわち、この単語列は分割ポイントを持つことが分かり、しかもブロック番号ペアを見ることによりそれぞれ並べ替え後にどの位置に移動するかも分かる。なお、この場合も分割ポイントは4個の単語の中央位置にあることを想定している。

[0059] こうした学習データを多数準備して分割モデルをトレーニングし、各単語列について並べ替え後のブロック番号ペアを推定することで、1個の分割モデルを用いて分割ポイントの位置と並べ替えパターンとを推定できる。

[0060] [コンピュータによる実現]

本発明の実施の形態に係る自動翻訳システム150及びその構成要素は、いずれもコンピュータハードウェアと、そのコンピュータハードウェア上で実行されるコンピュータプログラムとにより実現できる。図10はこのコンピュータシステム630の外観を示し、図11はコンピュータシステム630の内部構成を示す。

[0061] 図10を参照して、コンピュータシステム630は、メモリポート652

及びDVD (Digital Versatile Disk) ドライブ650を有するコンピュータ640と、いずれもコンピュータ640に接続されたキーボード646と、マウス648と、モニタ642とを含む。

[0062] 図11を参照して、コンピュータ640は、メモリポート652及びDVDドライブ650に加えて、CPU(中央処理装置)656と、CPU656、メモリポート652及びDVDドライブ650に接続されたバス666と、起動プログラム等を記憶する読出専用メモリ(ROM)658と、バス666に接続され、上記自動翻訳システム150の各部の機能を実現するプログラム命令、システムプログラム及び作業データ等を記憶するランダムアクセスメモリ(RAM)660と、ハードディスク654を含む。コンピュータシステム630はさらに、他端末との通信を可能とするネットワーク668への接続を提供するネットワークインターフェイス(I/F)644を含む。

[0063] コンピュータシステム630を上記した実施の形態に係る自動翻訳システム150及びその各機能部として機能させるためのコンピュータプログラムは、DVDドライブ650又はメモリポート652に装着されるDVD662又はリムーバブルメモリ664に記憶され、さらにハードディスク654に転送される。又は、プログラムはネットワーク668を通じてコンピュータ640に送信されハードディスク654に記憶されてもよい。プログラムは実行の際にRAM660にロードされる。DVD662から、リムーバブルメモリ664から又はネットワーク668を介して、直接にRAM660にプログラムをロードしてもよい。

[0064] このプログラムは、コンピュータ640を、上記実施の形態に係る自動翻訳システム150の各機能部として機能させるための複数の命令からなる命令列を含む。コンピュータ640にこの動作を行わせるのに必要な基本的機能のいくつかはコンピュータ640上で動作するオペレーティングシステム若しくはサードパーティのプログラム又はコンピュータ640にインストールされる、ダイナミックリンク可能な各種プログラミングツールキット又はプログラムライブラリにより提供される。したがって、このプログラム自体

はこの実施の形態のシステム、装置及び方法を実現するのに必要な機能全てを必ずしも含まなくてよい。このプログラムは、命令のうち、所望の結果が得られるように制御されたやり方で適切な機能又はプログラミングツールキット又はプログラムライブラリ内の適切なプログラムを実行時に動的に呼出すことにより、上記したシステム、装置又は方法としての機能を実現する命令のみを含んでいればよい。もちろん、独立したプログラムのみで必要な機能を全て提供してもよい。

[0065] 以上、実施の形態を例として本願発明を説明したが、本願発明が上記実施の形態に限定されるわけではない。例えば、既に述べたように、複数種類の分割モデルを並列に用い、入力文に対して最も高いスコアで分割ポイントを検出した分割モデルの結果にしたがって入力文を分割し、並べ替えるようにしてもよい。また、自動翻訳装置 172 として、上記実験では単語並べ替えを用いた場合に最もよい精度が得られることが分かったが、自動翻訳の仕組みが単語並べ替えを用いたものに限定されるわけではない。統計的な言語モデル又は翻訳モデルを用いて翻訳を行う自動翻訳装置であればどのようなものでも用いることができる。なお、この場合の分割モデルの学習データは、前述したとおり、並べ替えパターンにしたがった対訳のみを学習データとして選択してその分割ポイントを検出するようなプログラムにより行う必要がある。そうしたプログラムは、並べ替えパターンに伴う単語の対応付けを検討し、その特徴を捉えるようにすることで作成できる。

[0066] また、上記実施の形態では、分割モデルに入力する単語列の長さを 4 としたが、この値に限定されるわけではない。単語列は 2 以上であればどのような長さでもよい。また、分割ポイントの有無を判定する位置は単語列の中央としたが、それには限定されない。単語列の長さが奇数の場合も含め、分割ポイントの位置が中央以外で単語列内の単語の間であればどのような位置でもよい。もっとも、おそらくは分割ポイントの位置は中央又はその前後が最も望ましいと考えられる。

[0067] 上記実施の形態では、分割モデルとして SVM を用いたが、使用するモデ

ルがSVMに限定されるわけではない。例えばニューラルネットワークを用いてもよい。

[0068] また、上記実施の形態では、翻訳言語のペアとして日本語と英語とを扱った。しかし、本発明はそのような実施の形態には限定されず、任意の言語ペアに対して適用できる。ただし、日本語と英語、日本語と中国語等のように、互いの語順が大きく異なる言語ペアの間に本発明を適用することにより、語順が似通った言語ペアより大きな効果を得ることができる。

### 産業上の利用可能性

[0069] 本発明は、いわゆる翻訳産業に限らず、利用者が母語以外の言語に遭遇する可能性がある産業全般に利用可能である。例えば製造業において外国から装置、製品又は原料を輸入するための手続きを行う産業、そうした装置、製品又は原料を目的地に正しく配送する産業、配送された装置及び製品を正しくセットアップし、原料を正しく利用することが必要な産業等、幅広い産業において本発明は利用可能である。

[0070] 今回開示された実施の形態は単に例示であって、本発明が上記した実施の形態のみに制限されるわけではない。本発明の範囲は、発明の詳細な説明の記載を参酌した上で、請求の範囲の各請求項によって示され、そこに記載された文言と均等の意味及び範囲内での全ての変更を含む。

### 符号の説明

[0071] 114 第1のブロック  
116 第2のブロック  
118 第3のブロック  
120、122 分割ポイント  
150 自動翻訳システム  
160、162、194 記憶装置  
164 学習処理部  
166 分割モデル  
168、171 入力文

- 1 7 0 前処理部
- 1 7 2 自動翻訳装置
- 1 9 0 分割ポイント検出部
- 1 9 2 学習データ生成部
- 1 9 6 モデル学習部
- 2 0 0 分割部
- 2 0 2 並べ替え部

## 請求の範囲

### [請求項1]

第1の言語の文を第2の言語の文に自動翻訳により翻訳するために、前記第1の言語の文を複数の構造部品に分解して順序を並べ替えるパターンを学習する自動翻訳パターン学習装置であって、

前記第1の言語及び前記第2の言語の対訳を記憶するための対訳記憶手段と、

前記対訳記憶手段に記憶された対訳の各々について、前記第1の言語の文の構造部品の出現順序と、対応する前記第2の言語の文の構造部品の出現順序との対応関係が変化する位置を前記第1の言語の文及び前記第2の言語の文の各々の分割位置として検出するための検出手段と、

前記第1の言語の文の、連続する所定個数の単語列と、当該単語列内の所定位置に前記分割位置があるか否かを示す正解情報を付したデータを機械学習の学習データとして生成するための学習データ生成手段と、

前記学習データ生成手段により生成された学習データを用いて、機械学習により、前記第1の言語の文の前記所定個数の入力単語列が与えられたときに、当該入力単語列内の前記所定位置に前記分割位置があるか否かを判定するための判定モデルの学習を行うための学習手段とを含む、自動翻訳パターン学習装置。

### [請求項2]

前記学習データ生成手段は、前記第1の言語の文の、連続する所定の偶数個数の単語列と、当該単語列内の中央に前記分割位置があるか否かを示す正解情報を付したデータを前記機械学習の学習データとして生成するための手段を含む、請求項1に記載の自動翻訳パターン学習装置。

### [請求項3]

前記学習データ生成手段は、前記第1の言語の文の、連続する所定個数の単語列と、当該単語列内の所定位置に前記分割位置があるか否かを示す情報、及び前記分割位置の前後の単語列の各々について、前記



第2の言語の文における当該単語が属する構成部品の並べ替えパターンを示す情報とからなる正解情報を付したデータを機械学習の学習データとして生成するための手段を含む、請求項1に記載の自動翻訳パターン学習装置。

[請求項4]

第1の言語の文を第2の言語の文に自動翻訳により翻訳する前処理として、前記第1の言語の文を複数の構造部品列に分割し、その順序を並べ替える前処理装置であって、

前記第1の言語の文の所定個数の入力単語列が与えられたときに、当該入力単語列内の所定位置に、前記第1の言語の文の分割位置があるか否かを判定するように、予め学習済の判定手段と、

前記第1の言語の入力文が与えられたことに応答して、当該入力文を構成する、連続する前記所定個数の単語列を生成するための単語列生成手段と、

前記単語列生成手段により生成された前記所定個数の単語列の各々を前記判定手段に与えることにより、前記入力文を構成する構成部品の分割位置を特定する分割位置特定手段と、

前記分割位置特定手段により特定された分割位置で前記入力文を複数の構造部品に分割し、前記判定手段に対応して定められる並べ替えパターンにしたがって、当該複数の構成部品を並べ替える並べ替え手段とを含む、前処理装置。

[請求項5]

前記判定手段は、

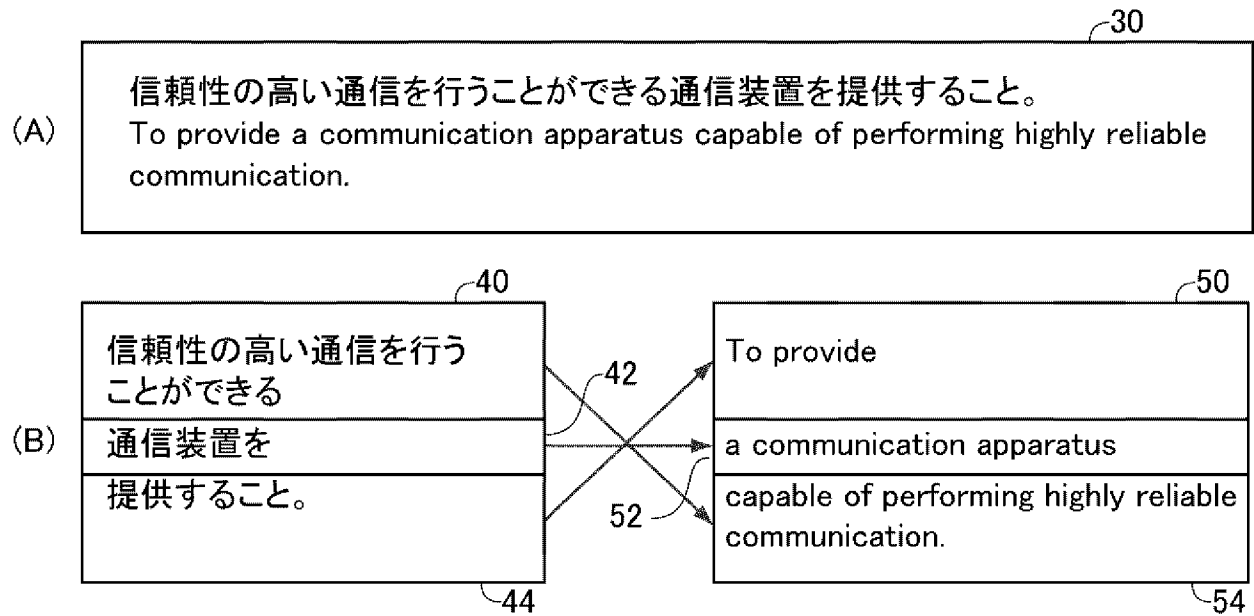
前記第1の言語の文の所定個数の入力単語列が与えられたときに、当該入力単語列内の所定位置に、前記第1の言語の文の分割位置があるか否かを判定し、当該入力単語列内の、前記分割位置前後の単語が属する構成部品が前記第2の言語の文のどの位置に配置されるかを示す分割位置情報を出力するように予め学習済の手段を含み、

前記並べ替え手段は、前記分割位置特定手段により判定された分割位置で前記入力文を複数の構造部品に分割し、前記判定手段により

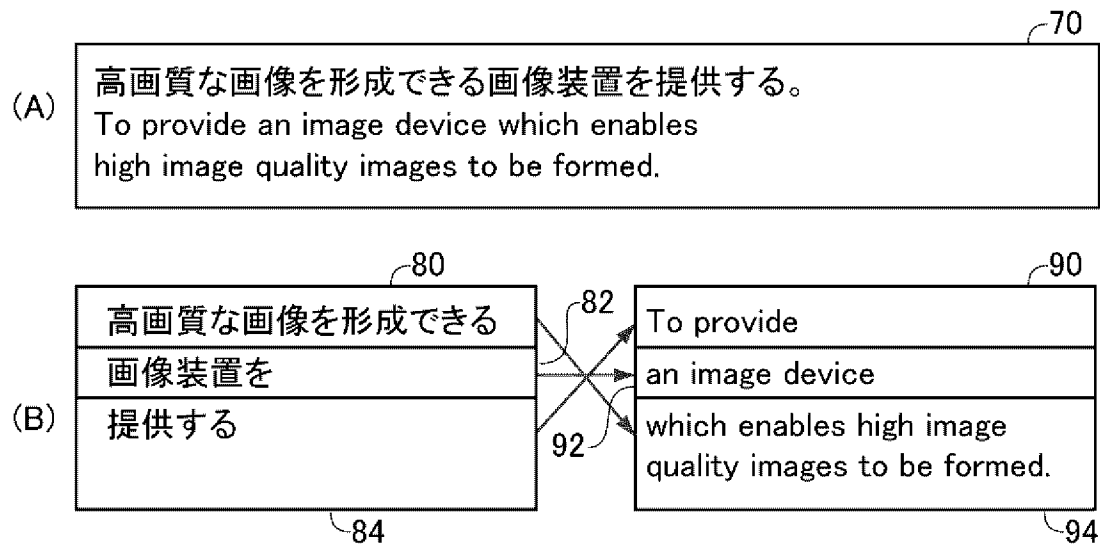
出力された前記分割位置情報にしたがって、当該複数個の構造部品を並べ替えるための手段を含む、請求項 4 に記載の前処理装置。

[請求項6]      コンピュータを、請求項 1 ～請求項 5 のいずれかに記載の各手段として動作させる、コンピュータプログラム。

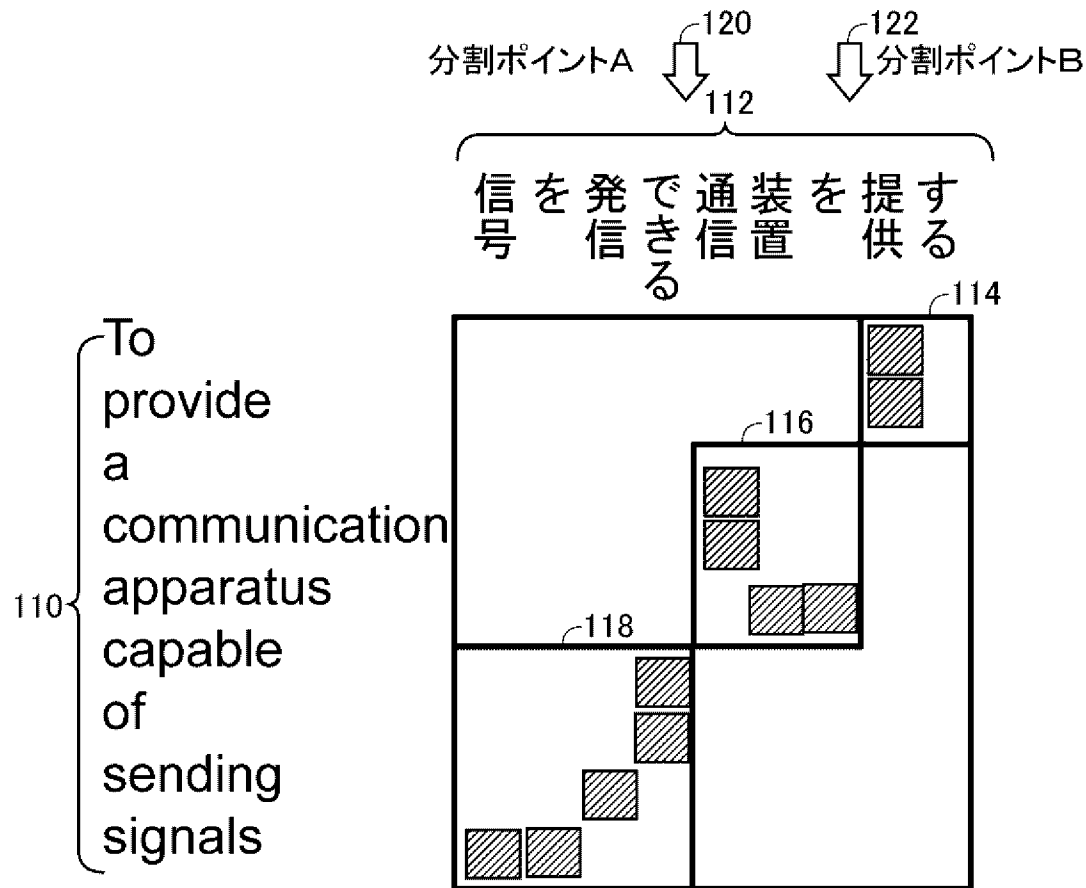
[図1]



[図2]



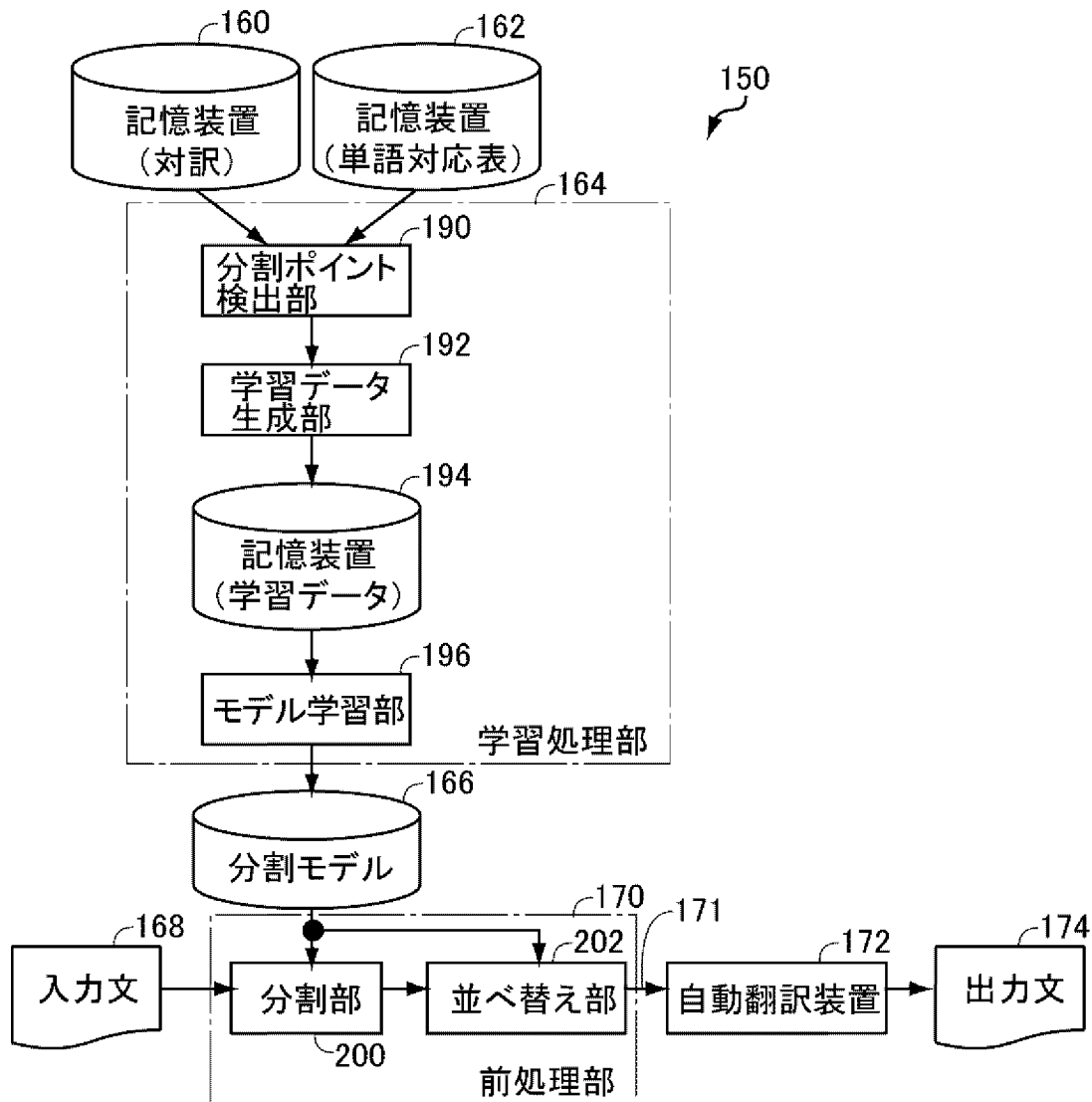
[図3]



[図4]

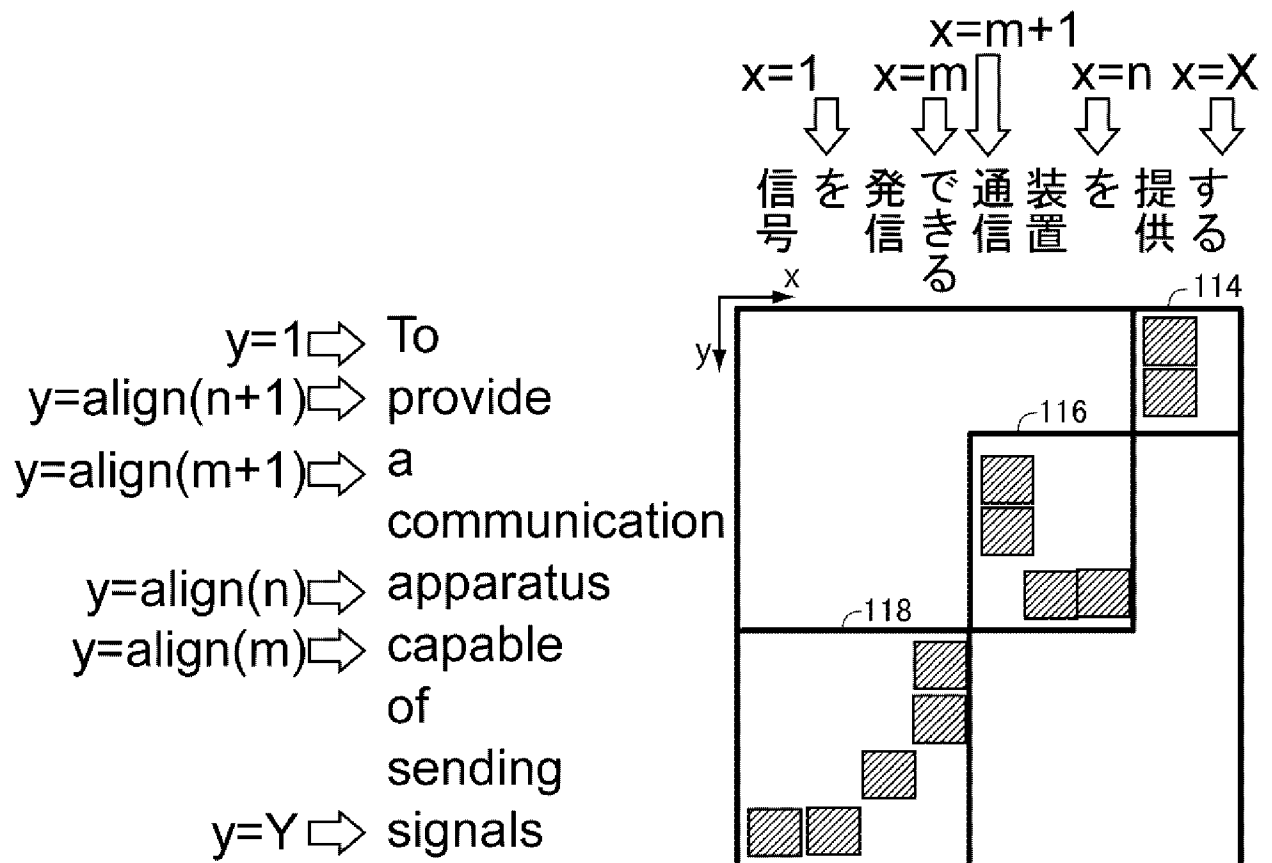


[図5]

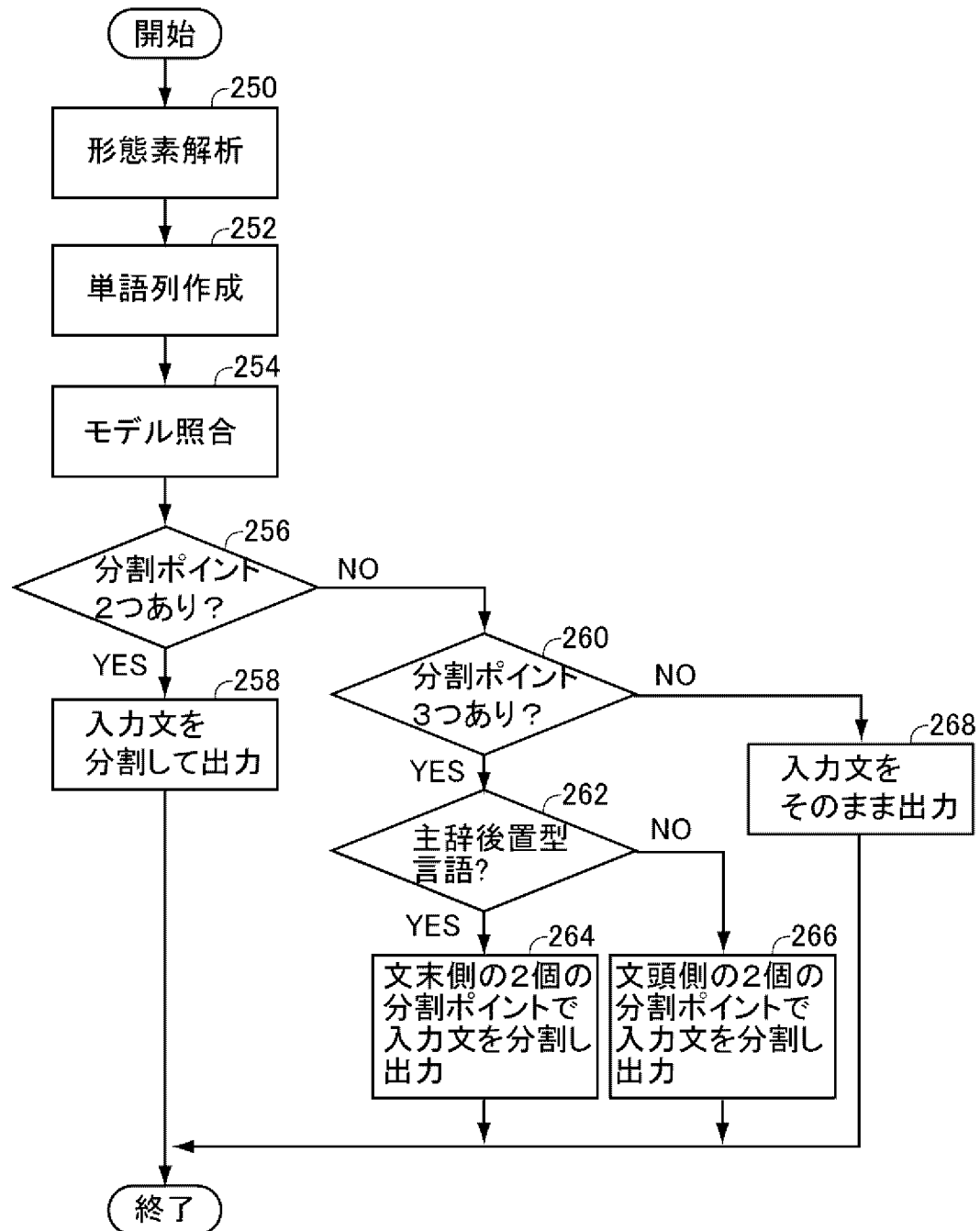


[図6]

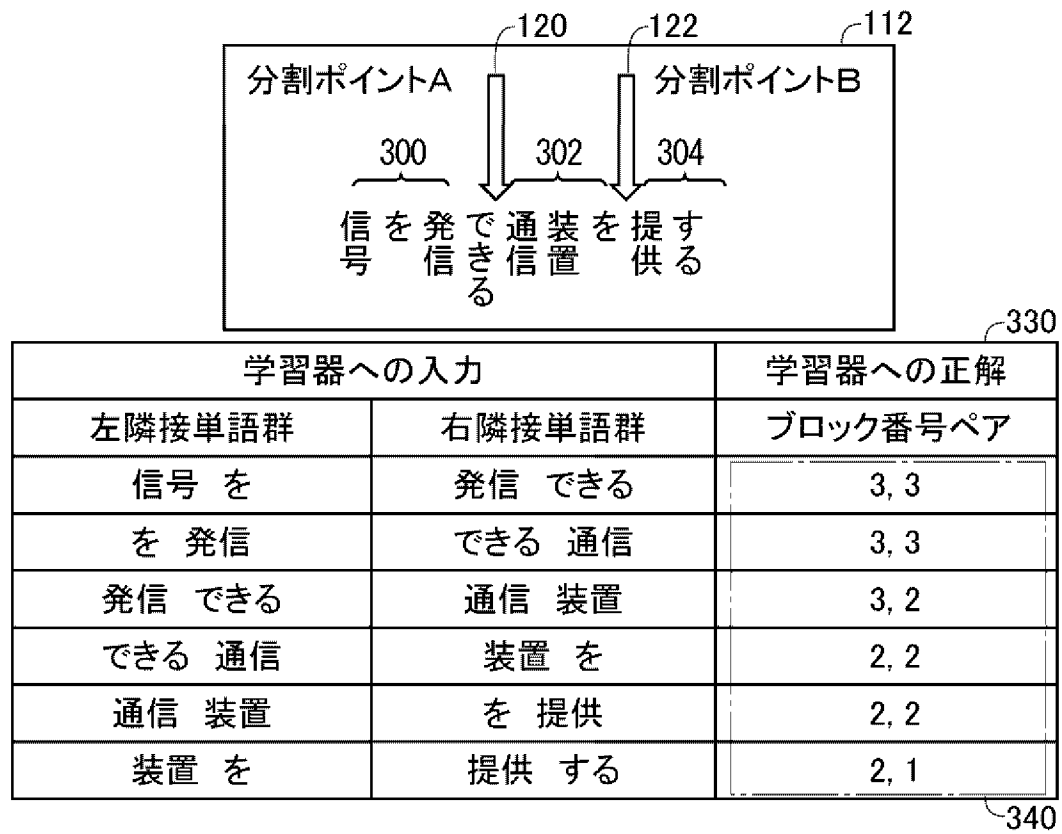
|      |                                         |
|------|-----------------------------------------|
| 001: | for j = 1 to X                          |
| 002: | for i = 1 to j-1                        |
| 003: | min1 = min y in [align(1)...align(i)]   |
| 004: | min2 = min y in [align(i+1)...align(j)] |
| 005: | if(align(i+1)...align(j) < min1 and     |
| 006: | align(j+1)...align(X) < min2 and        |
| 007: | align(i+1) < align(j)) then             |
| 008: | boundary1 = i                           |
| 009: | boundary2 = j                           |
| 010: | return boundary1 and boundary2          |
| 011: | end if                                  |
| 012: | end for                                 |
| 013: | end for                                 |



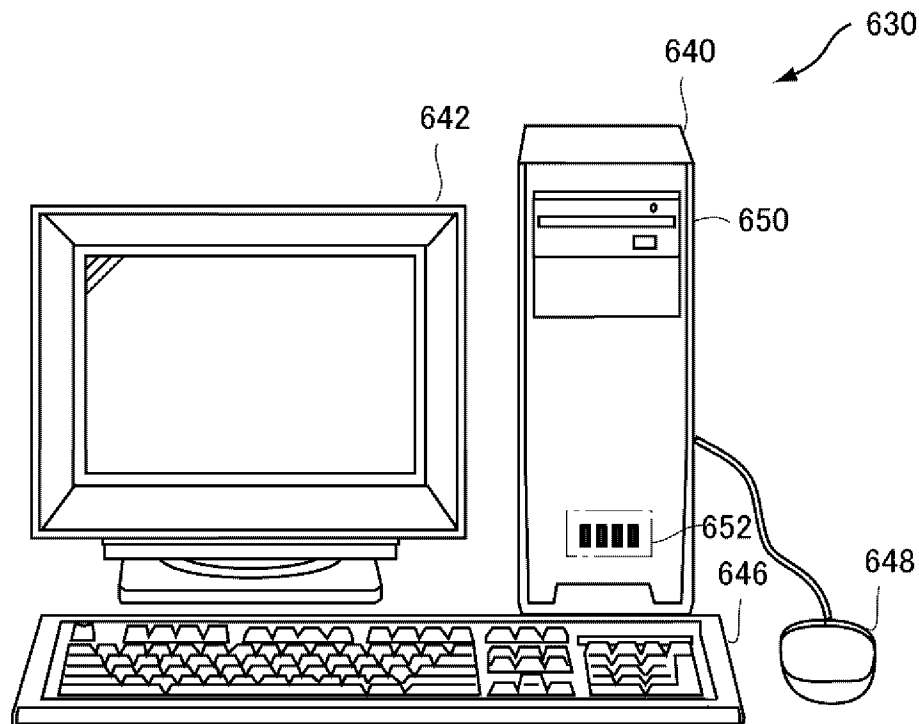
[図8]



[図9]

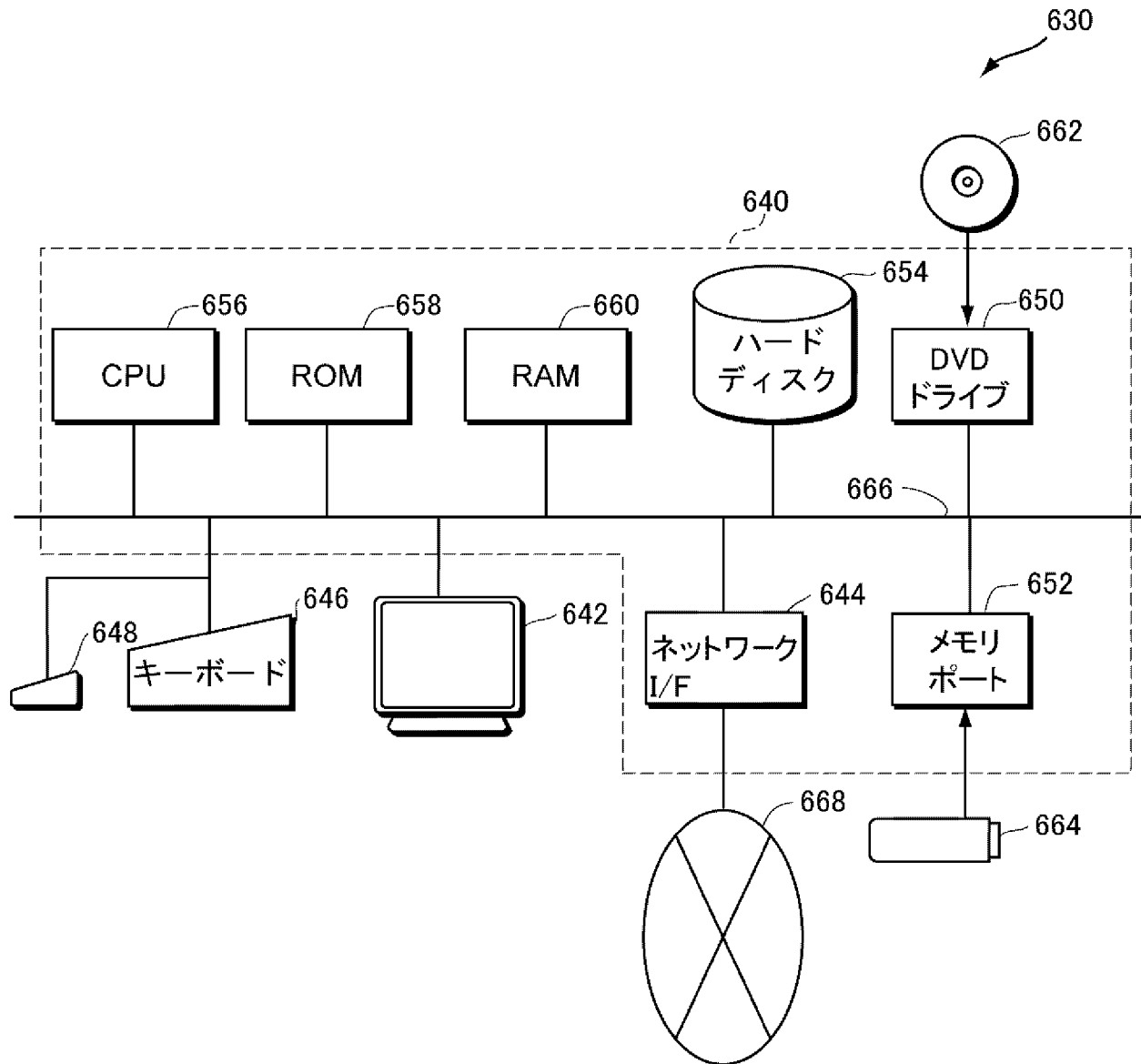


[図10]





[図11]



## INTERNATIONAL SEARCH REPORT

International application No.

PCT/JP2017/041249

## A. CLASSIFICATION OF SUBJECT MATTER

Int.Cl. G06F17/28 (2006.01) i

According to International Patent Classification (IPC) or to both national classification and IPC

## B. FIELDS SEARCHED

Minimum documentation searched (classification system followed by classification symbols)

Int.Cl. G06F17/27-17/28

Documentation searched other than minimum documentation to the extent that such documents are included in the fields searched

|                                                          |           |
|----------------------------------------------------------|-----------|
| Published examined utility model applications of Japan   | 1922-1996 |
| Published unexamined utility model applications of Japan | 1971-2018 |
| Registered utility model specifications of Japan         | 1996-2018 |
| Published registered utility model applications of Japan | 1994-2018 |

Electronic data base consulted during the international search (name of data base and, where practicable, search terms used)

## C. DOCUMENTS CONSIDERED TO BE RELEVANT

| Category* | Citation of document, with indication, where appropriate, of the relevant passages                                                                                                      | Relevant to claim No. |
|-----------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-----------------------|
| A         | 渡辺太郎, 統計的機械翻訳の現場, 人工知能学会誌, 01 May 2012, vol. 27 no. 3, pp. 288-295, (WATANABE, Taro, Field of Statistical Machine Translation, Journal of Japanese Society for Artificial Intelligence) | 1-6                   |
| A         | JP 2015-153182 A (NIPPON TELEGRAPH AND TELEPHONE CORP.)<br>24 August 2015, entire text, all drawings (Family: none)                                                                     | 1-6                   |



Further documents are listed in the continuation of Box C.



See patent family annex.

\* Special categories of cited documents:

"A" document defining the general state of the art which is not considered to be of particular relevance

"E" earlier application or patent but published on or after the international filing date

"L" document which may throw doubts on priority claim(s) or which is cited to establish the publication date of another citation or other special reason (as specified)

"O" document referring to an oral disclosure, use, exhibition or other means

"P" document published prior to the international filing date but later than the priority date claimed

"T" later document published after the international filing date or priority date and not in conflict with the application but cited to understand the principle or theory underlying the invention

"X" document of particular relevance; the claimed invention cannot be considered novel or cannot be considered to involve an inventive step when the document is taken alone

"Y" document of particular relevance; the claimed invention cannot be considered to involve an inventive step when the document is combined with one or more other such documents, such combination being obvious to a person skilled in the art

"&amp;" document member of the same patent family

Date of the actual completion of the international search

30 January 2018 (30.01.2018)

Date of mailing of the international search report

13 February 2018 (13.02.2018)

Name and mailing address of the ISA/

Japan Patent Office  
3-4-3, Kasumigaseki, Chiyoda-ku,  
Tokyo 100-8915, Japan

Authorized officer

Telephone No.

## A. 発明の属する分野の分類（国際特許分類（IPC））

Int.Cl. G06F17/28(2006.01)i

## B. 調査を行った分野

調査を行った最小限資料（国際特許分類（IPC））

Int.Cl. G06F17/27-17/28

最小限資料以外の資料で調査を行った分野に含まれるもの

|             |            |
|-------------|------------|
| 日本国実用新案公報   | 1922-1996年 |
| 日本国公開実用新案公報 | 1971-2018年 |
| 日本国実用新案登録公報 | 1996-2018年 |
| 日本国登録実用新案公報 | 1994-2018年 |

国際調査で使用した電子データベース（データベースの名称、調査に使用した用語）

## C. 関連すると認められる文献

| 引用文献の<br>カテゴリー* | 引用文献名 及び一部の箇所が関連するときは、その関連する箇所の表示                             | 関連する<br>請求項の番号 |
|-----------------|---------------------------------------------------------------|----------------|
| A               | 渡辺太郎, 統計的機械翻訳の現場, 人工知能学会誌, 2012.05.01, Vol.27 No.3, 288-295 頁 | 1-6            |
| A               | JP 2015-153182 A (日本電信電話株式会社) 2015.08.24, 全文・全図 (ファミリーなし)     | 1-6            |

☐ C欄の続きにも文献が列挙されている。☐ パテントファミリーに関する別紙を参照。

## \* 引用文献のカテゴリー

「A」特に関連のある文献ではなく、一般的技術水準を示すもの  
「E」国際出願日前の出願または特許であるが、国際出願日以後に公表されたもの  
「L」優先権主張に疑義を提起する文献又は他の文献の発行日若しくは他の特別な理由を確立するために引用する文献（理由を付す）  
「O」口頭による開示、使用、展示等に言及する文献  
「P」国際出願日前で、かつ優先権の主張の基礎となる出願

の日の後に公表された文献

「T」国際出願日又は優先日後に公表された文献であって出願と矛盾するものではなく、発明の原理又は理論の理解のために引用するもの  
「X」特に関連のある文献であって、当該文献のみで発明の新規性又は進歩性がないと考えられるもの  
「Y」特に関連のある文献であって、当該文献と他の1以上の文献との、当業者にとって自明である組合せによって進歩性がないと考えられるもの  
「&」同一パテントファミリー文献

国際調査を完了した日

30.01.2018

国際調査報告の発送日

13.02.2018

国際調査機関の名称及びあて先

日本国特許庁 (ISA/J P)

郵便番号 100-8915

東京都千代田区霞が関三丁目4番3号

特許庁審査官（権限のある職員）

成瀬 博之

5M

9192

電話番号 03-3581-1101 内線 3599