

(19) 日本国特許庁 (JP)

(12) 特 許 公 報 (B2)

(11) 特許番号
特許第4496347号
(P4496347)

(45) 発行日 平成22年7月7日 (2010.7.7)

(24) 登録日 平成22年4月23日 (2010.4.23)

(51) Int.Cl.

F I

GO 6 F 17/28 (2006.01)

GO 6 F 17/30 (2006.01)

GO 6 F 17/28 U

GO 6 F 17/30 2 1 O D

請求項の数 7 (全 27 頁)

(21) 出願番号	特願2003-418871 (P2003-418871)	(73) 特許権者	591040236
(22) 出願日	平成15年12月17日 (2003.12.17)		石川県
(65) 公開番号	特開2005-182218 (P2005-182218A)		石川県金沢市鞍月 1 丁目 1 番地
(43) 公開日	平成17年7月7日 (2005.7.7)	(74) 代理人	100111822
審査請求日	平成15年12月17日 (2003.12.17)		弁理士 渡部 章彦
審判番号	不服2007-4688 (P2007-4688/J1)	(72) 発明者	上田 芳弘
審判請求日	平成19年2月15日 (2007.2.15)		石川県金沢市鞍月 2 丁目 1 番地 石川県工
特許法第30条第1項適用 平成15年9月21日平成15年度電気関係学会北陸支部連合大会実行委員会事務局発行の「平成15年度電気関係学会北陸支部連合大会講演論文集」に発表			業試験場内
		(72) 発明者	加藤 直孝
			石川県金沢市鞍月 2 丁目 1 番地 石川県工
			業試験場内
		(72) 発明者	林 克明
			石川県金沢市鞍月 2 丁目 1 番地 石川県工
			業試験場内
最終頁に続く			

(54) 【発明の名称】 文書分類装置及びそのプログラム

(57) 【特許請求の範囲】

【請求項 1】

入力された文書に出現する単語を用いて、複数のカテゴリの各々について設けられる複数のカテゴリ辞書からなり、各々のカテゴリ辞書が、当該カテゴリにおいて出現した単語毎に、文書中に現れた単語の頻度を表すパラメータである単語の重要性を示すtf値と、単語がどの程度当該カテゴリに現れるかを表すパラメータである単語の特定性を示すidf値とに基づいて、単語の重要性と特定性とを加味した重みを表すパラメータである単語単独の重要度を表すtf・idf値を算出することにより作成された、当該カテゴリにおいて出現した単語毎に単語単独の重要度を表すtf・idf値を格納する重要性辞書と、当該カテゴリにおいて出現した単語についての複数の単語の組み合わせ毎に、各々の単語の前記idf値と、単語の共起性を示す確信度を表すパラメータであるconf値とに基づいて、特定性の高い単語間の共起性を示すidf/conf値を算出することにより作成された、当該カテゴリにおいて出現した単語についての第1単語と第2単語の組み合わせ毎に2単語間の共起性を表すidf/conf値を格納する同時出現性辞書とからなる辞書で、当該単語を照合して、前記辞書の単語毎のtf・idf値とidf/conf値とを求める手段と、

前記辞書の単語毎のtf・idf値とidf/conf値と入力された文書に出現する単語との一致率との積を算出して前記単語毎のスコアを算出する手段と、

前記単語毎のスコアに基づいて前記カテゴリ毎のスコアを算出する手段と、

これに基づいて前記入力された文書を前記複数のカテゴリのいずれかに分類する手段とを備える

ことを特徴とする文書分類装置。

【請求項 2】

前記文書は電子メールであり、前記カテゴリは文書を分配すべき担当者であり、前記文書分類装置は前記入力された文書を当該分類された担当者に配信する

ことを特徴とする請求項 1 記載の文書分類装置。

【請求項 3】

前記分類する手段は、前記カテゴリ毎のスコアに基づいて、前記入力された文書を、当該スコアが上位から所定の数のカテゴリに分類する

ことを特徴とする請求項 1 記載の文書分類装置。

【請求項 4】

当該文書分類装置が、更に、

所定の文書を前記文書分類装置により前記複数のカテゴリのいずれかに分類した結果に基づいて、前記重要性辞書における $tf \cdot idf$ 値を更新し、前記同時出現性辞書における $idf/conf$ 値を更新する学習装置とからなる

ことを特徴とする請求項 1 記載の文書分類装置。

【請求項 5】

前記学習装置は、

前記所定の文書が真のカテゴリに分類された場合、前記所定の文書に出現する単語であって当該真のカテゴリの辞書で照合された単語を重要語として、その重みを大きくする手段と、

前記所定の文書が真のカテゴリに分類されなかった場合、前記所定の文書に出現する単語の中で当該真のカテゴリについて特定性の高い単語を再評価することにより、当該真のカテゴリの特定語の重みを大きくする手段と、

前記所定の文書が真のカテゴリ以外のカテゴリに分類された場合、前記所定の文書に出現する単語であって当該誤って分類されたカテゴリの辞書で照合した単語を不要語として、当該誤って分類されたカテゴリの辞書における重みを小さくし、真のカテゴリ以外のカテゴリの辞書における重みを小さくする手段とを備える

ことを特徴とする請求項 4 記載の文書分類装置。

【請求項 6】

前記重要性辞書は、属するカテゴリの判っている文書に基づいて辞書を作成する辞書編集装置に設けられた手段であって、当該カテゴリにおいて出現した単語毎に、文書中に現れた単語の頻度を表すパラメータである単語の重要性を示す tf 値を算出し、単語がどの程度当該カテゴリに現れるかを表すパラメータである単語の特定性を示す idf 値を算出し、これらに基づいて、単語の重要性と特定性とを加味した重みを表すパラメータである単語単独の重要度を表す $tf \cdot idf$ 値を算出することにより、当該カテゴリにおいて出現した単語毎に $tf \cdot idf$ 値を格納する重要性辞書を作成する手段により作成され、

前記同時出現性辞書は、前記辞書編集装置に設けられた手段であって、当該カテゴリにおいて出現した単語についての複数の単語の組み合わせ毎に、各々の単語の前記 idf 値を算出し、単語の共起性を示す確信度を表すパラメータである $conf$ 値を算出し、これらに基づいて、特定性の高い単語間の共起性を示す $idf/conf$ 値を算出することにより、当該カテゴリにおいて出現した単語についての複数の単語の組み合わせ毎に当該単語間の共起性を表す $idf/conf$ 値を格納する同時出現性辞書を作成する手段により作成される

ことを特徴とする請求項 1 記載の文書分類装置

【請求項 7】

文書分類装置を実現するプログラムであって、

前記プログラムは、

入力された文書に出現する単語を用いて、複数のカテゴリの各々について設けられる複数のカテゴリ辞書からなり、各々のカテゴリ辞書が、当該カテゴリにおいて出現した単語毎に、文書中に現れた単語の頻度を表すパラメータである単語の重要性を示す tf 値と、単語がどの程度当該カテゴリに現れるかを表すパラメータである単語の特定性を示す idf 値

10

20

30

40

50

とに基づいて、単語の重要性和特定性を加味した重みを表すパラメータである単語単独の重要度を表す $tf \cdot idf$ 値を算出することにより作成された、当該カテゴリにおいて出現した単語毎に単語単独の重要度を表す $tf \cdot idf$ 値を格納する重要性辞書と、当該カテゴリにおいて出現した単語についての複数の単語の組み合わせ毎に、各々の単語の前記 idf 値と、単語の共起性を示す確信度を表すパラメータである $conf$ 値とに基づいて、特定性の高い単語間の共起性を示す $idf/conf$ 値を算出することにより作成された、当該カテゴリにおいて出現した単語についての第1単語と第2単語の組み合わせ毎に2単語間の共起性を表す $idf/conf$ 値を格納する同時出現性辞書とからなる辞書で、当該単語を照合して、前記辞書の単語毎の $tf \cdot idf$ 値と $idf/conf$ 値とを求める処理手段と、

前記辞書の単語毎の $tf \cdot idf$ 値と $idf/conf$ 値と入力された文書に出現する単語との一致率との積を算出して前記単語毎のスコアを算出する処理手段と、

前記単語毎のスコアに基づいて前記カテゴリ毎のスコアを算出する処理手段と、

前記カテゴリ毎のスコアに基づいて前記入力された文書を前記複数のカテゴリのいずれかに分類する処理手段、

の各手段としてコンピュータを機能させるためのプログラムである

ことを特徴とする文書分類プログラム。

【発明の詳細な説明】

【技術分野】

【0001】

本発明は、辞書編集装置、文書分類装置及びそのプログラムに関し、特に、電子メールやWeb等で入力した電子文書を、適切な担当者に自動的に正確に分類するための辞書編集装置、文書分類装置及びそのプログラムに関する。

【背景技術】

【0002】

最近、企業のコールセンター等では、電話やFAXに依らず、電子メールやWebで顧客からの問い合わせに対応することへのニーズが高まっている。しかし、電子メール等による問い合わせ（問い合わせメール）に少人数で対応するためには、適切な担当者へ問い合わせメールを分類し、回答しなければならない場合が多い。電子メールやWebにより充実したサービスを提供するためには、このような分類業務の効率化が重要であり、これを自動分類するシステムが強く望まれている。

【0003】

このような電子メールの自動分類の手段としては、電子メールを分類するための類推ルールをサンプルデータから学習して生成し、これに基づいて電子メールを自動分類できるようにした装置が知られている（例えば、特許文献1参照）。

【0004】

一方、文書分類に関する多くの研究において、単語単独の重要度である $tf \cdot idf$ 値($term\ frequency\ times\ inverse\ document\ frequency$)が応用されている。更に、本発明者等により、 $tf \cdot idf$ 値の他に、2単語間の共起性を表す $idf/conf$ 値を用いて、電子メール等の文書分類を行うことが提案されている（非特許文献1参照）。

【特許文献1】特開2001-256251号公報

【非特許文献1】成田他、(F-67)テキストマイニングと強化学習による電子メール自動分配、平成14年度電気関係学会北陸支部連合大会

【発明の開示】

【発明が解決しようとする課題】

【0005】

前述のように、本発明者等は、先に $tf \cdot idf$ 値及び $idf/conf$ 値を用いて電子メール等の文書分類を行うことを提案した。しかし、その後の本発明者の検討によれば、 $tf \cdot idf$ 値と $idf/conf$ 値を用いただけの文書分類では、電子メールの分類等に用いた場合、実用的な分類精度が得られないことが判った。即ち、 $tf \cdot idf$ 値及び $idf/conf$ 値を併せて1個のパラメータとして用いて1個の辞書を作成及び使用したのでは実用に耐え得る分類結果が得

10

20

30

40

50

られず、一方、 $tf \cdot idf$ 値と $idf/conf$ 値とを別々の独立した 2 個のパラメータとして用いて 2 個の辞書を作成及び使用すると、実用的な分類の実現に有効であることが判った。更に、辞書の作成後も単語の重み（ウェイト）を学習することが有効であるが、別々の独立した 2 個のパラメータである $tf \cdot idf$ 値と $idf/conf$ 値とについて個々に学習すると、実用的な分類の実現に有効であることが判った。そして、この学習の課程において、3 種類の学習（重要語、特定語、不要語）を行うことが有効であることが判った。

【0006】

本発明は、文書を適切な担当者に自動的に正確に分類するための辞書を作成する辞書編集装置を提供することを目的とする。

【0007】

また、本発明は、文書を適切な担当者に自動的に正確に分類する文書分類装置を提供することを目的とする。

【0008】

また、本発明は、文書を適切な担当者に自動的に正確に分類する文書分類プログラムを提供することを目的とする。

【課題を解決するための手段】

【0010】

本発明の文書分類装置は、入力された文書に出現する単語を用いて、複数のカテゴリの各々について設けられる複数のカテゴリ辞書からなり、各々のカテゴリ辞書が、当該カテゴリにおいて出現した単語毎に、文書中に現れた単語の頻度を表すパラメータである単語の重要性を示す tf 値と、単語がどの程度当該カテゴリに現れるかを表すパラメータである単語の特定性を示す idf 値とに基づいて、単語の重要性と特定性とを加味した重みを表すパラメータである単語単独の重要度を表す $tf \cdot idf$ 値を算出することにより作成された、当該カテゴリにおいて出現した単語毎に単語単独の重要度を表す $tf \cdot idf$ 値を格納する重要性辞書と、当該カテゴリにおいて出現した単語についての複数の単語の組み合わせ毎に、各々の単語の前記 idf 値と、単語の共起性を示す確信度を表すパラメータである $conf$ 値とに基づいて、特定性の高い単語間の共起性を示す $idf/conf$ 値を算出することにより作成された、当該カテゴリにおいて出現した単語についての第 1 単語と第 2 単語の組み合わせ毎に 2 単語間の共起性を表す $idf/conf$ 値を格納する同時出現性辞書とからなる辞書で、当該単語を照合して、前記辞書の単語毎の $tf \cdot idf$ 値と $idf/conf$ 値とを求める手段と、前記辞書の単語毎の $tf \cdot idf$ 値と $idf/conf$ 値と入力された文書に出現する単語との一致率との積を算出して前記単語毎のスコアを算出する手段と、前記単語毎のスコアに基づいて前記カテゴリ毎のスコアを算出する手段と、これに基づいて前記入力された文書を前記複数のカテゴリのいずれかに分類する手段とを備える。

【0011】

本発明のプログラムは、文書分類装置を実現するプログラムであって、前記プログラムは、入力された文書に出現する単語を用いて、複数のカテゴリの各々について設けられる複数のカテゴリ辞書からなり、各々のカテゴリ辞書が、当該カテゴリにおいて出現した単語毎に、文書中に現れた単語の頻度を表すパラメータである単語の重要性を示す tf 値と、単語がどの程度当該カテゴリに現れるかを表すパラメータである単語の特定性を示す idf 値とに基づいて、単語の重要性と特定性とを加味した重みを表すパラメータである単語単独の重要度を表す $tf \cdot idf$ 値を算出することにより作成された、当該カテゴリにおいて出現した単語毎に単語単独の重要度を表す $tf \cdot idf$ 値を格納する重要性辞書と、当該カテゴリにおいて出現した単語についての複数の単語の組み合わせ毎に、各々の単語の前記 idf 値と、単語の共起性を示す確信度を表すパラメータである $conf$ 値とに基づいて、特定性の高い単語間の共起性を示す $idf/conf$ 値を算出することにより作成された、当該カテゴリにおいて出現した単語についての第 1 単語と第 2 単語の組み合わせ毎に 2 単語間の共起性を表す $idf/conf$ 値を格納する同時出現性辞書とからなる辞書で、当該単語を照合して、前記辞書の単語毎の $tf \cdot idf$ 値と $idf/conf$ 値とを求める処理手段と、前記辞書の単語毎の $tf \cdot idf$ 値と $idf/conf$ 値と入力された文書に出現する単語との一致率との積を算出して前記単語

10

20

30

40

50

語毎のスコアを算出する処理手段と、前記単語毎のスコアに基づいて前記カテゴリ毎のスコアを算出する処理手段と、前記カテゴリ毎のスコアに基づいて前記入力された文書を前記複数のカテゴリのいずれかに分類する処理手段、の各手段としてコンピュータを機能させるためのプログラムである。

【発明の効果】

【0012】

本発明の辞書編集装置によれば、単語単独の重要度を表す $tf \cdot idf$ 値と複数の単語間の共起性を表す $idf/conf$ 値とを別々の独立した2個のパラメータとして用いて2個の辞書（ $tf \cdot idf$ 辞書、 $idf/conf$ 辞書）を作成することができるので、これを文書分類装置の辞書として用いることにより、電子メール等の文書の分類に用いた場合、実用的な分類精度を得ることができる。従って、文書を適切な担当者に自動的に正確に分類するための辞書を容易に作成することができる。

10

【0013】

本発明の文書分類装置によれば、前述の2個の辞書 $tf \cdot idf$ 辞書、 $idf/conf$ 辞書を用いることにより、入力した文書を基本的には単語単独の重要度を表す $tf \cdot idf$ 値と複数（2つ）の単語間の共起性を表す $idf/conf$ 値とに基づいてカテゴリに分類できるので、電子メール等の文書の分類に用いた場合、実用的な分類精度を得ることができる。従って、文書を適切な担当者に自動的に正確に分類することができる。

【0014】

本発明の文書分類プログラムによれば、これをフレキシブルディスク、CD-ROM、CD-R/W、DVD等の媒体に格納すること、又は、インターネット等のネットワークを介してダウンロードすることにより供給することができ、これにより前述の文書分類システムを容易に実現することができ、正確な文書分類を可能とすることができる。

20

【発明を実施するための最良の形態】

【0015】

図1は、文書分類システム構成図であり、本発明の文書分類システムの構成を示す。文書分類システムは、3個のサブシステム、即ち、辞書編集装置1、文書分類装置2、強化学習装置3からなる。これらの間は、例えばLAN（Local Area Network）により相互に接続される。この例では、文書分類装置2は、例えばメール解析装置2からなる。従って、以下の例においては、分類（解析）対象である前記入力された文書は例えば電子メール8であり、前記カテゴリは文書を分配すべき担当者（回答者）であり、前記文書分類装置2は入力された文書を当該分類された担当者に配信（分配）する。

30

【0016】

辞書編集装置1は、文書ファイル6に基づいて辞書を作成する辞書編集装置であって、当該担当者において出現した単語毎に、単語単独の重要度を表す $tf \cdot idf$ 値を格納する重要性辞書（ $tf \cdot idf$ 辞書）41を作成し、また、当該担当者において出現した単語についての複数の単語の組み合わせ毎に、当該単語間の共起性を表す $idf/conf$ 値を格納する同時出現性辞書（ $idf/conf$ 辞書）42を作成する。文書ファイル6は、各担当者が日常業務で作成した文書である（文書を格納している）ので、属するカテゴリ即ちその担当者の判っている文書である。従って、当該担当者において出現した単語とは、当該担当者の作成した文書ファイル6に出現した単語である。この例では、同時出現性辞書42は、当該担当者において出現した単語についての第1単語と第2単語の組み合わせ毎に2単語間の共起性を表す $idf/conf$ 値を格納する。従って、辞書4は、複数の担当者の各々について設けられる複数の担当者（の）辞書40からなり、各々の担当者辞書40が重要性辞書41と同時出現性辞書42とからなる。

40

【0017】

辞書編集装置1は、実際には、文書ファイル（担当者毎のフォルダ）6に保存された各種の文書を変換して得たテキストファイルを読み込み、読み込んだテキストデータから改行とスペースを取り除いた後に、周知の形態素解析処理により単語に分割することにより抽出した単語を、担当者毎の辞書40に登録する。抽出した単語において、品詞が助詞、

50

助動詞、接続詞、接頭詞、副詞、連体詞、感動詞、記号は不要語と考えられるので、辞書 4 には登録しない。

【 0 0 1 8 】

即ち、本発明では、辞書 4 の作成及び編集において、単語単独の重要度である $tf \cdot idf$ 値を用い、更に、単語間の共起性を $idf/conf$ 値 (idf divided by confidence) をも用いる。ここで、 tf 値は文書中に現れた単語の頻度 (即ち、単語の重要性) を示し、 idf 値は単語がどの程度特定の分類に現れるか (即ち、単語の特定性) を示し、 $tf \cdot idf$ 値は文書分類の研究で用いられているものであり単語の重要度と特定性とを加味した重み (即ち、単語単独の重要度) を示し、 $conf$ 値は単語の共起性を示す確信度を示し、 $idf/conf$ 値は特定性の高い単語間の共起性を示す。

10

【 0 0 1 9 】

図 2 (A) は $tf \cdot idf$ 値の辞書 (重要性辞書) 4 1 の一例を示す。重要性辞書 4 1 は、その担当者が使用した単語 a_1 、 a_2 、 a_3 、 $\dots$ 毎に、その tf 値、 idf 値、 $tf \cdot idf$ 値を格納する。 tf 値、 idf 値、 $tf \cdot idf$ 値の算出については後述する。図 2 (B) は $idf/conf$ 値の辞書 (同時出現性辞書) 4 2 の一例を示す。同時出現性辞書 4 2 は、その担当者が使用した単語 b_1 、 b_2 、 b_3 、 $\dots$ における第 1 単語と第 2 単語の組み合わせ毎に、当該第 1 単語の idf 値、第 1 単語と第 2 単語間の $conf$ 値、第 1 単語と第 2 単語間の $idf/conf$ 値を格納する。前述のように、この例における $idf/conf$ 値は、当該 2 単語間の特定性を考慮した共起性を示す。 $conf$ 値、 $idf/conf$ 値の算出については後述する。

20

【 0 0 2 0 】

メール解析装置 2 は、入力された文書である問い合わせの電子メール (問い合わせメール) 8 に出現する単語を用いて辞書 4 で当該単語を照合して、辞書 4 の単語毎の $tf \cdot idf$ 値と $idf/conf$ 値とを求め、これらに基づいて所定の演算を行って単語毎のスコアを算出し、単語毎のスコアに基づいて担当者毎のスコアを算出し、これに基づいて入力された文書を複数の担当者のいずれかに分類する。この例では、メール解析装置 2 は、前記所定の演算として、辞書 4 の単語毎の $tf \cdot idf$ 値と $idf/conf$ 値と入力された文書に出現する単語との一致率との積を算出して、単語毎のスコアを算出する。また、メール解析装置 2 は、担当者毎のスコアに基づいて、入力された文書である問い合わせメール (以下、単にメール) 8 を、当該スコアが上位から所定の数の担当者に分類する。

30

【 0 0 2 1 】

メール解析装置 2 は、辞書編集装置 1 と同様に、入力されたメール 8 の本文について、改行とスペースを取り除いた後に周知の形態素解析処理を行って、単語に分割することにより得られた単語と辞書 4 とを照合する。この時、メール解析装置 2 は、マッチした (照合できた) 単語又は単語の組み合わせに関する $tf \cdot idf$ 値と $idf/conf$ 値と単語の一致率とから、当該メール 8 に対する全ての担当者のスコアを算出して、最後にスコアの高い一人又は複数の担当者を当該メール 8 への回答者 (当該メール 8 に回答すべき担当者) として推定し、図 3 に示すように、その結果を例えばメール解析装置 2 に表示する。即ち、質問者からのメール 8 を入力として、これに回答すべき担当者の候補 (即ち、図 3) を出力として得る。また、この例では、メール解析装置 2 は、この分類の結果 (スコア) に従って、図 3 において「○」の付された担当者の担当者 (回答者) 端末 5 に、当該入力されたメール 8 を分配する (配信する)。

40

【 0 0 2 2 】

図 3 は、メール解析結果の一例を示す。当該結果は、担当者 d 毎に、回答者とするか否かの結果、スコア $Score(d)$ 、メール 8 と辞書とで照合できた単語 t 、当該単語毎のスコア (d, t) からなる。担当者 d 及び照合単語 t はスコアの高い順に表示される。最もスコアの高いもの (第 1 (i) 位) を第 1 (i) 位照合単語という。図 3 において、「○」の付された担当者が回答者とされる。この例では上位 2 人が回答者として推定される。入力された (受信された) メール 8 は、回答者とされた担当者にメール 8 として分類 (送信) される。一方、「×」の付された担当者は回答者と推定されない。人間の分類者を介在させる場合や分類を受けた担当者が他の担当者に再分類する場合は、図 3 に示す推定結果を参照して分類

50

を行うようにすれば良い。

【0023】

なお、図示はしないが、メール解析装置2は、例えばLANを介して回答者端末5に接続され、また、メールサーバに接続される。メールサーバは、インターネットに接続され、外部からのメール8を受信し、これを内部の各端末5に配信するために、メール解析装置2に送る。メール8は、例えば質問者が周知のWebブラウザを用いて、問い合わせ内容、氏名、返信用電子メールアドレス等を入力することにより作成(入力)する。

【0024】

本発明では、2個の辞書41、42に基づいて得た2個の重みをそのまま用いるのではなく、マッチした(照合できた)単語の一致率をも考慮してスコアを算出する。即ち、処理対象であるメール8に出現する単語と辞書の単語との照合を、完全一致ではなく、部分一致で行うと共に、当該部分一致の割合(単語の一致率)をスコアの算出に用いる。これにより、同一の用語についての僅少な個人差の影響等を排除することができる。

10

【0025】

強化学習装置3は、所定の文書(即ち、別文書ファイル7)をメール解析装置2により複数の担当者のいずれかに分類した結果に基づいて、一旦作成した重要性辞書41における $tf \cdot idf$ 値を更新し、同時出現性辞書42における $idf/conf$ 値を更新する(強化する)。即ち、強化学習装置3は、重要語学習、特定語学習、不要語学習を行う。重要語学習処理は、所定の文書が真の担当者に分類された場合、前記所定の文書に出現する単語であって当該真の担当者の辞書40で照合された単語を重要語として、その重みを大きくする。特定語学習の処理は、所定の文書が真の担当者に分類されなかった場合、前記所定の文書に出現する単語の中で当該真の担当者について特定性の高い単語を再評価することにより、当該真の担当者の特定語の重みを大きくする。不要語学習は、所定の文書が真の担当者以外の担当者に分類された場合、前記所定の文書に出現する単語であって当該誤って分類された担当者(真の担当者以外の担当者(カテゴリ))の辞書40で照合した単語を不要語として、当該誤って分類された担当者の辞書40における重みを小さくし、真の担当者以外の担当者の辞書40における重みを小さくする。これにより、メール8の分類精度を更に向上することができる。

20

【0026】

本発明では、一旦作成した辞書4における $tf \cdot idf$ 値と $idf/conf$ 値を強化するために、別文書ファイル7、即ち、分類を受ける各担当者自身が様々な業務の中で作成した文書を読み込んだファイルを用いる。別文書ファイル7は、文書ファイル6及びメール8とは別の文書(文書ファイル)であって、各担当者が作成した別文書である(別文書を格納する)。即ち、別文書ファイル7は、担当者が随時作成する文書を随時取り込んで格納したものである。別文書ファイル7は、文書ファイル6と同様にして、メール解析装置2により複数の担当者のいずれかに分類される。別文書ファイル7の担当者も明確であるのでこの結果は正確に評価することができ、また、担当者の業務内容の変化に応じて随時 $tf \cdot idf$ 値と $idf/conf$ 値を強化することができ、メール8の正確な分類に有効である。

30

【0027】

また、本発明では、重要語学習、特定語学習、不要語学習即ち、 $tf \cdot idf$ 値と $idf/conf$ 値の強化を、プロフィットシェアリング(以下PSと言う)により行う。即ち、別文書ファイル7についてその担当者(分類先)を推定し、その推定結果の適切さに応じた報酬(プロフィット)により $tf \cdot idf$ 値と $idf/conf$ 値を補正する。PSは強化学習において注目されている。これにより、分類の専門家と同等レベルの分類精度を得ることができる。PSは、周知のように、報酬を得たときに複数ルール(本発明では複数単語の重み)を一括して強化するので、効率的にメール8の分類精度を向上することができる。なお、このように強化学習を文書分類に取り入れたシステムは、これまでに開発されていない。

40

【0028】

以上のように、本発明によれば、2個の辞書41、42に基づいて得た2個の重みを用いてメール8等の文書を分類することにより、実用的な分類を実現することができる。即

50

ち、日常業務としてメール 8 を分類している専門家の分類結果（分類精度）が、実用上必要な精度と考えることができる。本発明のメール解析装置 2 による分類精度として、当該専門家の分類精度とほぼ同等の精度を得ることができる。従って、本発明のメール解析装置 2（による分類）は十分に実用に耐え得るものである。なお、本発明者の検討によれば、実際は、2 個の辞書 4 1、4 2 のみを用いた分類精度は当該専門家の精度を少し下回るが、 $tf \cdot idf$ 値と $idf/conf$ 値を別文書ファイル 7 を用いて強化学習することにより、分類の専門家と同等な精度でメール 8 を分類することができる。

【0029】

以下、処理フローを参照して、本発明の文書分類システムにおける処理について、詳細に説明する。図 4 は、文書分類処理フローであり、本発明の図 1 に示す文書分類システムにおける文書分類処理を示す。辞書編集装置 1 が、各担当者が作成した文書ファイル 6 を収集して、この中の出現単語の重みを、単語単独の重要度を表す $tf \cdot idf$ 値と 2 単語間の共起性を表す $idf/conf$ 値として算出し、この 2 種類の辞書 4 1、4 2 を担当者毎に作成する（ステップ S101）。強化学習装置 3 が、PS を応用して、これらの重み又はウェイト（を示す $tf \cdot idf$ 値及び $idf/conf$ 値）を強化学習する（ステップ S102）。即ち、メール解析装置 2 が、メール 8 の受信の有無を判断し、当該受信がない場合、別文書ファイル 7 の文書と 2 種類の辞書 4 1、4 2 を照合して、単語の重みと単語の一致率とから、担当者毎にスコアを算出し、このスコアが高い担当者を当該メール 8 への回答者として推定し、当該推定の結果を強化学習装置 3 に入力して $tf \cdot idf$ 値及び $idf/conf$ 値の重みを更新する。一方、メール 8 を受信すると、メール解析装置 2 は当該メール 8 を解析する（ステップ S103）。即ち、メール 8 と 2 種類の辞書 4 1、4 2 を照合して、単語の重みと単語の一致率とから、担当者毎にスコアを算出する。そして、メール解析装置 2 は、このスコアが高い担当者を当該メール 8 への回答者として推定し、これをメール解析装置 2 の端末（図示せず）に表示する（ステップ S104）。

【0030】

図 5 は、辞書編集処理フローであり、図 1 の辞書編集装置 1 がステップ S101 において実行する辞書編集処理を示す。辞書編集装置 1 が、全担当者の全文書ファイル 6 を読み込み（ステップ S111）、全文書について前処理を行う（ステップ S112）。即ち、全文書から、その改行とスペースとを除去し、残りの部分について周知の形態素解析を行い、分かち書きした単語を得る。そして、当該得られた全単語から不要な品詞の単語を削除する。次に、辞書編集装置 1 が、図 2（C）に示す担当者・単語テーブルを生成し（ステップ S113）、これに基づいて $tf \cdot idf$ 値を算出して（ステップ S114）、当該単語及び $tf \cdot idf$ 値を $tf \cdot idf$ 辞書（重要性辞書）4 1 に書き込む（ステップ S115）。これにより、 $tf \cdot idf$ 辞書 4 1 が作成される。担当者・単語テーブル及びその生成については、図 6 を参照して後述する。次に、辞書編集装置 1 が、図 8 に示す単語の組み合わせテーブルを生成し（ステップ S116）、これに基づいて $idf/conf$ 値を算出して（ステップ S117）、当該単語とその組み合わせ及び $idf/conf$ 値を $idf/conf$ 辞書（同時出現性辞書）4 2 に書き込む（ステップ S118）。これにより、 $idf/conf$ 辞書 4 2 が作成される。単語の組み合わせテーブル及びその生成については、図 7 を参照して後述する。

【0031】

図 6 は、担当者・単語テーブル生成処理フローであり、辞書編集装置 1 がステップ S113 において実行する担当者・単語テーブル生成処理を示す。辞書編集装置 1 が、d に最初の担当者を設定し（ステップ S121）、図 2（C）の担当者・単語テーブルの当該列に d を追加しその要素 $tf(t, d)$ を全て「0」に初期化し（ステップ S122）、t に最初の単語を設定し（ステップ S123）、担当者・単語テーブルの行 t_1 、 t_2 、 t_3 、・・・に t が存在するか否かを調べる（ステップ S124）。存在しない場合、辞書編集装置 1 は、担当者・単語テーブルの行に t を追加し、その要素 $tf(t, d)$ を全て「0」に初期化する（ステップ S125）。存在する場合、辞書編集装置 1 は、ステップ S125 を省略し、当該担当者 d 及び単語 t に対応する要素に「1」を加算する（ステップ S126）。即ち、 $tf(t, d) = tf(t, d) + 1$ を実行する。この後、辞書編集装置 1 は、担当者 d の全単語について処理が終了したか否かを

調べ（ステップS127）、終了していない場合、tに次の単語を設定し（ステップS128）、ステップS124以下を繰り返す。終了している場合、辞書編集装置1は、全担当者について処理が終了したか否かを調べ（ステップS129）、終了していない場合、dに次の担当者を設定し（ステップS1210）、ステップS122以下を繰り返す。終了している場合、当該処理を終了する。これにより、担当者・単語テーブルが生成される。

【0032】

今、 $tf \cdot idf$ 値は、式（1）のように定義される。即ち、

【0033】

【数1】

$$tf_t^d \cdot idf_t^d = \frac{tf(t, d)}{\sum_{t' \in T_d} tf(t', d)} \cdot \log \frac{N}{df(t)} \quad (1)$$

【0034】

ここで、右辺第1項の $tf(t, d)$ は、担当者dが作成した全ての文書（文書ファイル6）中における単語tの出現回数を表す。従って、右辺第1項は同一の担当者が何度も繰り返して使用する単語に大きい重みを与える。例えば、図2（C）において、担当者d1における単語t1の出現回数 $tf(t1, d1) = 1$ である。右辺第1項の分母は担当者dが使用した単語の頻度の総和を表している。例えば、図2（C）において、担当者d1についての当該値 $\sum tf(t', d1) = 8538$ である。右辺第2項のNは分類先である全ての担当者数を、 $df(t)$ は単語tを使用した担当者数を表す。従って、右辺第2項は特定少数の担当者が使用する単語に大きな重みを与え、各担当者の特徴付ける特定性の指標となる。以上から、担当者・単語テーブルに基づいて、 tf 値（ tf_t^d ）、 idf 値（ idf_t^d ）、 $tf \cdot idf$ 値を算出することができる。

【0035】

図7は、単語の組み合わせテーブル生成処理フローであり、辞書編集装置1がステップS116において実行する単語の組み合わせテーブル生成処理を示す。辞書編集装置1が、最小支持度及び最小確信度を読み込み（ステップS131）、dに最初の担当者を設定し（ステップS132）、句点の検索をしてセンテンス集合 S^d を生成し（ステップS133）、最小支持度とセンテンス集合 S^d の全センテンス数とからLSCを算出する（ステップS134）。次に、辞書編集装置1が、単語の組み合わせ数を「1」とし、即ち、単語組み合わせテーブル（以下、同じ） $C = C1$ とし（ステップS135）、（Cの単語の組み合わせ）＝（センテンス集合 S^d から抽出した各センテンスにおける重複のない単語）とし（ステップS136）、Cのカウント値（Count）をCの単語の組み合わせが出現するセンテンス集合 S^d のセンテンス数とする（ステップS137）。次に、辞書編集装置1が、Cにおける単語の組み合わせの全てについて、当該組み合わせについてのCのカウント値がLSC以上であるか否かを調べ、そうでない場合にはCから当該単語の組み合わせを除き、そうである場合にはCから当該単語の組み合わせを除かないようにする（ステップS138）。この後、辞書編集装置1が、単語の組み合わせ数が「2」か否かを調べる（ステップS139）。「2」でない場合、即ち、単語の組み合わせ数が「1」の場合、辞書編集装置1が、単語の組み合わせ数を「2」とし、即ち、 $C = C2$ とし（ステップS1310）、（Cの単語の組み合わせ）＝（C1から生成した2単語の組み合わせ）とし（ステップS1311）、ステップS137以下を繰り返す。

【0036】

ステップS139において単語の組み合わせ数が「2」である場合、辞書編集装置1が、t1をC2の単語の組み合わせの第1単語とし（ステップS1312）、t1のカウント値（t1_Count）をC1のt1のカウント値（Count）とし（ステップS1313）、最小確信度とt1のC1のカウント値とからLCCを算出する（ステップS1314）。次に、辞書編集装置1が、C

10

20

30

40

50

2における単語の組み合わせの全てについて、当該組み合わせについてのC2のカウント値がLC2以上か否かを調べ、そうでない場合にはC2から当該単語の組み合わせを除き、そうである場合にはC2から当該単語の組み合わせを除かないようにする(ステップS1315)。この後、辞書編集装置1が、全担当者について処理済みか否かを調べ(ステップS1316)、そうでない場合、dに次の担当者を設定し(ステップS1317)、ステップS133以下を繰り返す。全担当者について処理済みである場合、処理を終了する。

【0037】

ここで、idf/conf値は、関連ルールの抽出に用いられる周知の確信度に基づいている。関連ルールの手法を本発明に応用するために、担当者dが使用した単語を要素とする集合を $T^d = (t_1, t_2, t_3, \dots, t_m)$ とする。また、担当者dが作成した文書から抽出したセンテンス(句点で区切られる1文)の集合を $S^d = (s_1, s_2, s_3, \dots, s_n) (s_i \in T^d)$ とする。ここで、単語tの支持度support(t)は S^d 全体に対しtを含むセンテンスの割合を表す。また、関連ルールは $t \rightarrow t'$ で表現され、単語tが出現したセンテンスには単語 t' が出現する確率が高いこと、即ちtと t' の共起性が高いことを表す。関連ルールは支持度(support)及び確信度(confidence)の2つのパラメータを有し、これらの値により関連ルールの有意性を示す。ここで、support($t \rightarrow t'$)は S^d 全体に対しtと t' を共に含むセンテンスの割合、confidence($t \rightarrow t'$)はtを含むセンテンスの中で t' を含むセンテンスの割合と定義されている。

【0038】

本発明では、図8に示すように、最小支持度と最小確信度を設定して、共起性の高い単語の組み合わせを求め、この共起性を表す重みとして確信度を用いる。この例では、最小支持度及び最小確信度を、各々、0.3及び0.7とした。これらの値は経験的に定めることができる。なお、索引語抽出の研究で良く知られているように、頻度の低い単語は不要語であるが、頻度が上位の単語も特徴語ではなく一般語であることが多く、不要語となる。ここでも同様に、最小支持度と最小確信度を満足しない単語の組み合わせは不要であり、辞書4には登録しない。これと同時に、支持度もしくは確信度が上位の組み合わせは一般語の組み合わせとなり不要である。そこで、上述のように最小支持度と最小確信度を満足した単語の組み合わせの共起性を表す具体的な重みとして、確信度の逆数に単語の特定性を考慮して第1単語のidf値を積算した値を用いる。なお、更に、第1単語に加えて第2単語のidf値も積算した値を重みとしても良いし、確信度の代わりに支持度を用いても良い。

【0039】

例えば、図8に示すように、担当者dのセンテンス集合 S^d が求まるとする。単語の組み合わせを1個とすると、各々の出現回数Countが求まる(最初のC1)。これから、(最小支持度) × (S^d のセンテンス数) = 0.3 × 4 よりもCount値の小さい単語{単語4}を除く(2番目のC1)。これにより、各々の単語についての最小支持度を満足する値が定まる。次に、残りの単語について出現する2個の単語の組み合わせの全てについて、各々の出現回数Countを求め(最初のC2)、これから、0.3 × 4 よりもCount値の小さい単語の組み合わせ{単語1, 単語2}及び{単語1, 単語5}を除く(2番目のC2)。次に、残りの単語の組み合わせについて、(最小確信度) × (各々の第1単語についての最小支持度を満足する値) よりもCount値の小さい単語の組み合わせ{単語2, 単語3}及び{単語3, 単語5}を除く(3番目のC2)。即ち、{単語2, 単語3}については2(C2におけるCount値、以下同じ) < 0.7 × 3(2番目のC1の単語2のCount値、以下同じ)であり、{単語3, 単語5}については2 < 0.7 × 3であり、除かれる。一方、{単語1, 単語3}については2 > 0.7 × 2であり、{単語2, 単語5}については3 > 0.7 × 3であり、残される。

【0040】

なお、本発明者の検討によれば、単語の組み合わせ数を3以上にしても、メール8(文書)の分類精度は向上しないことが判った。従って、この例においては、計算量の低減のために、単語の組み合わせは「2」に制限される。従って、この例では、同時出現性辞書

4 2 は、第 1 単語と第 2 単語の組み合わせ毎にidf/conf値を格納する。

【 0 0 4 1 】

今、単語 t が出現したとき単語 t' が共起する指標であるidf/conf値を式 (2) で定義する。

【 0 0 4 2 】

【 数 2 】

$$idf_t^d / conf_t^d \Rightarrow t' = idf_t^d \cdot \frac{confidence_{max}^d}{confidence^d(t \Rightarrow t')} (2)$$

10

【 0 0 4 3 】

ここで、右辺第2項の分子はある担当者 d におけるconfidence値 (即ち、確信度) の最大値 (max) であり、担当者毎に大きさの異なるconfidence値を標準化している。以上から、図 8 の単語の組み合わせテーブルに基づいて、conf値、idf/conf値を算出することができる。conf値の定義は、式 (2) の両辺より idf_t^d を除くことにより明らかであろう。

【 0 0 4 4 】

図 9 は、メール (問い合わせメール) 解析処理フローであり、図 1 のメール解析装置 2 がステップS103及びS104において実行するメール解析処理を示す。メール解析装置 2 が、
処理対象である受信したメール (問い合わせメール) 8 を読み込み (ステップS141)、
これについて、その改行とスペースとを除去し、残りの部分について周知の形態素解析を行
い (ステップS142)、分かち書きした単語を得る。次に、メール解析装置 2 が、 d に最初
の担当者を設定し (ステップS143)、 t_1 にメール 8 の最初の単語を設定し (ステップS144)、
 $tf \cdot idf$ 値重み付き加算を行い (ステップS145)、idf/conf値重み付き加算を行い (ステップS146)、
メール 8 の全単語について処理を終了した否かを調べる (ステップS147)。
 $tf \cdot idf$ 値重み付き加算については図 10 を参照して後述し、idf/conf値重み付き加
算については図 11 を参照して後述する。

20

【 0 0 4 5 】

全単語について処理を終了していない場合、メール解析装置 2 は、 t_1 にメール 8 の次の
単語を設定し (ステップS148)、ステップS145以下を繰り返す。全単語について処理を終
了している場合、メール解析装置 2 は、担当者 d のスコアを当該メール 8 の全単語につい
てのスコアの総計、即ち、 $Score(d) = \sum Score(d, t_1)$ として算出し (ステップS149)、全
担当者について処理を終了した否かを調べる (ステップS1410)。全担当者について処理
を終了していない場合、メール解析装置 2 は、 d に次の担当者を設定し (ステップS1411)、
ステップS144以下を繰り返す。全担当者について処理を終了している場合、メール解
析装置 2 は、全担当者についての $Score(d)$ の平均及び標準偏差を算出して、これに基づい
て、当該メール 8 に回答すべき担当者 (回答者) の候補を決定し (ステップS1412)、こ
れを表示する (ステップS1413)。

30

【 0 0 4 6 】

図 10 は、 $tf \cdot idf$ 値重み付き加算処理フローであり、メール解析装置 2 がステップS1
45において実行する $tf \cdot idf$ 値重み付き加算処理を示す。メール解析装置 2 が、Max__ tf
__ idf __Ratio に「 0 」を設定 (代入) し、 $Score(d, t_1)$ に「 0 」を設定し、 m を読み込
み (ステップS151)、 t_dic に担当者 d の $tf \cdot idf$ 辞書 4 1 の最初の単語を設定し (ス
テップS152)、 t_1 と t_dic とが部分一致するか否かを調べる (ステップS153)。部分一
致する場合、メール解析装置 2 が、 $tf \cdot idf$ に t_dic の $tf \cdot idf$ 値を設定し、Matched__
Ratio に、 t_1 と t_dic との一致率の m 乗を設定し、 $tf \cdot idf$ と Matched__Ratio とから tf
__ idf __Ratio を算出し (ステップS154)、 $tf \cdot idf$ __Ratio が Max __ $tf \cdot idf$ __Ratio
よりも大きいかな否かを調べる (ステップS155)。大きい場合、メール解析装置 2 が、Max
__ $tf \cdot idf$ __Ratio に $tf \cdot idf$ __Ratio を設定し (ステップS156)、担当者 d の $tf \cdot idf$

40

50

辞書 4 1 の全単語について処理が終了したか否かを調べる (ステップS157)。

【 0 0 4 7 】

全単語について処理が終了していない場合、t_dic に担当者 d の tf・idf 辞書 4 1 の次の単語を設定し (ステップS158)、ステップS153以下を繰り返す。ステップS153において t1 と t_dic とが部分一致しない場合、ステップS154～ステップS156を省略して、ステップS157を実行する。ステップS155において tf・idf __Ratio が Max __tf・idf __Ratio よりも大きくない場合、ステップS156を省略して、ステップS157を実行する。ステップS157において全単語について処理が終了している場合、Score(d, t1) に Max __tf・idf __Ratio を設定して処理を終了する (ステップS159)。

【 0 0 4 8 】

図 1 1 は、idf/conf 値重み付き加算処理フローであり、メール解析装置 2 がステップS145において実行する idf/conf 値重み付き加算処理を示す。メール解析装置 2 が、Max __idf/conf __Ratio に「0」を設定し、m を読み込み (ステップS161)、t_dic1 に担当者 d の idf/conf 辞書 4 2 の最初の第 1 単語を設定し、t_dic2 に担当者 d の idf/conf 辞書 4 2 の最初の第 2 単語を設定し (ステップS162)、t1 と t_dic1 とが部分一致するか否かを調べる (ステップS163)。部分一致する場合、メール解析装置 2 が、t2 にメール 8 の t1 の次の単語を設定し (ステップS164)、t2 と t_dic2 とが部分一致するか否かを調べる (ステップS165)。部分一致する場合、メール解析装置 2 が、idf/conf に t_dic1 と t_dic2 の idf/conf 値を設定し、Matched __Ratio1 に t1 と t_dic1 の一致率の m 乗を設定し、Matched __Ratio2 に t2 と t_dic2 の一致率の m 乗を設定し、idf/conf と Matched __Ratio1 と Matched __Ratio2 とに基づいて idf/conf __Ratio を算出する (ステップS166)。

【 0 0 4 9 】

ステップS165において t2 と t_dic2 とが部分一致しない場合、メール解析装置 2 が、メール 8 の最後の単語まで処理済か否かを調べ (ステップS1611)、最後の単語まで処理を終了していない場合、t2 にメール 8 の t2 の次の単語を設定し (ステップS1612)、ステップS165以下を繰り返す。

【 0 0 5 0 】

ステップS166の後、メール解析装置 2 が、idf/conf __Ratio が Max __idf/conf __Ratio よりも大きい場合、Max __idf/conf __Ratio に idf/conf __Ratio を設定して (ステップS168)、担当者 d の idf/conf 辞書 4 2 の全単語について処理済か否かを調べる (ステップS169)。全単語について処理を終了していない場合、メール解析装置 2 が、t_dic1 に担当者 d の idf/conf 辞書 4 2 の次の第 1 単語を設定し、t_dic2 に担当者 d の idf/conf 辞書 4 2 の次の第 2 単語を設定し (ステップS1610)、ステップS163以下を繰り返す。全単語について処理を終了している場合、メール解析装置 2 が、定数 α を読み込み (ステップS1613)、Score(d, t1) に α と Max __idf/conf __Ratio とにより求まる値を加算することにより Score(d, t1) を算出する (ステップS1614)。

【 0 0 5 1 】

ステップS163において t1 と t_dic1 とが部分一致しない場合、メール解析装置 2 が、ステップS164～S168 (S1611 及び S1612 を含む) を省略して、ステップS169を実行する。ステップS167において idf/conf __Ratio が Max __idf/conf __Ratio よりも大きくない場合、メール解析装置 2 が、ステップS168を省略して、ステップS169を実行する。ステップS1611においてメール 8 の最後の単語まで処理を終了している場合、メール解析装置 2 が、ステップS1612 (S166～S168を含む) を省略して、ステップS169を実行する。

【 0 0 5 2 】

このように、本発明では、メール 8 中の全ての出現単語を辞書 4 と照合して、メール 8 に対する各担当者のスコアを算出する。このとき、形態素解析により分割された単語の単語長が変化することがある。また、質問者と担当者とは同一の単語を使用する保証はない。このため、メール 8 に含まれる単語と辞書 4 中の単語を完全一致を条件に照合するとマッチする単語が少なくなり、結局は分類の精度が向上しない。そこで、本発明では、単語

10

20

30

40

50

の照合は部分一致によって行い、この一致率を単語の重みに乗ずることによる重み付き加算を行って、メール 8 に対する担当者 d のスコアを算出する。

【 0 0 5 3 】

まず、メール 8 中のある単語 t の $tf \cdot idf$ 値と $idf/conf$ 値を加味したスコアを式 (3) で定義する。なお、単語 t 及びこれと共起性を示す単語 t' は、ともに上記の部分一致のために辞書 4 中の複数の単語又は単語の組み合わせに照合する可能性があるので、このうち、各々 $tf \cdot idf$ 値と $idf/conf$ 値に一致率を乗じた値が最大値となるものを加算する。

【 0 0 5 4 】

【数 3】

$$Score_t^d = tf_t^d \cdot idf_t^d \cdot Match_Ratio(t) + \alpha \cdot \max_{\substack{t' \in Mail, \\ t' \neq t}} [conf_{t \Rightarrow t'}^d \cdot idf_t^d \cdot Match_Ratio(t) \cdot Match_Ratio(t')] \quad (3)$$

10

【 0 0 5 5 】

ここで、 $Match_Ratio(t)$ は単語の一致率、即ちメール 8 中の単語 t と辞書 4 中の単語の単語長に対する一致した文字数の比率を表し、式 (4) で定義される。

【 0 0 5 6 】

【数 4】

$$Match_Ratio(t) = \left(\frac{Match_Length(t, t_{dic})^2}{Length(t) \cdot Length(t_{dic})} \right)^m \quad (4)$$

30

【 0 0 5 7 】

ここで、 t_{dic} は辞書 4 中の単語を表し、 $Length(t)$ は単語 t の文字数、 $Match_Length(t, t_{dic})$ は t と t_{dic} の一致した文字数を表している。また、式 (3) の α は、第 1 項と第 2 項の重みを決める係数である。更に、式 (4) の m は単語の一致率をどのオーダでスコアに反映するかを表す。この一致率は 1 以下であるため、 m を大きくするほどこの率によってスコアに差が生じる。なお、この α と m のいずれも適当な範囲で変動させて、経験則により、最大の分類精度を得る α と m を採用することができる。

【 0 0 5 8 】

最後に、式 (3) で与えられる単語 t のスコア $Score_t^d$ について、メール 8 本文に出現する単語で総和を取ることで、入力されたメール 8 に対する担当者 d のスコアを式 (5) で定義する。

【 0 0 5 9 】

【数 5】

$$Score^d = \sum_{t \in Mail} Score_t^d \quad (5)$$

40

50

【0060】

ここで、入力したメール8に対する回答者は、全ての担当者のScore^dの分布を正規分布と仮定して、その平均スコアより 2σ (σ は標準偏差)以上大きいスコアを持つ担当者と推定する。これは、実際のメール8において、特に問い合わせ内容を限定しない場合は、複数の事柄について複合して質問していることや複数の専門分野にまたがって質問していることが多く、複数の担当者が連携して回答する必要があるためである。

【0061】

図12は、強化学習処理フローであり、主として図1の強化学習装置3がステップS102において実行する単語の重みの強化学習処理を示す。強化学習装置3が、別文書ファイル7を読み込み(ステップS171)、真の回答者集合を抽出する(ステップS172)。なお、別文書ファイル7には、例えば学習用メール8が含まれる(以下同じ)。メール解析装置2が、別文書ファイル7について前述のようにメール解析を実行し、図3の結果を得る。強化学習装置3が、当該結果をメール解析結果テーブルとして取得し(ステップS173)、メール解析結果テーブル中の全ての回答者について、当該回答者が真の回答者集合に含まれるか否かを調べ(ステップS174)、含まれる場合には当該回答者について重要語学習処理を実行し(ステップS175)、含まれない場合には当該回答者について不要語学習処理を実行する(ステップS176)。なお、実際には、点線で示すように、一人の回答者についてステップS174とS175又はS176とを実行することを、各回答者について繰り返す。この後、強化学習装置3が、メール解析結果テーブル中の全ての非回答者について、当該非回答者が真の回答者集合に含まれるか否かを調べ(ステップS177)、含まれる場合には当該非回答者について特定語学習処理を実行し(ステップS178)、含まれない場合にはステップS178を省略する。なお、実際には、点線で示すように、一人の非回答者についてステップS177又はS177及びS178を実行することを、各非回答者について繰り返す。この後、強化学習装置3が、全ての別文書ファイル7について処理済か否かを調べ(ステップS179)、処理済でない場合にはステップS171以下を繰り返し、処理済である場合には処理を終了する。

【0062】

本発明では、別文書ファイル7の解析結果を受けて、PSにより、推定した回答者の正しさから担当者毎の辞書4のtf・idf値とidf/conf値を更新する。この例では、作成者即ち担当者が判明している別文書ファイル7をメール解析装置2に入力する。この場合、真の文書作成者が判っているので、当該システムの分類結果の正否を判断できる。なお、人間の分類者が参照して別文書ファイル7を分類した結果を正しいとして学習しても良い。

【0063】

本発明では、以下の3種類の強化学習を行う。即ち、重要語の学習は、真の回答者をシステムが回答者であると推定できた場合に、別文書ファイル7とこの回答者の辞書40で照合した単語を重要語として、その重みを大きくする。また、特定語の学習は、真の回答者をシステムが回答者であると推定できなかった場合に、別文書ファイル7に出現した単語の中で真の回答者について特定性の高い単語を再評価することにより、この回答者の特定語の重みを大きくする。更に、不要語の学習は、真の回答者以外をシステムが回答者であると誤推定した場合に、別文書ファイル7とこの担当者の辞書40で照合した単語を不要語として、この担当者の辞書40における不要語の重みを小さくする。これと同時に、この単語を真の回答者以外の他の担当者の辞書40でも、同様に不要語として学習する。

【0064】

図13は、重要語学習処理フローであり、強化学習装置3がステップS175において実行する重要語の学習処理を示す。強化学習装置3が、定数c1、L、Sを定義ファイルから読み込み(ステップS181)、単語tにメール解析結果テーブルにおいて対象回答者の第1位照合単語を設定する(ステップS182)。次に、強化学習装置3が、対象回答者のtf・idf辞書41でtを検索してtf・idfにtと照合した単語のtf・idf値を設定し、tf・idfとc1とから $f_tf \cdot idf$ を算出し(ステップS183)、対象回答者のidf/conf辞書42の第1単語でtを検索してidf/confにtと照合した単語のidf/conf値を設定し、idf/confとc1とから $f_idf / conf$ を算出し(ステップS184)、tf・idf辞書41及びidf/conf辞書42に

おける t の $tf \cdot idf$ 値及び $idf/conf$ 値を更新する (ステップ S185)。即ち、新たな $tf \cdot idf$ 値を、それまでの $tf \cdot idf$ 値に $f_tf \cdot idf$ を加算した値とする。新たな $idf/conf$ 値を、それまでの $idf/conf$ 値に $f_idf/conf$ を加算した値とする。

【 0 0 6 5 】

次に、強化学習装置 3 が、 $i = 2$ として (ステップ S186)、単語 t にメール解析結果テーブルにおいて対象回答者の第 i 位照合単語を設定する (ステップ S187)。次に、強化学習装置 3 が、対象回答者の $tf \cdot idf$ 辞書 4 1 で t を検索して $tf \cdot idf$ に t と照合した単語の $tf \cdot idf$ 値を設定し、 $f_tf \cdot idf$ と S とから新たな $f_tf \cdot idf$ を算出し (ステップ S188)、対象回答者の $idf/conf$ 辞書 4 2 の第 1 単語で t を検索して $idf/conf$ に t と照合した単語の $idf/conf$ 値を設定し、 $f_idf/conf$ と S とから新たな $f_idf/conf$ を算出し (ステップ S189)、 $tf \cdot idf$ 辞書 4 1 及び $idf/conf$ 辞書 4 2 における t の $tf \cdot idf$ 値及び $idf/conf$ 値を更新する (ステップ S1810)。即ち、新たな $tf \cdot idf$ 値を、それまでの $tf \cdot idf$ 値に $f_tf \cdot idf$ を加算した値とし、新たな $idf/conf$ 値を、それまでの $idf/conf$ 値に $f_idf/conf$ を加算した値とする。この後、強化学習装置 3 が、 i が L に等しいか否かを調べ (ステップ S1811)、等しくない場合、 i を $i + 1$ とし (ステップ S1812)、ステップ S187 以下を繰り返す。等しい場合、ステップ S1812 を省略して、処理を終了する。なお、メール 8 の文章が極めて短いと i の最大値が L より小さくなることもあるが、この場合、読み込んだ L に代えて当該 i の最大値が用いられる。即ち、ステップ S1811 において i が当該 i の最大値に等しい場合、処理を終了する (以下、ステップ S1911 及び S2120 において同じ)。

【 0 0 6 6 】

P S では、エピソード単位にルールに付加された重みを強化する。エピソードとは、初期状態あるいは報酬を得た直後から、次の報酬までのルール系列を表す。この強化には報酬からどれだけ過去かを引数とする強化関数が用いられる。長さ l のエピソード ($r_1, \dots, r_i, \dots, r_l, r_1$) に対して、ルール r_i の重み w_i は式 (6) のように強化関数 f_i で強化される。

【 0 0 6 7 】

【数 6】

$$w_{r_i} = w_{r_i} + f_i \quad (6)$$

【 0 0 6 8 】

本発明では、ルールは辞書 4 中の単語又は単語の組み合わせに相当し、エピソードのルール系列は式 (5) でスコア $Score^d$ を決定した単語 ($t_1, \dots, t_i, \dots, t_l, t_1$) である。即ち別文書ファイル 7 と担当者 d の辞書 4 0 中で照合した単語である。更に、系列中の順序 i は式 (5) のスコア $Score^d$ への寄与の順、即ち式 (3) の $Score_t^d$ の大きさの順とし、この $Score_t^d$ が最大となるスコア 1 位の単語 t (第 1 位照合単語) を t_1 とする。また、本発明では、単語の重みを $tf \cdot idf$ 値と $idf/conf$ 値の 2 種類で構成しているので、式 (6) はそれぞれ式 (7) 及び式 (8) のようになり、当該システムが推定した回答者と真の回答者が一致した場合に担当者 d の単語 t_i に、 i が上位であるほど大きな正の報酬を与える。

【 0 0 6 9 】

【数 7】

$$tf_{t_i}^d \cdot idf_{t_i}^d = tf_{t_i}^d \cdot idf_{t_i}^d + f_{tf,i}^d \quad (7)$$

$$idf_{t_i}^d / conf_{t_i \Rightarrow t'_i}^d = idf_{t_i}^d / conf_{t_i \Rightarrow t'_i}^d + f_{conf,i}^d \quad (8)$$

【0070】

10

但し、本発明の学習方法においては、 $tf \cdot idf$ 値と $idf/conf$ 値について、それぞれ別に強化するが、その方法は全く同様である（図 13 中の処理を参照）。

【0071】

PS では、有効な単語の重みが強化され、無効な単語の重みが抑制されることを保証しなければならない。この条件は合理性定理により式（9）を満足することであることが周知のように証明されている。

【0072】

【数 8】

$$\forall k = 2, \dots, W + 1, \quad L \sum_{i=k}^{W+1} f_i^d < f_{i-1}^d \quad (9)$$

20

【0073】

ここで、 W はエピソードの最大長、 L は同一感覚入力下に存在する有効ルールの最大数であり、本発明では学習を有効にする単語数を限定し、別文書ファイル 7 と辞書 4 で一致した単語でスコア $Score^d$ の上位何単語を学習するかを設定することとし、その単語数を L とした（例えば、 $L = 10$ なら上位 10 位の単語を学習する、以下同じ）。また、この式（9）の定理を満足する最も簡単な強化関数には、式（10）で表される等比減少関数を用

30

【0074】

【数 9】

$$f_n^d = \frac{1}{S} f_{n-1}^d, \quad n = 2, \dots, W \quad (10)$$

【0075】

40

ここで、周知のように、 $S = L + 1$ を満足しなければならない。また、式（10）の初期値 f_1^d 、即ちスコア $Score_t^d$ が最大となる単語の強化値は以下の式（11）とする。ここで、 c_1 は定数である。

【0076】

【数 10】

$$f_1^d = c_1 w_1^d \quad (11)$$

50

【 0 0 7 7 】

このように、重要語の学習は、推定した回答者が真の回答者であった場合に行い、スコア $Score_t^d$ が上位である単語の重みに正の報酬を与え、その重みを大きくする。

【 0 0 7 8 】

図 1 4 は、不要語学習処理フローであり、強化学習装置 3 がステップ S176 において実行する不要語の学習処理を示す。強化学習装置 3 が、定数 $c2$ 、 L 、 S を定義ファイルから読み込み（ステップ S191）、単語 t にメール解析結果テーブルにおいて対象回答者の第 1 位照合単語を設定する（ステップ S192）。次に、強化学習装置 3 が、対象回答者の $tf \cdot idf$ 辞書 4 1 で t を検索して $tf \cdot idf$ に t と照合した単語の $tf \cdot idf$ 値を設定し、 $tf \cdot idf$ と $c2$ とから $f_tf \cdot idf$ を算出し（ステップ S193）、対象回答者の $idf/conf$ 辞書 4 2 の第 1 10
単語で t を検索して $idf/conf$ に t と照合した単語の $idf/conf$ 値を設定し、 $idf/conf$ と $c2$ とから $f_idf/conf$ を算出し（ステップ S194）、不要語強化処理を実行する（ステップ S195）。不要語強化処理については図 1 5 を参照して後述する。

【 0 0 7 9 】

次に、強化学習装置 3 が、 $i = 2$ として（ステップ S196）、単語 t にメール解析結果テーブルにおいて対象回答者の第 i 位照合単語を設定する（ステップ S197）。次に、強化学習装置 3 が、対象回答者の $tf \cdot idf$ 辞書 4 1 で t を検索して $tf \cdot idf$ に t と照合した単語の $tf \cdot idf$ 値を設定し、 $f_tf \cdot idf$ と S とから新たな $f_tf \cdot idf$ を算出し（ステップ S198）、対象回答者の $idf/conf$ 辞書 4 2 の第 1 20
単語で t を検索して $idf/conf$ に t と照合した単語の $idf/conf$ 値を設定し、 $f_idf/conf$ と S とから新たな $f_idf/conf$ を算出し（ステップ S199）、ステップ S195 と同様の不要語強化処理を実行する（ステップ S1910）。この後、強化学習装置 3 が、 i が L に等しいか否かを調べ（ステップ S1911）、等しくない場合、 i を $i + 1$ とし（ステップ S1912）、ステップ S197 以下を繰り返す。等しい場合、ステップ S1912 を省略して、処理を終了する。

【 0 0 8 0 】

図 1 5 は、不要語強化処理フローであり、強化学習装置 3 がステップ S195 及び S1910 において実行する不要語の強化処理を示す。強化学習装置 3 が、担当者 d に対象回答者を設定し（ステップ S201）、 $tf \cdot idf - f_tf \cdot idf$ が「0」より大きいかなかを調べ（ステップ S202）、大きい場合、 $tf \cdot idf$ 辞書 4 1 の t の $tf \cdot idf$ 値を $tf \cdot idf - f_tf \cdot idf$ に更新し（ステップ S203）、大きくない場合、 $tf \cdot idf$ 辞書 4 1 の t の $tf \cdot idf$ 値を「0」30
に更新する（ステップ S204）。次に、強化学習装置 3 が、 $idf/conf - f_idf/conf$ が「0」より大きいかなかを調べ（ステップ S205）、大きい場合、 $idf/conf$ 辞書 4 2 の t の $idf/conf$ 値を $idf/conf - f_idf/conf$ に更新し（ステップ S206）、大きくない場合、 $idf/conf$ 辞書 4 2 の t の $idf/conf$ 値を「0」に更新する（ステップ S207）。次に、強化学習装置 3 が、真の回答者集合以外の全担当者について処理済かなかを調べ（ステップ S208）、処理済でない場合、担当者 d に次の担当者（対象回答者）を設定し（ステップ S209）、担当者 d の $tf \cdot idf$ 辞書 4 1 で t を検索して $tf \cdot idf$ に t と照合した単語の $tf \cdot idf$ 値を設定し（ステップ S2010）、担当者 d の $idf/conf$ 辞書 4 2 の第 1 単語で t を検索して $idf/conf$ に t と照合した単語の $idf/conf$ 値を設定し（ステップ S2011）、ステップ S202 以下を繰り返す。ステップ S208 において全担当者について処理済である場合、処理を終了する。 40

【 0 0 8 1 】

不要語の学習は、真の回答者以外をシステムが回答者であると誤推定した場合に行う。この場合は、推定した回答者の上位単語は本来不要語であるが、大きな重みが付与されているものと考えられる。そこで、このような単語の重みには負の報酬を与え、式（12）によりその重みを小さくする。

【 0 0 8 2 】

【数 1 1】

$$w_i^d = \begin{cases} w_i^d - f_i^d & (\text{if } w_i^d \geq f_i^d) \\ 0 & (\text{if } w_i^d < f_i^d) \end{cases} \quad (12)$$

【0083】

なお、式(12)では重みが負になることを防いでいる。また、強化値の等比関数は式(10)と同様で、その初期値は式(13)で与える。

10

【0084】

【数 1 2】

$$f_1^d = c_2 w_1^d \quad (13)$$

【0085】

ここで、 c_2 は定数である。更に、経験的にこのような不要語は一般用語であることが大多数であり、システムが回答者として推定した者だけではなく、真の回答者以外の辞書40でも負の報酬を与えることにより、不要語の学習効率が大きく向上する。そこで、この場合の学習は、システムが誤推定した回答者の上位単語について、真の回答者以外の他の担当者の辞書40にその単語が登録されている場合、同様に負の報酬を与え、その重みを小さくする。

20

【0086】

図16は、特定語学習処理フローであり、強化学習装置3がステップS178において実行する特定語の学習処理を示す。強化学習装置3が、定数 c_3 、 L 、 S 、 $\min_idf$ 、 $\max_ave_ratio$ を定義ファイルから読み込み(ステップS211)、 idf 値再評価処理を実行する(ステップS212)。 idf 値再評価処理については図17を参照して後述する。次に、強化学習装置3が、単語 t にメール解析結果テーブルにおいて対象非回答者の第1位照合単語を設定し(ステップS213)、対象非回答者の $tf \cdot idf$ 辞書41で t を検索して $tf \cdot idf$ に t と照合した単語の $tf \cdot idf$ 値を設定し(ステップS214)、 $\max_tf \cdot idf$ に対象非回答者の $tf \cdot idf$ 辞書41での最大 $tf \cdot idf$ 値を設定し(ステップS215)、 $tf \cdot idf$ が $\max_tf \cdot idf$ と $\max_ave_ratio$ との積以上であるか否かを調べる(ステップS216)。なお、 $\max_ave_ratio$ は全担当者の $tf \cdot idf$ 値及び $idf/conf$ 値の最大値で標準化した平均的な重みを表し、この $\max_ave_ratio$ と $\max_tf \cdot idf$ との積は後述する式(14)の w_p^d となり、対象非回答者のスコアに有効となる重みを表す。 $\max_ave_ratio$ の算出方法は(図18で)後述する。 $f_tf \cdot idf$ が $\max_tf \cdot idf$ と $\max_ave_ratio$ との積以上でない場合、 $f_tf \cdot idf$ に $\max_tf \cdot idf$ と $\max_ave_ratio$ との積を設定し(ステップS217)、積以上である場合、 $f_tf \cdot idf$ に $tf \cdot idf$ と c_3 とから求まる値を設定する(ステップS218)。

30

40

【0087】

この後、強化学習装置3が、対象非回答者の $idf/conf$ 辞書42の第1単語で t を検索して $idf/conf$ に t と照合した単語の $idf/conf$ 値を設定し(ステップS219)、 $\max_idf/conf$ に対象非回答者の $idf/conf$ 辞書42での最大 $idf/conf$ 値を設定し(ステップS2110)、 $idf/conf$ が $\max_idf/conf$ と $\max_ave_ratio$ との積以上であるか否かを調べる(ステップS2111)。なお、 $\max_ave_ratio$ については、図16において前述した通りである。 $f_idf/conf$ が $\max_idf/conf$ と $\max_ave_ratio$ との積以上でない場合、 $f_idf/conf$ に $\max_idf/conf$ と $\max_ave_ratio$ との積を設定し(ステップS2112)、積以上である場合、 $f_idf/conf$ に $idf/conf$ と c_3 とから求まる値を設定する(ステップS2113)。

50

【 0 0 8 8 】

この後、強化学習装置 3 が、tf・idf 辞書 4 1 及び idf/conf 辞書 4 2 における t の tf・idf 値及び idf/conf 値を更新する（ステップ S2114）。即ち、新たな tf・idf 値を、それまでの tf・idf 値に f_tf・idf を加算した値とし、新たな idf/conf 値を、それまでの idf/conf 値に f_idf/conf を加算した値とする。この後、強化学習装置 3 が、 $i = 2$ として（ステップ S2115）、単語 t にメール解析結果テーブルにおいて対象回答者の第 i 位照合単語を設定する（ステップ S2116）。次に、強化学習装置 3 が、対象回答者の tf・idf 辞書 4 1 で t を検索して tf・idf に t と照合した単語の tf・idf 値を設定し、f_tf・idf と S とから新たな f_tf・idf を算出し（ステップ S2117）、対象回答者の idf/conf 辞書 4 2 の第 1 単語で t を検索して idf/conf に t と照合した単語の idf/conf 値を設定し、f_idf/conf と S とから新たな f_idf/conf を算出し（ステップ S2118）、tf・idf 辞書 4 1 及び idf/conf 辞書 4 2 における t の tf・idf 値及び idf/conf 値を更新する（ステップ S2119）。即ち、新たな tf・idf 値を、それまでの tf・idf 値に f_tf・idf を加算した値とし、新たな idf/conf 値を、それまでの idf/conf 値に f_idf/conf を加算した値とする。この後、強化学習装置 3 が、 i が L に等しいか否かを調べ（ステップ S2120）、等しくない場合、 i を $i + 1$ とし（ステップ S2121）、ステップ S2116 以下を繰り返す。等しい場合、ステップ S2121 を省略して、処理を終了する。

【 0 0 8 9 】

図 1 7 は、idf 値再評価処理フローであり、強化学習装置 3 がステップ S212 において実行する idf 値再評価処理を示す。強化学習装置 3 が、メール解析結果テーブルにおいて全非回答者の上位 L 位までの単語集合で idf 値を算出し（ステップ S221）、メール解析結果テーブルにおいて対象非回答者の上位 L 位までの単語の Score(d, t) を idf 値に書き換え（ステップ S222）、メール解析結果テーブルから対象非回答者の min_idf 未満の idf 値の単語と、L + 1 位以下の単語とを削除し（ステップ S223）、メール解析結果テーブルにおいて対象非回答者の単語を idf 値で降順に整列させる（ステップ S224）。なお、ステップ S223 において、min_idf 未満の idf 値の単語、又は、L + 1 位以下の単語が存在しない場合には、削除すべき単語が存在しない場合がある。また、ステップ S223 における単語の削除により i の最大値が L より小さくなることがあるが、この場合、読み込んだ L に代えて当該 i の最大値が用いられる。即ち、ステップ S2120 において i が当該 i の最大値に等しい場合、処理を終了する。

【 0 0 9 0 】

特定語の学習は、真の回答者をシステムが回答者であると推定できなかった場合に行う。この場合は、真の回答者の辞書 4 0 内で重要語の重みが本来の値より小さくなっている恐れがある。これは、真の回答者がある専門用語を固定的に多く使用し、その同意語や類似語を減多に使用しない場合等では同意語や類似語の tf 値が極端に小さくなるために起こり得る。そこで、真の回答者のスコア上位単語と他の担当者の上位単語を比較して、単語の特定性を idf 値の算出方法と同様にして再評価する。即ち、回答者としては推定されなかった真の回答者の上位単語に特定性が認められれば特定語として正の報酬を与え、その重みを大きくする。ここでの強化方法は、idf 値で単語系列を決定する以外は式 (10) の真の回答者を推定できた場合と同様であるが、idf 値が最大となる単語の強化値は式 (11) の代わりに式 (14) を用いる。

【 0 0 9 1 】

【 数 1 3 】

$$f_1^d = \begin{cases} c_3 w_1^d & (\text{if } w_1^d \geq w_p^d) \\ w_p^d & (\text{if } w_1^d < w_p^d) \end{cases} \quad (14)$$

【 0 0 9 2 】

10

20

30

40

50

ここで、式(11)と同様に c_3 は定数である。なお、辞書4中の全ての単語について、標準化した順位と重みの関係は図18に示したような曲線となる。図18の横軸は、各担当者の単語の最大順位(即ち、登録単語数)で各単語の順位を割った値を表し、縦軸は、同様に重みの最大値で各単語の重みを割った値を表す。図18は横軸の各順位に該当する全担当者の平均重み($tf \cdot idf$ 値と $idf/conf$ 値の両方の平均)をプロットしたものである。この曲線の平均傾きを「-1」と仮定し、この曲線と、傾きが「-1」の直線との接点における重みの値を図16における $\max_ave_ratio$ とした。式(14)の w_p^d はこの $\max_ave_ratio$ と最大重みとの積を表す。真の回答者を当該システムが回答者であると推定(分類)できなかった場合は、真の回答者のスコア $Score_i^d$ の上位単語は当該担当者の辞書40内で低順位であることが多いと考えられる。このような低順位のもの重みは、高順位のもの重みと比較して何桁も小さな値となっている。このため、強化値を元の重みの定数倍として与えても全体のスコアに影響する大きさにならない。そこで、学習の効果が現れないことを受け、スコア1位の w_1^d がこの w_p^d より小さければ、一度に w_p^d まで引き上げ、強化することとした。なお、本発明者の検討によれば、単純な重みの平均値を w_p^d とすると、重みの小さい単語が非常に多いために、 w_p^d は小さな値となり、学習効果が上がらないことが判った。

10

【0093】

学習時の単語の照合は、この例では完全一致としている(従って、図13~図17においては単語は必ず照合されることになる)が、これは本発明者の検討により、別文書ファイル7を学習に用いる場合、メール解析のように部分一致としなくても学習効果が上がることが確かめられたためであり、学習にメール8を用いる場合には、部分一致として強化値に単語の一致率を掛けて、強化しても良い。

20

【0094】

以上、本発明をその実施の形態に従って説明したが、本発明は、その主旨の範囲内で種々の変形が可能である。例えば、本発明の文書分類システムは、メール8に限られず、広く文書の分類に用いることができる。更に、分類の対象である文書に、例えばホームページを含むことができる。更には、分類の対象である文書に、例えば音声入力されたデータを音声認識して得た電子データを含むことができる。本発明の文書分類システムは、自然言語(記号を含む)により構成された電子データであれば、これを分類することができる。また、文書の分類先は、個々の担当者に限られず、種々の組織内における担当部署であっても良い。本発明の文書分類システムは、3個のサブシステムの全てを備えなくとも良い。即ち、辞書編集装置1、文書分類装置2、強化学習装置3の各々を独立に設けても良く、文書分類装置2のみを設けても良く、文書分類装置2に辞書編集装置1又は強化学習装置3を併設しても良い。

30

【産業上の利用可能性】

【0095】

以上説明したように、本発明によれば、辞書編集装置において、 $tf \cdot idf$ 値と $idf/conf$ 値とを別々の独立した2個のパラメータとして用いて $tf \cdot idf$ 辞書、 $idf/conf$ 辞書を作成することができるので、これを文書分類装置の辞書として用いることにより、電子メール等の文書の分類に用いた場合、実用的な分類精度を得ることができ、文書を適切な担当者に自動的に正確に分類するための辞書を容易に作成することができる。

40

【0096】

また、本発明によれば、文書分類装置において、前述の2個の辞書 $tf \cdot idf$ 辞書、 $idf/conf$ 辞書を用いることにより、電子メール等の文書の分類に用いた場合、実用的な分類精度を得ることができ、文書を適切な担当者に自動的に正確に分類することができる。

【0097】

また、本発明によれば、文書分類システムプログラムを、フレキシブルディスク、CD-ROM、CD-R/W、DVD等の媒体に格納すること、又は、インターネット等のネットワークを介してダウンロードすることにより供給することができ、前述の文書分類装置を容易に実現することができ、正確な文書分類を可能とすることができる。

50

【図面の簡単な説明】

【0098】

【図1】文書分類システム構成図である。

【図2】辞書説明図である。

【図3】文書分類結果（メール解析結果）説明図である。

【図4】文書分類処理フローである。

【図5】辞書編集処理フローである。

【図6】担当者・単語テーブル生成処理フローである。

【図7】単語の組み合わせテーブル生成処理フローである。

【図8】単語の組み合わせテーブル説明図である。

10

【図9】問い合わせメール解析処理フローである。

【図10】tf・idf 値重み付き加算処理フローである。

【図11】idf/conf 値重み付き加算処理フローである。

【図12】強化学習処理フローである。

【図13】重要語学習処理フローである。

【図14】不要語学習処理フローである。

【図15】不要語強化処理フローである。

【図16】特定語学習処理フローである。

【図17】idf 値再評価処理フローである。

【図18】順位と重みの関係説明図である。

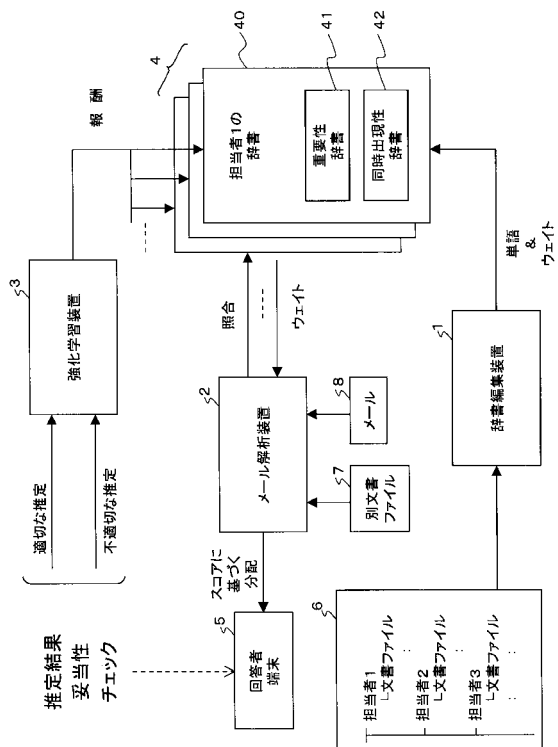
20

【符号の説明】

【0099】

- 1 辞書編集装置
- 2 文書分類装置（メール解析装置）
- 3 強化学習装置
- 4 辞書
- 40 カテゴリ辞書
- 41 重要性辞書（tf・idf 辞書）
- 42 同時出現性辞書（idf/conf 辞書）

【図 1】



【図 2】

(A)

単語	tf値	idf値	tf-idf値
単語a1	0.00424	1.838	0.00779
単語a2	0.00354	1.1451	0.00405
単語a3	0.0114	0.3502	0.00399
...	...	...	...

(B)

第1単語	第2単語	第1単語のidf値	conf値	idf/conf値
単語b1	単語b2	1.838	0.01207	152.3
単語b3	単語b4	1.1451	0.01263	90.67
単語b5	単語b6	1.8382	0.02241	82.03
...	...	...	...	...

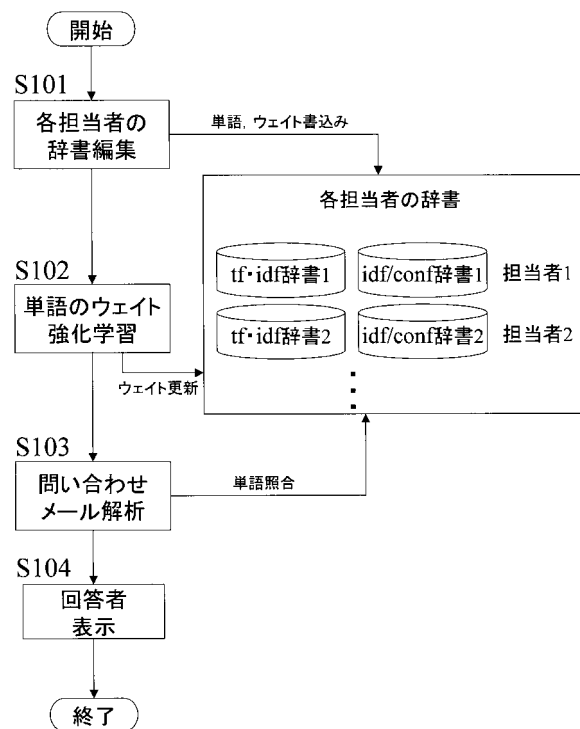
(C)

	d1	d2	d3...	df(t)
t1	1	11	5...	32
t2	0	25	2...	14
t3	33	0	0...	2
...	...	...	...	...
$\sum tf(t', d)$	8538	12945	35792	

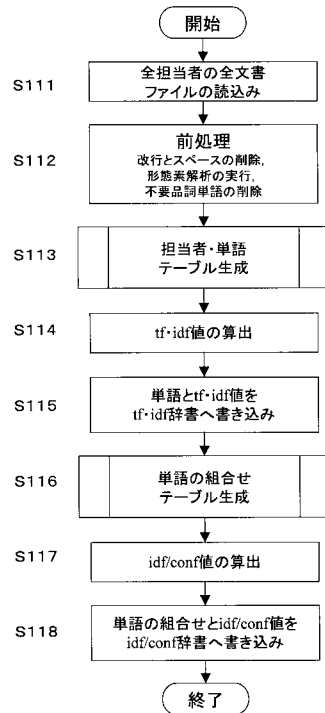
【図 3】

回答者	担当者名d	Score(d)	照合単語t	Score(d,t)
○	担当者1	0.008297	単語11 単語12 単語13 単語14 単語15 ...	0.003221 0.002776 0.000904 0.000586 0.000277 ...
○	担当者2	0.003147	単語21 単語22 単語23 単語24 単語25 ...	0.000917 0.000742 0.000447 0.000421 0.000183 ...
x	担当者3	0.003051	単語31 単語32 単語33 単語34 単語35 ...	0.001869 0.000291 0.000241 0.000122 0.000088 ...
...	...	...	...	...

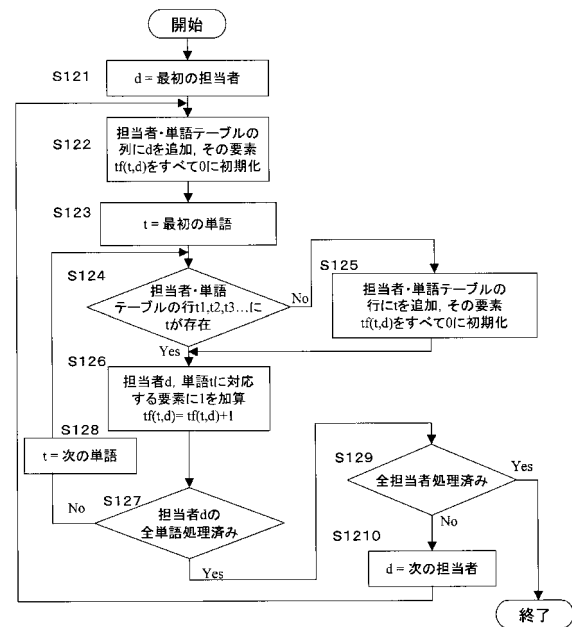
【図 4】



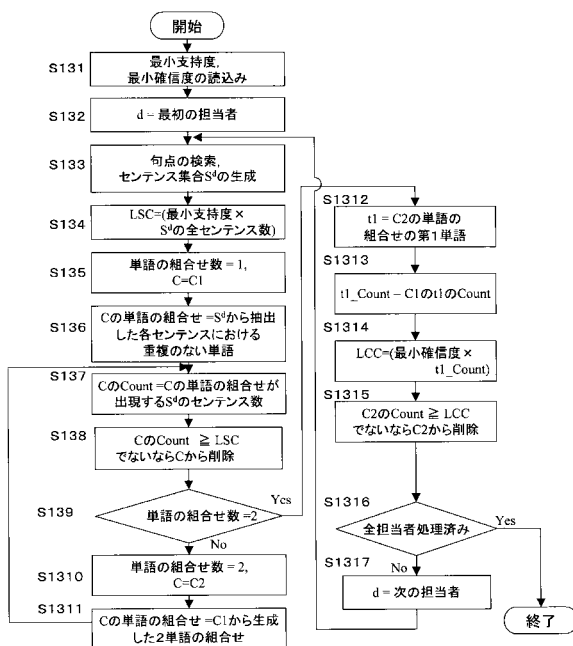
【図5】



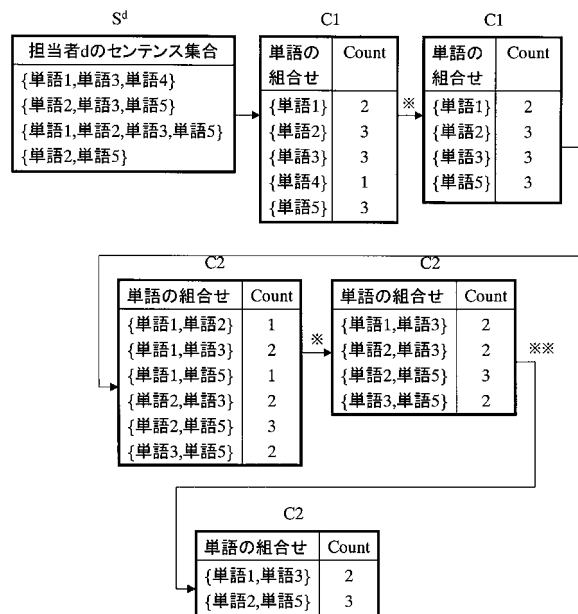
【図6】



【図7】

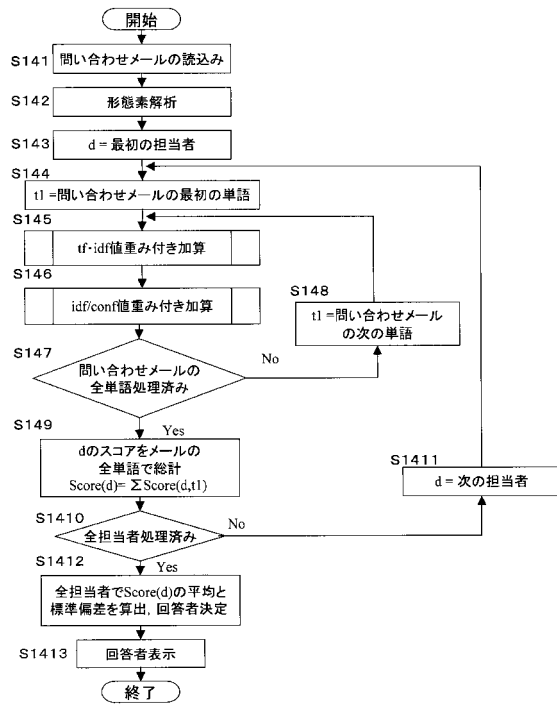


【図8】

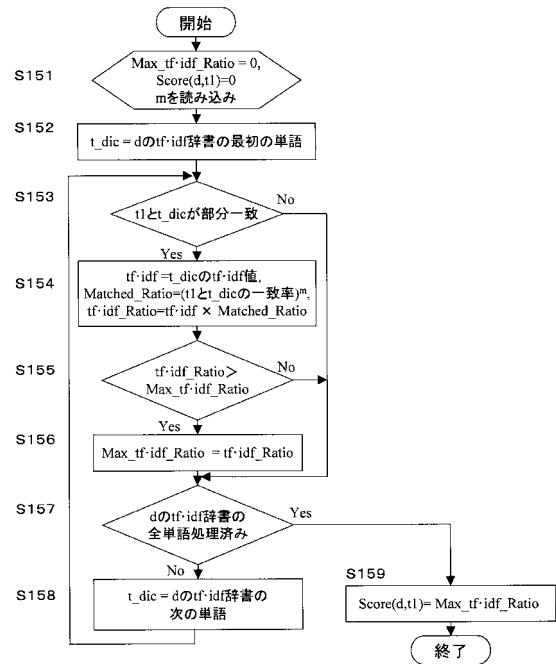


※: 最小支持度 = 0.3 ※※: 最小確信度 = 0.7

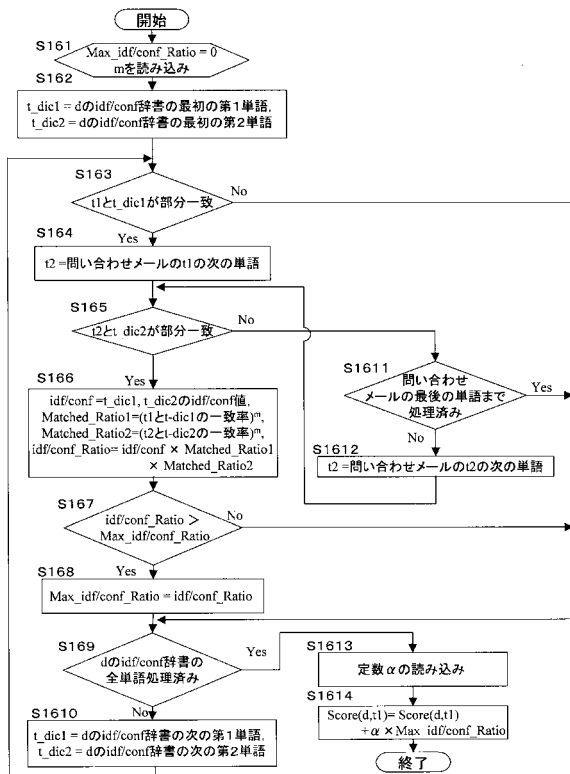
【図 9】



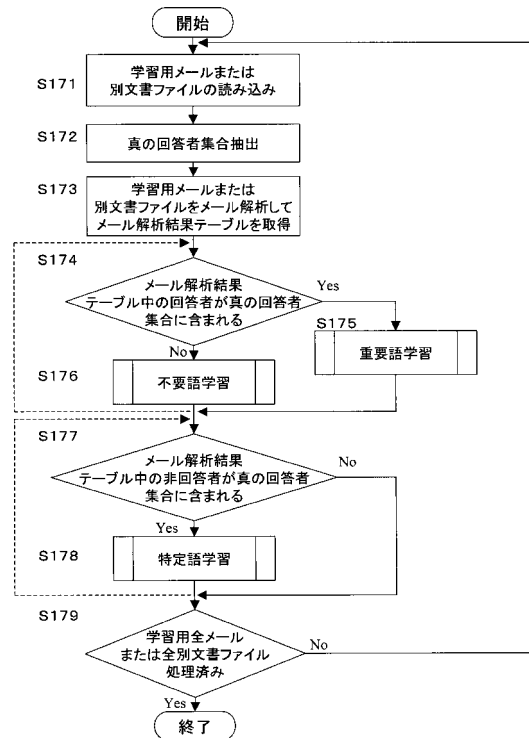
【図 10】



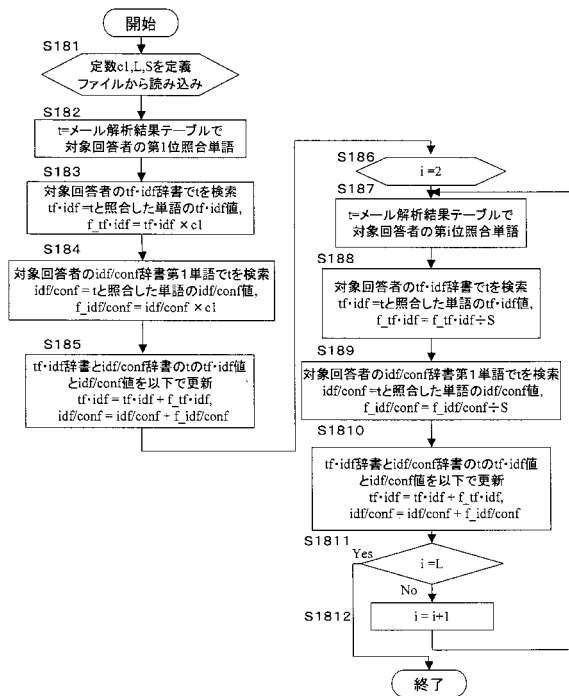
【図 11】



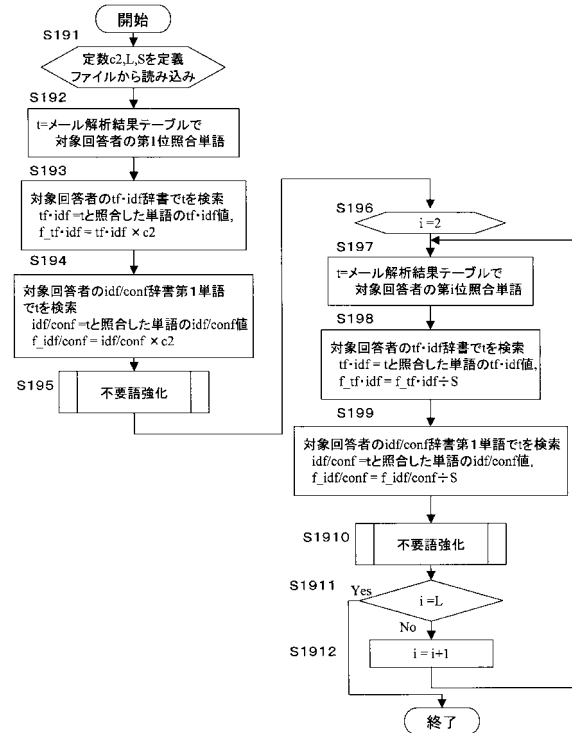
【図 12】



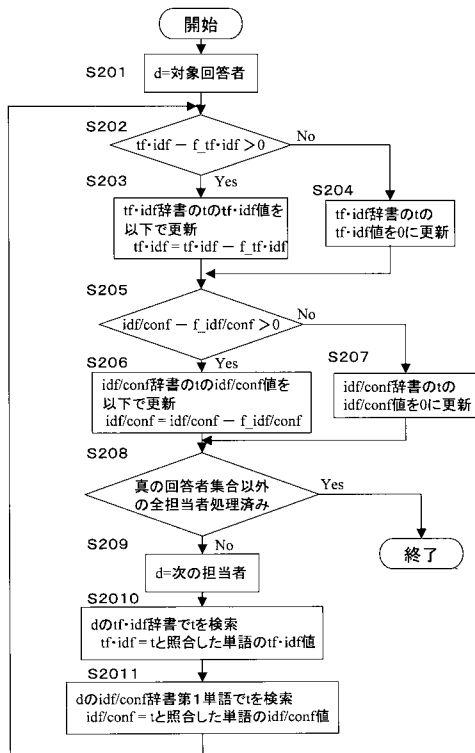
【図 13】



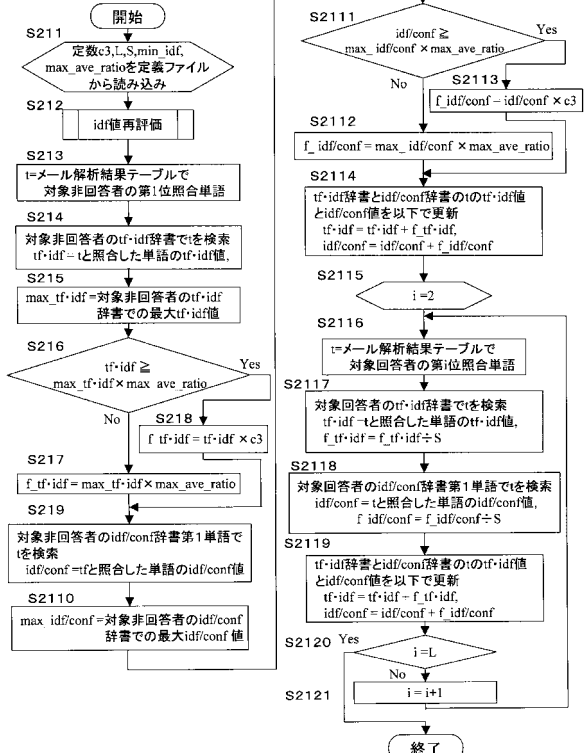
【図 14】



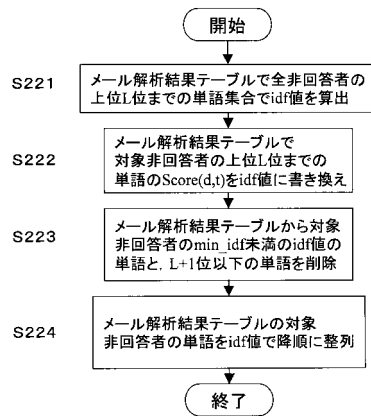
【図 15】



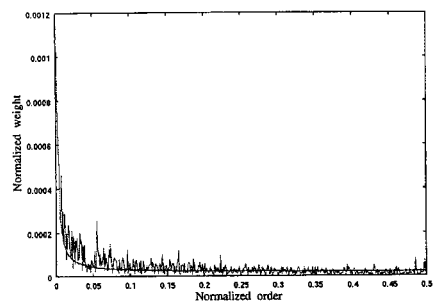
【図 16】



【図 17】



【図 18】



フロントページの続き

合議体

審判長 田口 英雄

審判官 飯田 清司

審判官 池田 聡史

(56)参考文献 特開平11-167581(JP,A)

特開2001-256251(JP,A)

特開2002-183175(JP,A)

上田他、テキストマイニングと強化学習を用いた電子メール自動分配、電子情報通信学会論文誌、D-I Vol. J87-D-I No. 10、2004年10月、p887-898

「特別企画 日本語完全対応 eメール自動分類・配信ソリューション-MatchMail-CallCenter-」、ビジネスコミュニケーション、Vol. 37 No. 5、2000年5月、p. 36-41