

機械翻訳を用いた英日・日英言語横断検索に関する一考察

酒井 哲也[†] 梶浦 正浩[†] 住田 一男[†]

Gareth Jones^{††} Nigel Collier^{†††}

本稿では、日本語テストコレクション BMIR-J2 およびこれを英訳したデータを用い、情報フィルタリングシステム NEAT と機械翻訳システム ASTRANSAC による日・英間の言語横断検索の検索精度評価実験を行う。英語検索要求による日本語文書の検索実験では、文書の日英機械翻訳と検索要求の英日機械翻訳のアプローチを比較し、前者のアプローチにより単言語検索の場合の 90% 以上、後者のアプローチにより 80% 以上の精度が実現できることを示す。さらに、日本語テストコレクションの検索要求を手で英訳して英日言語横断検索の模擬実験を行う場合に 2 人の翻訳者を起用することにより、検索精度が人手による翻訳結果に大きく左右されることを示す。一方、日本語検索要求による擬似英語文書の検索実験では、検索要求の日英機械翻訳の前後にローカルフィードバックを行うことにより、単言語検索の場合と同程度の精度を実現する。これは実際の英語文書を検索対象とした場合の精度の上限を示すものと考えられる。最後に、機械翻訳精度と検索精度との関係について考察を行う。

A Study on English-Japanese/Japanese-English Cross-language Information Retrieval Using Machine Translation

TETSUYA SAKAI,[†] MASAHIRO KAJIURA,[†] KAZUO SUMITA,[†]
GARETH JONES^{††} and NIGEL COLLIER^{†††}

We study a cross-language IR approach using the NEAT information filtering system and the ASTRANSAC machine translation system. The BMIR-J2 standard Japanese test collection and our own translated data are used for evaluation. In our *English-Japanese* experiments, we compare the retrieval performance of the *document translation* approach and the *query translation* approach, and show that they achieve over 90% and 80% of monolingual performance, respectively. Moreover, in our query translation experiments, we consider the case where the Japanese queries are translated into English by two different translators, and show that the manual translation can greatly affect the performance in traditional evaluation using a monolingual baseline. In our *Japanese-pseudo-English* experiments, we perform local feedback both before and after query translation, and achieve retrieval performance of the monolingual standard. This can be regarded as an upperbound of performance in Japanese-English cross-language IR using real English documents. Finally, we study the relationship between the quality of machine translation and retrieval performance.

1. はじめに

近年のインターネットの普及を背景に、世界中のユーザが自分の使いやすい言語で検索要求を表現し、これに適合する情報を世界中の多様な言語で表現された情報データベースから探し出すという多言語情報検

索 (multilingual information retrieval: MLIR) の環境が現実的なものとなりつつある。本稿では、MLIR のサブタスクとして、言語 L_1 で記述された検索要求を満たす言語 L_2 ($\neq L_1$) で記述された情報を検索するタスクと定義される言語横断検索 (cross-language information retrieval: CLIR) を扱う。CLIR は、たとえば以下の場合に有用である。

- 母国語が L_1 であるユーザ A が、言語 L_2 で記述された情報データベースから情報収集しようとしている。A 氏は、検索要求を言語 L_2 により自分でうまく記述する自信はないが、言語 L_2 で記述された情報を眺めて必要な情報を得る程度の語学

[†] 株式会社東芝研究開発センター

Toshiba R&D Center

^{††} Department of Computer Science, University of Exeter

^{†††} 東京大学工学部

Department of Information Science, University of Tokyo

力は有している。

- 母国語が L_1 であるユーザ B が、言語 L_2 で記述された情報データベースから情報収集しようとしている。B 氏は、言語 L_2 に関する知識をほとんど持たないため、検索された情報をシステムが選択的に言語 L_1 に翻訳してから提示する。

このようなニーズから、欧米ではここ数年 CLIR の研究がさかんである。たとえば、米国規格協会 (NIST) が 1992 年から開催している情報検索システムの評価会 TREC (Text REtrieval Conference) では、1997 年の TREC-6 から欧米語の検索要求および文書を対象とした CLIR が扱われるようになった¹⁾。これにならって、日本でも学術情報センターが日本語検索要求と英語文書を扱う CLIR のタスクを含む評価会 NTCIR (NACSIS Test Collection for Information Retrieval Systems) を 1999 年 9 月に開催する予定である²⁾。

CLIR では、以下のようなリソースが利用される。

- (a) 機械可読辞書 (machine-readable dictionaries: MRD) やシソーラス
- (b) parallel/comparable/unlinked/monolingual corpora^{*}
- (c) 機械翻訳システム (MT)

また、CLIR における検索要求と検索対象の関連性を求める手段、換言すれば CLIR を単言語検索 (monolingual IR: MIR) に帰着させる手段としては以下が考えられる。

- (i) 検索要求を翻訳する。
- (ii) 検索対象を翻訳する。
- (iii) 言語に依存しない中間言語表現を用いる。

CLIR における様々なアプローチにはそれぞれ一長一短がある。たとえば、多言語かつ多方向の CLIR を扱えるという長所を持つ (iii) は、一般に (b) の入手可能性に依存する。また、入手可能性が比較的高い (a) を利用する場合も、多義語の曖昧性解消のために (b) を併用せざるをえない場合が多い。逆に (b) からスタートして (a) の多言語シソーラスを構築し、これを CLIR に利用する事例もある。一方、(c) はもともと曖昧性解消の機構を有しているが、多言語および多方向への拡張のために新たな開発労力を要する。現状では、

(a), (b) と (i), (iii) の組合せが主流といえる^{3)~14)}。しかし、たとえば TREC-6 の比較評価では (c) の有効性が再認識されており、さらに (c) のもとでは (i) よりも (ii) のほうが有効であることが示唆されている^{1),15)}。(ii) は翻訳コストの観点から敬遠されることが多いが、たとえばユーザの母国語が既知である場合に、前処理の段階で検索対象を翻訳しておくことにより高精度な CLIR を実現するといった応用可能性がある。

本稿では、情報フィルタリングシステム NEAT^{16),17)} および機械翻訳システム ASTRANSAC^{18),19)} を用い、日本語と英語を扱う CLIR のうち (c) と (i), (ii) の可能性を探るための検索精度評価実験を行う。評価の題材としては、日本語情報検索システム評価用テストコレクション BMIR-J2²⁰⁾ およびこれを独自に英訳したものをを用いる。2 章で本研究の位置付けについて、3 章で CLIR の評価のベースラインとなる日本語単言語検索実験について、4 章で英語検索要求による日本語文書の検索実験について、5 章で日本語検索要求による英語文書の検索実験について述べる。さらに、6 章で考察を述べ、7 章でまとめを述べる。

2. 本研究の位置付け

MT を CLIR に利用している従来研究には、主として検索結果の提示のために MT を利用する試み¹⁴⁾ や、parallel/comparable corpora の代わりに monolingual corpora と MT を併用する試み^{3),12)} などがある。しかし、前章で述べたように、検索要求あるいは検索対象を MT により翻訳して CLIR を実現した事例は少ない。さらに、MRD やコーパス利用の研究の中には「現状の MT の性能では高精度な CLIR は期待できない」という主張さえしばしば見られる^{9),13)}。一方、MT を積極的に利用している研究者は「現状の MT は人が読んで満足するレベルには達していないが、CLIR には十分有用である」と主張する²¹⁾。我々もこの立場に立ち、MT の利用により高精度な CLIR が実現できることを示す。

本稿は以下の 4 つの特徴を持つ。

特徴 1: MT を用いた日英間の CLIR を扱い、かつその検索精度をテストコレクションを用いて定量的に評価している。

MT を用いた CLIR で日本語を扱い標準的な検索精

^{*} 制約が強い順、したがって実際の入手可能性が低い順に列挙してある。parallel corpora は文書とそれを直接翻訳したものが対応付けられているものであり、comparable corpora は同一トピックについての記述を含む複数言語の文書が対応付けられたものである。対応付けの単位には文書・文・語レベルがあり、対応関係の定義も様々である^{3),4)}。これに対し unlinked corpora は対応付けのされていない複数言語の monolingual corpora の集合である。

[☆] (社) 情報処理学会・データベースシステム研究会が、新情報処理開発機構との共同作業により、毎日新聞 CD-ROM'94 データ版を基に構築した情報検索システム評価用テストコレクション BMIR-J2 を利用。

度評価を行った従来研究は、我々の知るかぎり英語検索要求と日本語文書を扱った我々自身の研究²²⁾のみである。

特徴 2：日英間の双方向の CLIR を扱っている。

TREC-6 などでは欧米語間の双方向 CLIR を扱っている例が見られるが^{8),12)}、日本語を扱う CLIR では従来このような試みは見られなかった。ただし、日本語を扱う CLIR の評価用テストコレクションが現時点では公開されていないため、本稿では、日本語コレクション BMIR-J2 の全文書を MT で翻訳することにより、日本語検索要求と擬似英語文書を扱う CLIR の実験を可能にした（本稿ではこれを Je-CLIR と表記する。小文字の“e”は、純粋な自然言語ではなく日英 MT の出力であることを表している）。さらに、BMIR-J2 の検索要求を人手で英訳することにより、英語検索要求と日本語文書を扱った CLIR の実験を可能にした（これを EJ-CLIR と表記する）。

特徴 3：EJ-CLIR において、検索要求を翻訳した場合と文書を翻訳した場合を比較検討している。

Oard らは、TREC-6 の英語検索要求とドイツ語文書を扱った CLIR において検索要求を英独 MT により翻訳した場合と文書を独英 MT により翻訳した場合の性能比較を行っており、一般には後者のアプローチのほうが多くの文脈を活用できることから高精度を実現する可能性が高いことを示唆している¹⁵⁾。しかし、日本語を扱う CLIR においてはこのような比較はこれまでに報告されていない。

特徴 4：EJ-CLIR において、複数の翻訳者が日本語検索要求を英訳した場合を比較検討している。

通常、言語 L_1 の検索要求と言語 L_2 の検索対象を扱う CLIR の実験は以下の手順で行われる。

ステップ 1：言語 L_2 の検索要求 Q_2 および検索対象 D_2 を用いて MIR の実験を行う。

ステップ 2： Q_2 を人手により言語 L_1 の検索要求 Q_1 に翻訳する。

ステップ 3： Q_1 と D_2 を用いた CLIR の実験により、MIR の場合の何%の精度が達成できるかを示す。

しかし、上記手順で得られる評価値は、ステップ 2 の人手による検索要求の翻訳結果に大きく依存すると考えられる。そこで本稿では、2 人のバイリンガルにより検索要求を英訳し、それぞれを基に EJ-CLIR の実験を行う。このような試みは欧米でもこれまでに報告されていない。

なお、上記の CLIR 評価手順は、文書の言語を揃えて MIR と CLIR の比較を行うという TREC において採用された方式である。一方、NTCIR では検索対

象が parallel corpora に近いいため、検索要求の言語を揃えて MIR と CLIR の比較が行われる予定である²⁾。

3. ベースライン：日本語単言語検索 (J-MIR)

本章では CLIR のベースラインとなる日本語 MIR (以下、J-MIR) の実験について述べる。我々は文献 17) において、確率検索モデルと local feedback (以下、LF) を用いた高精度な J-MIR について報告した。本稿における J-MIR はこの実験を踏襲したものである。概要のみを説明する。なお、検索精度の評価尺度としては 11 点平均適合率²³⁾ (以下、11pt と表記する) を用いる。

- (1) BMIR-J2 の日本語検索要求 50 件 (Q_J と呼ぶ) を形態素解析し、名詞などを抽出して初期検索条件を生成する。
- (2) 式 (1) の検索語重み tw の総和に基づき NEAT により初期検索を行い、パラメータ K および b を最適化する。
- (3) LF を行う。すなわち、初期検索結果の上位 n 件を正解文書と見なし、この本文中から式 (2) の値の降順に新たな検索語 (展開語) を m 個選定して検索条件に付加する。
- (4) 再検索を行い n および m を最適化する。なお、本稿ではすべての実験において NEAT の text 条件のみを利用し、付加する検索条件の条件重み w は 0.2 で固定した^{16),17)}。

$$tw(t, d) = \frac{\log(|C|/df(t)) * tf(t, d) * (K + 1)}{K * \left((1-b) + b * \frac{l(d) * |C|}{\sum_{d \in C} l(d)} \right) + tf(t, d)} \quad (1)$$

$$rdf(t) * \log \frac{rdf(t) + 0.5}{\frac{|R| - rdf(t) + 0.5}{df(t) - rdf(t) + 0.5}} \quad (2)$$

ここで、

C : 検索対象となる文書集合、

R : 正解文書集合、

$l(d)$: 文書 d の文字数、

$tf(t, d)$: 文書 d 中における検索語 t の出現頻度、

$df(t)$: 文書集合 C の中で検索語 t を含む文書あるいは見出しの数、

K : $tf(t, d)$ の影響を調整するために経験的

表 1 J-MIR の結果
Table 1 J-MIR results.

名前	Q_J の説明	11pt
J-INIT	初期検索 ($K = 0.4, b = 0.5$)	0.463
J-LF	LF ($n = 4, m = 15$)	0.491
J-GLF	LF で n, m をグループごとに最適化	0.500

表 2 Q_J の初期検索語数に応じて最適化したパラメータ

Table 2 Parameters optimized for Q_J according to the number of initial search terms.

初期検索語数	1	2	3	4	5	6
擬似正解文書数 n	8	6	8	4	4	4
展開語数 m	15	5	15	20	30	5

に決定される定数 ($0 \leq K$),

$b: l(d)$ の影響を調整するために経験的に決

定される定数 ($0 \leq b \leq 1$),

$rdf(t)$: 検索語 t を含む正解文書の数.

今回の実験と文献 17) との違いは, 検索条件展開 (query expansion) に加えて検索語の重み調整 (term reweighting) も行った点である. これは具体的には, 式 (1) の $\log$ の項を式 (2) の $\log$ の項で置き換えることにより実現される²⁴⁾. 検索条件展開のみの効果は初期検索精度の 5% 程度であったが¹⁷⁾, これと重み調整の併用による効果は 6% 程度になることを予備実験により確認した.

表 1 に J-MIR の結果を示す. ここで, J- は J-MIR を意味し, INIT は K, b を最適化した初期検索結果を, LF は LF の n, m を検索要求全体について最適化した結果を表す. また GLF は, まず初期検索語数に応じて検索要求をグループ化し, 表 2 のように各グループごとに n, m の最適化を行ったものである. これは, LF によりどの程度文脈を補う必要があるかは初期検索条件中にどの程度の情報が含まれているかに依存するという仮説に基づくものである. なお, 符号検定²⁵⁾で有意差が見られたのは J-INIT と J-GLF の間 ($\alpha = 0.01$) のみであった.

本稿では, CLIR 評価のベースラインとして J-INIT および J-GLF を用いる. なお, BMIR-J2 を用いた J-MIR においてこのような高精度を達成した例はこれまでに報告されていない¹⁷⁾. また, 2 章で述べたように, EJ-CLIR のベースラインとして J-MIR を用いることは文書の言語を揃えた TREC 方式といえとともに, Je-CLIR のベースラインとして J-MIR を用いることは検索要求の言語を揃えた NTCIR 方式といえる. 図 1 はこの関係を図示したものである.

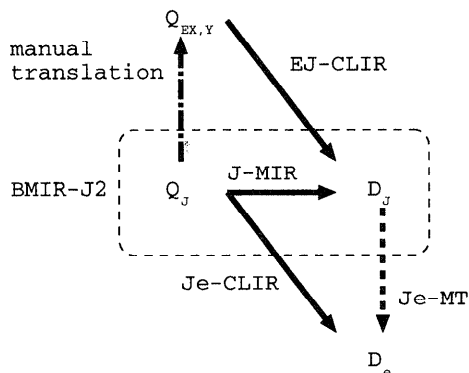


図 1 BMIR-J2 を用いた J-MIR, EJ-CLIR および Je-CLIR
Fig. 1 J-MIR, EJ-CLIR and Je-CLIR using BMIR-J2.

表 3 検索要求文の例
Table 3 Example search requests.

(1) Q_J	行政機関が関係する不況対策
(2) Q_e	Measure against depression relating with a governmental agency
(3) Q_{EX}	Anti-depression measures related to governmental agencies
(4) Q_{EY}	Specific measures taken by government administrators to combat the recession
(5) Q_{JX}	反低下法案は政治の機関に関係があった
(6) Q_{JY}	後退と格闘するために政府行政官によって得られた特定の法案

4. 英語検索要求による日本語文書の検索 (EJ-CLIR)

4.1 EJ-CLIR のための英語検索要求の作成

まず, Q_J をバイリンガルの日本人 X および Y が独自に英訳した. 各英語検索要求セットを Q_{EX} および Q_{EY} と呼ぶ. このうち X と Y の訳が同一になったものは 9 件で, これらに関しては以後 X と Y のすべての検索結果が同一となる. 表 3 の (1), (3), (4) に Q_J およびこれに対応する Q_{EX} と Q_{EY} の例を示す. 検索要求の表現は翻訳者によりかなりばらつきがあり, これは CLIR の精度に直接影響を与えたと考えられる.

4.2 文書の翻訳による EJ-CLIR

本節では, 日本語文書の機械翻訳による EJ-CLIR (EJD と呼ぶ) の実験について述べる. 実験手順は以下のとおりである.

- (1) BMIR-J2 の文書 5,080 件 (D_J と呼ぶ) を日英 ASTRANSAC により擬似英語文書 (D_e と呼ぶ) に翻訳し, NEAT で検索可能にする. なお, 本稿のすべての実験において, CLIR の実験の

text :1, agenc, depress,
government, measur;
text :0.2, busi, recoveri,
bottom, was, econom, plan,
hang, aspect, cycl,
research, april, consumpt,
postwar, still, low,
monthli, februaryi, reduct,
adjust, escap, index,
consider, yen, world,
diffusion, recess, real,
late, longest, uptrend;

図2 展開後の Q_{EX} の検索条件の例Fig. 2 An example expanded query from Q_{EX} .

ための ASTRANSAC の辞書の変更は英日・日英ともにいっさい行っておらず、専門分野辞書も使用していない。

- (2) Q_{EX} および Q_{EY} から生成した初期検索条件を用いて、 D_c に対して J-MIR と同様に初期検索および LF を行う。

このように、EJD は EJ-CLIR を英語検索に帰着させる。なお、NEAT では各英単語を stemming して検索を行う。図 2 に、表 3 の (3) に対応する初期検索条件に対して LF を行った後の検索条件を示す。「text: 1」の後に続いているのが初期検索語、「text: 0.2」の後に続いているのが展開語である¹⁷⁾。

表 4 に EJD の実験結果を示す。INIT, LF, GLF の意味は表 1 と同様である。 Q_{EX} については $K = 0.8$, $b = 0.5$, $n = 10$, $m = 25$, Q_{EY} については $K = 0.6$, $b = 0.7$, $n = 8$, $m = 25$ とした。表中の「%base」の 2 つの値はそれぞれ 11 点平均適合率が J-INIT/J-GLF の何%であるかを表している。この表の縦方向の対に関する符号検定で有意差が見られたのは、 Q_{EX} については EJD-INIT と EJD-LF/EJD-GLF の間 ($\alpha = 0.05$)、 Q_{EY} については EJD-INIT と EJD-LF の間 ($\alpha = 0.05$)、EJD-INIT と EJD-GLF の間 ($\alpha = 0.01$) であった。一方、横方向の対、すなわち Q_{EX} と Q_{EY} の対はすべて有意差なしであった。 Q_{EX} と Q_{EY} の平均精度に著しく差があるにもかかわらず有意差が検出されないのは、表 5 の (1) に示したように、ほとんどの検索要求について X が Y より勝っているわけではなく、X が勝っているものと Y が勝っているものがほぼ同数あるためである。実際、第 1 筆者（翻訳者 X）が Q_{EX} と Q_{EY} の英文を比較した結果、X よりも Y のほうが

表 4 EJD の結果

Table 4 EJD results.

	Q_{EX}		Q_{EY}	
	11 pt	%base	11 pt	%base
EJD-INIT	0.364	79/73	0.323	70/65
EJD-LF	0.416	90/83	0.353	76/71
EJD-GLF	0.427	92/85	0.361	78/72

表 5 X と Y の検索要求ごとの比較

Table 5 Per-query comparisons between X and Y.

11 pt	(1) EJD-		(2) EJQ-	
	INIT	GLF	INIT	GLF
X > Y	26	27	22	28
X = Y	4	0	10	3
X < Y	20	23	18	19
計	50	50	50	50

や口語的である傾向は見られたものの、前者の翻訳の質が後者より全体的に勝っているとはいえなかった。

TREC などにおける欧米言語どうしを扱う CLIR は、日本語と英語を扱う CLIR よりも一般には容易であると考えられる。たとえば、同一文字セットの欧米言語対においては固有名詞などが同一の表記になることが多いので、翻訳に失敗した語は原語のまま検索語として用いることにより検索に成功する可能性がある。それにもかかわらず TREC-6 における CLIR の精度が単言語検索の 50~75% 程度であったことを考慮すると¹⁾、我々の EJ-CLIR は高精度であるといえる。とくに、 Q_{EX} の EJD-LF/EJD-GLF は J-INIT の 90% 以上の精度を実現している。ただし、 Q_{EX} と Q_{EY} の比較から明らかなように、MIR をベースラインとした評価値は翻訳者に大きく依存するため、この数値は 1 つの目安としてとらえるべきであろう。1 人の翻訳者を用いた CLIR の多くの従来研究においても、翻訳者を差し替えることにより MIR をベースラインとした評価値は大きく変わることが予想される。

4.3 検索要求の翻訳による EJ-CLIR

本節では、英語検索要求の機械翻訳による EJ-CLIR (EJQ と呼ぶ) の実験について述べる。実験手順は以下のとおりである。

- (1) Q_{EX} および Q_{EY} の自然言語文を英日 ASTRANSAC によりそれぞれ擬似日本語文 Q_{jX} および Q_{jY} に翻訳する。小文字の “j” は、英日 MT の出力であることを表している。なお、 Q_j の検索要求が完全に復元されたものは Q_{jX} で 13 件、 Q_{jY} で 10 件あった。よって同じ K , b の値を用いればこれらの検索精度は J-MIR と同一 (100%) となる。

text :1, 関係, 機関, 政治, 低下, 反, 法案;
 text :0.2, 成立, 本会議, 改革, 参院, 否決, 可決,
 通過, 与野党, 関連, 審議, 同法案, 野党, 予算, 幹
 事, 護照;

図3 展開後の Q_{jX} の検索条件の例

Fig. 3 An example expanded query from Q_{jX} .

表6 EJQの結果
 Table 6 EJQ results.

	Q_{jX}		Q_{jY}	
	11 pt	%base	11 pt	%base
EJQ-INIT	0.335	72/67	0.262	57/52
EJQ-LF	0.379	82/76	0.313	68/63
EJQ-GLF	0.387	84/77	0.326	70/65

(2) Q_{jX} および Q_{jY} を基に, D_J に対して J-MIR
 と同様に初期検索および LF を行う。

このように, EJQ は EJ-CLIR を日本語検索に帰
 着させる。表3の(5), (6)に, 同表の(3), (4)に対
 応する Q_{jX} および Q_{jY} の例を示す。この例では,
 “depression” や “recession” といった検索要求中の重
 要語の誤訳(それぞれ「低下」「後退」)が見られ, ま
 た Q_{jX} の例では “related” を主動詞と解釈したこと
 による構文解析誤りが見られる。図3に, 表3の(5)
 に対応する初期検索条件に対して LF を行った後の検
 索条件を示す。

表6に EJQ の実験結果を示す。 Q_{jX} については
 $K=1.2$, $b=0.3$, $n=8$, $m=15$, Q_{jY} については
 $K=1.6$, $b=0.2$, $n=6$, $m=25$ とした。この表
 の縦方向の対に関する符号検定で有意差が見られたの
 は, Q_{jX} については EJQ-INIT と EJQ-LF/EJQ-
 GLF の間 ($\alpha=0.01$), Q_{jY} については EJQ-INIT
 と EJQ-LF の間 ($\alpha=0.05$) および EJQ-INIT
 と EJQ-GLF の間 ($\alpha=0.01$) であった。また 4.2 節
 同様, Q_{jX} のほうが Q_{jY} よりも全体的に精度は高
 いが, 横方向の対はすべて有意差なしであった。なお
 Q_{jX} と Q_{jY} との検索要求ごとの優劣は表5の(2)の
 とおりである。

表4と表6の比較から, Oard らの結果¹⁵⁾同様,
 EJQ よりも EJD のほうが全体的に精度が高いこと
 が分かる。一般に, 翻訳時には EJD のほうが多くの
 文脈が活用できるため翻訳精度が高くなり, さらに検
 索時には EJD のほうが誤訳された語の影響が文書中
 の他の語により緩和され検索精度が高くなると考えら
 れる。たとえば, 翻訳者 X の検索要求 ID = 124 に
 ついて EJQ と EJD を比較してみると, EJQ の場合,
 “Regulation relaxation regarding electric communi-

cation”(もともとは「電気通信に関する規制緩和」)
 が「電氣的なコミュニケーションに関する規則弛緩」
 と翻訳されてしまい, 「電氣的」「コミュニケーション」
 「規則」「弛緩」が初期検索語となり, 初期検索精度は
 0.000 となってしまう。一方, EJD では文書全
 体を翻訳するため, このように訳語が全滅してしまう
 確率はさきめて低く, 対応する初期検索精度は 0.262
 となっている。ただし, 比較的小規模な BMIR-J2 を
 用いた今回の実験では, 表4と表6で対応する結果
 の対はすべて有意差なしであった。

なお, MT を用いた CLIR において検索要求翻訳
 のアプローチと文書翻訳のアプローチを比較する際
 には, MT の非対称性を考慮する必要がある。我々の実
 験では, EJD において日英 ASTRANSAC を, EJQ
 において英日 ASTRANSAC を用いたが, 現状では日
 英に比べ英日のほうが翻訳品質が良いことが分かって
 いる。たとえば, 一般用語辞書の規模を比較すると,
 英日は約 24 万語, 日英は約 14 万語である。もしも日
 英と英口の品質が同程度の MT により EJD と EJQ
 の比較実験を行えば, BMIR-J2 の規模でも有意な差
 が検出できる可能性はある。

5. 日本語検索要求による擬似英語文書の検索 (Je-CLIR)

本章では, 日本語検索要求を MT により英訳し, 擬
 似英語文書を検索する Je-CLIR (JeQ と呼ぶ) の実
 験について述べる^{*}。これは, 1999 年 9 月に開催予
 定である NTCIR の JE-CLIR タスクの予備実験とし
 て行ったものである。本来ならば予備実験でも英語文
 書 (D_E と呼ぶ) を用いて JE-CLIR の実験を行い,
 J-MIR や EJ-CLIR との比較を行うべきである。しか
 し, 我々には D_E の持ち合わせがなく, かつこれまで
 に NEAT の英語文書に対する検索精度評価を行った
 ことがなかったので, 代わりに以下の手順で Je-CLIR
 の実験を行うことにした。

- (1) Q_J を日英 MT により擬似英語文 Q_e に翻訳
 し, 初期検索条件を作成する。
- (2) Q_e を基に, 4.2 節で作成した D_e に対して J-
 MIR と同様に初期検索および LF を行う。

表3の(2)に, 同表の(1)に対応する Q_e の例を
 示す。

上記 JeQ の検索精度は, 検索要求の翻訳による JE-
 CLIR (JEQ) の精度の上限と見なすことができる。な

^{*} 自然言語の英語文書を用いていないため, 文書翻訳によるア
 プロチ (JeD) を考慮することはできない。

ぜなら, JeQ が Q_e による D_e の検索, すなわち同一の日英 MT が出力した擬似英語どうしのマッチングを意味するのに対し, JEQ は Q_e による自然言語文書 D_E の検索を意味し, 一般には前者のほうが容易であると推測されるからである。

さらに本実験では, 通常の LF に加えて, 我々が日本経済新聞 CD-ROM1995 年版を利用して独自に作成した別の日本語コレクション TCIR-N1^{16),17)}を用いた翻訳前の LF (pre-translation local feedback)⁵⁾も試みた。CLIR における通常の LF が検索要求の翻訳後 (post-translation) に検索対象そのものを用いて行うのに対し, 翻訳前の LF は, 検索要求と言語が同じでかつ検索対象とは独立な monolingual corpus を用いる。すなわち, これは 1 章で述べた (b) のリソース利用の一例である。コーパスを用いた CLIR において, 翻訳前の LF は翻訳精度をあげるための文脈を補う効果がある。

本実験での翻訳前の LF の手順は以下のとおりである。

- (1) Q_J により TCIR-N1 の文書 5,048 件 (D'_J と呼ぶ) に対して初期検索を行う。 K, b には表 1 の BMIR-J2 における最適値 0.4 および 0.5 を流用する。
- (2) 通常の LF と同様に, 新しい検索語 (T'_J と呼ぶ) を抽出する。ただし, 簡単のために擬似正解文書数 n' は表 1 の最適値 4 で固定し, 展開語数 m' のみを最適化する。
- (3) T'_J の各語を日英 MT により擬似英語に翻訳し (T'_e と呼ぶ), 前述の Q_e に付加する。

図 4 に JeQ の実験の流れを示す。図中で, e-MIR は擬似英語の MIR を, Je-MT は日英 MT を表し, 上半分が翻訳前の LF に, 下半分が翻訳後の LF に対応する。ここで注意すべきは, コーパス利用の CLIR とは違い, 翻訳前の LF は Q_J から Q_e への翻訳精度になんら寄与しないという点である。すなわち, Q_J および T'_J はそれぞれ独立に日英 MT により Q_e および T'_e に翻訳される。また, T'_J は文ではなく語であるので, MT は MRD としてしか機能しない。

表 7 に JeQ の結果を示す。 Q_e については, $K = 0.2, b = 0.9, n = 10, m = 20$ とした。PRE は翻訳前の LF の結果 ($m' = 10$), PRE+GLF は GLF と PRE の検索条件をマージした場合の結果を表している。すなわち, 図 4 の上半分の処理のみ施したものが PRE, 下半分の処理のみ施したものが LF および GLF, そして両方の処理を施したものが PRE+GLF である。符号検定で有意差が見ら

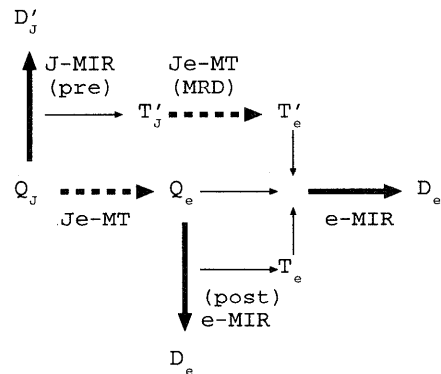


図 4 翻訳前後の LF をともなう JeQ
Fig. 4 JeQ with pre/post-translation LF.

表 7 JeQ の結果
Table 7 JeQ results.

	Q_e	
	11 pt	%base
JeQ-INIT	0.409	88/82
JeQ-PRE	0.426	92/85
JeQ-LF	0.457	99/91
JeQ-GLF	0.461	100/92
JeQ-PRE+GLF	0.462	100/92

れたのは JeQ -INIT と JeQ -PRE/ JeQ -LF の間 ($\alpha = 0.05$) および JeQ -INIT と JeQ -GLF/ JeQ -PRE+GLF の間 ($\alpha = 0.01$) であった。全体的に精度が高いのは, 前述のとおり D_E の代わりに D_e を用いているためである。Je-CLIR により MIR と同水準の精度が実現できたことから, もしも ASTRANSAC の対称性がほぼ成り立つならば, 実際の JE-CLIR の精度は EJ-CLIR と同程度になることが予想される。

JeQ -LF には及ばないものの, JeQ -PRE が有効であることは興味深い。翻訳前の LF の効果は, 利用するコーパス D'_J が検索対象 D_e とどの程度似た性質を持つか, および展開語の翻訳精度に依存すると考えられるが, 本実験の D'_J (TCIR-N1) は 1995 年の日本経済新聞の記事, D_e は 1994 年の毎日新聞の記事を日英 MT により翻訳したものであり, D'_J と D_e との関連性はそれほど高くはないと考えられる。さらに, 前述のとおり, 展開語の翻訳は単に MRD の訳語の第 1 候補を用いていることに相当するため,それほど高精度であるとは考えられない。以上にもかかわらず JeQ -PRE が有効であったことは, 日本語を扱う CLIR において, また MT 利用の CLIR において, 翻訳前の LF が有望であることを示していると考えられる。また, 翻訳前後の LF の併用 JeQ -PRE+GLF と JeQ -GLF との間には有意差がなかったが, 翻訳前

の LF においても n の最適化や JeQ-GLF と同様の最適化を行えば、併用の効果が得られる可能性はある。

6. 考 察

6.1 機械翻訳精度と検索精度の関係について

4 章および 5 章では、従来の CLIR 研究の評価方法に従い、単言語検索の精度の何%を実現できるかという総合的な観点から議論を行った。しかし、本研究における検索精度は、機械翻訳システムの翻訳精度に大きく依存していると考えられる。そこで、本節では機械翻訳精度と検索精度の関係について考察を行う。本研究では検索要求翻訳と文書翻訳とを扱っているため、本来は双方の翻訳精度について考慮する必要があるが、文書集合全体に関する翻訳精度を評価することは難しいため、以後、検索要求翻訳に限定して分析を行う。

まず、EJQ における検索要求の英日機械翻訳結果である Q_{jX} と、JeQ における検索要求の日英機械翻訳結果である Q_e に対し、第 1 筆者が以下のような簡単な基準により各検索要求の翻訳精度を 3 段階で評価した。

ランク A 完璧な翻訳結果であるもの。あるいは、若干言い回しが本来の検索要求と異なっているもの、その内容は完全に訳出されているもの。

ランク B ランク A でもランク C でもないもの。具体的には、検索要求の内容は推定できるが、自然言語として不自然なものがほぼ対応する。

ランク C 検索要求中の主要な概念が明らかに誤訳されているもの、あるいはその翻訳が完全に失敗して原語のまま残ってしまっているもの。あるいは、自然言語として意味不明のもの。

この結果、 Q_{jX} についてはランク A = 20 件、B = 9 件、C = 21 件、 Q_e についてはランク A = 21 件、B = 15 件、C = 14 件となった。

図 5 および図 6 に、それぞれ Q_{jX} および Q_e の各検索要求の初期検索精度をランク別にプロットしたものを示す。横軸は検索要求 ID = 101~150 であり、縦軸は 11 点平均適合率である。黒い四角形がランク A、白い四角形がランク B、白い丸がランク C の検索要求を表している。また、表 8 は、 Q_{jX} および Q_e についてランク別に初期検索精度の平均値を求めたものである。

図 6 の解釈は比較的容易である。まず、黒い四角形が図の上のほうに多く、白い丸が下のほうに多く、白い四角形がその間に多い傾向が見られることから、翻訳精度の良し悪しが初期検索精度の良し悪しに直接影

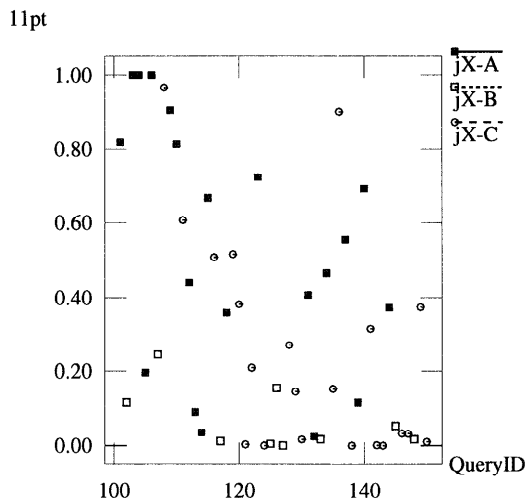


図 5 翻訳精度と EJQ-INIT の関係 (Q_{jX})

Fig. 5 Relationship between translation accuracy and EJQ-INIT (Q_{jX}).

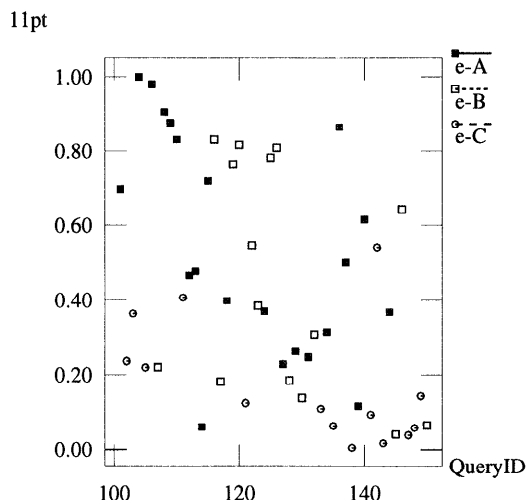


図 6 翻訳精度と JeQ-INIT の関係 (Q_e)

Fig. 6 Relationship between translation accuracy and JeQ-INIT (Q_e).

響を及ぼす場合が多いことが分かる。このことは表 8 の Q_e の欄において、ランクが下がるにつれて平均精度が下がっていることから分かる。さらに、黒い四角形が図の左のほうに多く、白い丸が右のほうに多く、白い四角形がその間に多い傾向が見られるが、これは BMIR-J2 の検索要求が構文的・意味的な難易度のほぼ昇順に並べられているため²⁰⁾、この難易度の増大にともなって機械翻訳精度が低下したものと考えられる。

図 5 では、図 6 ほど翻訳精度と初期検索精度との対応関係が顕著ではない。たとえば、翻訳精度がランク

表 8 翻訳精度と初期検索精度の平均値の関係

Table 8 Relationship between translation accuracy and mean initial performance.

	EJQ-INIT (Q_{jX})	JeQ-INIT (Q_e)
ランク A	0.534	0.538
ランク B	0.069	0.448
ランク C	0.259	0.173

C (白丸) だが初期検索精度が 0.5 以上である検索要求が 5 件も存在し、表 8 の Q_{jX} の欄において、ランク B の平均精度よりもランク C の平均精度がかえって高くなっている。この第 1 の原因としては、図 5 が擬似日本語検索要求 Q_{jX} と日本語文書 D_j のマッチングという現実的な検索の実験結果であるのに対して、図 6 は 5 章で述べたように擬似英語どうしのマッチングという一種の理想的な検索の実験結果であることが考えられる。すなわち、図 6 の Q_e の場合でも、もし D_e の代わりに自然言語文書 D_E を用いて実験を行えば、図 5 と同様にばらつきの多い結果が得られる可能性がある。

第 2 の原因としては、日本語検索に固有の問題により翻訳精度が初期検索精度に反映されない場合があるということが考えられる。たとえば、図 5 における ID = 111 (初期検索精度が 0.607 の白丸) は本来「賃貸住宅」と翻訳されるべきものが「賃貸料用の家」と翻訳され、これは意味をなさないのでランク C と判定されたものである。しかしこの翻訳結果を形態素解析することにより「賃貸」「料」「家」という検索語が抽出され、初期検索はある程度成功している。ここで、仮に「賃貸」ではなく「賃貸料」が形態素解析の結果として得られたとすると、「賃貸住宅」という語は含むが「賃貸料」という語は含まない正解文書が検索できなくなり、検索精度は低下する。このような正解は、実際にブール条件により調べたところ 4 件あることが分かった。また、4.3 節で示した ID = 124 はこの逆の例で、「電気」でなく「電氣的」が検索語として得られてしまったために検索に失敗したものである。以上は、語と語の間の明確な境界がない日本語に対する検索においては、どのような単位を検索語とするかにより検索精度が変わることを示す実例である。このことから、日本語を扱う CLIR は、欧米言語間の CLIR にはない複雑さがかかえていることがうかがえる*。

図 7 は、JeQ-GLF により JeQ-INIT に対して何 % の精度向上が得られたかを図 6 と対応させてプロッ

%Gain in 11pt

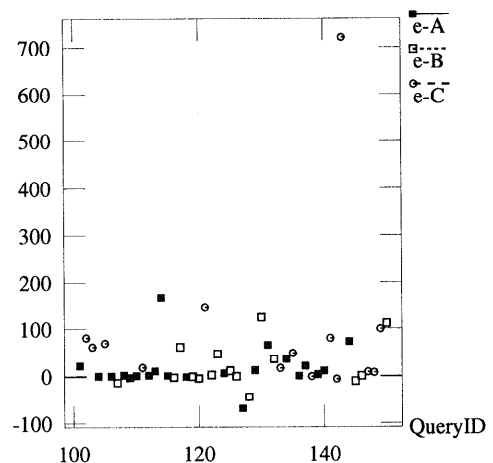


図 7 JeQ-GLF の JeQ-INIT に対する向上率 (Q_e)
Fig. 7 Percentage gain of JeQ-GLF over JeQ-INIT (Q_e).

トしたものである。黒い四角形、白い四角形、白い丸ともに 0% より上にプロットされているものが多いことから、翻訳精度の良し悪しにかかわらず、LF により精度向上が得られることが多いことが分かる。一方、0% より下にプロットされている検索要求が存在することから、LF が悪影響を及ぼす場合があることが分かる。この点は、ノイズを含むデータからの教師なし学習である LF の限界であると考えられる。なお、図中の異常値といえる ID = 143 (向上率が 700% 以上の白丸) は、検索要求「経営陣刷新」の英訳が失敗し、「刷新」という語が原語のまま残ってしまった場合を表している。欧米言語間の検索とは違い、文字セットの異なる原語は検索にまったく役立たないので、初期検索語としては「経営」の訳語である“management”を stemming した“manag”のみが用いられた。このため、初期検索精度は 0.017 と非常に低かったが、LF により検索精度は 0.143 まで向上している。このように、MT の失敗は文字セットの異なる言語間の検索において致命的であるが、その影響を LF が緩和できる場合があることが分かる。

6.2 LF の効果の符号検定結果について

本研究では、J-MIR においては J-INIT と J-GLF の間のみ有意差が見られたが、EJ-CLIR においては X と Y ともに EJD-INIT と EJD-LF/EJD-GLF の間および EJQ-INIT と EJQ-LF/EJD-GLF の間に有意差が見られ、Je-CLIR においては JeQ-INIT と JeQ-PRE/Jeq-LF/Jeq-GLF/Jeq-PRE+GLF の間に有意差が見られた。

* さらに、4.3 節の ID = 124 の例は、“communication”という語を「通信」と訳すか「コミュニケーション」というカタカナ外来語に訳すかという複雑さを示す例でもある。

すなわち、LFの効果は、MIRよりもCLIRにおいてより顕著であった。これは主に、MIRよりもCLIRの初期検索精度が平均的に低いことに起因すると考えられる。我々はこれまでに、文献17)における検索要求グループ²⁰⁾別の評価により、初期検索精度がある程度高い場合には、LFによる精度向上の余地がなくなる傾向があることを示したが、今回の実験でも同様の現象が起こったものと考えられる。

7. ま と め

本稿では、機械翻訳システムとローカルフィードバックを利用した日・英間のCLIRの検索精度評価を行い、その有効性を示した。英語検索要求による日本語文書の検索実験では、文書翻訳のアプローチによりMIRの90%以上の精度を、検索要求翻訳のアプローチによりMIRの80%以上の精度を実現し、またこれらの評価値は検索要求の翻訳者に依存することを示した。日本語検索要求による擬似英語文書の検索実験では、検索要求の翻訳前後のローカルフィードバックを扱い、MIRと同程度の精度を実現した。これは、日本語検索要求と英語文書を扱うCLIRにおける我々のアプローチの上限を示すものと考えられる。いずれの実験においても、ローカルフィードバックが検索精度向上のための有効な手段であることが分かった。さらに、機械翻訳精度と検索精度の関係について、具体例を見ながら考察を行った。

図8は、J-GLF、JeQ-PRE+GLF、EJD-GLF (Q_{EX})、EJQ-GLF (Q_{jX})の再現率・適合率曲線²³⁾である。また表9は、これらのCLIRの結果を検索要求ごとにJ-GLFと比較した結果である。この表から明らかなように、全体の3分の1程度の検索要求についてはCLIRのほうがMIRよりもかえって精度が高い。これは2章で述べた下記の主張を裏付けるものであると考えられる。

「現状のMTは人が読んで満足するレベルには達していないが、CLIRには十分有用である」

たとえば、表3の(2)は、翻訳結果としてはいま一歩であるが、これを初期検索条件に変換すると、表3の(3)の人手で作成した検索要求から作成した初期検索条件と同一になる。すなわち、初期検索語“agenc, depress, government, measur”を含む図2の1番目の条件に変換される。現状のMIRの性能自体が必ずしも十分であるわけではないが、少なくともMTとCLIRとの親和性は高いといえる。

CLIRではMIRに増して検索結果の提示方法が重要

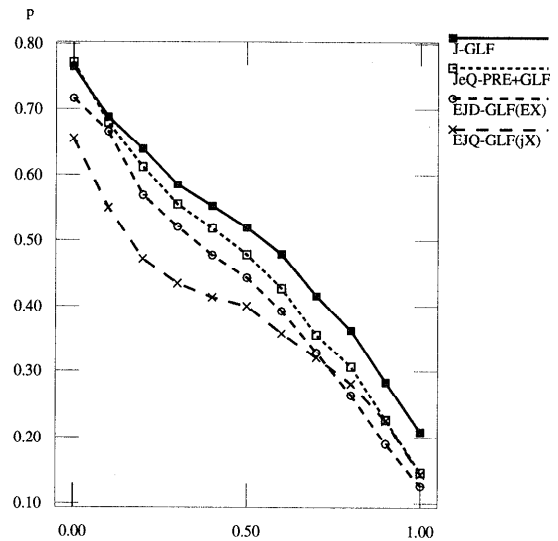


図8 再現率・適合率曲線

Fig. 8 Recall/precision curves.

表9 J-GLFとの検索要求ごとの比較

Table 9 Per-query comparisons with J-GLF.

J-GLF と比べ 11 pt が	(1) EJD-GLF		(2) EJQ-GLF		(3) JeQ- PRE+ GLF
	Q_{EX}	Q_{EY}	Q_{jX}	Q_{jY}	Q_e
高い	13	11	16	13	16
同じ	0	1	2	1	0
低い	37	38	32	36	34
計	50	50	50	50	50

となる。たとえば、1章で述べたB氏のようなユーザに対しては、言語障壁を隠蔽するユーザインタフェースを考慮する必要があると考えられる^{11),14)}。今後は、検索精度・翻訳精度の向上にとどまらず、ユーザが検索結果から実際に情報を入手するまでのインタラクションを含めた研究を進める予定である。

参 考 文 献

- 1) Schauble, P., et al.: Cross-Language Information Retrieval (CLIR) Track Overview, TREC6, <http://trec.nist.gov/pubs/trec6/t6-proceedings.html> (1997).
- 2) NTCIR: <http://www.rd.nacsis.ac.jp/~ntcadm>
- 3) Mateev, B., et al.: ETH TREC-6: Routing, Chinese, Cross-Language and Spoken Document Retrieval, TREC6, <http://trec.nist.gov/pubs/trec6/t6-proceedings.html> (1997).
- 4) Sheridan, P., et al.: Experiments in Multi-

- lingual Information Retrieval using the SPI-
DER system, *Proc. ACM SIGIR '96*, pp. 58-65
(1996).
- 5) Ballesteros, L., et al.: Resolving Ambiguity for
Cross-language Retrieval, *Proc. ACM SIGIR
'98*, pp.64-71 (1998).
 - 6) Carbonell, J.G., et al.: Translingual Informa-
tion Retrieval: A Comparative Evaluation, *IJ-
CAI'97*, pp.708-714 (1997).
 - 7) 藤井ほか：多言語情報検索のための複合語翻訳、
情処学会情報学基礎研究会, 98-FI-51, pp.55-62
(1998).
 - 8) Gaussier, E., et al.: Xerox TREC-6 Site Re-
port: Cross Language Text Retrieval, TREC6,
[http://trec.nist.gov/pubs/trec6/
t6.proceedings.html](http://trec.nist.gov/pubs/trec6/t6.proceedings.html) (1997).
 - 9) Hull, D., et al.: Querying Across Languages:
A Dictionary-Based Approach to Multilingual
Information Retrieval, *Proc. ACM SIGIR '96*,
pp.49-57 (1996).
 - 10) Kando, N., et al.: Cross-lingual Information
Retrieval using Automatically Generated Mul-
tilingual Keyword Clusters, *Proc. 3rd Interna-
tional Workshop on Information Retrieval with
Asian Languages*, pp.86-94 (1998).
 - 11) 菊井ほか：インターネット情報ナビゲーション
における多言語機能, 情処学会自然言語処理の応
用に関するシンポジウム,
[http://titan.mcnet.ne.jp/titan-help/jp/
3.html](http://titan.mcnet.ne.jp/titan-help/jp/3.html) (1995).
 - 12) Littman, M.L., et al.: Automatic Cross-
language Information Retrieval using Latent
Semantic Indexing, *Cross-Language Informa-
tion Retrieval*, Grefenstette, G. (Ed.), pp.51-
62, Kluwer Academic Publishers, Norwell, Mas-
sachusetts (1998).
 - 13) Pirkola, A.: The Effects of Query Structure
and Dictionary Setups in Dictionary-Based
Cross-language Information Retrieval, *Proc.
ACM SIGIR '98*, pp.55-63 (1998).
 - 14) Yamabana, K., et al.: A Language Con-
version Front-End for Cross-language Infor-
mation Retrieval, *Cross-Language Information
Retrieval* Grefenstette, G. (Ed.), pp.93-104,
Kluwer Academic Publishers, Norwell, Mas-
sachusetts (1998).
 - 15) Oard, D.W., et al.: Document Translation for
Cross-Language Text Retrieval at the Univer-
sity of Maryland, TREC6,
[http://trec.nist.gov/pubs/trec6/
t6.proceedings.html](http://trec.nist.gov/pubs/trec6/t6.proceedings.html) (1997).
 - 16) 酒井ほか：情報フィルタリングのためのプー
ル式と文書構造を利用した検索条件生成と検索精
度評価, 情報処理学会論文誌, Vol.39, No.11,
pp.3076-3083 (1998).
 - 17) 酒井ほか：確率モデルに基づく日本語情報フイ
ルタリングにおけるフィードバックによる検索条
件展開および検索精度評価, 情報処理学会論文誌,
Vol.40, No.5, pp.2429-2438 (1999).
 - 18) Hirakawa, H., et al.: EJ/JE Machine Tran-
slation System ASTRANSAC - Extensions to-
wards Personalization, *Proc. MT Summit III*,
pp.73-80 (1991).
 - 19) Kinoshita, S., et al.: ASTRANSAC - Toshiba
Machine Translation System, *Proc. MT Sum-
mit VI*, p.284 (1997).
 - 20) 酒井ほか：情報検索システム評価のためのテス
トコレクション, *Computer Today*, Vol.9, No.87,
pp.31-35, サイエンス社 (1998).
 - 21) Gachot, D.A., et al.: The SYSTRAN NLP
Browser: An Application of Machine Tran-
slation Technology in Cross-Language Infor-
mation Retrieval, *Cross-Language Information
Retrieval*, Grefenstette, G. (Ed.), pp.105-118,
Kluwer Academic Publishers, Norwell, Mas-
sachusetts (1998).
 - 22) Jones, G., et al.: Cross-Language Information
Access: a case study for English and Japanese,
情処学会情報学基礎研究会, 98-FI-51, pp.47-54
(1998).
 - 23) Witten, I.H., et al.: *Managing Gigabytes:
Compressing and Indexing Documents and Im-
ages*, pp.148-151, Van Nostrand Reinhold,
(1994).
 - 24) Robertson, S.E., et al.: *Simple, Proven Ap-
proaches to Text Retrieval*, Computer Labora-
tory, University of Cambridge (1994).
 - 25) Hull, D.: Using Statistical Testing in the Eval-
uation of Retrieval Experiments, *Proc. ACM
SIGIR '93*, pp.329-338 (1993).

(平成 10 年 12 月 17 日受付)

(平成 11 年 9 月 2 日採録)



酒井 哲也 (正会員)

昭和 43 年生。平成 5 年早稲田大
学大学院理工学研究科工業経営学専
門分野修士課程修了。同年 (株) 東
芝入社。自然言語処理, 情報検索・
フィルタリングの研究開発に従事。

**梶浦 正浩 (正会員)**

昭和 41 年生。平成 3 年慶應義塾大学大学院理工学研究科計算機科学専攻修士課程修了。平成 6 年同専攻後期博士課程単位取得退学。同年 (株) 東芝入社。情報検索・情報フィルタ

リング, EC の研究開発に従事。

**住田 一男 (正会員)**

昭和 32 年生。昭和 57 年東京工業大学大学院総合理工学研究科物理情報工学専攻修士課程修了。同年 (株) 東芝入社。自然言語処理, 情報検索・フィルタリングの研究開発に従事。

工学博士。

**Gareth Jones**

Gareth Jones received a B.Eng in Electrical and Electronic Engineering from the University of Bristol, UK in 1989, and a PhD in 1994 from the same institution

for his thesis examining the application of linguistic models in automatic speech recognition. From 1993 to 1996 he was a Research Associate at the University of Cambridge, UK carrying out research on multimedia information retrieval. In 1996 he was appointed as a Lecturer in Computer Science at the University of Exeter, UK. From 1997 to 1998 he was a Toshiba Fellow at the Toshiba R&D Center. His research interests include: information retrieval, information access, speech recognition, natural language engineering, and virtual technology. He is a member of the UK IEE and British Computer Society IRSG.

**Nigel Collier**

Nigel Collier received his BSc. in computer science from Leeds University in the UK in 1992, an MSc. in machine translation and a Ph.D. in language engineering

from UMIST in the UK in 1994 and 1996 respectively. From 1990 to 1991 he was a visitor at the University of Tokyo. From 1996 to 1998 he was a Toshiba Fellow at Toshiba R&D Center and is currently a JSPS research associate at Tokyo University investigating information extraction from genome domain texts. His current research interests are in machine translation, cross-language information retrieval, and knowledge acquisition. He is a member of the ACM and SIGART as well as the British Computer Society.