

低品質印刷文字を高精度に識別する複合認識アルゴリズム

進 藤 宣 博[†] 阿 曾 弘 具^{††} 木 村 正 行^{†††}

ワープロやパソコンのプリンタで印刷された文書は潰れによる低品質な文字が多く、イメージスキャナによる文書読み取り時には、潰れやかすれが発生し低品質になる。このように低品質文字は実用的な文書に多く現れ、その認識技術の確立は文字認識装置の実用化のために必要なことである。本論文は、低品質な文字のうち潰れ文字を含んだ印刷文字の高精度な認識アルゴリズムの開発を目的とする。そのために、まず、従来の認識手法のいくつかについて潰れ文字に対する性能を実験的に調べる。その結果に基づいて、潰れ文字を高精度に認識する複合認識アルゴリズムを提案する。認識実験によりその評価を行い、確かに潰れ文字に対して有効であることを示す。本論文で実験的検討の対象にした特徴量は、高品質な印刷文字の認識で有効性が確かめられている方向線素特徴量とメッシュ特徴量、文字パターン自体を特徴量とみなすドットパターン特徴量の3つである。提案する複合認識アルゴリズムは、方向線素特徴量とメッシュ特徴量を併用して初段認識を行い、ドットパターン特徴量を用いてその正しさの検証をするものである。潰れ文字も含んだ印刷文字を対象に認識実験を行い、99.9% 台の認識率が達成できることを示した。すなわち、複数の特徴量を併用して、それぞれの特徴量の特性を生かした複合認識アルゴリズムにより全体の認識精度が向上することを示した。

A Compound Recognition Algorithm Good for Printed Characters of Bad Quality

NOBUHIRO SHINDO,[†] HIROTOMO ASO^{††} and MASAYUKI KIMURA^{†††}

Documents printed by dot-printers of word-processors involve many blurred or thinned characters. Since there are such characters of bad quality in practical documents, to develop a character recognition system for practical use, it is necessary to establish a technique to recognize such blurred characters of bad quality. In this paper, performances of pattern matching algorithms using oriented-dot feature (originally called directional element feature), mesh feature and dot-pattern feature are investigated for recognizing characters of bad quality. The results tell that each algorithm has own good points. Using the complementary characteristics, a compound recognition algorithm is proposed, which consists of first recognition and verification steps for characters. First recognition step is a compound algorithm of one with oriented-dot feature and one with mesh feature. Verification step is a matching in dot-pattern feature. The algorithm is evaluated by experiments for printed characters including blurred ones. The result is that the recognition rate is more than 99.9%.

1. はじめに

計算機が身近に利用できるようになり、印刷物や書籍等の電子化が要求され、それらを自動的に行うための文字認識装置が注目されている。現在、身近にある印刷物はワープロやパソコンのプリンタで印刷されていることが多く、そのような場合、ドットプリンタ等による印字では潰れが生じやすく、レーザプリンタでも通常のサイズの文字は高品質であるがサイズの小さな文字の印字では潰れが生じ、低品質な文字が多数あ

るものも少なくない。イメージスキャナによる文書読み取り時には潰れやかすれが現れることが多い。また、新聞の文字は紙質のせいでも低品質になってしまうことが多い。このように低品質文字は実用的な文書に多く現れ、その認識技術の確立は文字認識装置の実用化のための重要な課題である。

高品質な印刷文字については、いろいろな認識手法が提案され、実験的な検討もなされ、高い認識精度を得ている^{2), 3), 7), 8), 11), 12)}。しかし、低品質な文字に関する実験的検証は少ない¹¹⁾。そこで、本論文では、潰れのある低品質な文字を含んだ印刷文字の高精度な認識アルゴリズムの開発を目的として、従来の認識手法のいくつかについて低品質文字に対する性能を実験的に

[†] (株)リコー ソフトウェア事業部

^{††} 東北大学工学部通信工学科

^{†††} 北陸先端科学技術大学院大学情報科学研究科

調べる。その結果に基づいて、低品質文字を高精度に認識する複合認識アルゴリズムを提案する。認識実験によりその評価を行い、確かに低品質文字に対して有効であることを示す。

低品質文字は潰れとかすれ文字に大別できるが、プリンタ印字ではかすれはほとんど生じない。また、通常の印字に現われるかすれ文字については特徴量を求める段階でその影響が軽減される。そのため特にかすれさせた文字を対象にする必要がないと判断し、本論文では低品質文字として潰れ文字に注目している。

本論文で実験的検討の対象にした特徴量は、高品質な印刷文字の認識で有効性が確かめられている方向線素特徴量とメッシュ特徴量、文字パターン自体を特徴量とみなすドットパターン特徴量の3つである。

方向線素特徴量^{(11),(12)}は、線素整合法⁴⁾、輪郭線方向パターンマッチング法⁶⁾、加重方向指数ヒストグラム法¹⁰⁾における特徴量と同様の特徴量である。方向線素特徴量以外は手書き文字に対する特徴量として提案され、手書き文字認識において比較的高精度な認識を実現している。それぞれ、方向線素化処理が細線化パターンまたは輪郭線ターンのどちらをベースに得られているか、線素の数を計数するときの重み（ぼけフィルタ）をどのように設定しているか、計数のための部分パターンの大きさをどうするか（次元数が決まる）、などで特徴付けられる。特徴抽出が比較的簡単であること、印刷文字に対する実験結果が知られていることなどの理由で、方向線素特徴量を選んで実験を行った。

方向線素特徴量については、潰れのある低品質な文字について認識精度が低下し、その原因が特徴抽出時に細線化処理を行うことにあと指摘されている¹¹⁾。その理由は、潰れ文字に対して細線化処理を施したとき、細線化の目的である文字の骨格の抽出ができず、元の文字が推測できないようなパターンになってしまうことであり、前処理として細線化処理を必要とする特徴量において一般的に言えることと思われる。

そこで本研究では、潰れの影響の少ない特徴量として細線化処理を行わない特徴量を選び、認識に併用することを考えた。まず、印刷文字の線分の位置等が極端に変動しないことに注目し、文字パターン自体を特徴量とするドットパターン特徴量を採用することにした。ドットパターン特徴量は文字認識研究の初期の頃から考えられていたものであるが、それを用いた整合性の判定に長い時間を必要とする。この処理時間の短縮のために、階層的パターンマッチング法²⁾が考案さ

れ、メッシュ特徴量が導入された。その後、メッシュ特徴量はマルチフォント漢字認識にも適用されその有効性が検証されている^{(3),(7),(8)}。ペリフェラル特徴量³⁾などこのほかにもパターン整合法のための特徴量が提案されているが、処理時間の短縮化の観点からメッシュ特徴量、認識の正確さの点からドットパターン特徴量を検討の対象とした。

提案する複合認識アルゴリズムは、方向線素特徴量とメッシュ特徴量を併用して初段認識を行い、ドットパターン特徴量を用いて初段の認識結果の正しさの検証をするものである。同様な手法は、大分類後に細分類をするという形のものとして、階層的パターンマッチング法^{(2),(5),(7),(8),(11)}、段階的整合方式⁹⁾など、既にいくつか提案されている。それらは、細分類で用いる固有の特徴量の認識に時間がかかることから、大分類用の特徴量を考え、細分類のための候補を絞り込むために大分類を考えている。また、それらの実験による検証は、文字サイズが8mm以上で高品質な文字パターンを対象に、字種もJIS第1水準の漢字をすべて含んでいない当用漢字の文字セットに対して行われている^{(2),(3),(7),(8)}（これは当時の計算機的能力によるところが大きいと思われる）。本論文では、従来から知られているドットパターン特徴量、メッシュ特徴量についても、第1水準の全漢字を対象に通常のワープロ用プリンタ出力文字（約4mm角で低品質なものを含む）を用いて評価実験を行い、その性能を再評価する。また、複合認識アルゴリズムが、大分類—細分類によるものとは性質が異なることが示される。一方、複合認識アルゴリズムは、ニューラルネットによる文字認識で用いられている相互補完的な出力評決法と方式としては類似しているが、3千字種弱という漢字のための高精度な具体的な認識手法として提案し、その有効性を明らかにすることを課題としている。

以下、2章で、認識実験に用いた認識システムの概要を示し、実験データの説明をする。3章で、3つの特徴量の定義を示し、認識実験によってそれぞれの性能を確認する。4章では、複合認識アルゴリズムを提案し、その性能を実験によって確かめる。

2. 認識システムと文字データ

認識システムは、入力部、前処理部、特徴抽出部、識別部の4つの部分からなる。

入力部では、イメージスキャナ（分解能は12ドット/mm）で文書を読み取り、2値パターンに変換した

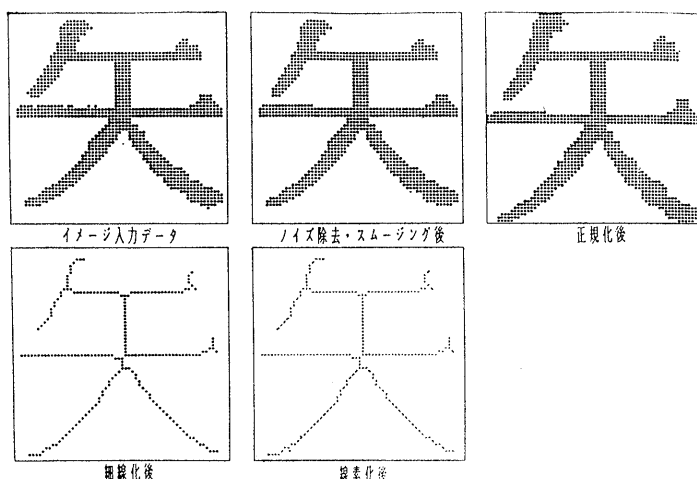


図 1 前処理による文字パターンの変化

Fig. 1 Transformation of character pattern by pre-processing (image, noise reduction, normalization, thinning, dot orientation).

後、文字単位に切り出す処理を行う。

前処理部では、ノイズ除去、平滑化、正規化、細線化、線素化を行う。最後の2つは方向線素特微量を求めるための処理である。この前処理における文字パターンの変化の例を図1に示す。

正規化までの処理は、各特微量に対して共通で、その詳細は割愛するが、切り出された1文字単位のイメージは、 64×64 ドットの枠いっぱいに広がる2値データ（正規化文字パターンと呼ぶ）に変換される。

細線化処理は、幅を持った文字の線分を幅1の線分にする処理であり、Hilditchの方法⁹⁾を採用した。この方法では、文字の線の輪郭を1ドット分割している、線幅が1ドットになると処理を終了する。終了するまでの輪郭を削る処理の回数を細線化回数と呼ぶ。

線素化処理は、細線化された文字イメージに対して、各黒画素の方向を定める処理である（詳細は文献11)参照）。方向の定まった画素を方向線素または単に線素と呼ぶ。方向線素としては、縦線素「|」、横線素「-」、右上斜め線素「/」、左上斜め線素「\」の4つを考える。線素化処理により細線化パターンは方向線素パターンに変換される。

特徴抽出部は、正規化文字パターンないし線素パターンから数値ベクトルである特微量を求める処理で、各特微量ごとに後述する。

識別部では、あらかじめ作成しておいた辞書内の文字ごとの標準特微量と入力文字の特微量の整合性を比較して、よりよく整合する（ほぼ一致する）ものを文

字候補と決定する。整合性は特微量を定める数値ベクトルに対する距離尺度で計る。辞書作成は文字セットの学習ともいわれるが、その詳細は特微量ごとに後述する。

本論文の実験で使用した文字データは、2種類のワイヤドット方式プリンタと1種類のレーザショット方式プリンタで印字したものである。対象文字集合は、JIS第1水準漢字の2965字種で、4倍角、縦倍角、横倍角、全角の4種類の字形をもつものである。認識実験用のサンプルは、各字種ごとに字形4種類の各1文字を3種類のプリンタで印字したもので、全部で12セット35580文字である。識別のための辞書は、各

字種ごとに字形2種類（4倍角、全角）の各2文字を3種類のプリンタで印字したもの、計12文字を使用して作成した。これは、辞書が高品質文字と低品質文字の両方の傾向を標準特微量に反映できていればよいという観点から、中間的傾向をもつ縦倍角と横倍角は不要と考えたためである。辞書用データと認識実験用データは別々に印字したものである。

文字サンプルの品質について、正規化後の文字パターンの例を図2に示す。一般に、ワイヤドット方式のプリンタで印字した文字の潰れがひどいが、全角文字についてはレーザショット方式のプリンタで印字した文字についても潰れが見られる。文字の潰れの程度は、文字パターンの線幅を一定値と考え、細線化回数で評価できる。すなわち、細線化回数は、線分が太くなるほど増加し、潰れが存在するとかなり増加するはずである。表1に、認識対象の文字サンプルについて字形ごとの細線化回数別文字数を示す。表1から、4倍角に比べ、他の字形、特に全角について、回数の多い部分に文字数の分布が移動していて、潰れのある文字が多数あることがわかる。

3. 各特微量の性能

3.1 方向線素特微量

本論文で実験対象とする文字データセットに関して方向線素特微量が文字の潰れの程度に応じてどの程度有効かを調べる。

方向線素特微量¹¹⁾は、線素化された文字パターンか

4倍角 (レーザショットプリンタ)

爵 駕 寡 竈 贗

全角 (レーザショットプリンタ)

爵 駕 寡 竈 贗

全角 (リボンプリンタ)

爵 駕 寡 竈 贗

図 2 入力文字パターンの例

Fig. 2 Samples of character patterns.

表 1 細線化回数別文字数

Table 1 The number of characters against the number of thinnings.

字 形	細線化回数 (回)									
	~4	5	6	7	8	9	10	11	12~	
4 倍角	439	4210	3451	738	54	3				
縦倍角	145	2732	3903	1590	420	84	18	3		
横倍角	193	3768	3690	951	145	57	33	23	35	
全 角	25	1206	3810	2382	919	314	115	59	65	
Total	802	11916	14854	5661	1538	458	166	85	100	

文字データ: 35580 文字

(文字)

ら求める。まず、 64×64 ドットの文字領域を8ドット間隔に分割し、 8×8 ドットの基本領域に分ける。 2×2 の隣接する基本領域を統合することにより、オーバーラップしている 16×16 ドットの小領域が $7 \times 7 = 49$ 個できる。この 49 個の小領域ごとに、4種類の線素の数を重みをつけて計数することにより、 $49(\text{領域}) \times 4(\text{線素}) = 196$ 次元の数値ベクトルを得る。このベクトルを方向線素特徴量と呼ぶ。なお、重みは小領域のまん中が大きく周辺にいくほど小さくなるようにし、線分の位置変動に対して、計数値の変化が少ないようになっている (ぼけ変換に相当する)。

認識実験において、辞書は前述した辞書用データを用いて、各字種ごとに12個の文字パターンの方向線素特徴量の平均ベクトルとして作成した。距離尺度は、ユークリッド距離を使用している。実験は、全字種を一文字ずつもつセット (使用プリンタと字形が同一) ごとに行い、その1位認識率の平均 (平均認識率)

は、99.29% であった。

潰れの程度が細線化回数で評価できるということから、整合性を示す距離値の順位について細線化回数別に集計した。順位は1位、2~10位、および11位以下の3つに分類し、この分類に対する文字の分布を表2に示す。細線化回数の少ない文字については、認識率はかなり高く、全文字数の93.4%を占める7回以下の文字については、99.7%の認識率である。しかし、細線化回数が多い文字ほど認識率は下がり、12回以上の文字は、認識率が57%である。

方向線素特徴量による認識性能が高品質な文字で高く潰れのある低品質な文字で低下することについて、本論文で実験対象とする文字データセットに関して定量的な確認をした。

3.2 ドットパターン特徴量

細線化を行わない特徴量として、文字パターン自体を特徴量とみなすものを選び、性能を調べる。

ドットパターン特徴量は、正規化した文字パターン自身を特徴量と考えた $64 \times 64 = 4096$ 次元の2値ベクトルである。識別のための辞書の標準特徴量は辞書用データの各文字ごとの特徴量の平均ベクトルとして作成し、平均ベクトルの各要素は計算の高速化のために0から $M (=16)$ の整数値に変換した。文字データに潰れがあることやプリンタが異なることから、各文字の全文字パターンで安定して0または1となる画素とそうでない画素とが生ずる。辞書用の各文字の標準特徴量は、安定している画素で0または M となり、他の画素でその中間値をとる。この辞書データに応じ

表 2 細線化回数別各順位文字数

Table 2 The number of characters ordered in recognition against the number of thinnings.

認識順位	細線化回数（回）									
	～4	5	6	7	8	9	10	11	12～	
1 位	799	11881	14809	5636	1518	433	147	64	57	
2～10位	3	31	38	19	14	20	15	12	15	
11位以下	0	4	7	6	6	5	4	9	28	
TOTAL	802	11916	14854	5661	1538	458	166	85	100	

文字データ: 35580 文字

(文字)

て、入力文字のドットパターン特徴量の各次元でとる値は0と M の2値とした。距離尺度としては、重み付きユークリッド距離を使用した。重みは各文字ごとにあらかじめ定義し、次元 i の重み w_i は、辞書ベクトルを作成するために用いたデータの分散 σ_i の逆数に比例するものを用いた。ただし、分散が0のことがあるので逆数計算には正の小さな値 ε を加えたものを用いた。すなわち、 $\rho_i = 1/(\sigma_i + \varepsilon)$ として、

$$w_i = A \cdot \rho_i / \sum_j \rho_j$$

ただし、 A は整数化のための係数である。実験では、 $A=50000$ 、 $\varepsilon=1.61$ とした。 $(\varepsilon$ は種々の値で予備実験をし、認識率が最もよかったものを選んだ。 $M > \mu_i \geq 1$ のとき $\sigma_i \geq M-1$ で、 ε は十分小さい。)分散がすべて等しい場合の重みは $12 (= A/4096)$ となり、分散が等しくない場合はその前後の値をとる。

辞書ベクトルの値(平均値)を μ_i とすると

$$\sigma_i = \mu_i(M - \mu_i)$$

となり、重みは、文字部分(値 M)および背景部分(値0)が安定に観測される画素を強調するようになっている。文字パターンの安定な画素に文字を識別する情報があると考えられ、この重みは直観的に妥当なものであると思われる。その実際の妥当性は実験によって確認できると考える。

ドットパターン特徴量の分類能力を調べるための認識実験では、各入力文字に対する候補字種を、方向線素特徴量を用いた認識実験の上位10字種に絞って、整合性を計算するようにした。また、その上位10字種の中に正解の字種が入っているものだけを入力文字とした。これによる認識率(正解文字数/全対象文字数)は全字種を候補とするときの認識率の近似値を与えたと考えられる。候補を10個に限った実験を計画した理由は、ドットパターン特徴量の次元数が4096次元と多いため、距離計算に時間がかかるためである。実際、全字種を候補とする処理時間をSUN 4/260で計測したところ、方向線素特徴量を用いた認識実験では約2.6秒/文字であるのに対し、ドットパターン特徴量を用いた実験では約28.1秒/文字(従って、23時間/2965文字)であった。

実験結果は、すべての文字セットで、その認識率は99.9%以上の認識率であり、平均認識率は99.95%であった。字形別の認識率を方向線素特徴量の場合(ドットパターン特徴量の実験対象となった文字集合に限定した認識率である)と比較して図3に示す。認識率は、字形に依らず安定しており、方向線素特徴量を用

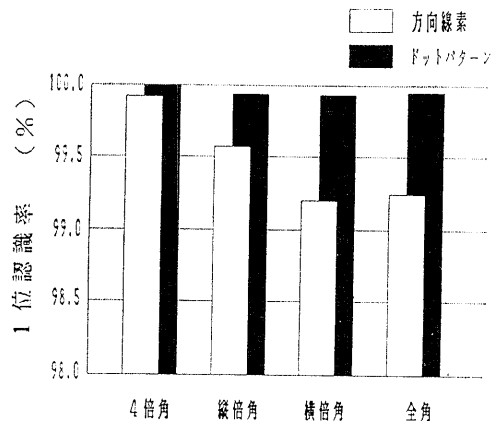


図3 字形別認識率(ドットパターン対方向線素)

Fig. 3 Recognition rate (dot pattern [black] vs. oriented dot [white]).

いた場合よりかなり高い。すなわち、ドットパターン特徴量が潰れの影響を受けにくいことを示している。

3.3 メッシュ特徴量

ドットパターン特徴量は処理時間がかかり過ぎることが欠点であり、その情報を圧縮して低次元にしたとみなせるメッシュ特徴量が考えられた。その性能を同一条件で調べる。

本論文で用いたメッシュ特徴量は、正規化した文字パターンについて、文字領域の縦横を8分割し、それぞれの分割された領域(8×8ドット)内での黒画素の割合を計測したもの(黒画素数/64)である。領域の数は8×8であり、メッシュ特徴量は64次元となる。また、整数で計算を行うため、各領域の黒画素の割合は、0から $M (=16)$ の整数値で表現している。

メッシュ特徴量を用いた認識実験における辞書は辞書用データの特徴量の平均ベクトルを要素として作成し、距離尺度はユークリッド距離を用いた。

実験の結果は、平均認識率で99.49%であり、10位までの平均累積認識率が99.99%と高く、大分類方法として優れている。字形ごとの認識率を、方向線素特徴量の場合(全データセットを対象とした認識率である)と比較して、図4に示す。方向線素特徴量に比べて、メッシュ特徴量を用いた場合の認識率は、各字形で安定しており潰れの影響が少なく、全角、横倍角の潰れの多い字形でもよい性能を示している。

しかし、4倍角、縦倍角などの比較的高品質な文字の多い字形において方向線素特徴量を用いる方が優れている。これは、両者の認識方法の特徴を示していると考えられる。この特徴を定量的に示したものが表3

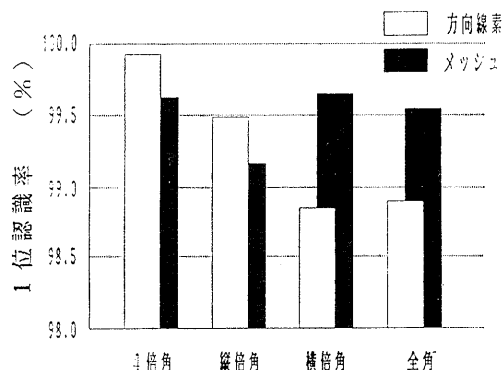


図 4 字形別認識率 (メッシュ対方向線素)

Fig. 4 Recognition rate (mesh [black] vs. oriented dot [white]).

表 3 正・誤認識文字数 (方向線素対メッシュ)

Table 3 The number of correct and error characters in recognition (oriented dot vs. mesh).

		メッシュ特徴量	
		正	誤
方向線素特徴量	正	35163文字 (98.83%)	165文字 (0.46%)
	誤	237文字 (0.67%)	15文字 (0.04%)

全対象文字数: 35580 文字

である。表 3 は、認識対象の各文字について、方向線素特徴量とメッシュ特徴量それぞれによる識別結果について、正誤文字数の分布を示したものである。この表から、どちらの特徴量を用いても認って認識される文字が全体の 0.04% のみであることがわかる。従って、これは両特徴量を併用することで両者の識別能力を相互に補完すべきことを示唆している。

4. 複合認識アルゴリズム

4.1 方 法

低品質文字が存在する文書に対して、方向線素特徴量とメッシュ特徴量を併用すべきことが、両者の識別実験から明らかになった。この実験結果をふまえ、速くかつ正確な認識を実現するため、図 5 の複合認識アルゴリズムを考案した。

本アルゴリズムは、方向線素特徴量とメッシュ特徴量をもつ高い認識性能を相互補完的に発揮させるために、同一未知文字に対して両者で独立に認識させる初段認識とそれにより得られる候補に対してドットパターン特徴量を用いて正しいものを選ぶ検証段階とからなる。検証方法は、時間がかかっても 100% 正しく認識できる方法が望ましいが、現在そのような方法は

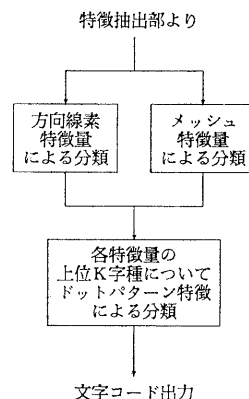


表 5 複合認識アルゴリズム

Fig. 5 Compound recognition algorithm.

存在せず、次善の策として方向線素特徴量やメッシュ特徴量より良い認識率を与えることがわかったドットパターン特徴量を採用している。本方法はそれぞれの初段認識が他方によって訂正されうることを根拠にしておき、単独のものより高い認識性能を与えると考えられる。

4.2 認識実験

初段認識で選出する候補の数 K の値をいくつか変えて認識実験を行った。対象文字パターンは、3章の実験で用いた 12 セットに加えて、新しい 12 組のデータセットを用意して用いた。実験結果を表 4 に示す。ここで、「 n 位以内誤り」の欄は n 位までに正解の候補が入っていなかった場合の入力パターン数である。

既存の大分類-細分類の手法では、大分類で候補数を大きくする方がその中に正解が入りやすくなり高精度になっていたが、本複合認識アルゴリズムでは逆に初段認識での候補数を多くすると誤認識が増えている。これは、ドットパターン特徴量が 100% の認識率を与える方法ではないため、初段認識では 2 位以下になった間違った候補を 1 位候補にしてしまうことが起きているためである。ただし、2・3 位累積認識率をみると、劇的に認識率が向上しており、 K を大きくすると $2K$ 個の候補の中に正解が含まれる場合がほとんどで、含まれない場合は 7 万余りの文字パターンのうち 5 個以下であることがわかる。従って、正解を選ぶ確実な方法があればそれをドットパターン特徴量による認識の代わりに用いることで、初段認識で選ばれた $2K$ 個 ($K \geq 2$) 以下の候補の中にある正解を確実に選ぶことができ、正解がない場合は間違ってしまうがその状況になるのは 5 個以下の文字パターンだけであ

表 4 複合認識アルゴリズム実験結果

(初段認識候補数 K と累積認識率・誤認識文字数)Table 4 Experimental results by compound recognition algorithm (accumulated recognition rate and the number of misrecognized characters for each candidate number K).

	1 位認識率	2 位累積認識率	3 位累積認識率
	1 位候補誤り	2 位以内誤り	3 位以内誤り
$K=1$	99.94%	99.97%	—
	42文字	24文字	—
$K=2$	99.92%	99.99%	99.99%
	59文字	6文字	5文字
$K=5$	99.91%	99.99%	99.99%
	62文字	7文字	4文字
$K=10$	99.91%	99.99%	99.99%
	63文字	8文字	4文字

全対象文字数: 71160 文字

り、誤認識率を1万分の1以下にできることがわかる。しかし、ドットパターン特徴量による整合法を使って1位候補を認識結果とする場合は、 $K \geq 2$ とすることは認識率を下げ、 $K=1$ が最善であることがわかった。すなわち、方向線素特徴量とメッシュ特徴量とで異なる候補が選出された場合、どちらの候補が正しいかを検証するためだけにドットパターン特徴量を用いることが妥当であると考えられる(初段認識の候補が一致すれば、そのまま出力し検証は行わない)。

$K=1$ の場合について認識結果の詳細を述べる。2章で述べた入力データ(その大部分についてドットパターン特徴量だけで認識率 99.95% (3.2 節)を与えている)について、35580 の文字サンプル中誤認識文字数は 18 文字であり、平均認識率で 99.95% を得ている。誤認識の原因は初段認識で両者が誤る場合と検証段階で誤る場合とに分類できる。この 18 文字は 15:3 に分類され、初段認識で検証が必要となる2文字の候補が生ずるのは表3から 402 文字であり、検証段階で 399 文字を正しく判定したことを示している。

処理時間は、約 3.2 秒/文字(初段認識で一致)、ないしは、約 3.3 秒/文字(検証段階も実施)であった。

新しい 12 組のデータセットについては、誤認識が 24 文字 (9:15) で、平均認識率 99.93% であった。この文字セットは最初の実験用データセットに比べて潰れ文字が多かったが、同程度の誤認識ですんでいるといえる。

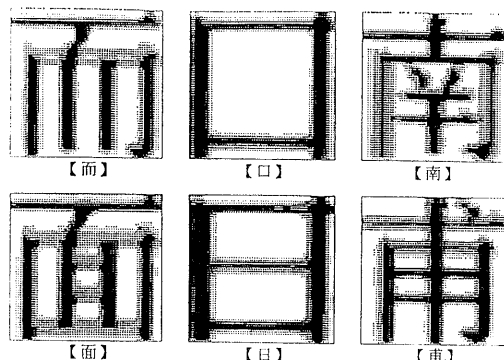


図 6 誤認識文字の標準特徴量

Fig. 6 Standard features of misrecognized characters.

検証段階における誤りは、「南」と「甫」、「口」と「日」、「而」と「面」などの文字パターンに多かった。その標準特徴量を図6に示す。標準特徴量がよく似ていることが理解されよう。すなわち、文字パターンを構成している線分の位置ずれが少ないことを期待したが、実は、少数の文字に微妙な位置ずれがあり、安定した線分だけをみると区別がつかない特徴量になってしまっているといえる。より高精度を求めるには線分の位置ずれにつよい認識手法が必要であることを示唆している。

5. おわりに

潰れのある低品質な印刷文字に対する識別性能を、方向線素特徴量、メッシュ特徴量、ドットパターン特徴量の3つについて調べ、それぞれの特性を明らかにした。その結果から、それらを統合した複合認識アルゴリズムを提案し、その性能を確かめた。

メッシュ特徴量が潰れ文字に対しても有効であることが示され、しかも、方向線素特徴量とは誤認識の仕方が異なることがわかった。また、ドットパターン特徴量も潰れ文字に極めて有効であることが示された。

複合認識アルゴリズム($K=1$)は、潰れ文字を含むワープロ用プリンタ印字文字に対して認識率 99.9% 台を与え、方向線素特徴量とメッシュ特徴量のそれぞれのよい点を生かした高精度な認識アルゴリズムであるといえる。すなわち、細線化したパターンに基づく方向線素特徴量が高品質な文字を正確に識別し、細線化をしないで求まるメッシュ特徴量が比較的潰れに強いという両者の長所を生かし、互いに補完している。

階層的認識アルゴリズムでは、大分類で誤ると細分類では救いようがなく、今回提案の複合認識アルゴリ

ズム ($K=1$) でも大分類に相当する初段認識で間違えるものが間違いの半分以上であり、真の自動認識のためには $K \geq 2$ とすることが必要である。すなわち、初段認識による上位 K 個の候補に正解が必ず入るようにして、検証段階でより高精度な認識アルゴリズムを用いることがよいと思われる。高精度なアルゴリズムとして、手書き文字認識の分野で知られる文脈情報を用いるものや構造解析法などを併用することが考えられる。すなわち、複数の認識アルゴリズムを考え、それぞれの特性を生かすような複合的な利用によりより高精度化できるとと思われる。本論文で示した初段認識がそのときでも有効であることは、それぞれの性能確認実験から明らかであり、低品質文字を含む印刷文字の自動認識のための基礎的データを与えたともいえる。

かなやアルファベットは通常潰れることがなく今回は実験しなかったが、それらに対する性能を確認する必要がある。また、プリンタや字形をかえたマルチフォントで実験したが、雑誌や新聞、本などのマルチフォント活字に対する性能を確認することも重要である。これらのことは今後に残された課題である。

謝辞 本研究を進めるにあたり御討論頂いた北陸先端科学技術大学院大学堀口進教授、下平博助教授、並びに東北大学大型計算機センター孫寧助手に深く感謝する。本研究の一部は文部省科学研究費補助金特別推進研究 (No. 63060001) の援助による。ここに記して感謝する。

参 考 文 献

- 1) Hilditch, C. J.: Linear Skeleton from Square Cupboards, *Machine Intelligence 6*, Meltzer, B. and Michie, D. eds., pp. 403-420, Univ. Press Edinburgh (1969).
- 2) 山本真司, 中田和男: 階層的パターンマッチングによる漢字認識の基礎—印刷漢字認識の研究—, 信学論 D, Vol. 56-D, No. 6, pp. 365-372 (1973).
- 3) 梅田三千雄: マルチフォント印刷漢字の分類, 信学論 D, Vol. J62-D, No. 2, pp. 133-140 (1979).
- 4) 安田道夫, 藤沢浩道: 文字認識のための相関法の一改良, 信学論 D, Vol. J 62-D, No. 3, pp. 217-224 (1979).
- 5) 萩田紀博, 梅田三千雄, 増田 功: 三つの概形特徴を用いた手書き漢字の分類, 信学論 D, Vol. J 63-D, No. 12, pp. 1096-1102 (1980).
- 6) 齊藤泰一, 山田博三, 山本和彦: 手書き漢字の方向パターンマッチング法による解析, 信学論 D, Vol. J 65-D, No. 5, pp. 550-557 (1982).
- 7) 目黒真一, 梅田三千雄: マルチフォント印刷漢字の認識, 信学論 D, Vol. J 65-D, No. 8, pp. 1026-1033 (1982).
- 8) 北島宗雄: 教科書体を含むマルチフォント印刷文字認識, 信学論 D, Vol. J 66-D, No. 4, pp. 400-

406 (1983).

- 9) 馬場口登, 国弘秀人, 相原恒博: 漢字認識における段階的整合方式の効果, 信学論 D, Vol. J 68-D, No. 11, pp. 1918-1925 (1985).
- 10) 鶴岡信治, 栗田昌徳, 原田智夫, 木村文隆, 三宅康二: 加重方向指数ヒストグラム法による手書き漢字・ひらがな認識, 信学論 D, Vol. J 70-D, No. 7, pp. 1390-1397 (1987).
- 11) 孫 寧, 田原 透, 阿曾弘具, 木村正行: 方向線素特徴量を用いた高精度文字認識, 信学論 D-II, Vol. J 74-D-II, No. 3, pp. 330-339 (1991).
- 12) 孫 寧, 阿曾弘具, 木村正行: 連想整合法に基づく高速文字認識アルゴリズム, 情報処理学会論文誌, Vol. 32, No. 3, pp. 404-413 (1991).

(平成 4 年 12 月 28 日受付)

(平成 6 年 5 月 12 日採録)



進藤 宣博

1965 年生。1988 年東北大学工学部情報工学科卒業。1990 年同大学院大学院博士課程前期 2 年の課程 (情報工学専攻) 修了。同年, (株) リコーに入社。在学中, 文字認識の研究に従事, 現在, プリンターのソフトウェア設計に従事。



阿曾 弘具 (正会員)

1946 年生。1973 年東北大学大学院博士課程修了。工学博士。同年同大工学部助手, 1979 年名古屋大学工学部講師, 助教授を経て, 1986 年東北大学工学部助教授, 1991 年同教授。学習オートマトン, セル構造オートマトン, 並行処理理論, シストリックアルゴリズム設計論, 文字認識システムの開発などの研究に従事。IEEE, ACM, EATCS, 電子情報通信学会, 人工知能学会, LA 各会員。



木村 正行 (正会員)

1927 年生。1959 年東北大学大学院博士課程修了。工学博士。1962 年東北大学助教授, 1970 年同工学部教授, 1991 年同大学名誉教授。同年 4 月北陸先端科学技術大学院大学情報科学研究科教授, 現在に至る。システム工学および情報システム工学分野の研究, パターン認識, 学習自己組織化等の研究を行い, 最近では, 遺伝的アルゴリズム, 知識情報処理の分野にも興味を持つ。著書「自己組織系構成論」, 「システム工学基礎論」, 「しきい値論理とその応用」等。電子情報通信学会会員。