

CycleGAN を用いた近代書籍風文字の生成と そのデータ拡張への応用

角張凜¹ 飯田紗也香¹ 城和貴¹

概要: 明治から昭和初期にかけて刊行された近代書籍には活版印刷手法が使用されており、文字にはにじみやかすれが存在する。そのため、文字認識には現代印字文字を対象とした一般的な OCR ではなく、近代書籍に特化したモデルを使用する。モデルの認識精度を向上させるには、学習データとして近代書籍文字画像が大量に必要となる。しかし、手作業で書籍画像から文字を切り出し収集することは大きな手間がかかる。また、書籍内で文字の出現頻度に差があるため、低頻出文字はデータが集まりにくいといった難点もある。そこで、本稿では、CycleGAN を用いて「近代書籍風文字」を生成することによるデータ拡張手法を提案する。生成した近代書籍風文字画像を学習データに含めて近代書籍の学習と文字認識を行い、その結果を示す。

1. はじめに

近代書籍とは、明治から昭和初期にかけて刊行された書籍を指す。国立国会図書館デジタルコレクション[1]では、帝国議会会議録から雑誌、文学作品等、多岐にわたる近代書籍が保存・公開されている。その書籍データはすべて画像形式である。テキストデータ化ができれば、利便性が高まり、活用の幅が広がる。現在、国立国会図書館デジタルコレクションで公開されている近代書籍は約 35 万点である。それらの書籍画像を 1 点ずつ手作業でテキストデータ化することは容易でなく、光学文字認識 (Optical Character Recognition, OCR) 技術の活用が求められている。一方で、近代書籍には活版印刷手法が使用されており、時代・出版者により印刷様式が異なるなど、現代の書籍画像とは異なる点が多く存在する。そのため、一般的に使用されている OCR 技術では文字認識を正確に行うことはできない。したがって、近代書籍に特化した文字認識手法・モデルの開発が試みられている[2][3]。

高精度な近代書籍文字認識を実現するために、現在の課題とされているのが、学習データ収集である。一般的に文字認識モデルの学習では、認識対象と同じ種類の文字画像を大量に学習データとして用いる。しかし、現在の近代書籍画像は資料 1 ページごとの画像であるため、そこから手作業で文字画像を何千、何万文字と切り出すためには大変な手間と時間がかかる。また、文字の出現頻度の差も存在する。資料の中で高頻出である文字は集まりやすいが、低出現頻度文字種では獲得の難易度が極めて高くなることが報告されている[4]。そのため、近代書籍の文字画像を大量に、かつまんべんなく収集することは容易ではない。このような状況が、近代書籍文字認識の精度向上を困難にする原因となっている。

本稿では、近代書籍文字の特徴を持つ「近代書籍風文字」を生成することによって、これらの問題の解決を試みる。

提案手法では、CycleGAN を使用して近代書籍文字の特徴を学習し、現代印字文字と組み合わせた文字画像を生成する。生成される文字を近代書籍風文字と呼ぶ。これを学習データとして近代書籍文字を補い、文字認識実験を行う。

以下、本稿における構成について示す。2 章で、関連研究として近代書籍の文字認識についてと、CycleGAN に関する研究について述べる。3 章で、提案する CycleGAN を用いた近代書籍風文字の生成手法について説明する。4 章で、提案手法を用いた文字認識モデルの学習と認識実験について述べる。

2. 関連研究

2.1 近代書籍文字の文字認識

1 章で述べたとおり、近代書籍の文字認識には、近代書籍に特化した文字認識手法およびモデルの開発が必要であり、これまでに様々な研究が行われている。近代書籍に特化した文字認識手法として初めに提案されているのが、多フォント活字認識法[5]である。近代書籍文字が持つ、出版者ごとに規格が異なるという特性は、現代印字文字よりも手書き文字に近い。このことを踏まえ、オフライン手書き文字認識に用いる手法を利用することによって、近代書籍の多フォント活字認識を実現する。実験で示されている認識対象は 10 文字のみだが、99.0%の認識率を記録している。その後も、実用的な近代書籍文字認識の実現を目標として、いくつかの手法が提案されている[2][6]。CNN を利用して、近代書籍文字に加えて現代印字文字を学習データに用いる近代書籍文字認識[3]では、約 97%の精度が報告されている。この数値は、実用的な近代書籍文字認識モデルを提案する先行研究の中では最高の認識率となっている。本稿では、この CNN モデルを用いて文字認識実験を行う。近代書籍文字と現代印字文字を学習データとする先行研究が近代書籍文字画像を 5 種類使用して学習していることに対し、本稿ではそれよりも少ない近代書籍文字画像データセットによる、高精度な文字認識の実現を試みる。

¹ 奈良女子大学
Nara Women's University

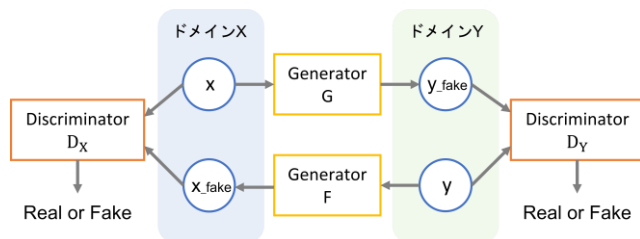


図1 CycleGANの構成

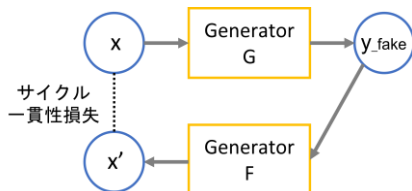


図2 サイクルの流れ

本稿と同様に、学習データ拡張を目的として、近代書籍の特徴を持つ文字画像を生成する手法も提案されている[7]。本稿と異なる点は、教師あり学習として現代印字文字と近代書籍文字を1対1で学習させる点である。したがって、生成される画像の一致率が評価とされている。本提案手法では、文字画像同士をペアにするのではなく、文字群全体の特徴を学習する。そのため、文字画像生成において、特定の近代書籍文字画像を完璧に再現することは求めない。

2.2 CycleGAN

CycleGANは、GAN (Generative Adversarial Network, 敵対的生成ネットワーク) を用いる画像変換手法である[8]。一般的に、GANはGeneratorとDiscriminatorと呼ばれる2つのニューラルネットワークを持つ。Generatorは、真のデータセットからその生成分布を学習し、真のデータをまねた偽のデータを生成する。Discriminatorは、入力されるデータが真のデータか、Generatorによって生成される偽のデータか識別する。2つのニューラルネットワークが互いに競い合いながら最適化を行うことで、本物のような高精度の画像を生成することが可能となる。CycleGANでは、2つのGANを組み合わせ、2つのデータセットについてドメイン間の関係を捉える。ドメインとは、一般的に分野や領域という意味の言葉である。本稿では、画像が持つ固有の特徴によってまとめられる集合のことをドメインとして指す。CycleGANは、別のドメインである2つの画像データセットの特徴をそれぞれ学習することで、一方のドメインの画像を、他方のドメインの特徴を持つ画像へ変換することが可能となる。なお、GANは教師なし学習であり、1枚の画像ごとにペアとなる正解データが存在する必要はない。図1に一般的なCycleGANの構成を示す。データセットが属するドメインとして、ドメインXとドメインYが存在する。ドメインXとドメインYは、それぞれ相互に変換される画像データセットである。ドメインXのデー

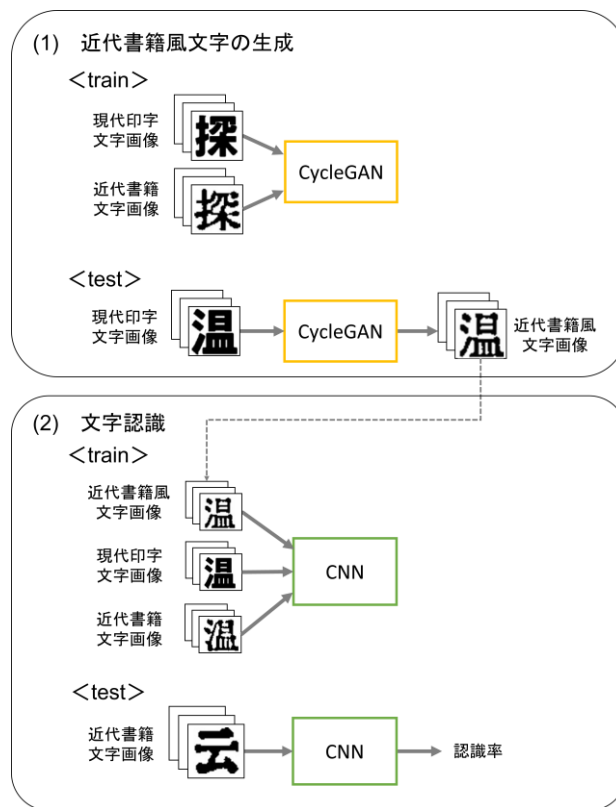


図3 提案手法の流れ

タxは、Generator GによってドメインYのデータを模したy_fakeへと変換される。一方で、ドメインYのデータyはGenerator FによってドメインXのデータを模したx_fakeへと変換される。ドメイン内の真のデータと偽のデータは、それぞれのドメインのDiscriminatorによって識別される。その結果をもとに損失関数を計算する。この損失関数の値は、敵対的損失と呼ばれる。加えて、CycleGANにはサイクル一貫性損失が導入されている。図2に、CycleGAN内に設けられるサイクルの流れを示す。Generator GによってドメインYへと変換されているデータxは、Generator FによってドメインXへと逆変換される。xと逆変換の結果x'の差から、サイクル一貫性損失を計算する。ここではデータxの場合のみを示しているが、データyの変換も同様である。サイクル一貫性損失の導入によって、Generatorは互いに独立することなく、教師なし学習でも相互変換を行うことが可能なモデルを学習する。CycleGANはこれらの構成によって、高精度な画像変換を実現している。この特徴が注目され、CycleGANを用いたデータ拡張は様々な分野で行われている。文字認識にまつわる研究では、手書き漢字画像の生成[9]や、スタイル変換[10]などが提案されている。また、学習にペアデータを必要としない点は、データ収集が困難な分野の深層学習に非常に適しており、多くの研究でその有効性が報告されている[11][12]。本稿では、近代書籍文字画像のデータ拡張を目的として、CycleGANを用いる。

	学習データ		テスト データ	生成画像 (近代書籍風文字)	対応する 近代書籍文字
	現代印字文字	近代書籍文字			
例1	探	探	温	温	温
例2	探		温	温	
失敗例	探		温	温	

図4 近代書籍風文字画像生成に関わる文字画像

3. 提案手法

本章では、CycleGAN を用いて近代書籍文字の特徴を持つ「近代書籍風文字」を生成し、文字認識の学習データとして活用する手法について述べる。提案手法の流れを図3に示す。提案手法を2つの手順に分けて説明する。

3.1 近代書籍風文字画像の生成

1つ目の手順である、近代書籍風文字画像の生成について述べる。CycleGAN による近代書籍風文字画像生成の流れを図3の(1)に示す。図3(1)に示す train フェーズでは、現代印字文字画像データセットと近代書籍文字画像データセットを用意する。これらのデータセットを CycleGAN に与え、敵対的損失とサイクルー貫性損失の最適化によって学習を行う。CycleGAN は、現代印字文字というドメインと、近代書籍文字というドメインを学習することで、それぞれの画像の特徴を捉えることが可能となる。同時に、ドメイン間の関係を把握し、一方のドメインが持つ特徴をどのように別ドメインの画像に変換できるか見いだす。図3(1)に示す test フェーズでは、学習済みの CycleGAN に現代印字文字画像を入力する。その出力として、現代印字文字の特徴が近代書籍文字の特徴に変換されている画像が生成される。これを近代書籍風文字画像と呼ぶ。近代書籍風文字画像生成に関わる文字画像の例を図4に示す。学習データは train フェーズで、テストデータは test フェーズで用いる。近代書籍風文字画像の右列に、テストデータの文字に対応する近代書籍文字画像を示している。これは、近代書籍風文字画像との比較のためである。近代書籍風文字は入力される現代印字文字の形を保ちつつ、近代書籍文字特有のにじみやかすれなどの特徴を持つことが確認できる。テストデータの文字に対応している近代書籍文字の画像と比較すると、近代書籍風文字は特定の近代書籍文字を完全に再現しているわけではないことも分かる。

CycleGAN の学習に用いるデータセットは、それぞれ 64x64 ピクセルのグレースケール文字画像 1200 文字ずつで構成される。CycleGAN は教師なし学習を行うため、データセット間で文字同士が 1 対 1 で結びつくことはない。ただし、本稿で行う実験ではデータセット作成を簡略化するため、両文字画像データセットの文字は共通としている。

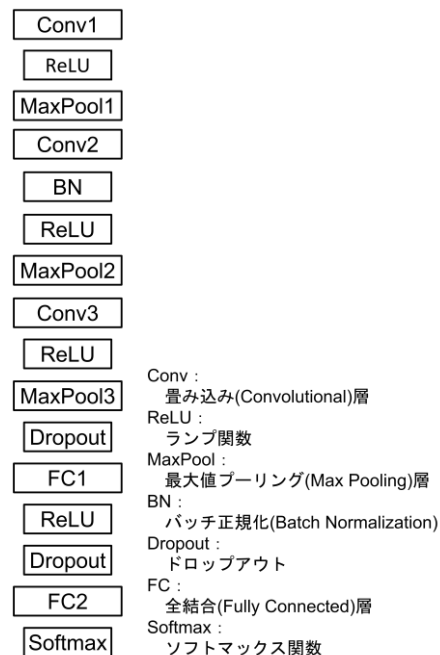


図5 近代書籍文字認識を行う CNN

なお、両方のデータセットの文字が全く同じ場合とすべて異なる場合のそれぞれでの学習結果として、生成される画像全体が持つ特徴に大きな変化がないことが確認できている。

3.2 文字認識

2つ目の手順である、文字認識について述べる。図3の(2)にその流れを示す。(1)で生成する近代書籍風文字を学習データに含めてモデルを学習し、近代書籍の文字認識を行う。文字認識を行うモデルは、2.1 節で述べたように先行研究[3]に用いられている CNN と同一のものを使用する。図5にその構成を示す。畳み込み層は3層で、それぞれのフィルタ数は (160, 320, 640) である。ただし、フィルタサイズは全層で 7x7 を用いるのではなく、層ごとに (7x7, 5x5, 3x3) とする。これは、アルゴリズムの挙動を制限するハイパーパラメータの調節を目的とする予備実験で、全層 7x7 とする場合よりも層ごとにフィルタサイズを変更する場合に高い認識精度が得られることが確認されているためである。

CNN の学習では、近代書籍風文字画像、現代印字文字画像、近代書籍文字画像のデータセットを組み合わせる学習データセットを作成する。このデータセットを入力して、それぞれの文字画像から文字を認識できるようにモデルを学習させる。学習完了後、CNN にテストデータを入力して近代書籍文字画像の認識を行う。

4. 近代書籍風文字を用いた文字認識実験

本章では、提案手法の有効性を示すため、近代書籍風文

表1 現代印字文字一覧

A1 ゴシック B	リュウミン U-KL
Arial Unicode MS	見出しゴ MB1
AR 丸ゴシック体	見出しゴ MB31
BIZ UD ゴシック	黒体
HGS ゴシック E	秀英にじみ角ゴシック金 B
HGS 創英プレゼンス EB	秀英にじみ角ゴシック銀 B
HGS 創英角ポップ	秀英にじみ明朝 L
HGS 明朝 E	秀英横太明朝 M
HG 丸ゴシック M-PRO	秀英角ゴシック銀 B
MigMix	秀英丸ゴシック B
Migu	秀英明朝 M
MS ゴシック	新ゴ U
UD デジタル教科書体 NP-B	新丸ゴ H
ゴシック MB101U	太ゴ B101
こぶりなゴシック W6	中ゴシック BBB
じゅん 34	凸版文久ゴシック DB
ソフトゴシック U	凸版文久見出しゴシック EB
ヒラギノ角ゴ W9	凸版文久見出し明朝 EB
ヒラギノ角ゴオールド W6	梅ゴシック
ヒラギノ角ゴオールド W9	黎ミン U
ヒラギノ角ゴシック Std	黎ミン Y10U
ヒラギノ丸ゴ W8	黎ミン Y20U
ヒラギノ丸ゴシック ProN	黎ミン Y30U
ヒラギノ明朝 W8	黎ミン Y40B
メイリオ	黎ミン Y40U

字を用いた文字認識実験を行う。初めに、CNN の文字認識学習に用いる学習データセットについて述べる。その後、学習データに近代書籍風文字を含める場合と、近代書籍風文字を含めずに現代印字文字や近代書籍文字のみで学習する場合の実験結果を比較する。指標とするのは、学習済みモデルに近代書籍文字を入力して出力される認識率である。実験によって、提案手法が CNN を用いた近代書籍文字認識の学習データ拡張手法として有効であることを確認する。

4.1 データセット

学習データには近代書籍文字、現代印字文字、近代書籍風文字の 3 種類の文字画像を使用する。以下、それぞれの文字画像データセットについて述べる。

現代印字文字画像データセットは、50 種類の現代印字文字で構成される。使用する現代印字文字の一覧を表 1 に示す。データセット構築時には、設定されるデータセット数に合わせて 50 種類の中から最大 40 種類がランダムに選択される。

近代書籍文字画像データセットは、6 種類の近代書籍文字で構成される。ただし、そのうち学習に用いるのは 4 種

表2 文字認識結果

学習用文字画像データセット (種類)			認識率 (%)
近代書籍	現代印字	近代書籍風	
4	0	0	89.6
	40	0	97.7
	20	20	97.7
	0	40	97.6
1	40	0	93.4
	20	20	95.0
	0	40	95.5
0	40	0	83.3
	20	20	89.8
	0	40	91.1

類のみである。残りの 2 種類は、それぞれテストデータとバリデーションデータに使用する。バリデーションデータは、重みづけなどモデルの学習には関わらないが、ハイパーパラメータの調整に用いられる。

近代書籍風文字画像データセットは、50 種類の近代書籍風文字で構成される。これは、1 種類の近代書籍文字画像データを、50 種類の現代印字文字画像データにそれぞれ組み合わせ生成されるものである。生成される画像の中には、うまく学習を行うことができず、図 4 に示す失敗例のように文字の形が崩れているものも一部存在する。これらを除く 40 種類を、文字認識の学習データとして使用する。なお、一部の漢字には、同一 JIS 文字コードを持つ一方で形が異なる「異字体」が存在するが、今回は極力取り除いてデータセットを作成する。例として、「高」と「高 (はしごだか)」などが挙げられる。

すべての文字画像データセットは、1 種類につきそれぞれ JIS 第 1・第 2 水準レベルの漢字 2000 文字で構成されている。文字画像は 64 x 64 ピクセルで、グレースケールである。

本稿では、次のような場合に分けて実験を行う。

- 近代書籍文字画像が十分存在する場合
(近代書籍文字画像を 4 種類使用)
- 近代書籍文字画像が少ない場合
(近代書籍文字画像を 1 種類使用)
- 近代書籍文字画像が存在しない場合
(近代書籍文字画像を 0 種類使用)

上記それぞれの場合に、現代印字文字と近代書籍風文字を合計 40 種類加えて学習データセットとする。40 種類のうち、現代印字文字と近代書籍風文字の割合を変更し、近代書籍風文字の有効性について確認する。

4.2 実験結果

実験結果を表2に示す。表2の1行目から4行目が、近代書籍文字画像データセットが4種類存在する場合の結果である。近代書籍のみで学習する場合には認識率が89.6%であるのに対し、現代印字文字画像を40種類追加する場合、現代印字文字画像と近代書籍風文字画像を20種類ずつ追加する場合、近代書籍風文字画像を40種類追加する場合のいずれにおいても97.7%ほどの認識率となっている。このことから、近代書籍風文字画像が近代書籍文字認識の学習データセットとして、現代印字文字画像と少なくとも同等の有効性を持つことを確認できる。一方で、学習データセットにおける現代印字文字画像と近代書籍風文字画像の割合を変更しても、認識率は0.1%程度しか変化していない。学習データに近代書籍文字が4種類存在している場合には、その他のデータセット構成は大きな影響を与えないということが分かる。すなわち、近代書籍文字画像が十分な量存在し、かつドメインを問わずに文字ごとのデータ数を十分に確保することができれば、近代書籍文字が持つ特徴の学習は、近代書籍文字自身のデータのみで賄うことができると言える。表2の5行目から7行目は、近代書籍文字画像データセットが1種類のみ存在する場合の認識結果である。現代印字文字画像データセットのみ40種類を追加する場合の認識率は93.4%となっている。対して、現代印字文字画像を近代書籍風文字画像に半分置き換える場合、認識率は95.0%となる。全部置き換える場合では、認識率は95.5%まで上昇している。表2の8行目から10行目は、近代書籍文字画像が存在しない場合の文字認識結果である。現代印字文字画像のみで学習する場合、83.3%の認識率である。半分の20種類を近代書籍風文字画像データセットで置き換える場合、認識率は89.8%となり、全部の40種を置き換える場合には91.1%となる。これらの結果から、近代書籍文字画像データセットが少ない、もしくは存在しない場合には、近代書籍風文字画像データセットを利用することで認識率が高くなることが確認できる。また、近代書籍文字画像データセットが少なくなるにつれて、現代印字文字画像データセットのみで学習する場合と、近代書籍風文字画像データセットを加えて学習する場合の認識率の差が大きくなるということが言える。近代書籍文字画像データセットが1種類のとき、近代書籍文字画像と現代印字文字画像のみで構成されるデータセットで学習する結果と、近代書籍風文字画像を用いるデータセットで学習する結果の差は最大で約2.1%である。近代書籍文字画像データセットが存在しない場合には、その差は最大で約7.8%となる。これらの結果から、データ量が不十分な近代書籍文字の学習において、近代書籍風文字画像の有無が結果に大きく作用していることが分かる。近代書籍文字の認識を行うためにはその特徴の学習が不可欠である。近代書籍文字の特徴を持って生成される近代書籍風文字が、不足している近代書

誤認識される 近代書籍文字	対応する 現代印字文字
惑	惑
感	感
畜	畜

図6 誤認識される文字の例

籍文字を補っていると考えられる。このことから、低出現頻度文字種など収集が困難な文字の文字認識に対して、現代印字文字のみを学習するよりも近代書籍風文字を活用することが効果的であると言える。

また、近代書籍文字画像データセットが何種類存在する場合でも、学習データに現代印字文字画像データセットを全く使用しない結果と、学習データに現代印字文字画像と近代書籍風文字画像が両方含まれている結果には、大きな差がないことが確認できる。したがって、必ずしも近代書籍風文字画像データセットのみを使用すればいいのではなく、他の文字画像データセットとの組み合わせで認識率が向上する可能性があると考えられる。

課題として挙げられるのは、近代書籍文字画像データが十分に確保できない場合における認識率のさらなる向上である。近代書籍文字画像データが存在しない学習データセットを用いる場合における文字認識精度は、近代書籍風文字画像の活用によって90%を超えるなど、大きく改善している。しかし、提案手法の実用化のためにはより高精度な文字認識が求められる。近代書籍文字画像が1種類含まれる、もしくは全く含まれない学習データセットを用いる場合に、複数の学習済みモデルで誤認識されている近代書籍文字画像の例を図6に示す。これらの誤認識されている文字の多くに共通する特徴は、異字体ではないが、近代書籍文字と現代印字文字で異なる形状をしているという点である。この差異が、文字認識率低下の一因であると考えられる。この問題における解決策の1つとして、近代書籍文字の形状を正確に真似た現代印字文字画像を作成し、それを基に近代書籍風文字画像を生成する手法が挙げられる。しかし、すべての文字に対してその手法を適用することは手間がかかる。そのため、文字形状の微妙な違いに対して寛容である文字認識モデルを開発することが、課題解決において有効な手段であると考えられる。さらに、出版者ごとに文字の特徴を再現するなど、CycleGANの特性を活かして新たなアプローチに取り組むことで、近代書籍文字画像データの有無にかかわらず高精度な文字認識の実現を目指す。

5. まとめ

本稿では、CycleGAN を用いて近代書籍風文字を生成し、近代書籍の文字認識学習に用いる手法を提案する。近代書籍には活版印刷手法が用いられており、一般的な OCR 技術では文字認識精度が不十分である。そのため、近代書籍に特化した文字認識手法が開発され、高精度な近代書籍文字認識の実現が試みられている。同時に、学習に用いる文字画像データの不足が大きな課題となっている。本提案手法では、近代書籍文字認識の学習データを補うことを目的として、CycleGAN を用いて近代書籍風文字画像を作成する。CycleGAN の学習で必要となる近代書籍文字画像は、その文字種を問わない。近代書籍風文字画像を生成する際に必要となるのは、現代印字文字画像のみである。そのため、さまざまな現代印字文字を用いて近代書籍風文字を生成することで、近代書籍が存在しない文字についても近代書籍文字認識の学習データを補うことができる。近代書籍文字画像と現代印字文字画像を CycleGAN に与えることで、CycleGAN は2つのデータセットの特徴と関係を学習する。学習完了後の CycleGAN を用いて、現代印字文字と近代書籍文字の特徴を組み合わせた近代書籍風文字画像を生成することができる。作成される近代書籍風文字画像の有効性を確認するために、CNN を用いた近代書籍文字認識実験を行う。近代書籍文字、現代印字文字、近代書籍風文字のそれぞれの画像データセットから数種類を学習データに用いて、CNN を学習させる。学習済みの CNN を用いて近代書籍文字認識を行い、認識率を比較する。その結果、近代書籍風文字画像を学習データに用いるすべての場合で、用いない場合と同等、もしくは上回る認識率が確認できる。特に、学習データセットに近代書籍文字画像が存在しない場合に、認識率の差が大きくなることが実験結果から示されている。現代印字文字画像のみで構成されるデータセットの結果と、近代書籍風文字画像のみで構成されるデータセットの認識結果を比較すると、後者で認識率が約 7.8% 上昇している。これらの結果から、提案手法が近代書籍文字認識の学習データ拡張として有効であることが示される。CycleGAN によって近代書籍文字の特徴を再現する近代書籍風文字の生成が可能であるということに加え、近代書籍風文字画像によって近代書籍文字の認識率が向上できることが言える。特に、データ収集が困難な低出現頻度の文字に有効である。

今後は、近代書籍文字画像が十分に確保できない場合や存在しない場合にも、十分に高精度な文字認識が行えるようにモデルの改善や新たな手法の開発を行う。

謝辞

本研究の一部は文部科学省科学研究費補助金(20H04483)による。

参考文献

- [1] “国立国会図書館デジタルコレクション”. <http://dl.ndl.go.jp>, (参照 2021-11-15).
- [2] Fujimoto, K., Ishikawa, Y., Takata, M. and Joe, K., “Early-Modern Printed Character Recognition using Ensemble Learning”, PDPTA2017, pp.288-294 (2017).
- [3] Yasunami, S., Koiso, N., Takemoto Y., Ishikawa, Y., Takata, M. and Joe, K., “Applying CNNs to Early-Modern Japanese Printed Character Recognition”, PDPTA2019, pp.189-195 (2019).
- [4] 藤田未希, 竹本有紀, 石川由羽, 高田雅美, 城和貴. 近代書籍における低出現頻度文字種の獲得. 情報処理学会研究報告. 2019, Vol.2019-MPS-126, No.6, p.1-6
- [5] Ishikawa, C., Ashida, N., Enomoto, Y., Takata, M., Kimesawa, T. and Joe, K., “Recognition of Multi-Fonts Character in Early-Modern Printed Books”, PDPTA2009, Vol.2, pp.728-734 (2009).
- [6] Kosaka, K., Fujimoto, K., Ishikawa, Y., Takata, M. and Joe, K., “Comparison of Feature Extraction Methods for Early-Modern Japanese Printed Character Recognition”, PDPTA2016 Final Edition, pp.408-414 (2016).
- [7] Takemoto, Y., Ishikawa, Y., Takata, M. and Joe, K., “Structure of Neural Network Automatically Generating Fonts for Early-Modern Japanese Printed Books”, PDPTA2019, accepted (2019).
- [8] Zhu, J. Y., Park, T., Isola, P., & Efros, A. A., “Unpaired image-to-image translation using cycle-consistent adversarial networks”, Proceedings of the IEEE international conference on computer vision, pp. 2223-2232(2017).
- [9] Chang, Bo, et al., “Generating handwritten chinese characters using cyclegan”, 2018 IEEE Winter Conference on Applications of Computer Vision (WACV). IEEE, 2018.
- [10] 成沢淳史, 柳井啓司. スタイル転移によるフォント画像変換. 研究報告コンピュータビジョンとイメージメディア (CVIM), 2018(1), p.1-6.
- [11] Sandfort, V., Yan, K., Pickhardt, P. J., & Summers, R. M. “Data augmentation using generative adversarial networks (CycleGAN) to improve generalizability in CT segmentation tasks”, Scientific reports, 9(1), pp.1-9 (2019).
- [12] Bao, Fang, Michael Neumann, and Ngoc Thang Vu. “CycleGAN-Based Emotion Style Transfer as Data Augmentation for Speech Emotion Recognition.” INTERSPEECH. (2019).