

(405) パターン認識

電子部品検査における画像データ拡張のための自己符号化器の性能把握

Performance comprehension of auto-encoder for image data augmentation in electronic component inspection

田中 智裕[†] 大井 健太郎[†] 青戸 勇太^{††} 松林 幹大^{††} 前田 俊二[†]
Tomohiro Tanaka[†] Kentarou Ooi[†] Yuta Aoto^{††} Kanta Matsubayashi^{††} Shunji Maeda[†]
[†]Hiroshima Institute of Technology ^{††}Hiroshima Institute of Technology, Graduate School

1 概要

精密な電子部品を製造する際、電子部品に欠陥がないか検査装置により自動識別されている。電子部品の微細化、多機能化に伴い、検査の高感度化が要求されており、微細欠陥の検出は、正常部を欠陥として誤って認識する過検出を誘発してしまう。そのため、目視による再検査などが行われており、製造工程効率化のため再検査の自動化が望まれている。

目視検査自動化の手法には、Deep Learning を用いる方法があるが、Deep Learning を用いた検査では、学習に多数の画像が必要となるため、製造現場で学習用画像が確保できない場合、低正答率、低再現性を引き起こす。

本研究では、学習用欠陥画像の自動生成(1)(2)の可能性を検討している(3)。本報告では、低次元モデルで表現可能な Auto-Encoder(4)(自己符号化器。以下 AE と略す)による欠陥画像の自動生成手法を提案する。

本報告では、微細な欠陥の画像を複数生成することを狙い、AE の潜在変数の演算により、内挿画像を生成する。このとき、妥当な内挿画像が得られる条件を模擬パターンにより明らかにし、AE の安定的な適用を進める。

なお、正常データのみを用いて AE を学習させ、欠陥をこれからの乖離として認識する手法(5)が報告されているが、本報告は 2 クラス分類器の性能向上に寄与する学習用パターンの生成を目的としている。

2 欠陥画像の自動生成

欠陥画像の自動生成手法を図 1 に示す。文献 3 にて紹介した手法である。電子部品の配線回路パターンを用いて説明する。取得済みの欠陥画像と、これと同一となるよう形成された、対応する良品部の画像から差画像を求め、欠陥部を抽出する。この欠陥抽出画像と、対象とする貼り付け先の配線回路パターンの画像を、AE にそれぞれ入力し、Encoder 部で次元削減を行って求めた各潜在変数間で演算を行う。演算結果を Decoder 部で復元して、欠陥がコピーされた画像を得る。潜在変数間

の演算は、加減算が可能である。さらに、抽出欠陥の割合を変えることにより、欠陥の寸法を制御することも可能である。文献 3 では、この手法によりデータ拡張を効果的に行い、高感度な検査方法が実現できることを示した。

上記が欠陥画像生成の基本的な考えであるが、AE の適用限界などが不明瞭と考え、本報告では潜在変数を用いた内挿演算が良好に動く範囲を把握すべく、評価を行うことにした。

次に AE を説明する。AE は、図 2 に示すように、Encoder と Decoder で構成される。入力画像を低次元特徴へ変換する部分を Encoder と呼び、その低次元特徴量から元の入力画像を復元する部分を Decoder と呼ぶ。中間層の潜在変数の次元は中央で最も低くなり、入力画像 x の情報を圧縮表現した特徴が現れる。図 3 に潜在変数のベクトル波形を示す。Encoder に x を入力したとき、潜在変数 h は、以下の次式によって与えられる。

$$h = \sigma(x * W + b) \dots\dots\dots (1)$$

W は Encoder 側の結合重みであり、 b は対応するバイアスである。そして、 σ は Encoder 側の活性化関数であり、 $*$ は 2D 畳込み演算を示す。

畳み込んだ特徴マップを元の形状に復元するには、次式が行われる。

$$y = \tilde{\sigma}(h * \tilde{W} + c) \dots\dots\dots (2)$$

$\tilde{W}$ は Decoder 側の結合重みであり、 c は対応するバイアスである。 $\tilde{\sigma}$ は Decoder 側の活性化関数である。AE は入力画像のみで演算が行える教師なし学習である。

上記 AE により得られる潜在変数に対する内挿演算は次式により行われる。

$$V_i = l_s + \alpha(l_e - l_s) = \alpha \cdot l_e + l_s(1 - \alpha) \dots\dots\dots (3)$$

ここで、 l_s は良品の潜在変数、 l_e は欠陥の潜在変数であり、 α はパラメータである。このパラメータ α を変えて内挿演算が有効に動く範囲を評価する。図 4 に内挿の例を示す。以下に述べる評価実験ではパラメータ α を変えて、内挿の良否を把握する。

3 評価条件と評価結果

3.1 評価条件

電子部品等に使われているプリント基板を検査対象と想定する。対象とする画像は、8bit 階調の 24×24 画素の一層模擬パターンとする。図 5 にその 1 例を示す。模擬パターンは 1 種類とし、これに欠陥などを作り込み、配線パターンの明るさの違いなども考慮することとし、計 22 の場合を評価することにした。

評価条件を表 1 に示す。まずは 2 値画像の良品画像と欠陥画像を各 1 枚準備し、以下の制御項目について、 α をパラメータとして評価する。

- (1) 欠陥の寸法の制御 (画像の基本的な類似性)
- (2) 配線パターンの明るさの制御
- (3) フィルタとの畳込みによる画像のパターンエッジ勾配の制御

(フィルタのサイズは $2 \times 2, 3 \times 3$ とし、これを繰り返し作用させる。重みの合計は 1 とした)

これらの評価実験を通して、AE による内挿演算が良好な濃淡画像を生成可能な α の範囲を求め、画像データ拡張が信頼性高く実施可能な条件を明らかにする。

AE の特徴抽出に用いる Encoder 部は、図 6 に示すように 5 層とし、バッチサイズは 1、抽出する特徴は 100 次元、学習回数は 3000 回とした。

3.2 評価結果

提案手法を用いて生成した内挿画像の例を図 7 に示す。画像中央部の水平波形も併せて表示した。同図はフィルタごとの内挿結果を示す、同図(a)

(b)(c)は欠陥のサイズが小中大的場合を示す。

欠陥のサイズは 3×3 (小), 5×5 (中), 7×7 (大)画素である。波形が単調増加でない場合は内挿演算が誤っていると考えた。誤った箇所は背景を橙色にした。 3×3 (小)の欠陥サイズの濃淡画像の場合、 α に依存せず内挿により良好な画像が生成できた。

図 7(d)は配線パターンの明るさごとの内挿結果である。同様に図 7(e)はパターンエッジ勾配ごとの内挿結果である。

図 7(a)(b)(c)から欠陥の寸法が大きくなると、潜在変数の内挿演算が誤った動作をすることを確認した。図 7(d)から配線パターンの明るさの違いの影響も小さいことが確認できた。図 7(e)からパターンエッジ勾配の違いについても、内挿はあまり左右されないことを確認した。すなわち、 3×3 (小)の欠陥サイズの濃淡画像の場合、配線パターンの明るさの違いやエッジ勾配の違いには影響を受けず、かつ α の値に関係なく良好な内挿画像が生成できている。

図 7(a)(b)(c)に対応する潜在空間の 1 例を図 8 に示す。欠陥の寸法に応じて二つの画像の距離が遠

くなり、内挿精度が劣化すると考える。

これらの特性は、AE による画像生成の動作条件を与えるものであり、有益な結果となった。

上記特性に基づき、微細な欠陥が学習データに追加できるようになり、有効にデータを拡張でき、それによって高感度な欠陥検査を実現できるようになる。例えば文献(3)で述べたようにオートエンコーダで得られる潜在変数をサポートベクトルマシンに入力して欠陥判定が可能である。その一例を図 9 に示す。横軸がパラメータ α 、縦軸が検出率である。同図に示すように欠陥を配線パターンのエッジ、配線パターン間、配線パターン上に埋め込んだ。欠陥検出感度を評価できることが分かる。

4 結論

Deep Learning を用いた電子部品検査において、AE を用いた信頼性高い画像データ拡張について述べた。AE に与える画像の類似性をパラメータにして、内挿による欠陥画像生成が有効に動作する範囲を求めた。評価を通して高感度検査に有効な知見が得られたと考える。

参考文献

- (1) 伊藤秀将ほか:画像認識精度を向上させるディープラーニング技術を用いた学習用画像の自動生成,東芝レビュー, 72.4(2017)
- (2) 片山隼多ほか:DNN による外観検査の自動化における学習画像生成の検討,ViEW(2017)
- (3) 柳部正樹ほか, 電子部品検査における欠陥検査感度制御のための Deep Convolutional Auto-encoder を用いた学習画像生成の検討,精密工学会誌, Vol.85, No.8, pp.733-740 (2019.8)
- (4) Francois Chollet, Deep Learning with Python, Manning Publications (2017/12/22)
- (5) 梅田弘之:エンジニアなら知っておきたい AI のキホン(2019)

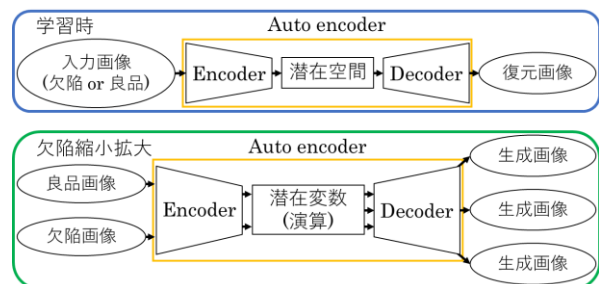


図1 オートエンコーダによる画像生成の全体構成

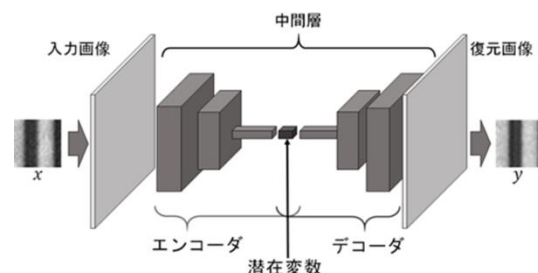


図2 オートエンコーダの構成

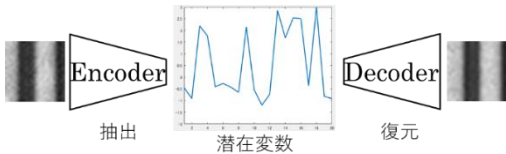


図3 オートエンコーダの潜在変数

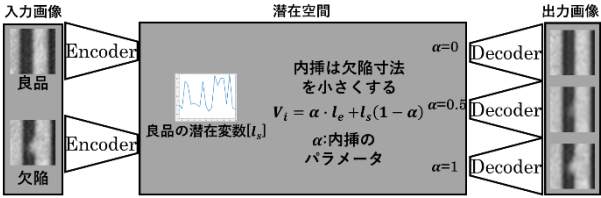


図4 内挿による画像の生成手法

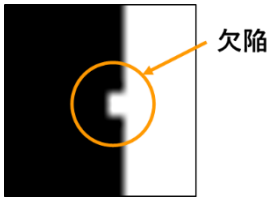


図5 模擬パターンの一例
(画像のサイズ:24×24 画素)

表1 生成条件

Index		Conditions
入力画像	良品	1枚
	欠陥	1枚
入力画像寸法		24×24画像
Epoch数		3000
Batch size		1
潜在変数次元数		100

Input	None,24,24,1
Convolution2D	None,24,24,32
MaxPooling2D	None,12,12,32
Convolution2D	None,12,12,16
MaxPooling2D	None,6,6,16
Flatten	None,576
Dense	None,100
Dense	None,576
Reshape	None,6,6,16
UpSampling2D	None,12,12,16
Convolution2D	None,12,12,32
UpSampling2D	None,24,24,32
Convolution2D	None,24,24,1
Output	None,24,24,1

図6 オートエンコーダの層構成

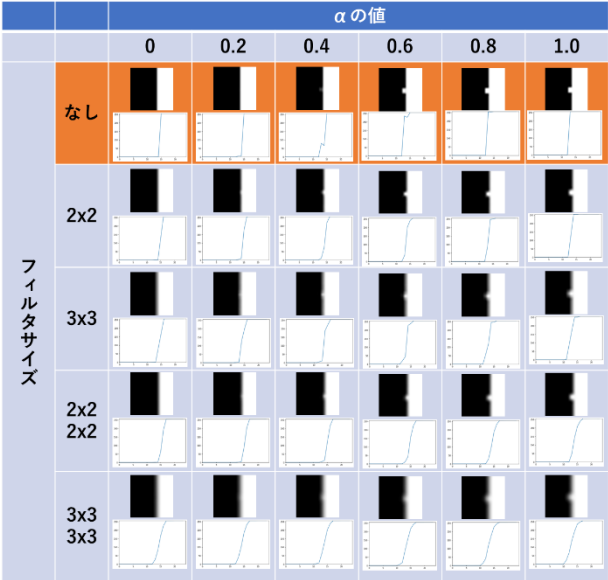


図7(a)フィルタごとの内挿結果(欠陥サイズ小)

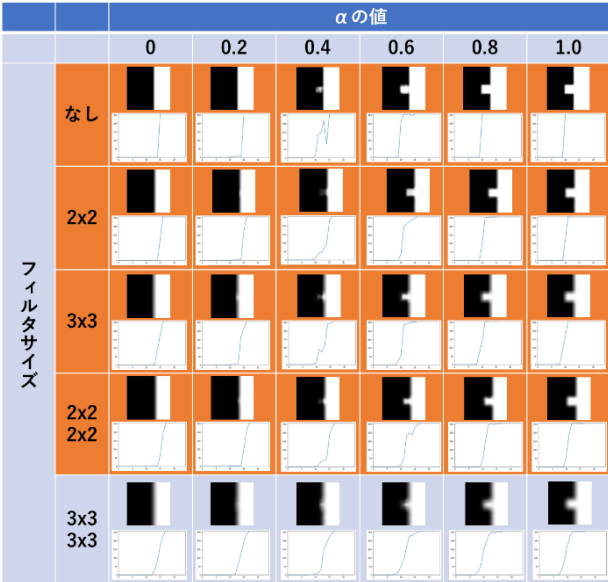


図7(b) フィルタごとの内挿結果(欠陥サイズ中)

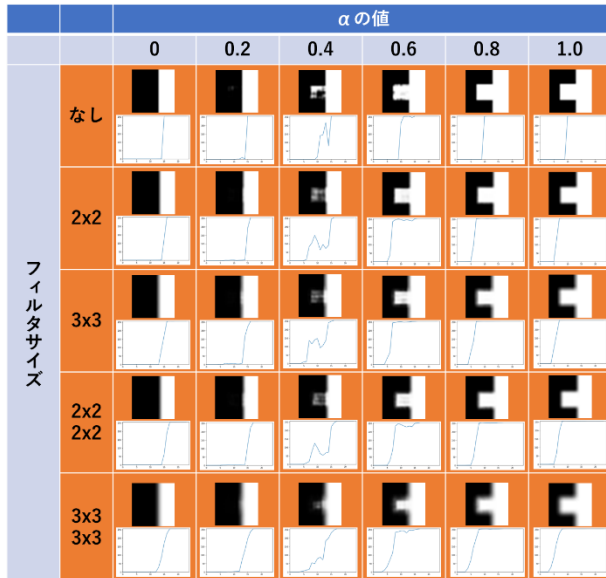


図 7(c) フィルタごとの内挿結果(欠陥サイズ大)

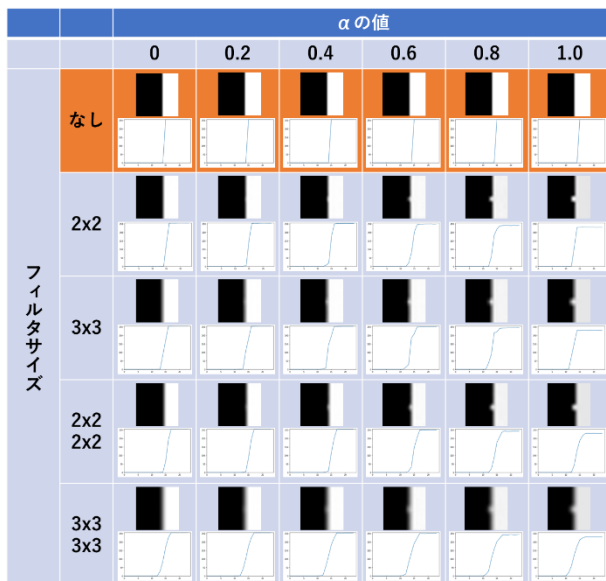


図 7(d) 配線パターンの明るさごとの内挿結果
; 2 値の場合の上段は学習時の loss が
下がらなかった
(欠陥サイズ小)

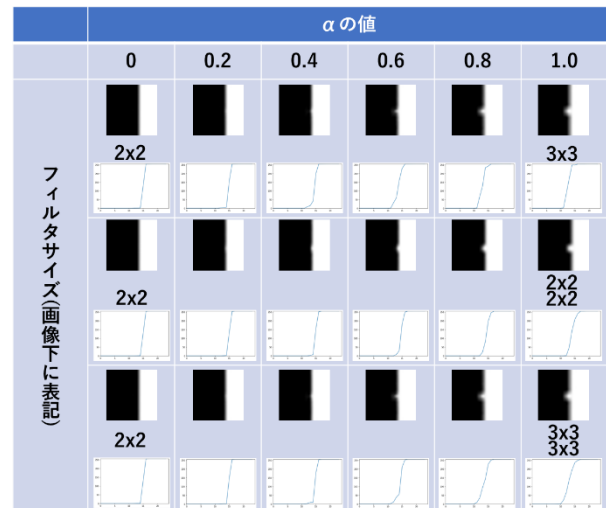


図 7(e) パターンエッジ勾配ごとの内挿結果
(欠陥サイズ小)

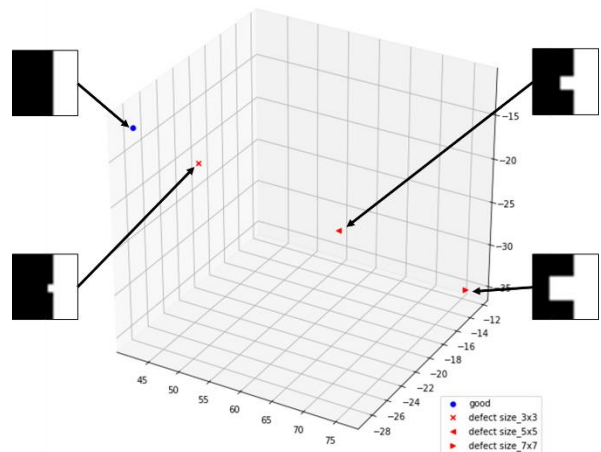


図 8 3次元潜在空間における良品と欠陥の分布
(欠陥サイズ大中小)

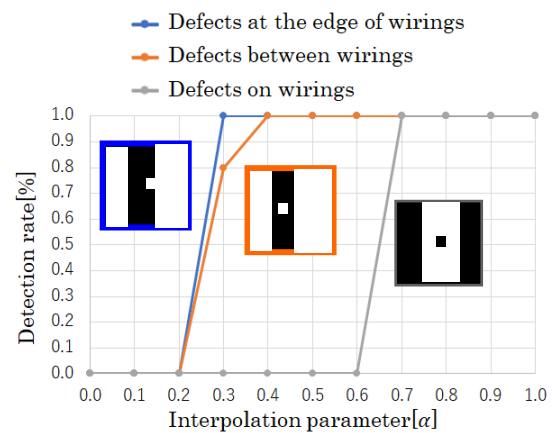


図 9 欠陥検出感度の評価結果
(実際には実物パターンを対象)