

NAIST-IS-MT1551048

修士論文

フェイルダイ特性を利用した LSI テスト結果予測の精度 向上に関する研究

佐藤 敬済

2017 年 3 月 16 日

奈良先端科学技術大学院大学
情報科学研究科 情報科学専攻

本論文は奈良先端科学技術大学院大学情報科学研究科に
修士（工学）授与の要件として提出した修士論文である。

佐藤 敬済

審査委員：

井上 美智子 教授 （主指導教員）

池田 和司 教授 （副指導教員）

大下 福仁 准教授 （副指導教員）

フェイルダイ特性を利用した LSI テスト結果予測の精度向上に関する研究*

佐藤 敬済

内容梗概

今日, LSI(集積回路) は鉄道や ATM などの社会インフラに組み込まれ, 求められる信頼性がますます高まっている. そのため, 様々なテストを行い, すべてのテストをパスした製品のみを出荷している. しかし, 大量の製品に対して, 数多くのテストを行うことは, 製造コストの半分以上を占めると言われるほど高いコストがかかる. そこで近年, データマイニング手法を用いて, 前段のテスト計測値から最終的な結果を予測し, テスト工程を省略することで, テストコストを削減する手法が注目されている.

本研究では, LSI のフェイルダイの持つ特性に注目し, テスト結果予測精度を向上する. 具体的には, 機械学習による LSI テスト結果予測において, フェイルダイをクラスタリングする手法と, ある特定のフェイルダイを抽出するため, 学習結果から特徴的なダイを抽出する手法を組み合わせ用いた. テスト結果予測は, サポートベクターマシン (Support Vector Machine, SVM) によって行う. 予測モデルの性能は, 受信者動作特性 (Receiver Operating Characteristic, ROC) 曲線に基づく, 曲線下面積 (Area Under Curve, AUC) によって評価する. 本研究では, 提案手法を用いることで, フェイル特性を考慮しない場合と比べて, AUC が 0.05 向上した.

キーワード

データマイニング, バーンインテスト, LSI テスト, フェイル特性, サポートベクターマシン

*奈良先端科学技術大学院大学 情報科学研究科 情報科学専攻 修士論文, NAIST-IS-MT1551048, 2017 年 3 月 16 日.

Pass/Fail Prediction in LSI Test Utilizing Fail Die Characteristics*

Takazumi Sato

Abstract

Today, VLSI (Very Large Scale Integrated Circuit) is embedded into social infrastructure such as railroad and ATM, the required reliability is increasing more and more. Therefore, we apply various kinds of tests and ship only fully reliable products. However, various kinds of tests for a lot of products is costly. Test costs is said to be approaching a half of manufacturing cost. Therefore, research on Pass/Fail prediction of test result by data mining has been conducted. This method enables reduction of test cost by omitting some test processes.

In this research, we focus on the fail die characteristics and effectively predict Pass/Fail. Specifically, we utilize the method of clustering and extraction of fail dies those are sensitive to be predicted for a part of products as fail. Pass/Fail prediction is done by Support Vector Machine. The performance of the prediction model is evaluated by the Area Under Curve based on the Receiver Operating Characteristic curve. In this research, by using the proposed method, the AUC was improved by at most 0.05 compared to when not considering the fail die characteristics.

Keywords:

data mining, burn-in test, LSI test, fail die characteristics, support vector machine

*Master's Thesis, Department of Information Science, Graduate School of Information Science, Nara Institute of Science and Technology, NAIST-IS-MT1551048, March 16, 2017.

目次

図目次	iv
1. はじめに	1
1.1 LSI の製造	2
1.2 故障の種別とバーンイン	3
1.3 データマイニングを用いた LSI テストコスト削減	4
1.4 本研究の目的	6
2. 対象データ	8
2.1 対象データの概要	8
2.2 テストフロー	9
3. 対象データの解析	11
3.1 フェイルカテゴリ	12
3.2 フェイルダイ特性	13
4. 提案手法	17
4.1 フェイルダイ特性を考慮したテスト結果予測の概要	17
4.2 フェイルダイ特性を活用した学習フェイルダイ選択法	18
4.2.1 k-means 法によるフェイルダイのクラスタリング	18
4.2.2 フェイルダイの抽出	19
4.3 テスト結果予測モデルの統合	21
5. 実験	23
5.1 実験データ	23
5.2 実験結果	24
5.2.1 k-means 法によるクラスタリングを用いた手法	24
5.2.2 フェイルダイの抽出を用いた手法	25
5.2.3 両手法 (抽出後にクラスタリング) を適用した結果	27

6.	まとめと今後の課題	29
	謝辞	31
	参考文献	32
	発表リスト	33

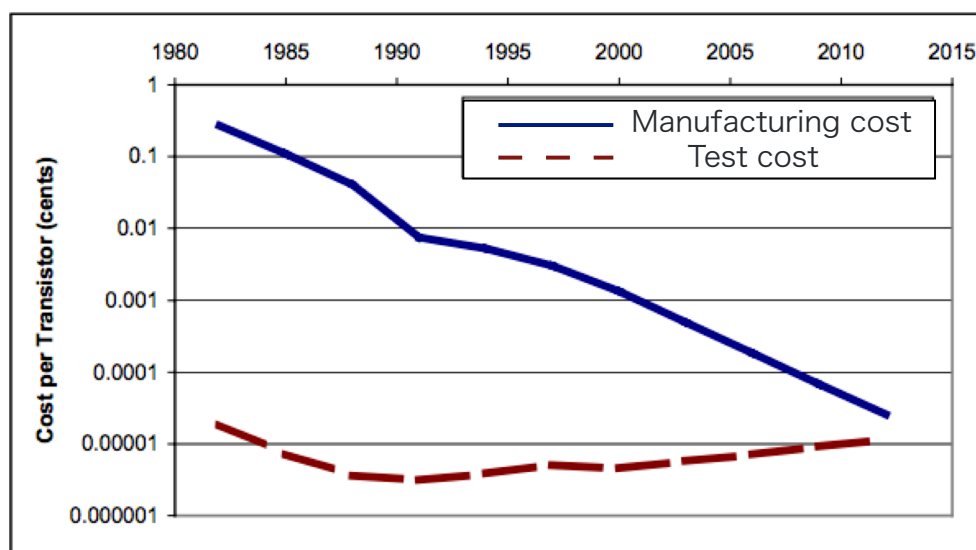
図目次

1.1	LSI の製造コストとテストコストの推移	1
1.2	製造フロー	3
1.3	故障率曲線	4
1.4	一般的なテストフロー	6
1.5	提供されているデータのテストフローとデータマイニングを利用 したテストフロー	6
2.1	テストフロー	9
3.1	ROC 曲線の例	12
3.2	学習・検証ダイによって予測精度にばらつきがある例	14
3.3	学習に用いたダイのクラスによって予測精度にばらつきがある例 (ウェハ X, Y からウェハ Z を予測)	15
3.4	学習に用いたダイのクラスによって予測精度にばらつきがある例 (ウェハ X, Z からウェハ Y を予測)	15
3.5	学習に用いたダイのクラスによって予測精度にばらつきがある例 (ウェハ Y, Z からウェハ X を予測)	16
4.1	提案手法のフロー	18
4.2	予測の容易なフェイル特性ダイの抽出フロー	20
4.3	代表ダイの抽出例	21
6.1	フェイル特性を考慮しない場合と提案手法を用いた場合の ROC 曲線	29

1. はじめに

半導体デバイスのコストには、設計に関わるコスト、製造に関わるコスト、テストコストが存在する。デバイスの微細化が進み、デバイスを構成する半導体素子の集積度が上がると、図 1.1 に示す通り、素子一つあたりの製造コストは低減する。一方で、テストコストは集積度が上がり素子数が増加しても試験対象となる素子数が増えるため、デバイスのテストコストを決める主要因である時間は増加する。そのため、製造コストと違い、素子一つあたりのテストコストは低減しない。また、同一のデバイスの生産量が増えてくるとデバイス一つあたりの設計コストは低減するが、デバイス一つあたりのテスト時間は変わらないので、テストコストは低減しない [8]。

そこで、テストコストを削減する手法として、データマイニング技術が注目されている。この節では、まず第 1.1 節において一般的な LSI の製造工程を示す。次に第 1.2 節では、LSI の故障の種別とバーンインテストについて述べる。そして、第 1.3 節では、データマイニングを用いた LSI テストコスト削減に関する研究について述べる。最後に、第 1.4 節では本研究の目的を述べる。



*引用:ITRS 2002 Update Conference

図 1.1 LSI の製造コストとテストコストの推移

1.1 LSI の製造

LSI は一般的に図 1.2 のフローのように製造される。フローに示すように、LSI の製造は大きく前工程と後工程に分かれている。

前工程では、まずウェハ処理が行われる。高純度のシリコン単結晶棒を薄い円盤状に切り出したものをウェハと呼び、このときに LSI 製造用のウェハのサイズを決定する。次は洗浄工程である。半導体素子は非常に微細で汚れによる影響を受けるため、このときにウェハ表面を鏡面状に研磨、洗浄する。次が成膜工程である。成膜工程では、トランジスタを構成するためにシリコンウェハ上に絶縁膜や配線層を成形する。その後、リソグラフィが行われる。リソグラフィでは薄膜に対して感光性の物質を塗布してパターン上に露光させることで、露光している部分と露光していない部分を構成する。そうすることで、ウェハ上に絶縁膜で覆われている部分とそうでない部分を作り、導電する部分とそうでない部分を構成する。次が不純物拡散工程である。シリコン上に P 型や N 型の半導体領域を形成することでトランジスタができる。この工程ではホウ素やリンなどの不純物をウェハに添加し、熱処理などにより、不純物を拡散させる。以上の工程でウェハ上にトランジスタが構成される。前工程の最後にウェハ上に構成されたダイに対してテストを行う。この工程をウェハテストと呼ぶ。

次は後工程である。後工程ではウェハ上に形成されたダイを切り分けるダイシングが行われる。その後、切り取ったダイをパッケージングすることでチップが構成される。パッケージングの工程では、各ダイはリードフレームと呼ばれる基盤に接続され、ゴミや水分などからデバイスを守るために樹脂によってパッケージされる。

最後に後工程によりパッケージ化されたチップに対してテストが行われる。この工程をパッケージテストと呼び、テストをパスすればチップは市場へと出荷される。

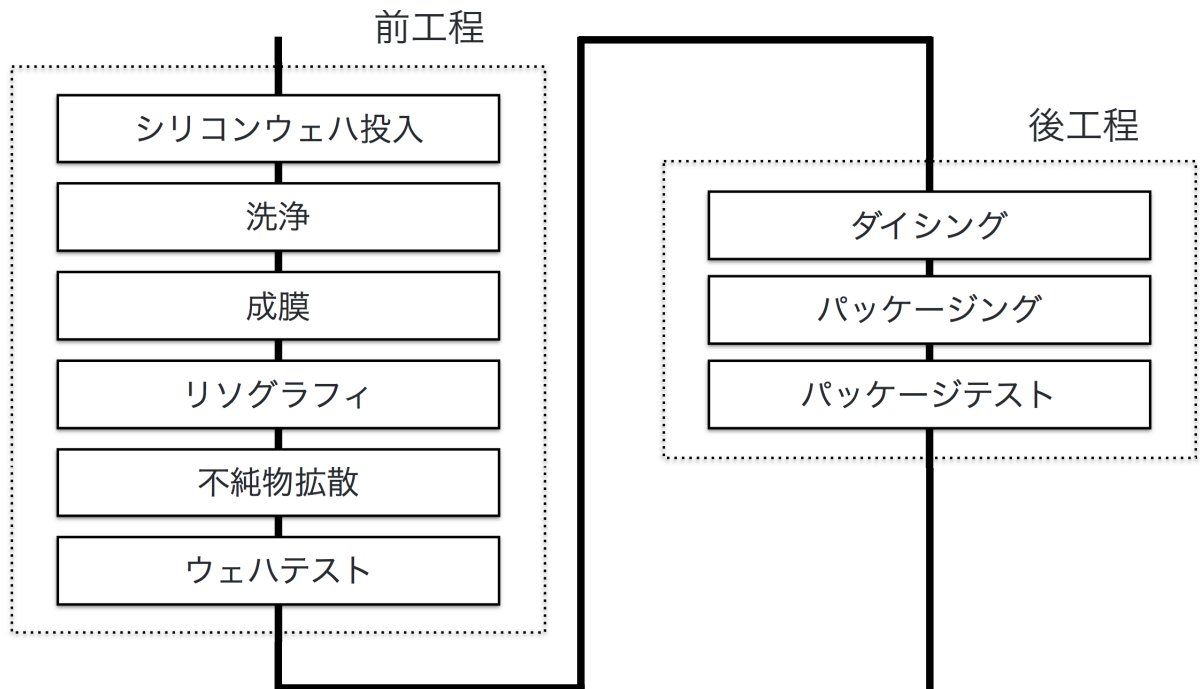


図 1.2 製造フロー

1.2 故障の種別とバーンイン

LSI の故障は、初期不良、偶発故障、磨耗故障の 3 つに大きく分類される。故障率と時間軸との関係は図 1.3 のようにバスタブカーブを描く。初期不良は 1.1 節で示した製造工程に起因した潜在的な欠陥が顕在化することで発生し、生産直後はこの初期不良により高い故障率となる。その後低い故障率で推移し、さらに時間が経過すると、LSI の寿命による磨耗故障によって再び故障率は上昇する。

パッケージテスト工程の 1 つに、バーンイン (BI: Burn-In) がある。この工程は、LSI に高温や高電圧といったストレスを与えて劣化を加速させることで、初期不良となる LSI を検出するために行われる。よって、このバーンインは初期不良となる LSI を取り除き、高い品質を保証するために有効な手法と知られている。しかし、バーンインの工程には、大きな問題点が二つ存在する。第一に、多くのコストがかかるという点である。バーンインを行う装置は高価な上、バーンイン工程は、数時間から数十時間に及ぶ時間が必要とされる。そのため、長時間テスト装置を占有することになり、高いコストが発生する。バーンインのコストは製造する製品によっては

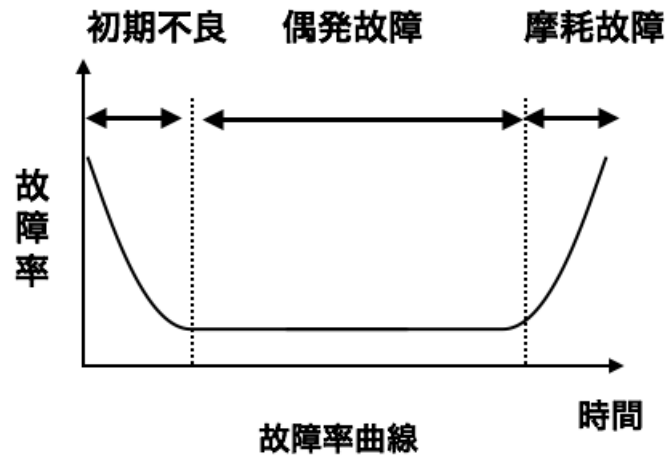


図 1.3 故障率曲線

全体のコストの 40% に達するとの報告 [1] もある。第二に、バーンイン工程自体が製品を劣化させてしまうという点である。バーンイン工程では、製品にストレスを与えることにより劣化を加速させ、初期不良を出荷前に顕在化させることで検出する。しかし、その反面、バーンインは全ての製品に行われるため、最終的に出荷される良品の寿命も縮めてしまう。

1.3 データマイニングを用いた LSI テストコスト削減

バーンインの問題は、データマイニング技術を用いて、初期不良と判明するダイを予測することができれば省略し、改善することができる。そのため、バーンイン結果を予測する研究が近年盛んに取り組まれている。

一般的なテスト工程では、図 1.4 の上図に示すようにテストが行われ、全てのテストにパスした製品のみが出荷される。このようなテストフローにおいて、テストコストを削減するために、データマイニングを用いたテスト結果の予測手法が考えられている。本研究では、図 1.5 中の上図のように一般的なテスト工程を区切って、少しずつ実施したデータが提供されている。このような区切られたテストフローにおいて、前段の各テスト工程で観測された物理的特性などの計測結果を用いて、後段のテスト結果を予測することで、図 1.5 中の下図のようなテストフローを取ることができる。このテストフローでは、良品と予測された LSI には実施する後段のテスト

工程を減らし、それ以外の製品には通常の工程を行う。この方法により、LSI に対する信頼性を十分保持したまま、テスト時間とテストコストの削減が可能となる。また、良品に対するバーンイン工程を減らすことで、良品の寿命を縮めてしまうという問題点も避けることが可能となる。

しかし、バーンインによりフェイルするダイを正確に予測することは、とても困難とされている。一般的にテストにフェイルしたダイは、コスト削減のため、それ以降はテストされない。そのため、テスト工程の初期段階でフェイルしたダイに対しては、後段のテスト結果が得られず解析に利用することができない。また、バーンインにより検出されるフェイルは、製品やテクノロジーにも依存するが、パスダイのわずか数 % 程度と非常に少数のため、学習に用いるパスダイとフェイルダイの個数に不均衡が生じ、予測が困難である。

近年、データマイニングの手法を用いて、テスト間の関連や、出荷後の不良により返品された製品と製造テストの関連を解明する研究がなされている [2-6]。特にバーンインテストの結果を予測することで LSI テストコストの効率化を行う研究として、Sumikawa らの研究 [2]、や、野々山らの研究 [6] がある。Sumikawa らの研究 [2] では、サポートベクターマシン (Support Vector Machine, SVM) を用いてバーンイン後の結果を予測している。この際、計測機器によるばらつきやフェイルの原因ごとに予測モデルを構築することで、85% のバーンインコストの削減を行っている。野々山らの研究 [6] では、標準偏差によるクラスタリングや、サイト間ばらつきの補正を用いた SVM による予測が提案されている。さらに Sumikawa らの研究 [3] では、製品 100 万個に 1 つ程度の割合で発生する、出荷後に顧客のもとで不具合が発生し返品される可能性のある製品の特特定を、LSI に行う多変量テストの結果から予測している。また、Ye らの研究 [5] では、LSI の故障製品のテストや修理といった複雑な工程のコストを SVM を用いた手法により、削減している。小河らの研究 [7] では、計測機器やパッケージ間のばらつきを補正することで、ばらつきによってフェイルの特徴が隠蔽されることを防ぐ手法が提案されている。このように、近年データマイニングを用いて、LSI テストコストの削減を目指す研究はさかんに行われている。

通常のテストフロー



図 1.4 一般的なテストフロー

本研究において提供されたデータのテストフロー



データマイニングを用いたときのテストフロー

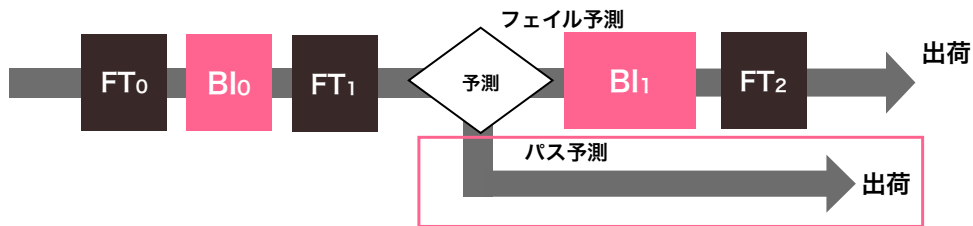


図 1.5 提供されているデータのテストフローとデータマイニングを利用したテストフロー

1.4 本研究の目的

本研究では、テストにフェイルしたダイのもつフェイル特性を利用することで、フェイル特性の似たフェイルダイを高精度に検出する。具体的には、全テスト工程のうち、前段部分のパッケージテストのテスト結果を使用し、SVM を用いて最終的なテスト結果を予測する。SVM によるテスト結果予測の際、フェイルダイの特性を活かした学習データの選択を行うことにより、予測精度を向上させる手法を提案

する。

以降の本論文の構成を示す。まず第 2 節で、本研究の対象とするデータについて述べる。次に第 3 節では、対象データの解析より明らかになったフェイルダイ特性について述べる。第 4 節でフェイルダイ特性を利用した学習手法について説明し、第 5 節で実験結果と考察を示す。最後に第 6 節でまとめと今後の課題を述べる。

2. 対象データ

本節では、本研究で対象とするデータについて説明する。まず第 2.1 節では、本研究で扱うデータの概要について述べる。次に第 2.2 節では、ダイのデータが計測されたテストフローについて述べる。

2.1 対象データの概要

本研究で対象とするデータについて説明をする。対象データに含まれる情報を表 2.1 に示す。

表 2.1 対象データの概要

項目	概要
X 座標	ウェハ上のダイの X 座標
Y 座標	ウェハ上のダイの Y 座標
パッケージロット	パッケージロットの番号
フェイルプロセス	フェイルしたテスト工程
フェイルテスト	フェイルしたテスト項目
フェイルカテゴリ	フェイルしたカテゴリ
テスト項目	計測値
テストタ	テスト番号
テストサイト	サイト番号

まず各ウェハには番号が付与されており、各ダイが各ウェハ内に所属する位置の X 座標と Y 座標が示されている。また、1 つのウェハから複数のパッケージが作成されるため、各ダイがどのパッケージに所属するかという情報も提供されている。フェイルプロセス、フェイルテストはそれぞれ各フェイルダイがフェイルしたテスト工程やテスト項目を表している。フェイルカテゴリは各フェイルダイをフェイルの原因によって大きく分類したときのカテゴリであり、本研究では特にフェイルカテゴリによるフェイル特性の違いに注目している。また、テストを行ったときの計

測値, その際に利用したテストの番号やテストサイトの番号も示されている.

2.2 テストフロー

本節では, 対象データのダイがどのようにテストされたか説明する. 提供されたデータは図 2.1 に示すようなフローでテストが行われた.

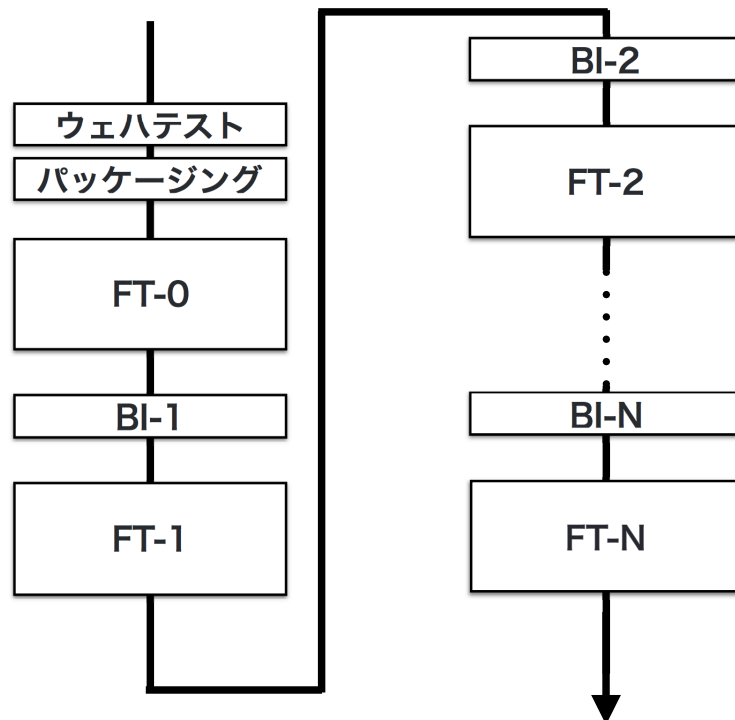


図 2.1 テストフロー

まず, ウェハテストが行われる. 次に, ウェハテストをパスしたダイのみパッケージングが実施される. その後, パッケージテストとバーンイン (BI: Burn-In) が行われる. パッケージテストはファイナルテスト (FT: Final Test) と呼称される. 本研究で用いた対象データでは, 図 2.1 に示したようにバーンインとファイナルテストを繰り返し実施する. 1 つのテスト工程 (FT-0 や FT-1 など) は, 電流, 電圧, 周波数など数百種類のテスト項目で構成されており, 温度などのテスト環境を変えて繰り返し, 同じテスト項目をテストしている.

本研究では, 図 2.1 における FT-1 までのテスト計測値を用いて, それ以降のテス

トでフェイルするダイを特定することを目的とする。これは、バーンインによって本来後段で検出される劣化性のフェイルを、前段で検出することで、テストコストの削減を図るためである。

3. 対象データの解析

本節では、提供されたデータについて実験に用いたデータと解析から明らかになった知見を紹介する。

まず、実験に用いたデータとして、前処理後のデータを表 3.1 に示す。テスト項目数は、パッケージテスト工程の前段部分である、FT1 以前の電流や電圧、周波数などを各ダイごとに計測するテストの前処理後の数を表し、これらのテストの計測値から、最終的なテスト結果を予測する。各ウェハのダイ数とフェイルダイ数は、実験に用いた 3 枚のウェハ X, Y, Z のそれぞれに所属する前処理後のダイ数とフェイルダイ数を表している。

表 3.1 実験データの概要

テスト項目数		771
各ウェハのダイ数 (フェイルダイ数)	ウェハ X	1817 (Fail :24)
	ウェハ Y	2718 (Fail :45)
	ウェハ Z	3127 (Fail :46)

次に、解析を行う際に用いたテスト結果予測精度の評価方法を示す。予測精度の評価は、受信者動作特性 (Receiver Operating Characteristic, ROC) 曲線に基づく、曲線下面積 (Area Under Curve, AUC) を用いて行った。本研究における ROC 曲線の例を、図 3.1 に示す。ROC 曲線は、パスとフェイルを決める閾値を変化させ、縦軸に全フェイルダイの内フェイルダイと予測できた割合、横軸に全パスダイの内フェイルダイと予測してしまったダイの割合を示す。理想的な予測ができていれば、ROC 曲線は (0, 0) の位置から (0, 1) に上昇し、さらに水平に (1, 1) へ推移する。AUC は、ROC 曲線の下側の面積を表し、こちらは理想的な予測ができていれば、1 となる。

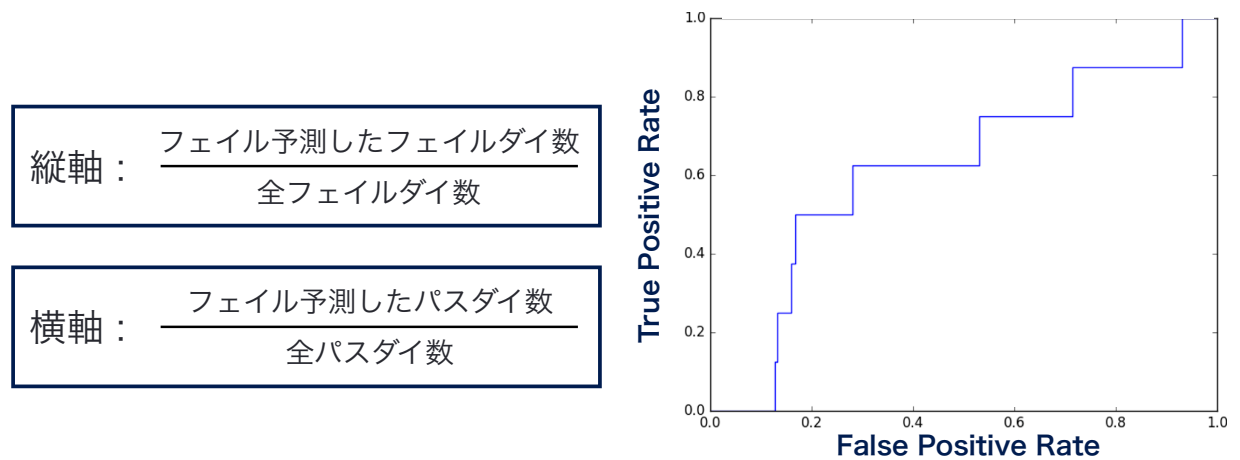


図 3.1 ROC 曲線の例

以降では、3.1 節ではフェイルカテゴリに関して、フェイルカテゴリごとにテスト結果予測を行うことが有効なカテゴリの存在や、一部のカテゴリは予測しやすいことを紹介し、本研究の予測対象とするフェイルカテゴリについて述べる。3.2 節では、提供されたデータにおけるフェイルダイ特性について述べる。

3.1 フェイルカテゴリ

本節では、フェイルカテゴリの概要と、全てのフェイルダイを対象に予測を行うよりも、特定のカテゴリのフェイルを対象として予測を行った方が有効な例を示す。表 3.2 に実験データのフェイルダイの属するカテゴリを示した。フェイルカテゴリは各フェイルダイをフェイルの原因によって大きく分類したときのカテゴリであり、実験データには 6 種類存在する。フェイルカテゴリに属するダイの数にはばらつきがあり、カテゴリによっては、あるウェハにはほとんど存在しないようなカテゴリもある。

表 3.2 各カテゴリのフェイルダイ数

	C2	C9	C11	C12	C14	C15
フェイルダイ数	1	13	14	73	12	2

表 3.3 にウェハ X, Y, Z について, ウェハ 2 枚分のダイを学習用に, ウェハ 1 枚分を検証用に用いてサポートベクターマシン (Support Vector Machine, SVM) によるテスト結果予測を 3 通り行った結果の AUC を示す. 表 3.3 の AUC は, 予測対象とするフェイルカテゴリを考えない場合と, 特定のフェイルカテゴリのみにした場合の結果を示す. それぞれの学習・検証の組み合わせにおいて, カテゴリ 12 やカテゴリ 14 はフェイルカテゴリごとに予測することによって, フェイルカテゴリを考えない場合と比べて高い AUC を得ていることがわかる. このように, カテゴリごとの予測を行うことが有効なカテゴリが存在し, カテゴリによって予測の難度が異なる. 本研究では, カテゴリ 12 のフェイルを予測の対象とする. カテゴリ 12 のフェイルを予測対象とする理由は, AUC の改善の余地が大きく, また全てのウェハに多く存在するからである.

表 3.3 AUC(検証ウェハのフェイルダイ数)

	X, Y $\rightarrow$ Z	X, Z $\rightarrow$ Y	Y, Z $\rightarrow$ X
カテゴリ分けなし	0.553(46)	0.567(45)	0.604(24)
C9	0.385(11)		0.613(2)
C11	0.483(8)	0.453(4)	0.192(2)
C12	0.615(21)	0.665(37)	0.641(15)
C14	0.678(5)	0.706(3)	0.712(4)

3.2 フェイルダイ特性

本節では, 実験の結果から, フェイルダイの持つ特性に注目した経緯を述べる. まず複数のウェハを混ぜ, ランダムに SVM による学習と検証を 15 回試行した実験結果の ROC 曲線を図 3.2 に示す. ほとんどの試行では, およそ同じような曲線を描いているが, 二回大きく異なる曲線を描いている試行が確認できる. この結果より, 学習と検証に用いるフェイルダイの関係性に注目し, さらに顕著な関係性を確認するため, クラスタリングを行った.

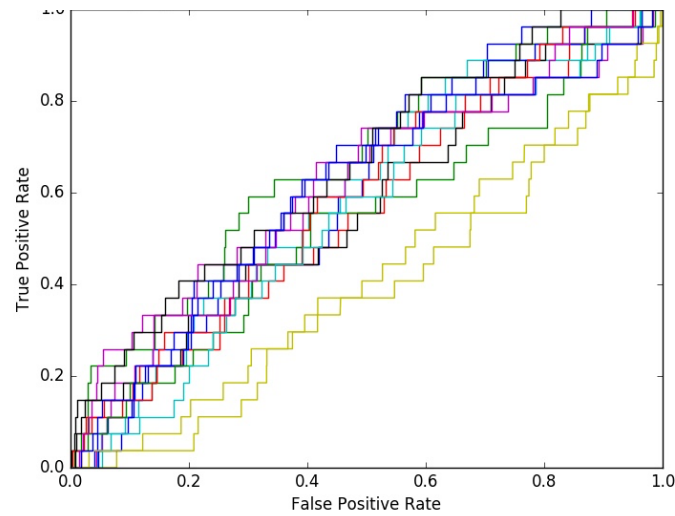


図 3.2 学習・検証ダイによって予測精度にばらつきがある例

クラスタリングの実験では、ウェハ X, Y, Z について、ウェハ 2 枚分のダイを学習に、ウェハ 1 枚分を検証に用いて SVM によるテスト結果予測を 3 通り行った。その結果を示す。ここで通常の予測に加えて、学習に用いるフェイルダイを k-means 法を用いて 2 クラスにクラスタリングし、それぞれのフェイルクラスターのダイとパスダイで学習し、予測を行った結果の ROC 曲線を図 3.3, 図 3.4, 図 3.5 に示す。図中にそれぞれの予測モデルによる AUC を示した。k-means 法によるフェイルダイのクラスタリングから作成したモデルに関しては、モデルの作成に用いたフェイルダイ数を示した。いずれの結果においても、学習に用いたダイのクラスターによって予測精度が大きく異なることが確認できる。また、この予測精度は、学習に用いるフェイルダイの数に依存しないことも確認できる。

これらの実験結果から、フェイルダイのもつ特性に注目した。学習に用いたフェイルダイと近い特性を持つフェイルダイは、予測しやすいと考える。そこで、本研究では、フェイルダイの特性を活かし、学習に用いることのできるフェイルダイから、代表的なフェイルダイの選択し、学習に用いることで、予測を高精度化する。具体的な代表フェイルダイの選択法に関して、次節で述べる。

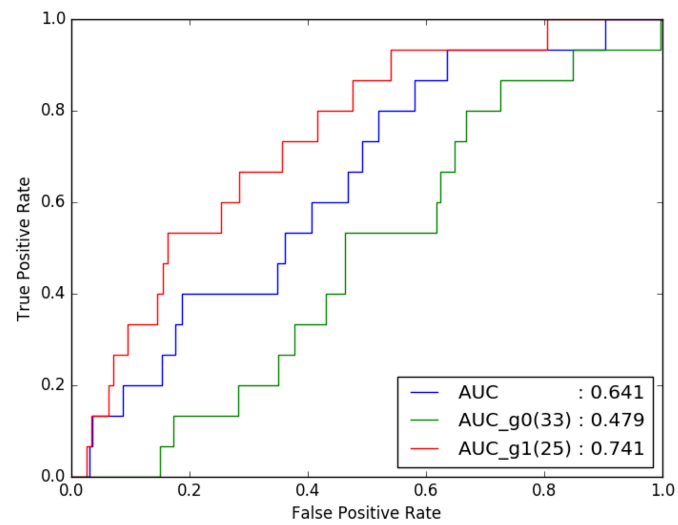


図 3.3 学習に用いたダイのクラスによって予測精度にばらつきがある例 (ウェハ X, Y からウェハ Z を予測)

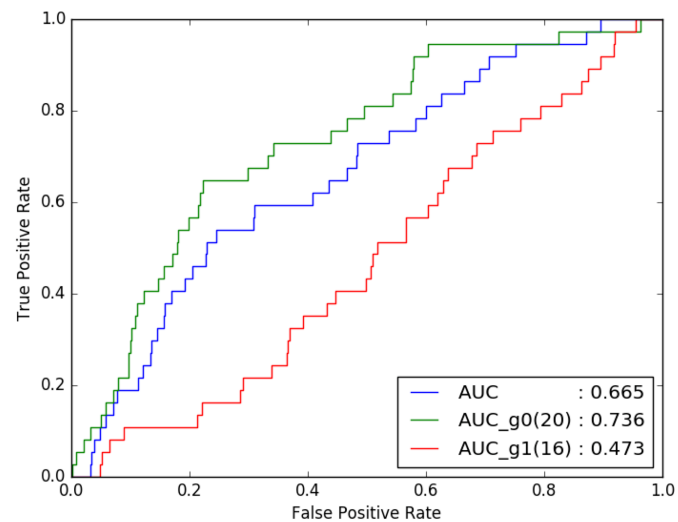


図 3.4 学習に用いたダイのクラスによって予測精度にばらつきがある例 (ウェハ X, Z からウェハ Y を予測)

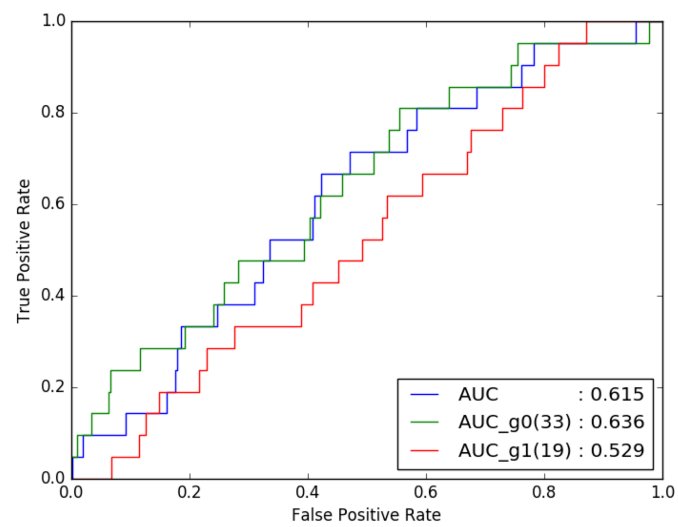


図 3.5 学習に用いたダイのクラスによって予測精度にばらつきがある例 (ウェハ Y, Z からウェハ X を予測)

4. 提案手法

3.2 節で示したように、予測対象のフェイルダイと学習に用いたフェイルダイが似た特性を持つとき、予測精度が高くなる傾向がある。そこで本節では、フェイルダイ特性を利用し、学習に用いるフェイルダイから代表的なフェイルダイを選択し、それらを用いてテスト結果予測を行う手法を提案する。まず第 4.1 節では提案する手法の全体像を示す。提案する手法では、SVM を用いてテスト結果の予測を行う。次に第 4.2 節において、具体的にフェイルダイ特性を利用する方法として、フェイルダイ特性によってフェイルダイをクラスタリングする手法を提案する。最後に第 4.3 節では、実験で構築される複数の予測モデルによるテスト結果予測の統合方法について述べる。

4.1 フェイルダイ特性を考慮したテスト結果予測の概要

第 3 節で述べたように、フェイルダイはそれぞれフェイルダイ特性を持つ。そこで本研究では、データマイニングによるテスト結果予測を行う際、フェイルダイ特性を考慮して予測モデルを構築することで、予測精度の向上を行う手法を提案する。

提案手法全体の流れを図 4.1 に示す。提案手法の詳細をフェーズごとに述べる。

まず、学習データをパスダイとフェイルダイに分割する。このときフェイルダイに対して、フェイルダイ特性による学習フェイルダイの選択を行う。具体的なフェイルダイ特性の活用法は、4.2 節において述べる。

次に、パスダイと 1 つのフェイルダイ集合より、SVM による予測モデルを構築する。ここで、フェイルダイ特性によって選択されたフェイルダイによって、それぞれ予測モデルを構築する。これらの予測モデルは、学習データからそれぞれ選択されたフェイルダイ集合の特性に似たフェイルダイに感度の高い予測モデルであると考えられる。

最後に、構築したそれぞれの予測モデルを用いて、検証データのテスト結果の予測を行う。この際、予測モデルごとにテスト結果の予測結果が異なるため、フェイル確度の大きいモデルを採用することで、予測結果を統合する。モデルの統合方法に関しては、第 4.3 節に詳細を示した。予測モデルの性能は、第 3 節で示した通り、ROC

曲線に基づく, 曲線下面積 AUC によって評価する.

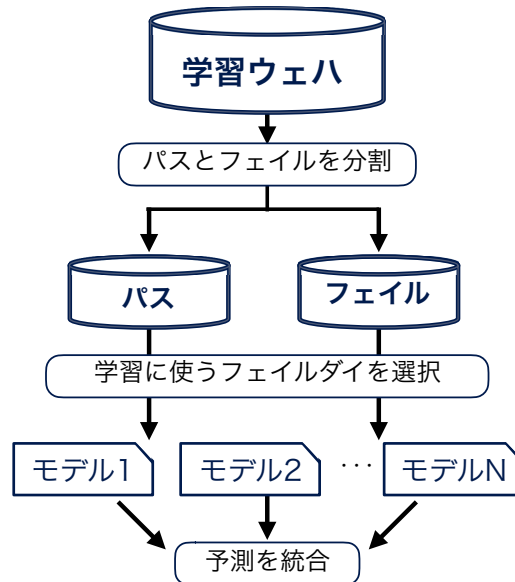


図 4.1 提案手法のフロー

4.2 フェイルダイ特性を活用した学習フェイルダイ選択法

第 4.1 節では, フェイルダイ特性を活用し, テスト結果予測を行う手法の概要を示した. 本節では, 具体的にフェイルダイ特性活用の手法を示す. 本研究では, 以下の 2 種類の方法でフェイルダイ特性を扱った. それぞれの手法について示す.

4.2.1 k-means 法によるフェイルダイのクラスタリング

この手法では, 学習データのフェイルダイを k-means 法によりクラスタリングすることで, フェイル特性によるダイの分割を行う. クラスタリングによって, 学習データのフェイルダイクラスに似たフェイル特性を持つフェイルダイの予測を効果的に行う. k-means 法によって分割するクラスタ数は, 実験結果の AUC から検討し, いずれのウェハの組み合わせのときもクラスタ数は 2 とした.

4.2.2 フェイルダイの抽出

まず、学習データのフェイルダイから、結果予測が容易なフェイルダイを抽出する。結果予測が容易なフェイルダイは、多くのフェイルダイに共通したフェイル特性を持っているため、予測が容易であると考えられる。そこで、結果予測が特に容易なフェイルダイを用いて学習を行うことで、検証データにおいて共通するフェイル特性を持つフェイルダイを効果的に予測する。具体的には、テスト結果が既知であるデータを用いて交差検証を行い、予測が困難なダイを特定する。予測が困難であったフェイルダイは、学習データから除くことで、予測が容易なフェイルダイを抽出する。

予測の容易なダイ抽出方法の詳細を述べる。本稿で実施した代表ダイの抽出の流れを図 4.2 に示す。まず、テスト結果が既知であるデータをランダムに 2:1 に分割し、2/3 を学習用データ、1/3 を検証用データとし、SVM によるテスト結果予測を行う。そして検証の際、一定の FPR(FPR: False Positive Rate) を閾値として、閾値までにフェイルを予測できなかったフェイルダイは、予測が困難であるとし、学習データから除く。この工程における検証結果の一例の ROC 曲線を図 4.3 に表す。この例において、一定の閾値をフェイルとしても予測できなかったフェイルダイ数を表 4.1 に示す。表 4.1 に示した、一定の閾値までのダイをフェイルとしても予測できなかったフェイルダイをフェイル特性が異なる予測の困難なダイと定義する。本研究では、閾値とする FPR を 0.8, 0.6, 0.4 と変化させ、実験を行った。例えば、FPR が 0.8 のとき、フェイルと予測できなかったフェイルダイは、全パスダイの 80% をフェイルとしても予測できなかったフェイルダイであるため、予測は困難である。本研究では、閾値とする FPR を 0.8, 0.6, 0.4 と変化させ、この工程を 15 回試行することで、予測の容易なフェイルダイの抽出を行った。試行の際、複数回に渡ってフェイル予測が困難であったフェイルダイは、その回数も記録する。これは、複数回予測が困難であったフェイルダイを優先して除くべきと考えるためである。このような手順でテスト結果が既知であるデータから、フェイル特性によってダイを抽出する。なお、学習データと検証データの割合や、試行回数はフェイルダイの数や、実験結果から最良のものを選択した。

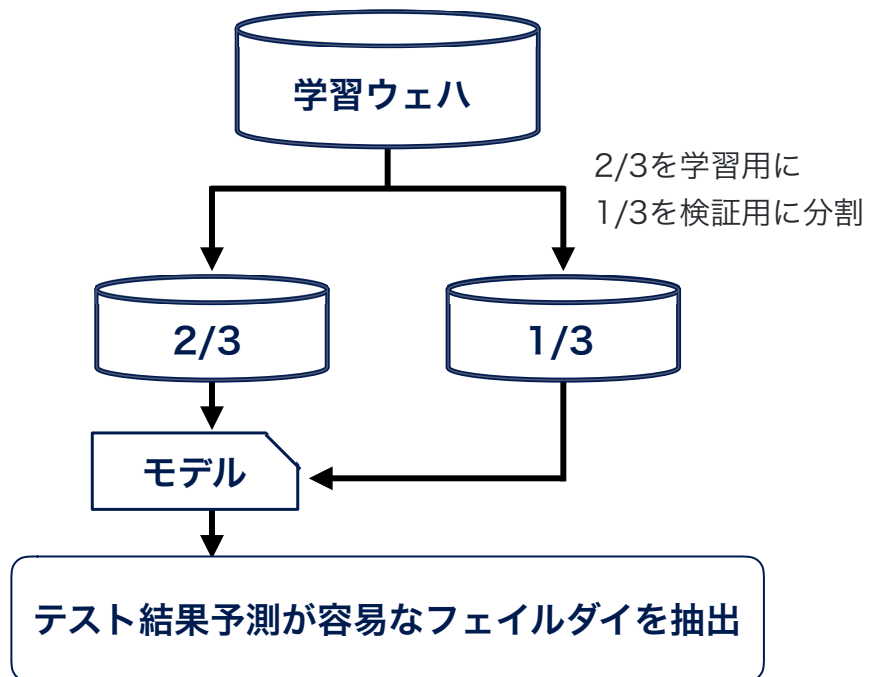


図 4.2 予測の容易なフェイル特性ダイの抽出フロー

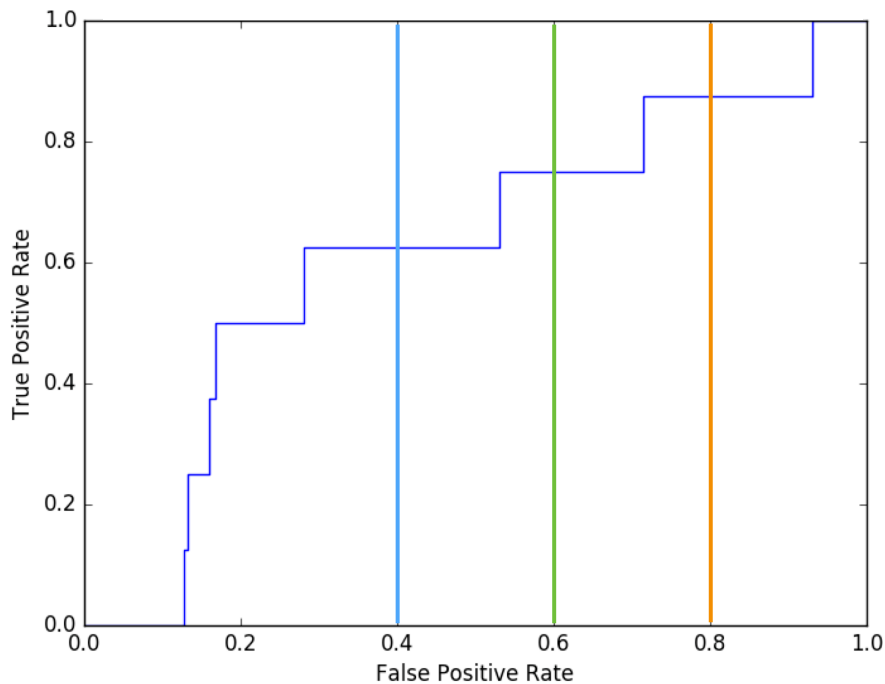


図 4.3 代表ダイの抽出例

表 4.1 閾値以下をフェイルとしたとき予測できなかったフェイルダイの数

閾値 (FPR)	0.8	0.6	0.4
予測できなかったダイ数	1	2	3

4.3 テスト結果予測モデルの統合

第 4.2 節では、フェイルダイ特性をクラスタリングする手法と抽出する手法を提案した。いずれの手法であっても、複数の予測モデルが構築されるため、モデルによる予測を統合する必要がある。そこで本研究では、各モデルによる各ダイへのテスト結果予測のうち、最もフェイル確度が高いものを採用することで、結果予測を統合する。複数のモデルによる結果予測モデルの統合例を表 4.2 に示す。表 4.2 は、予測モデル 1, 2, 3 とこれらを統合したモデルによる各ダイに対するフェイル確度を表している。この例に示したように、各モデルの最も高いフェイル確度を採用する

ことでモデルの統合を行う。

表 4.2 テスト結果予測モデルの統合例

	モデル 1	モデル 2	モデル 3	統合モデル
ダイ A	0.2	0.3	0.4	0.4
ダイ B	0.5	0.8	0.2	0.8
ダイ C	0.1	0.1	0.2	0.2

5. 実験

本節では、第 4 節で述べた方法によって、フェイルダイ特性を考慮したテスト結果予測を提供されたデータに対して行う。そして、提案手法の予測精度を AUC により評価する。以降、第 5.1 節では、実験に用いるデータを示す。また、第 5.2 節では、フェイル特性のクラスタリングを行ったテスト結果予測の実験結果を示す。

5.1 実験データ

実験で用いるデータについて説明する。実験データの概要を表 5.1 に示す。まず、提供されているデータのうち、ウェハ 3 枚分のダイを実験に用いる。実験では、ウェハ X, Y, Z のうち、ウェハ 2 枚分のダイを予測モデル構築のための学習に利用し、ウェハ 1 枚分のダイを検証に用いる。予測モデルの構築には、テストフローの前段部分に属するテスト項目のうち 771 テスト項目を利用する。これは、前処理後の電流や電圧、周波数などのテスト項目を合わせた数である。ウェハ X, Y, Z には、それぞれ数千個程度のダイが所属しており、検出対象とするフェイルはそれぞれ 15 個, 37 個, 21 個存在する。

表 5.1 実験データの概要

学習に用いるダイ数		ウェハ 2 枚分
検証に用いるダイ数		ウェハ 1 枚分
テスト項目数		771
ウェハのダイ数 (検出対象フェイルダイ数)	ウェハ X	1817(Fail : 15)
	ウェハ Y	2718(Fail : 37)
	ウェハ Z	3127 (Fail : 21)

5.2 実験結果

SVM によるフェイルダイ特性を考慮したテスト結果予測を行った結果を示す。実験は、以下の 3 通りの組み合わせで行い、それぞれの検証時の AUC によって予測モデルの精度を評価する。

1. ウェハ X, Y を用いて学習, ウェハ Z で検証
2. ウェハ X, Z を用いて学習, ウェハ Y で検証
3. ウェハ Y, Z を用いて学習, ウェハ X で検証

第 5.2.1 節では、k-means 法によってフェイル特性をクラスタリングを行った実験結果を述べる。また第 5.2.2 節では、フェイルダイの抽出を行った実験結果を述べる。最後に、第 5.2.3 節において、フェイルダイの抽出を行った後に、k-means 法によってフェイル特性をクラスタリングした実験結果を述べる。

5.2.1 k-means 法によるクラスタリングを用いた手法

第 4.2.1 節において示したように、k-means 法によるクラスタリングを用いてテスト結果予測を行った。表 5.2 にそれぞれのウェハの組み合わせごとに本手法の実験結果を示す。original は、フェイル特性を考慮せずに SVM による予測を行った際の AUC を表す。k-means Integration は、2 クラスのフェイルダイから作成した 2 つの予測モデルの結果を統合した際の AUC を表す。k-means class0, および class1 は、2 クラスのフェイルダイから作成した 2 つの予測モデルのうち、片方を用いた際の AUC を表し、括弧中には学習フェイルダイのうち、そのクラスに属するダイの割合を表す。

本研究で提案する k-means Integration は original と比較して、ウェハ X, Z で学習し、ウェハ Y を検証する際に限り、AUC が 0.18 向上したが、提案手法により、全体の予測精度が改善されたとは言えない。しかし、フェイルダイの片方のクラスから構成されるモデル、k-means class0, および class1 による予測では、学習ウェハ XY のとき 0.021, 学習ウェハ XZ のとき 0.070, 学習ウェハ YZ のとき 0.100 の AUC の改善が確認できる。この結果は、予測対象となるウェハの最終的なテスト結果がこの段階では不明であるため、実際の工程に適用することはできないが、学習

データをクラスタリングした片方のクラスは、似たフェイルダイ特性を持ったダイを高感度に予測している。

表 5.2 k-means 法によるクラスタリングを行った際の AUC(モデルの構築に用いたフェイルダイの割合)

	X, Y $\rightarrow$ Z	X, Z $\rightarrow$ Y	Y, Z $\rightarrow$ X
original	0.615	0.665	0.641
k-means Integration	0.612	0.682	0.630
k-means class0	0.636(63%)	0.736(56%)	0.479(57%)
k-means class1	0.529(37%)	0.473(44%)	0.741(43%)

5.2.2 フェイルダイの抽出を用いた手法

続いて、第 4.2.2 節において示した、予測の容易なダイの抽出し、テスト結果予測を行った結果を示す。まず、学習用ウェハの組み合わせごとに、閾値とする FPR を 0.8, 0.6, 0.4 と変化させ、代表的なフェイルダイ特性と異なる特性を持つダイを特定した。閾値ごとに予測できなかったフェイルダイを除いて抽出した代表的なフェイルダイ特性を持つダイを表 5.3, 表 5.4, 表 5.5 に示す。表はそれぞれの FPR を閾値とすると、予測できなかったフェイルダイ数を表す。また 15 回の試行時に、複数回に渡って予測できなかったフェイルダイは、その回数ごとに区別して除き、予測の容易なフェイルダイの抽出を行った。

表 5.3 閾値を 0.8 としたとき抽出したフェイルダイの割合

閾値内で予測できなかった回数	X, Y	X, Z	Y, Z
8 回以上			98%
4 回以上	98%		93%
3 回以上	96%	94%	
2 回以上	90%	75%	91%
1 回以上	76%	55%	75%

表 5.4 閾値を 0.6 としたとき抽出したフェイルダイの割合

閾値内で予測できなかった回数	X, Y	X, Z	Y, Z
8 回以上			98%
6 回以上	98%		96%
5 回以上	92%		93%
4 回以上	88%	89%	91%
3 回以上	86%	75%	87%
2 回以上	73%	55%	74%
1 回以上	56%	16%	35%

表 5.5 閾値を 0.4 としたとき抽出したフェイルダイの割合

閾値内で予測できなかった回数	X, Y	X, Z	Y, Z
7 回以上	98%		
6 回以上	92%	94%	93%
5 回以上	88%	89%	86%
4 回以上	71%	69%	78%
3 回以上	63%	39%	72%
2 回以上	48%	30%	55%
1 回以上	30%	8%	36%

表 5.3, 表 5.4, 表 5.5 に示した FPR の閾値と予測が困難であった回数ごとに抽出したフェイルダイ集合を用いてそれぞれ予測モデルを構築する。最後に、全学習フェイルダイから抽出したフェイルダイの割合を条件とし、予測モデルを統合する。条件ごとに統合した際の AUC を表 5.6 示す。表 5.6 において、original は、フェイル特性を考慮せずに SVM による予測を行った際の AUC を表す。以降は、全フェイルダイのうち学習に使ったダイの割合を条件に予測モデルの統合を行った際の AUC を表す。例えば、ウェハ XY において 90% の条件でモデルを統合する場合、表 5.3, 表 5.4, 表 5.5 における、フェイル 52 個中 46 個以上を学習に用いたモデルを統合した予測モデルによる AUC を表す。いずれのウェハの組み合わせ、モデル

統合条件においても全体的に original と比べて AUC が向上していることが確認できる. AUC は各ウェハの組み合わせにおいて, 学習ウェハ X, Y のとき 0.030, 学習ウェハ X, Z のとき 0.003, 学習ウェハ Y, Z のとき 0.025 改善した. AUC は, 抽出フェイルダイの割合が小さい予測モデルを統合した際に高い傾向がある. これは, 学習データにおいて予測が容易なフェイルダイに特化した予測モデルによって, 似た特性をもつフェイルダイのフェイル確度が上昇したためと考えられる.

表 5.6 モデルの統合条件と AUC

	X, Y(Fail: 52) $\rightarrow$ Z	X, Z(Fail: 36) $\rightarrow$ Y	Y, Z(Fail: 58) $\rightarrow$ X
original	0.615	0.665	0.641
90% 以上	0.631	0.661	0.646
80% 以上	0.638	0.663	0.659
70% 以上	0.641	0.665	0.653
60% 以上	0.641	0.665	0.653
50% 以上	0.643	0.668	0.662
40% 以上	0.644	0.668	0.662
30% 以上	0.645	0.663	0.666
20% 以上	0.645	0.663	0.666
10% 以上	0.645	0.663	0.666

5.2.3 両手法 (抽出後にクラスタリング) を適用した結果

最後に, 第 5.2.2 節において述べた, ダイの抽出を行った後に, 抽出したダイを k-means 法によってクラスタリングし, 予測モデルの統合を行った実験結果を表 5.7 に示す. 表 5.6 同様, 表 5.7 において, original は, フェイル特性を考慮せずに SVM による予測を行った際の AUC を表す. 以降は, 全フェイルダイのうち学習に使ったダイの割合を条件に予測モデルの統合を行った際の AUC を表す. 実験 5.5.2 同様, いずれのウェハの組み合わせ, モデル統合条件においても全体的に original と比べて AUC が向上していることが確認できる. AUC は最高で各ウェハの組み

合わせにおいて, 学習ウェハ X, Y のとき 0.048, 学習ウェハ X, Z のとき 0.022, 学習ウェハ Y, Z のとき 0.050 改善した. この結果は, 学習データにおいて予測が容易なフェイルダイを抽出, さらにクラスタリングすることで, 学習データにおいて予測が容易であったフェイルダイと似た特性をもつフェイルダイのフェイル確度が上昇したためと考えられる. 実験結果より, 実際の工程では, 代表的な特性をもつダイの選択をフェイルダイの 30% 以上を用いて予測モデルを構築すると, 安定した予測精度の向上を望むことができる.

表 5.7 モデルの統合条件と AUC

	X, Y(Fail: 52) $\rightarrow$ Z	X, Z(Fail: 36) $\rightarrow$ Y	Y, Z(Fail: 58) $\rightarrow$ X
original	0.615	0.665	0.641
k-means Integration	0.612	0.682	0.630
90% 以上	0.630	0.675	0.687
80% 以上	0.632	0.675	0.689
70% 以上	0.648	0.673	0.685
60% 以上	0.648	0.675	0.685
50% 以上	0.649	0.684	0.691
40% 以上	0.654	0.684	0.691
30% 以上	0.663	0.686	0.687
20% 以上	0.663	0.686	0.687
10% 以上	0.663	0.687	0.687

6. まとめと今後の課題

本論文では, LSI テストにおけるテスト結果予測のため, フェイルダイのもつ特性を利用することで, 予測精度を向上する方法について提案した. 学習データに含まれるフェイルダイと特性の似たフェイルダイは予測しやすいことを活かし, 学習データのフェイルダイをフェイルダイ特性によって選択した予測モデルを構築することで, 一部のフェイルダイの予測に特化した予測モデルを構築した. 本研究における提案手法である, k-means 法とフェイルダイ抽出を組み合わせた手法では, フェイル特性を考慮しない場合に比べて, 最高で AUC を 0.05 向上させることに成功した. 提案手法と, フェイル特性を考慮しない場合のテスト結果予測の ROC 曲線の一例を図 6.1 に示した. 図 6.1 に示すように序盤から中盤で予測されていたフェイルが, 提案手法によってより早期に予測できていることが確認できる.

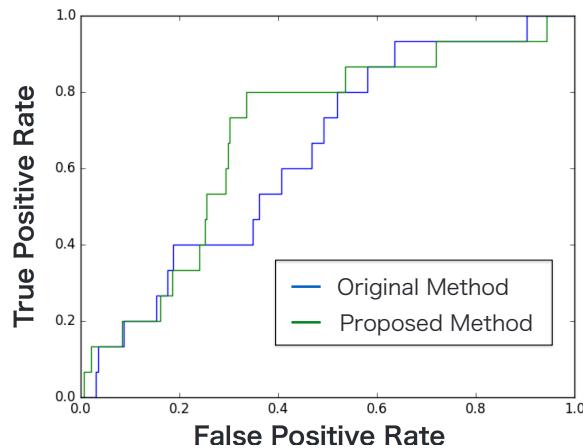


図 6.1 フェイル特性を考慮しない場合と提案手法を用いた場合の ROC 曲線

今後の課題は, 学習データにおいて予測難度の高いフェイルダイの予測と, ウェハ間のばらつきの補正を考えることである. まず, 学習データにおいて予測難度の高いフェイルダイの予測について, 本稿では一部のフェイルダイに特化した予測を行うことで予測精度の向上を行った. そのため, 学習データにおいて予測難度の高いフェイルダイのもつフェイル特性と似たフェイル特性をもつダイが, 未知データに多数存在した場合, それらを予測することは難しい.

次にウェハ間のばらつきの補正に関して, 実験では学習と検証に用いるウェハに

よって、予測の難易度や、提案手法を適用した際の効果が異なった。実際の工程では、同様の環境であっても、製造やテストの時期が異なるといった理由でよりウェハ間のばらつきは大きくなり、予測は困難になると考えられる。しかし、ウェハ間のばらつきを補正することができれば、提案手法による一定の効果が期待できると考えられる。

謝辞

本研究は奈良先端科学技術大学院大学にて、井上美智子教授の御指導の下で行ったものである。非常に熱心かつ丁寧にご指導賜ったことについて深く感謝いたします。また、本研究を行う上で丁寧に御助言くださった同大学院の池田和司教授、大下福仁准教授に心より感謝いたします。九州工業大学の梶原誠司教授、大分大学の大竹哲史准教授には、本研究について有益な議論や御助言をくださったことについて感謝いたします。本研究はルネサス セミコンダクタ パッケージ & テストソリューションズ株式会社の協力の元で行われたものである。本研究を進めるにあたって、内容について深く議論し、有益な御助言をくださった同社の中村芳行様 (現:ルネサスシステムデザイン) に感謝いたします。

参考文献

- [1] A.Righter, C.F.Hawkins, J.M.Soden, P.Maxwell, “CMOS IC Reliability Indicators and Burn-In Economics,” International Test Conference, 1998.
- [2] N.Sumikawa, L.Wang, M.S.Abadir, “An Experiment of Burn-in Time Reduction Based on Parametric Test Analysis, ” IEEE International Test Conference, 2012.
- [3] N.Sumikawa, J.Tikkanen, L.Wang, “Screening Customer Returns With Multivariate Test Analysis,” IEEE International Test Conference, 2012.
- [4] N.Sumikawa, D.Drmanac, L.Wang, L.Winemberg, M.S.Abadir, “Forward prediction based on wafer sort data,” IEEE International Test Conference, 2011.
- [5] Y.Fangming, Z.Zhaobo, C.Krishnendu, G.Xinli, “Board-Level Functional Fault Diagnosis Using Learning Based on Incremental SupportVector Machines, ” IEEE 21st Asian Test Symposium, 2012.
- [6] 野々山聡, 佐藤康夫, 梶原誠司, 中村芳行, “データマイニング手法によるバーンインテスト結果予測の検討,” 信学技報, vol.113, no.321, DC2013-57, 2013.
- [7] 小河亮, 中村芳行, 井上美智子, “LSI テストにおけるフェイル予測のためのばらつき補正に関する検討,” 研究報告システムと LSI の設計技術 (SLDM), 2016(46), 1-6.
- [8] 浅田邦彦, パワーデバイス・イネーブリング協会監修, “はかる × わかる半導体,” 日経 BP コンサルティング, Japan, 2013.

発表リスト

- [1] 佐藤 敬済, 中村 芳行, 井上 美智子, “フェイルダイ特性を考慮した LSI テスト結果予測の精度向上に関する研究,” 信学技報, 2017 年 3 月 (to appear).