

スペクトラムに基づく欠陥限局に適したプログラム構造の再調査

久保 光生[†] 肥後 芳樹[†] 楠本 真二[†]

[†] 大阪大学大学院情報科学研究科

E-mail: †{hkr-kubo,higo,kusumoto}@ist.osaka-u.ac.jp

あらまし プログラムに含まれる欠陥箇所を自動で推測する方法として、スペクトラムに基づく欠陥限局 (SBFL) がある。SBFL では、各テストケースの成否と実行経路情報をもとに欠陥箇所を特定する。同じ機能を持つプログラムでも、その構造によって SBFL の欠陥限局精度には差が生じる。よって、プログラム構造を SBFL に適する形に変更することで、機能を保ちつつ SBFL の精度向上が期待される。先行研究では SBFL に適するプログラム構造を発見するために、SBFL スコアが提案された。SBFL スコアはプログラムが SBFL にどの程度適しているかを評価する指標の 1 つである。また、先行研究では、同じ機能を持つが構造の異なるプログラムペアを対象として SBFL スコアを計測し、SBFL スコアが高くなるプログラム構造を得ている。しかし、実験対象のプログラム数が 10 個、ミューテーション演算子が 11 種類と少ないことが課題であった。そこで本研究では、実験対象のプログラム数を約 36 倍、ミューテーション演算子の種類数を約 2.5 倍にして実験を行った。実験の結果、新たに SBFL スコアを高めるプログラム構造が 3 つ特定できた。

キーワード SBFL, ミューテーションテスト, ソフトウェア品質モデル

1. はじめに

ソフトウェア開発において、デバッグは多くの労力とコストを必要とする作業である。ソフトウェア開発に必要なコストのうち、半分以上をデバッグ作業が占めているという報告もある [1]。そのため、デバッグを支援するための研究が盛んに行われている。デバッグ支援に関する研究分野の 1 つに欠陥限局がある。欠陥限局とは、プログラム中の欠陥箇所を推測する技術である。中でも近年、スペクトラムに基づく欠陥限局 (Spectrum-Based Fault Localization, 以降 SBFL) に関する研究が盛んに行われている [2]。SBFL では、テストがどの文を実行し、どの文を実行しなかったかという情報を用いて、プログラムの欠陥と疑われる箇所を自動的に特定する。SBFL の基本的なアイデアは、多くの失敗テストで実行される文ほど欠陥である可能性が高く、成功テストで多く実行される文ほど欠陥である可能性が低いと判断することである。

SBFL の精度は、欠陥自体の性質やテストの内容など、様々な要因に左右される [3]。その中でも、佐々木らによる先行研究 [4] では、プログラム構造に着目した。この先行研究では、プログラムがどの程度 SBFL に適するかを表す SBFL 適合性が提案されている。また、SBFL 適合性の評価指標として、SBFL スコアが提案されている。SBFL スコア計測の基本的なアイデアは、すべてのテストを通過するプログラムについて、ミューテーションテストを用いて、様々な箇所に意図的に欠陥を発生させることである。これらの欠陥が SBFL によってどの程度正確に特定できたか計測することにより、元のプログラムの SBFL 適合性を評価できる。先行研究では同じ機能を持つが構造の異なる 5 組のプログラムペアを用いて SBFL スコアを計測することにより、SBFL に適するプログラム構造が調査された。しかし、計測対象のプログラム数が 10 個と少なく、SBFL に適するプログラム構造が十分に明らかになっ

たとはいえない。また、プログラムに意図的に欠陥を発生させるためのミューテーション演算子の種類が 11 個と少なく、SBFL スコアの信頼性が低いことが課題であった。

そこで本研究では、365 個のプログラムを用いて SBFL スコアを計測することにより、SBFL に適する新たなプログラム構造の発見を目指す。また、SBFL スコアの信頼性を高めるために、16 種類の新たなミューテーション演算子を定義する。これにより、合計で 27 種類のミューテーション演算子をプログラムに適用する。

2. 準備

2.1 スペクトラムに基づく欠陥限局 (SBFL)

デバッグを支援する技術の 1 つに、欠陥限局がある。欠陥限局とは、プログラム中の欠陥箇所を推測する技術である。テストを用いた自動的な欠陥限局方法の 1 つに、スペクトラムに基づく欠陥限局 (Spectrum-Based Fault Localization, SBFL) がある。SBFL におけるスペクトラムとは、テスト実行時にどの文が実行されたかという実行経路情報である。SBFL の基本的なアイデアは、多くの失敗テストで実行される文ほど欠陥である可能性が高く、成功テストで多く実行される文ほど欠陥である可能性が低いと判断することである。

SBFL による欠陥箇所の特定方法について説明する。まず、すべてのテストを実行し、テストの成否と実行経路情報を記録する。次に、これらの情報を利用して、疑惑値と呼ばれる、欠陥である可能性の高さを示す値を文ごとに算出する。疑惑値の算出方法には様々なあるが、Abreu らが SBFL で用いられる計算式の有効性を評価した結果、Ochiai の計算式 [5] が最も優れていると結論づけている [6]。

Ochiai の計算式における疑惑値 $susp(s)$ の算出方法を (1) に示す。ここで、 $fail(s)$ は文 s を実行した失敗テストの数、 $pass(s)$

は文 s を実行した成功テストの数, $totalFail$ はすべての失敗テストの総数である.

$$susp(s) = \frac{fail(s)}{\sqrt{totalFail \times (fail(s) + pass(s))}} \quad (1)$$

この $susp(s)$ をすべての文 s に対して算出し, その値が高い文ほど, 欠陥の原因箇所である可能性が高いと推測する.

なお, SBFL の疑惑値の算出においては, 失敗テストの実行経路情報が最も重要な要素となる. 失敗テストでどのような文が実行され, どのような文が実行されなかったかが, テストの失敗, すなわち欠陥の原因箇所に対する大きな手がかりとなるためである. Ochiai の計算式では, 分子が $fail(s)$ であることから, 失敗テストの実行経路情報の重要性が見て取れる.

2.2 SBFL 適合性

先行研究 [4] で提案された SBFL 適合性について述べる. SBFL 適合性とは, プログラムが持つ品質特性の 1 つであり, プログラム自体が SBFL にどの程度適しているかを表す. プログラムの機能やテストスイートが同じでも, プログラム構造が異なれば SBFL を用いた欠陥限局の精度に違いが生じることがある.

プログラム構造の違いによる SBFL 適合性の変化について, 例を用いて説明する. 図 1 に示すプログラム (a) とプログラム (b) は, 機能が同じだが, 構造が異なる. 図 1 に示すテスト (c) を用いて両方のプログラムに SBFL を実行すると, 各文の疑惑値が算出される. プログラム (a) では, 欠陥がある箇所と同じ疑惑値を持つ文が 4 個存在する. 一方, プログラム (b) では, 欠陥がある箇所と同じ疑惑値を持つ文は 1 個である. 欠陥がある箇所と同じ疑惑値を持つ文が少ないほど, 確認しなくてはならない文が少なくなるため, SBFL による欠陥限局の精度が高いといえる. よって, この場合, (b) の方が SBFL 適合性が高い.

2.3 SBFL スコア

SBFL 適合性と同時に先行研究 [4] で提案された SBFL スコアについて述べる. SBFL スコアは, SBFL 適合性の評価指標の 1 つとして提案された. SBFL スコアを計測するための基本的なアイデアは, 与えられたプログラムに意図的に変更を加え, 人工的な欠陥を生成することである. これらの人工的な欠陥を SBFL でどの程度正確に特定できるかを計測することにより, プログラム全体がどの程度 SBFL に適するかを測定できる.

2.3.1 SBFL スコアの計測方法

先行研究 [4] で提案された, SBFL スコアの計測方法の概要を示す. プログラム P の SBFL スコア計測時の入力には以下の 2 つである.

- ミュータント生成器 G
- テストスイート T

計測の大まかな流れを図 2 に示す. SBFL スコアの計測は以下の 3 つのステップで構成される.

- (1) プログラムに対してミュータントを生成
- (2) ミュータントに対して SBFL を実行
- (3) SBFL スコアを算出

以降では, 各ステップの詳細を述べる.

プログラム (入力: a, b)	susp	t ₁	t ₂	t ₃	t ₄
s ₁ : boolean result = false;	0.50	✓	✓	✓	✓
s ₂ : if (0 < a)	0.50	✓	✓	✓	✓
s ₃ : result = true;	0.00	✓	✓		
s ₄ : if (0 <= b) //correct: 0 < b	0.50	✓	✓	✓	✓
s ₅ : result = true;	0.50	✓	✓	✓	✓
s ₆ : return result;	0.50	✓	✓	✓	✓
テスト結果: P P P F					

(a) リファクタリング前

プログラム (入力: a, b)	susp	t ₁	t ₂	t ₃	t ₄
s' ₂ : if (0 < a)	0.50	✓	✓	✓	✓
s' ₃ : return true;	0.00	✓	✓		
s' ₄ : if (0 <= b) //correct: 0 < b	0.71			✓	✓
s' ₅ : return true;	0.71			✓	✓
s' ₆ : return false;	-				
テスト結果: P P P F					

(b) リファクタリング後

テストケース	入力 (a, b)	期待値	実際の値
t ₁ :	(1, 1)	true	true
t ₂ :	(1, 0)	true	true
t ₃ :	(0, 1)	true	true
t ₄ :	(0, 0)	false	true

(c) テストスイート

図 1 プログラム構造の違いにより SBFL 適合性が変化する例

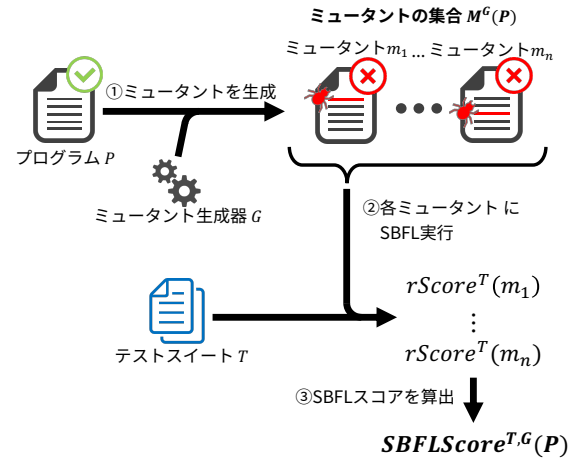


図 2 SBFL スコアの計測方法の概要

Step1. プログラムに対してミュータントを生成

対象プログラムとテストスイートを用意する. プログラムはすべてのテストに通過することを前提とする. このプログラムに対して, ミュータント生成器を用いてミュータントを生成する. このとき, 各ミュータントは元のプログラムに対して 1 箇所だけが変更されるように生成する. プログラムには変更可能な箇所が複数存在するため, ミュータントは複数生成される. ミュータント生成器 G によって生成されたミュータントの集合を $M^G(P)$ と定義する.

Step2. ミュータントに対して SBFL を実行

$M^G(P)$ に含まれる各ミュータントに対して, テストスイート T を用いて SBFL を実行し, 文ごとの疑惑値を算出する. ここで, あるミュータント $m \in M^G(P)$ に含まれる各文 s について, 以下を定義する.

- $susp^T(s)$: 文 s の疑惑値

- $rank^T(s)$: 文 s の疑惑値の順位
- $rScore^T(s)$: 文 s の疑惑値の正規化順位

疑惑値の算出には Ochiai の計算式を用いる。疑惑値の順位は、疑惑値の高い順に文を並べた際に、欠陥が存在する文を発見するまでに最大で確認しなければならない文の総数とする。例えば、疑惑値 1.0 の文が 2 つ、0.8 の文が 1 つ存在する場合は、疑惑値 1.0 の文はどちらも 2 位と扱い、疑惑値 0.8 の文は 3 位とする。疑惑値の順位は、文の総数により異なる価値を持つ。例えば、10 個の文のうちの 10 位と、100 個の文のうちの 10 位とでは、後者の方が順位としての価値が高い。そこで、各文が文全体の中でどの程度上位に位置するかを表すため、順位を 0 以上 1 以下の範囲で線形に正規化する。文 s の正規化順位 $rScore^T(s)$ を以下の (2) の通り算出する。1 が最も価値が高く、0 が最も価値が低いことを表す。(2) において、 $totalStatements^T$ はテストスイート T によって実行される文の数である。

$$rScore^T(s) = 1 - \frac{rank^T(s) - 1}{totalStatements^T - 1} \quad (2)$$

また、ミュータント m に含まれる欠陥の疑惑値の正規化順位を $rScore^T(m)$ と定義する。各ミュータントに含まれる欠陥は 1 箇所であるため、この値はミュータントごとに一意である。ミュータント m の欠陥を含む文を s_{fault}^m とすると、 $rScore^T(m)$ は文 s_{fault}^m の疑惑値の正規化順位である。ミュータントの $rScore$ が高いほど、欠陥箇所を正確に特定できることを意味する。

$$rScore^T(m) = rScore^T(s_{fault}^m)$$

Step3. SBFL スコアを算出

対象プログラムから生成された各ミュータントの $rScore$ の平均をとることにより、SBFL スコアが算出される。なお、 $|M^G(P)|$ は、生成されたミュータントの総数を表す。

$$SBFLScore^{T,G}(P) = \frac{1}{|M^G(P)|} \sum_{m \in M^G(P)} rScore^T(m)$$

2.4 先行研究の調査方法

先行研究 [4] では、Java で実装されたプログラムを対象として、プログラム構造の違いにより SBFL スコアがどのように変化するか確認するための実験が行われた。

2.4.1 実験対象プログラムとテストスイート

先行研究 [4] では、単一のメソッドで構成されるプログラムに対して SBFL スコアを計測する実験が行われた。実験対象は、同じ入力を与えられると同じ出力を返すが、構造の異なる 5 種類のプログラムペアである。以降、同じ入力を与えられると同じ出力を返すが、構造の異なるメソッドを同機能異構造メソッドと呼ぶ。5 種類のプログラムペアは、佐々木らによってリファクタリングを題材として作成された。リファクタリングには多くのパターンが存在するが、Fowler によって「条件記述の単純化」に分類されたリファクタリングパターン [7] の中から選定した単一のメソッド内で完結する 5 種類のリファクタリングが題材となった。

また、各プログラムのテストスイートは佐々木らによって、以下を満たすように作成された。

- すべてのミュータントはいずれかのテストに失敗する。
- 条件網羅率が 100% になる。

2.4.2 ミュータント生成器

先行研究 [4] では、表 1 に示す 11 種類のミューテーション演算子が用いられた。これらは、オープンソースのミューテーションテストツールである PIT [8] の基本ミューテーション演算子を参考に実装された。PIT は、ミューテーションテストの分野でミュータントの生成に広く用いられている [9]。

2.5 先行研究の調査結果

先行研究 [4] での調査の結果、同一条件分岐先で実行される文が少ないほど SBFL 適合性が高いことが明らかになった。また、SBFL 適合性を高めるプログラム構造の変換方法として、以下の 2 つが得られている。

- return 文を用いて条件分岐を早期終了させる
- 同一の条件分岐先の文の数が多い方から少ない方へ文を移動する

3. 先行研究の課題とその解決策

先行研究 [4] には、以下の 2 つの課題が存在する。

- ミューテーション演算子の種類が少ない。
- 調査対象プログラム数が少ない。

以下では、課題の詳細と、その解決策について述べる。

3.1 ミューテーション演算子の種類が少ない

先行研究で使用されたミューテーション演算子は 11 種類であった。先行研究で実験対象となったプログラムは、これらのミューテーション演算子が多くの箇所に適用できるように実装されたため、十分な量のミュータントが生成できていた。しかし、それ以外のプログラムを用いる場合、十分な量のミュータントを生成できない場合がある。この状態で計測された SBFL スコアは、プログラム全体の SBFL 適合性を表さないため、信頼性が低い。

この問題の解決のために、本研究では、新たに表 2 に示す 16 種類のミューテーション演算子を実装した。なお、新たに追加したミューテーション演算子は、既存のミューテーション演算子と同じミュータントを生成する可能性がある。例えば、プログラムに

表 1 先行研究で用いられたミューテーション演算子

	変換前	変換後
Conditional Boundary	<code>a<b</code>	<code>a<=b</code>
Increments	<code>n++</code>	<code>n-</code>
Invert Negatives	<code>-n</code>	<code>n</code>
Math	<code>a+b</code>	<code>a-b</code>
Negate Conditionals	<code>a==b</code>	<code>a!=b</code>
Void Method Calls	<code>method();</code>	<code>;</code>
Primitive Returns	<code>return 5;</code>	<code>return 0;</code>
Empty Returns	<code>return "str";</code>	<code>return "";</code>
False Returns	<code>return true;</code>	<code>return false;</code>
True Returns	<code>return false;</code>	<code>return true;</code>
Null Returns	<code>return object;</code>	<code>return null;</code>

プログラム (入力: a)	t_1	t_2	t_3	プログラム (入力: a)	t_1	t_2	t_3
s_1 : if (a == 0)	✓	✓	✓	s'_1 : if (a > 0)	✓	✓	✓
s_2 : return 0;			✓	s'_2 : return 1;		✓	✓
s_3 : else if (a > 0)	✓	✓		s'_3 : else if (a < 0)		✓	✓
s_4 : return 1;	✓			s'_4 : return -1;		✓	
s_5 : return -1;		✓	✓	s'_5 : return 0;			✓

図 3 実行経路が異なるプログラムペアの例

プログラム (入力: a,b)	t_1	t_2	t_3	プログラム (入力: a,b)	t_1	t_2	t_3
s_1 : if (a > 0)	✓	✓	✓	s'_1 : if (a > 0)	✓	✓	✓
s_2 : return 1;	✓			s'_2 : return 1;		✓	
s_3 : if (b < 0 && a != 0)	✓	✓		s'_3 : else if (a < 0 && b < 0)		✓	✓
s_4 : return -1;		✓		s'_4 : return -1;		✓	
s_5 : return 0;			✓	s'_5 : return 0;			✓

図 4 実行経路が同じであるプログラムペアの例

return false; という文がある場合、Change Boolean Literal と、False Return が同じミュータントを生成する。よって、同じミュータントは複数生成されないようにミュータント生成器を実装した。

また、ミュータントの増加数を検証するために、ミューテーション演算子の追加前と追加後のそれぞれについて、実際のプログラムで生成されたミュータントの個数を比較する。このとき、プログラムの規模を考慮する必要がある。なぜなら、プログラムの規模が大きくなると、ミュータントの個数は増加する傾向にあるからである。規模に対するミュータントの増加数を評価するための指標として、IMPL (Increased Mutants Per LOC) を求める。IMPL は式 (3) で定義される。なお、論理 LOC とは、プログラムのうち、空行や括弧のみの行、コメントのみの行を除いた行数を指す。IMPL の値が大きいほど、SBFL スコアの精度が高まったことを表す。

$$IMPL = \frac{\text{ミュータントの増加数}}{\text{論理 LOC}} \quad (3)$$

表 2 本研究で新たに用いるミューテーション演算子

ミューテーション演算子	変換前	変換後
Change String Literal	"String"	"String1"
Change Instanceof	a instanceof A	a instanceof B
Nonvoid Method Calls	a = method();	a = null;
Constructor Calls	a = new A();	a = null;
Compund Operator	a += 1	a -= 1
Change Numeric Literal	if(x<0)	if(x<1)
Change Boolean Literal	true	false
Change Unary Operator	!ismethod()	ismethod()
Add Not Operator	if(b)	if(!b)
More Specific If	a&& b	a, b, a b
Less Specific If	a b	a, b, a&&b
Break and Continue	break;	continue;
Null Assignment	o1 = o2;	o1 = null;
Empty Assignment	s = a.toString();	s = "";
Primitive Assignment	int a = b;	int a = 0;
Change Throw Exception	Throw new A;	Throw new B;

3.2 実験対象プログラム数が少ない

実験対象プログラムは 5 組のプログラムペアに含まれる 10 個のプログラムであった。よって、調査結果の一般化可能性に欠ける。また、SBFL スコアが高くなるプログラム構造を十分に調査しきれていない。

この問題の解決のために、本研究では、Java で実装された同機能メソッドを含むデータセット [10] を用いる。本データセットは、単一のメソッドからなるプログラムを 728 個含む。これらはオープンソースソフトウェアから収集された。また、メソッドは機能によって 276 組のグループに分類されている。各グループは、同機能メソッドを 2 個から 12 個含む^(注1)。各プログラムのテストスイートは Evosuite [11] によって自動生成された。

なお、本研究で用いるデータセットの中には、記述方法は異なるが、テストスイートを実行した際の実行経路が同じになるグループが含まれる可能性がある。「実行経路が同じ」とは、2 つのプログラムが以下の 2 つの条件を満たすことを指す。ここで、 s_i は 1 目目のプログラムの上から i 番目の文を、 s'_i は 2 目目のプログラムの上から i 番目の文を表す。

- テストスイートで実行される文の総数 m が同じである。
- テストスイートに含まれるすべてのテストにおいて、 $1 \leq i \leq m$ について、文 s_i における実行の有無が、文 s'_i における実行の有無と等しい。

実行経路が異なるプログラムペアの例を図 3 に、実行経路が同じであるプログラムペアの例を図 4 に示す。2 つのプログラムの実行経路が同じになる場合、適用されるミューテーション演算子の違いのみによって、SBFL スコアに差が生じてしまう。このようなプログラムは実験対象として適切でないと考えたため、実験対象から除外する必要がある。よって、本研究においては、プログラム構造を「テストスイートを実行した際の実行経路」と定義する。記述方法が異なっても、テストスイートを実行した際の実行経路が同じ場合には、プログラム構造は同じであるとみなす。

4. 実験

実験では、どのようなプログラム構造が SBFL に適するか明らかにする。

4.1 実験対象プログラムとテストスイート

実験対象プログラムは、3.2 節で示したデータセットに含まれる 728 個のプログラムのうち、以下の 3 つの条件を満たす 133 組のグループに含まれる 365 個のプログラムである。

- グループ内のプログラムをテストスイートに通したときの条件網羅率が 100% になる。
- グループ内のプログラムが同機能である。
- グループ内のプログラム間で、プログラム構造が互いに異なる。

条件の確認は以下の 2 つのステップで構成される。

(1) テストスイートを手動で修正

(2) 条件を満たさないプログラムを除外

以下、これらの手順の詳細を述べる。

(注1)：実際には、同機能でないメソッドが存在する。

Step1. テストスイートを手動で修正

各プログラムのテストスイートは Evosuite [11] によって自動生成された。自動生成されたテストスイートでは、条件網羅率が 100% にならない場合がある。よって、テストスイートを手動で修正し、条件網羅率が 100% になるようにした。

Step2. 条件を満たさないプログラムを除外

すべてのプログラムをテストスイートに通し、網羅率や等価性、実行経路情報を確認した。その結果、合計で 143 組のグループと 363 個のプログラムが実験対象から除外された。除外された要因とそれに当てはまるプログラムの数を表 3 に示す。

4.2 実験方法

実験は以下の 3 つのステップで構成される。

(1) 実験対象プログラムの SBFL スコアを計測

(2) SBFL スコアに差が生じないグループを考察対象から除外

(3) SBFL スコアが高くなる要因を考察

以下では、これらの詳細を述べる。

Step1. 実験対象プログラムの SBFL スコアを計測

すべての実験対象プログラムについて、計測ツールを用いて SBFL スコアを計測する。計測ツールは SBFL スコアを計測するために、自動欠陥修正ツールである kGenProg [12] を用いて実行経路情報を取得する。kGenProg はプラグインとして JaCoCo^(注2) を呼び出し、実行経路情報を得ている。各ミュータントに対して SBFL を実行した際の実行経路情報はファイルに保存され、後から確認できる。

また、計測ツール内のミュータント生成器に、2.4.2 節で示した 11 種類のミューテーション演算子と、3.1 節で示した 16 種類のミューテーション演算子を実装した。さらに、使用するミューテーション演算子を切り替えられるようにした。これにより、演算子の追加前後におけるミュータント数の変化を記録できる。

Step2. SBFL スコアに差が生じないグループを考察対象から除外

プログラム構造が異なっても、SBFL スコアが同じになる場合が考えられる。この場合、SBFL スコアが高くなる要因を考察することはできない。よって、プログラム構造が異なるグループのうち、SBFL スコアに差が生じないグループを考察対象から除外する。

Step3. SBFL スコアが高くなる要因を考察

各ミュータントに対して SBFL を実行した際の実行経路情報を目視で確認し、SBFL スコアが高くなる要因を考察する。

4.3 先行研究からの改善点の評価

ミュータントの増加数

各プログラムについて IMPL を計測した結果を図 5 に示す。IMPL が 0 より大きいプログラムは 365 個中 358 個であった。よって、ほとんどのプログラムにおいて、SBFL スコアの信頼性は

表 3 条件にあてはまらないプログラム数

条件	プログラム数
条件網羅率が 100% にならない	18
グループ内のプログラムが等価でない	90
グループ内のプログラムの実行経路が同じである	255

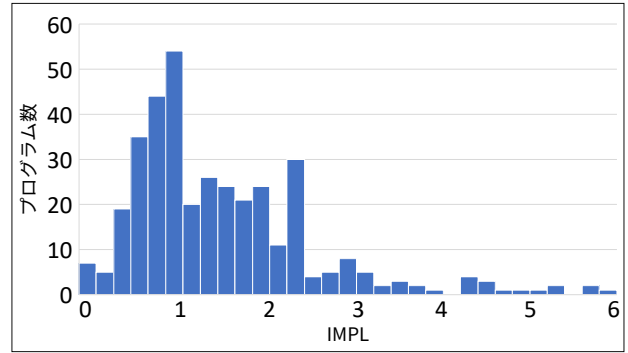


図 5 各プログラムの IMPL の分布

表 4 SBFL スコアが高くなる要因とグループ数

要因	グループ数
分岐がない処理の代わりに if 文や for 文を用いている	21
early return のための if 文を追加している	15
if 文の条件式に論理演算子を用いず、別の文に分けている	12
同一条件分岐先で実行される文が少ない	62
不明	10

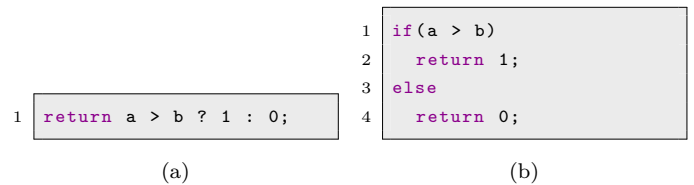


図 6 if 文や for 文の有無を示すプログラムペアの例

向上した。

考察対象プログラムの増加数

実験対象である 133 組のグループのうち、120 組のグループで SBFL スコアに変化が生じた。よって、考察対象プログラム数は先行研究から大幅に増加した。

SBFL スコアに変化が生じないグループは 11 組存在した。変化が生じない理由は、以下の 2 つであった。

- グループ内のいずれのプログラムにも分岐が全く存在しない。
- 実行経路が入れ替わっているのみで、実行される文の数は同じである。

4.4 実験結果と考察

SBFL スコアが高くなる要因として、新たに以下の 3 つが特定できた。

- 分岐がない処理の代わりに if 文や for 文を用いている
- early return [13] のための if 文を追加している
- if 文の条件式に論理演算子を用いず、別の文に分けている

また、先行研究で発見された要因である「同一条件分岐先で実行される文が少ない」に当てはまるグループも存在した。SBFL スコアが高くなる要因とそれに当てはまるグループの数を表 4 に示す。

以下では、新たに発見した要因について、具体例を用いて説明する。

分岐がない処理の代わりに if 文や for 文を用いている

図 6 の (b) のプログラムでは、if 文が用いられている。よって、テストを実行した際の実行経路に差が生じるため、SBFL スコア

(注2) : <https://www.jacoco.org/jacoco/>

<pre> 1 int s = 0; 2 for (int i=0; i<1; i++) 3 s += p[i]; 4 return s; </pre>	<pre> 1 if (1 < 1) 2 return 0; 3 int s = 0; 4 for (int i=0; i<1; i++) 5 s += p[i]; 6 return s; </pre>
(a)	(b)

図 7 if 文による early return の有無を示すプログラムペアの例

<pre> 1 if (i < s i > e) { 2 return -1; 3 return 0; </pre>	<pre> 1 if (i < s) 2 return -1; 3 else if (i > e) 4 return -1; 5 return 0; </pre>
(a)	(b)

図 8 論理演算子の有無を示すプログラムペアの例

が 0 よりも大きくなる。一方、(a) のプログラムでは、三項演算子が用いられており、if 文と for 文が存在しない。if 文や for 文がプログラム中になく、各文の疑惑値に差が生じないため、各ミュータントの $rScore$ が 0 になり、SBFL スコアも 0 になってしまう。よって、分岐がない処理の代わりに if 文や for 文を用いることで、SBFL スコアは向上する。

early return のための if 文を追加している

プログラムの入力によっては、長い処理を行う前に出力が分かる場合がある。例えば、図 7 の (a) では、配列の内容の処理に for 文が用いられている。一方、(b) では if 文を追加することにより、配列の長さが 0 である場合には for 文による処理を行わずに 0 を返す。if 文がない場合、テストスイートの実行経路に差が生じにくい。一方で、if 文による early return を追加することにより、テストスイートの実行経路を分散させることができる。これにより、その後続く文に欠陥がある場合に疑惑値を上昇させられるため、SBFL スコアが向上する。

if 文の条件式に論理演算子を用いず、別の文に分けている

図 8 の (a) のプログラムでは、条件式を || でつなぎ、1 つの if 文を用いている。一方、(b) のプログラムでは、条件式を分け、2 つの if 文を用いている。条件式を || でつないだ場合、条件式に欠陥がある際に、その文が成功テストで多く実行される。成功テストで多く実行される文の疑惑値は低下するため、他の文の疑惑値の順位が上昇する。その結果、 $rScore$ が低下し、SBFL スコアも低下する。反対に、条件式を別の if 文に分けると、SBFL スコアは向上する。

5. 妥当性への脅威

SBFL スコアの計測結果は、テストスイート及びミュータント生成器に影響を受ける。よって、異なるテストスイート及びミュータント生成器を使用すると、SBFL スコアの値が変化し、異なる結果が現れる可能性がある。

また、本研究とは異なるプログラムを実験対象とすると、今回の実験で得られた結果を否定する事例が現れる可能性がある。

6. おわりに

本研究では、先行研究の課題であるミューテーション演算子の種類の少なさと実験対象プログラム数の少なさを解決した。また、SBFL スコアの計測実験を行い、SBFL 適合性が高いプログラム構造を調査した。調査の結果、新たに 3 つの SBFL 適合性が高いプログラム構造が得られた。

今後の課題としては、SBFL 適合性を高める変換方法に基づいてプログラムを自動変換するツールの作成が挙げられる。また、SBFL 適合性が高くなる一方で、プログラムの保守性が低下する場合がある。よって、SBFL 適合性と他の品質特性との関連を明らかにすることは今後の重要な課題である。

謝辞 本研究は JSPS 科研費 (JP20H04166, JP21K18302, JP21K11820, JP21H04877, JP22H03567, JP22K11985) の助成を得て行われた。

文 献

- [1] B. Hailpern and P. Santhanam, “Software debugging, testing, and verification,” IBM Systems Journal, vol.41, no.1, pp.4–12, 2002.
- [2] W.E. Wong, R. Gao, Y. Li, R. Abreu, and F. Wotawa, “A survey on software fault localization,” IEEE Transactions on Software Engineering, vol.42, no.8, pp.707–740, 2016.
- [3] S. Ali, J.H. Andrews, T. Dhandapani, and W. Wang, “Evaluating the accuracy of fault localization techniques,” Proc. International Conference on Automated Software Engineering, pp.76–87, 2009.
- [4] 佐々木唯, 肥後芳樹, 松本真佑, 楠本真二, “プログラムに対する欠陥限局の適合性計測,” 情報処理学会論文誌, vol.62, no.4, pp.1029–1038, 2021.
- [5] R. Abreu, P. Zoetewij, and A.J. Van Gemund, “An Evaluation of Similarity Coefficients for Software Fault Localization,” Proc. Pacific Rim International Symposium on Dependable Computing, pp.39–46, 2006.
- [6] R. Abreu, P. Zoetewij, R. Golsteijn, and A.J.C. vanGemund, “A Practical Evaluation of Spectrum-based Fault Localization,” Journal of Systems and Software, vol.82, no.11, pp.1780–1792, 2009.
- [7] F. Martin, Refactoring: Improving the Design of Existing Code, Addison-Wesley, 1999.
- [8] H. Coles, T. Laurent, C. Henard, M. Papadakis, and A. Ventrresque, “PIT: A Practical Mutation Testing Tool for Java (Demo),” Proc. International Symposium on Software Testing and Analysis, pp.449–452, 2016.
- [9] J. Xuan and M. Monperrus, “Test case purification for improving fault localization,” Proc. International Symposium on Foundations of Software Engineering, pp.52–63, 2014.
- [10] Y. Higo, S. Matsumoto, S. Kusumoto, and K. Yasuda, “Constructing dataset of functionally equivalent java methods,” Proc. the International Conference on Mining Software Repositories, pp.682–686, 2022.
- [11] G. Fraser and A. Arcuri, “EvoSuite: automatic test suite generation for object-oriented software,” Proc. Symposium and the European Conference on Foundations of Software Engineering, pp.416–419, 2011.
- [12] Y. Higo, S. Matsumoto, R. Arima, A. Tanikado, K. Naitou, J. Matsumoto, Y. Tomida, and S. Kusumoto, “kgenprog: A high-performance, high-extensibility and high-portability apr system,” Proc. the Asia-Pacific Software Engineering Conference, pp.697–698, 2018.
- [13] M.A. Saca, “Refactoring improving the design of existing code,” Proc. IEEE Central America and Panama Convention, pp.1–3, 2017.