

修士学位論文

題目

複数の工数予測結果の期待される精度比較による予測値決定支援

指導教員

楠本 真二 教授

報告者

生方 克馬

平成 23 年 2 月 7 日

大阪大学 大学院情報科学研究科
コンピュータサイエンス専攻

内容梗概

ソフトウェア開発において、適切なプロジェクト計画を立案し、プロジェクトを円滑に遂行するためには工数予測が重要である。正確に工数を知ることにより、適切なスケジュール管理や開発資源の割当てが可能となり、プロジェクト失敗のリスクを抑えることができる。そのため、現在までに様々な工数予測手法が提案されており、近年では過去のプロジェクトの実績データをもとに工数を予測する工数予測手法が注目されている。しかしながら、実績データに基づく工数予測手法の予測精度は使用するデータセット（過去のプロジェクトの集合）の性質に大きく影響され、ひとつの工数予測手法で常に高い精度を得ることは難しい。そこで本研究では、1 件の予測対象のプロジェクトに対して複数の工数予測手法を適用し、データセットに応じてその都度工数予測手法を選択する方法を提案する。工数予測手法の選択基準には、**Leave-One-Out** を用いてデータセット中のプロジェクトの工数を予測したときの誤差をもとに算出した、期待される精度を用いる。また、複数の工数予測手法を同時に適用可能な工数予測ツールが公開されていないことから、提案手法を実現する工数予測ツール「 e^3 」の開発を行った。 e^3 は、1 件の予測対象のプロジェクトに対して複数の工数予測手法を同時に適用し、各工数予測手法による工数の予測値と期待される精度を表示し、状況に応じた予測値の決定を支援するツールである。これにより、特定の工数予測手法にとらわれない、柔軟な工数予測が可能となる。提案手法の効果を確認する実験としては、**ISBSG** のデータセットとその他 2 社のデータセットを用い、特定の工数予測手法を使用し続けた場合と、ツールの表示する期待される精度の最も高い工数予測手法をその都度選択して使用した場合とで精度の比較を行った。その結果、特定の工数予測手法を使用し続けた場合はデータセットと工数予測手法の組合せによって精度にばらつきがあったが、期待される精度の最も高い工数予測手法をその都度選択して使用した場合は、データセットごとに使用し続けたときに最も精度の高かった工数予測手法とほぼ同程度の精度で予測が可能であることがわかった。

主な用語

工数予測, 予測精度, ソフトウェア開発, **Leave-One-Out**

目次

1	まえがき	1
2	関連研究	3
2.1	工数予測とデータセットの関係	3
2.2	工数予測ツール	3
3	工数予測手法	5
3.1	類似性に基づく手法	5
3.2	回帰分析	6
3.2.1	直線単回帰分析	6
3.2.2	累乗単回帰分析	7
3.2.3	重回帰分析	7
4	提案手法	9
4.1	概要	9
4.2	期待される精度	9
4.3	工数予測手法の選択	10
5	工数予測ツール「e^3」	12
5.1	概要	12
5.2	入出力	13
5.2.1	入力画面	13
5.2.2	プロジェクトデータファイル	14
5.2.3	プロジェクトデータの編集	15
5.2.4	出力画面	15
5.3	実装済みの工数予測手法	16
6	実験	18
6.1	概要	18
6.2	使用したデータセット	18
6.3	手順	19
6.4	結果	21
7	考察	24
7.1	実験結果の考察	24
7.2	ツールの評価	25

8	あとがき	28
	謝辞	29
	参考文献	30
A	工数予測手法の追加	33
A.1	概要	33
A.2	命名規則	33
A.3	ファイルの配置	33
A.4	メソッドの引数	35
A.5	記述例	37

図目次

1	LOO を用いた過去プロジェクトの工数予測	9
2	ツールの構成図	12
3	入力画面	13
4	CSV ファイルを Excel 等で開いたイメージ	14
5	出力画面	16
6	ISBSG : MMRE の比較	22
7	ISBSG : MMER の比較	22
8	A 社 : MMRE の比較	22
9	A 社 : MMER の比較	23
10	B 社 : MMRE の比較	23
11	B 社 : MMER の比較	23
12	ISBSG : 各試行で最小の誤差で予測が行えた工数予測手法の割合	25
13	ファイルの配置場所と記述例	34
14	CSV ファイルの例	35
15	メトリクス情報の配列	36
16	プロジェクト名情報の配列	36
17	プロジェクト情報の配列	36

表 目 次

1	ISBSG：使用メトリクスと欠損値の割合	19
2	実験で収集されるデータ	20
3	選択された工数予測手法の割合	25
4	ツールに関するアンケート結果	26
5	ツールの実行時間	26

1 まえがき

ソフトウェア開発において、適切なプロジェクト計画を立案し、プロジェクトを円滑に遂行するためには工数予測が重要である。工数とはプロジェクトの要員と工期の積によって表される数値であり、スケジュール管理や開発資源の割当てなど、プロジェクトマネジメントの重要な参考資料となる。正確に工数を予測することで、結果として納期遅れやコスト超過といったプロジェクトの失敗のリスクを抑えることが可能である。そのため、現在までに様々な工数予測手法が提案され [1][2][3][4], 予測精度の向上が図られている。

工数予測手法は、大きく、パラメトリック法、積上げ法、類推法の 3 種類に分類できる [5]。パラメトリック法は、プロジェクトの工数を目的変数、その他のプロジェクトの特性を表す項目（メトリクス）を説明変数とした関数を用い、工数を予測したいプロジェクト（現行プロジェクト）の工数を求める手法である。使用する関数としては COCOMO[6] などの予め定義された関数も存在するが、回帰分析やニューラルネット [4] などを用いて過去に実施したプロジェクト（過去プロジェクト）をもとに開発組織ごとに関数を作成する手法もある。積上げ法では、プロジェクトの成果物を階層的に細かい構成要素に分解し、構成要素ごとの工数を計算してそれらを積上げることで最終的な予測工数を得る。類推法は、過去プロジェクトの中から現行プロジェクトと類似しているプロジェクト（類似プロジェクト）を探し、類似プロジェクトをもとに工数を予測する手法である。このように、多くの工数予測手法では関数の作成時や類似プロジェクトの探索時に過去プロジェクトの実績データを用いる。しかしながら、実績データに基づく工数予測手法の場合、予測精度はデータセット（過去プロジェクトの集合）に大きく依存する [7] ため、1 つの工数予測手法で常に高い精度を得ることは難しいという問題がある。例えば、特異なプロジェクトが少なく、目的変数と説明変数の相関が強いような「きれいな」データセットに対してはパラメトリック法が有効であり、逆に様々な環境のプロジェクトが多く、特定の関数では目的変数を説明しにくいような「乱雑な」データセットに対しては類推法が有効であるといえる [8]。

そこで、本研究では 1 件の現行プロジェクトに対して複数の工数予測手法を同時に適用し、その都度データセットに合った工数予測手法を選択する方法を提案する。工数予測手法の選択基準には、Leave-One-Out (LOO) [9] を用いてデータセット中の過去プロジェクトの工数を予測したときの誤差をもとに算出した、現行プロジェクトの工数の予測値に対する期待される精度を用いる。また、複数の工数予測手法を同時に適用可能な工数予測ツールが公開されていないことから、提案手法を実現し、期待される精度に応じた予測値の決定を支援する工数予測ツール「 e^3 」の開発を行った [10][11]。本ツールの表示する期待される精度の高い工数予測手法をその都度選択することにより、特定の工数予測手法にとらわれることなく、使用するデータセットの特徴に応じた柔軟な予測値の決定が可能となる。なお、本ツールはユーザがプログラミングにより比較的簡単に工数予測手法を追加実装できるよう設計されており、組織独自の工数予測手法と既存の工数予測手法の比較実験を行う用途にも利用できる。

提案手法の評価実験としては、特定の工数予測手法を使用し続けた場合と、ツールの表示する期待される精度の最も高い工数予測手法をその都度選択して使用した場合とで予測誤差の比較を行う。実験は、ISBSG のデータセットから抽出した 122 件のプロジェクトのデータ、ある企業の実施した特定の業界へ向けた 53 件のプロジェクトのデータ、別のある企業の実施した特定のアプリケーションに関する 37 件のプロジェクトのデータを使用し、LOO により行う。また、ツールの実用性の評価として、ツールを試用して頂いた企業の方へのアンケートも行う。

以降、2 では関連研究について述べ、3 で既存の工数予測手法を説明する。続いて 4 にて提案手法の説明を行い、5 で提案手法を実現する工数予測ツール「 e^3 」について説明する。次に 6 で実験とその結果について説明し、7 で実験の考察とツールの評価について述べる。最後に 8 でまとめる。

2 関連研究

2.1 工数予測とデータセットの関係

使用するデータセットによる予測精度の違いに関する既存研究としては, Jeffery ら [12], Kitchenham ら [13], Mendes ら [14] の研究がある. これらの研究では, 複数の組織のプロジェクトで構成される (cross-company) データセットと, 現行プロジェクトと同一の組織のプロジェクトで構成される (within-company) データセットによる工数予測の精度を比較している. 実験の結果, それぞれの研究で用いたデータセットの場合, within-company のデータセットを用いた方がより高い精度で予測が可能であることを示した. ただし, 同様の研究を行った Wieczorek ら [15] は, within-company のデータセットを用いた場合でも, cross-company のデータセットを用いた場合と比べて有意に高い精度で予測が行えたとはいえないとの実験結果を報告していることから, 使用するデータセットなどの条件によって結果は変化し得るといえる.

一方, データセットの特徴に応じた有効な工数予測手法に関する研究として, Shepperd ら [7] の研究がある. Shepperd らは, 特徴の異なる以下の 4 種類のデータセットを人工的に作成し, パラメトリック法である重回帰分析や, 類推法である類似性に基づく手法 [1][16] など複数の工数予測手法を用いたときのそれぞれの精度を比較した.

(A) 各メトリクスの値が正規分布に従うデータセット

(B) メトリクスの値として一部特異な値を持つデータセット

(C) 互いに相関の強いメトリクスを含むデータセット

(D) データセット (B) と (C) の条件を複合したデータセット

その結果, データセット (A) や (B) については重回帰分析が有効であり, データセット (C) や (D) については類似性に基づく手法が有効であった. また, Mendes ら [14] は, cross-company と within-company それぞれのデータセットにおいて有効であった工数予測手法についても言及しており, cross-company のデータセットを用いた場合は類似性に基づく手法が, within-company のデータセットを用いた場合は重回帰分析がそれぞれより有効であったと述べている.

2.2 工数予測ツール

既存の工数予測ツールとしては, 奈良先端科学技術大学院大学の Magi[17], 情報処理推進機構ソフトウェア・エンジニアリング・センターの CoBRA 法に基づく見積り支援ツール [18], 構造計画研究所の KnowledgePLAN[19] がある.

Magi は, 類似性に基づく手法を用いて工数予測を行うツールである. 類似性に基づく手法に特化し, 工数の予測値のほか, 類似プロジェクトの工数の実測値のばらつきや, 現行プロジェクトと類似プロジェクトとの比較などが出力される. また, データセットの品質診断機能も備えており, 工数予

測にどの程度役立つかという観点から、プロジェクト件数、期待される精度など 8 項目についての診断結果も出力される。

CoBRA 法に基づく見積り支援ツールは、CoBRA 法 [20] を用いた工数予測を支援するツールである。CoBRA 法に必要な変動要因関係図を作成する機能など、CoBRA 法に特化した機能を備えている。

KnowledgePLAN は、世界中から集められた過去プロジェクトと独自の予測アルゴリズムによって予測を行うツールである。ユーザが過去プロジェクトを保有していなくても、内蔵の 10000 件以上の過去プロジェクトをもとに世界標準の予測が可能である。

いずれのツールも特定の工数予測手法に特化した細かい出力が特徴であり、複数の工数予測手法には対応していない。そのため、複数の工数予測手法の結果から予測値を決定したい場合にはそれぞれのツールを実行し、予測値を得る必要がある。

3 工数予測手法

3.1 類似性に基づく手法

類似性に基づく手法 (Estimation by Analogy, EbA) は、過去プロジェクトの中から現行プロジェクトと類似しているプロジェクト (類似プロジェクト) を探し、類似プロジェクトの工数の実測値をもとに現行プロジェクトの工数を予測する手法である。これは、メトリクスの値の類似しているプロジェクト同士は工数の値も類似しているという仮説を前提としている。類似プロジェクトのみを用いて予測を行うため、過去プロジェクトの中に特異なプロジェクトが存在したとしても、そのプロジェクトの影響を受けずに予測が可能という長所がある。

EbA は、ダミー変数化、正規化、類似度計算、予測値計算の 4 つの手順から構成される。各手順で用いるアルゴリズムはそれぞれいくつか提案されているが、ここでは本研究で用いたアルゴリズム [21] のみを示す。

手順 1 ダミー変数化

データセットに名義尺度のメトリクスが含まれる場合、他の比例尺度のメトリクスと同様に扱えるようにするため、カテゴリごとにダミー変数に置き換える。プロジェクト p_i が名義尺度のメトリクス m_j においてカテゴリ c に属するかどうかを表すダミー変数 $d_{ij}(c)$ は式 (1) で定義される。

$$d_{ij}(c) = \begin{cases} 1 & \dots \text{カテゴリ } c \text{ に属する} \\ 0 & \dots \text{カテゴリ } c \text{ に属さない} \end{cases} \quad (1)$$

手順 2 正規化

一般に、データセット中の各メトリクスの値域にはばらつきがある。そこで、メトリクスごとの類似度への影響度を均等にするため、すべてのメトリクスの値を 0 から 1 の範囲で正規化する。メトリクス m_j のデータセット中での最大値を $\max(m_j)$ 、最小値を $\min(m_j)$ とすると、プロジェクト p_i のメトリクス m_j の値 v_{ij} を正規化した値 v'_{ij} は式 (2) で定義される。

$$v'_{ij} = \frac{v_{ij} - \min(m_j)}{\max(m_j) - \min(m_j)} \quad (2)$$

手順 3 類似度計算

各過去プロジェクトについて、現行プロジェクトとの類似度を計算する。プロジェクト p_a とプロジェクト p_b の類似度 $\text{sim}(p_a, p_b)$ は式 (3) で定義される。

$$\text{sim}(p_a, p_b) = \frac{1}{\text{dist}(p_a, p_b)} \quad (3)$$

ここで, $dist(p_a, p_b)$ はプロジェクト p_a とプロジェクト p_b のユークリッド距離を表し, 正規化されたメトリクスの値 v'_{ij} を用いて式 (4) で定義される.

$$dist(p_a, p_b) = \sqrt{\sum_{j=1}^n (v'_{aj} - v'_{bj})^2} \quad (4)$$

なお, n はメトリクス数を表す.

手順 4 予測値計算

類似プロジェクトの工数の実測値をもとに, 現行プロジェクトの工数の予測値を算出する. 予測値計算のアルゴリズムとしては, **Weighted Sum**[22] を用いた. プロジェクト p_i の工数の実測値を e_i とすると, 現行プロジェクト p_c の工数の予測値 $\hat{e}_c$ は式 (5) で定義される.

$$\hat{e}_c = \frac{\sum_{i \in k-nearestProjects} (e_i \times sim(p_c, p_i))}{\sum_{i \in k-nearestProjects} sim(p_c, p_i)} \quad (5)$$

ここで, $k-nearestProjects$ は現行プロジェクトとの類似度の高い上位 k 件の過去プロジェクト (類似プロジェクト) の集合を表す. k の値は別途実験的に求める必要がある.

3.2 回帰分析

回帰分析は, プロジェクトの工数を目的変数, その他のメトリクスを説明変数として, 過去プロジェクトをもとに回帰式を導出し, そこに現行プロジェクトのメトリクスの値を代入することで工数の予測値を得る手法である. 本研究では, 回帰モデルとして式 (6) の直線モデル, 式 (7) の累乗モデル, 式 (8) の線形モデルを用いて回帰分析を行った.

$$y = a + bx \quad (6)$$

$$y = ax^b \quad (7)$$

$$y = a_0 + a_1x_1 + a_2x_2 + \dots + a_kx_k \quad (8)$$

3.2.1 直線単回帰分析

直線単回帰分析では, 目的変数 y を工数, 説明変数 x をプロジェクトの規模とし, 式 (6) の直線モデルを用いて回帰分析を行う. 最小二乗法を用いることで, 式 (6) 中の係数 a と b はそれぞれ式 (9), (10) で求めることができる.

$$a = \bar{y} - b\bar{x} \quad (9)$$

$$b = \frac{s_{xy}}{s_x^2} \quad (10)$$

ここで, s_{xy} は x と y の共分散, s_x^2 は x の分散を表し, 標本数 (過去プロジェクト数) を n としてそれぞれ式 (11), (12) で定義される.

$$s_{xy} = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y}) \quad (11)$$

$$s_x^2 = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2 \quad (12)$$

3.2.2 累乗単回帰分析

累乗単回帰分析では, 目的変数 y を工数, 説明変数 x をプロジェクトの規模とし, 式 (7) の累乗モデルを用いて回帰分析を行う. 式 (7) は, 両辺を自然対数化すると式 (13) で表すことができる.

$$\ln(y) = \ln(a) + b \ln(x) \quad (13)$$

したがって, 直線単回帰分析と同様に, 式 (13) 中の係数 $\ln(a)$ と b はそれぞれ式 (14), (15) で求めることができる.

$$\ln(a) = \overline{\ln(y)} - b \overline{\ln(x)} \quad (14)$$

$$b = \frac{s_{\ln(x) \ln(y)}}{s_{\ln(x)}^2} \quad (15)$$

また, a は $\ln(a)$ と自然対数の底 e を用いて式 (16) によって表される.

$$a = e^{\ln(a)} \quad (16)$$

3.2.3 重回帰分析

重回帰分析では, 目的変数 y を工数, 説明変数 x_j をその他のメトリクスとし, 式 (8) の線形モデルを用いて回帰分析を行う. 最小二乗法を用いることで, 式 (8) 中の係数 a_0 と $a_1, a_2, \dots, a_k$ はそれぞれ式 (17), (18) で求めることができる.

$$a_0 = \bar{y} - \begin{pmatrix} a_1 & a_2 & \cdots & a_k \end{pmatrix} \begin{pmatrix} \overline{x_1} \\ \overline{x_2} \\ \vdots \\ \overline{x_k} \end{pmatrix} \quad (17)$$

$$\begin{pmatrix} a_1 \\ a_2 \\ \vdots \\ a_k \end{pmatrix} = S^{-1}M \quad (18)$$

ここで、 S と M は変数の分散や変数間の共分散を用いて、それぞれ式 (19), (20) で定義される行列である。

$$S = \begin{pmatrix} s_{x_1}^2 & s_{x_1 x_2} & \cdots & s_{x_1 x_k} \\ s_{x_2 x_1} & s_{x_2}^2 & \cdots & s_{x_2 x_k} \\ \vdots & \vdots & \ddots & \vdots \\ s_{x_k x_1} & s_{x_k x_2} & \cdots & s_{x_k}^2 \end{pmatrix} \quad (19)$$

$$M = \begin{pmatrix} s_{x_1 y} \\ s_{x_2 y} \\ \vdots \\ s_{x_k y} \end{pmatrix} \quad (20)$$

重回帰分析の結果は使用するメトリクスによって変化するが、変数増加法 [23] を用いることで機械的にメトリクスを選択することも可能である。変数増加法では、以下の手順で使用するメトリクスを選択する。

手順 1 使用するメトリクスの集合 X を空に、 X に含まれるメトリクスを用いて重回帰分析を行ったときの自由度調整済み決定係数 $adjR_X^2$ を 0 に初期化する。

手順 2 X に含まれないメトリクスを 1 つずつ X に追加したと想定し、それぞれ重回帰分析を行い、自由度調整済み決定係数 $adjR^2$ と、追加したメトリクスのトレランスを計算する。

手順 3 手順 2 で求めた $adjR^2$ がその時点の $adjR_X^2$ より大きく、かつトレランスが十分大きいメトリクスのうち、最も $adjR^2$ の大きいメトリクスを X に追加する。 $adjR_X^2$ を $adjR^2$ で更新する。もしそのようなメトリクスが無い場合は、その時点での X を確定し、終了する。

手順 4 手順 2 から繰り返す。

ここで、トレランスとは多重共線性を検出するための指標であり、トレランスが小さいとき多重共線性があるといえる。本研究では、トレランスの値が 0.1 以下であったときに多重共線性があると判断している。したがって、手順 3 において「トレランスが十分大きい」とは、具体的にはトレランスが 0.1 より大きいことである。

4 提案手法

4.1 概要

現行プロジェクトの工数を予測するとき、最も小さな誤差で予測ができる工数予測手法を使用することが望ましい。しかしながら、工数予測の時点では現行プロジェクトの工数の実測値は未知であり、事前に実測値と予測値の誤差を知ることはできない。そこで本研究では、現行プロジェクトの工数予測とは別に、データセットを構成する過去プロジェクトの工数を予測し、そのときの誤差を「期待される精度」として工数予測手法の選択基準にすることを提案する。

4.2 期待される精度

ある工数予測手法によって得られた現行プロジェクトの工数の予測値に対する期待される精度は、同一の工数予測手法を用いて過去プロジェクトの工数を予測したときの実測値と予測値の誤差で表現される。現行プロジェクトと違い、過去プロジェクトの工数の実測値は既知であるため、実測値と予測値の誤差を計算することが可能である。そのときの誤差が小さいほど、同一の工数予測手法で現行プロジェクトの工数を予測したときの期待される精度は高く、小さい誤差で現行プロジェクトの工数を予測できることが期待される。

過去プロジェクトの工数を予測する手順としては、Leave-One-Out (LOO) [9] を用いる。LOO は、過去プロジェクト n 件によって構成されるデータセットから 1 件を予測対象プロジェクトとして抽出し、残りの $n-1$ 件の過去プロジェクトをもとに予測対象プロジェクトの工数を予測して誤差を求めることを n 件の全過去プロジェクト分繰り返して工数予測手法の精度を計算する方法である。LOO を用いて、1 件の過去プロジェクトを抽出し、精度計算のための予測を行う場面を図 1 に示す。ある工数予測手法について、LOO を用いて現行プロジェクトの工数の予測値に対する期待される精度を求める具体的な手順は以下のとおりである。

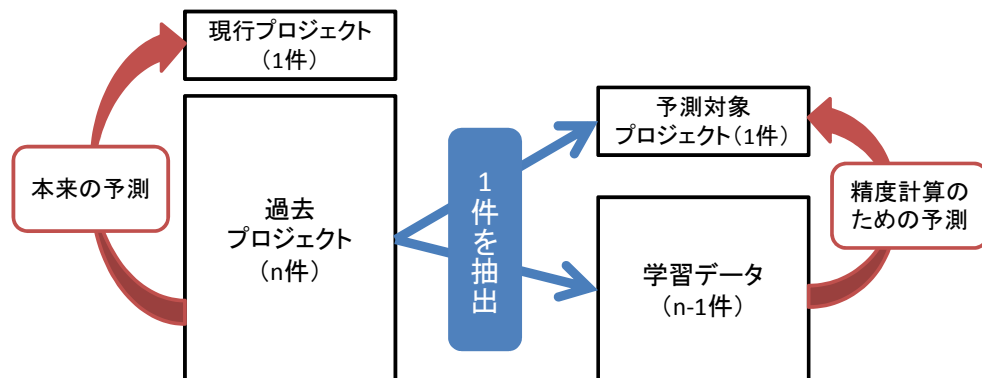


図 1: LOO を用いた過去プロジェクトの工数予測

手順 1 n 件の全過去プロジェクトから 1 件のプロジェクトを抽出する.

手順 2 抽出した 1 件を予測対象プロジェクト, 残りの $n-1$ 件を学習データ (過去プロジェクト) と
して, 当該工数予測手法にて学習データをもとに予測対象プロジェクトの工数を予測する.

手順 3 予測対象プロジェクトの工数の実測値と予測値の誤差を計算する.

手順 4 予測対象プロジェクトとして抽出するプロジェクトを変更し, n 件の全過去プロジェクトに
ついて手順 1 から手順 3 を繰り返す.

手順 5 n 回繰り返された手順 3 によって求められた n 個の誤差をもとに, 当該工数予測手法を使用
した際の現行プロジェクトの工数の予測値に対する期待される精度を算出する.

手順 3 で計算する実測値と予測値の誤差の尺度としては, 実測値からみた予測値の相対誤差である
MRE (Magnitude of Relative Error) と, 予測値からみた実測値の相対誤差である MER (Magnitude
of Error Relative) [24] を用いる. MRE と MER は, 実測値を $real$, 予測値を est としてそれぞれ式
(21), (22) で定義される.

$$MRE = \frac{|real - est|}{real} \quad (21)$$

$$MER = \frac{|real - est|}{est} \quad (22)$$

また, 手順 5 で求める期待される精度としては, MRE の平均値である MMRE (Mean MRE), 中央
値である MdMRE (Median MRE), MRE が 0.25 以下であった予測の割合を示す Pred(0.25), MER
の平均値である MMER (Mean MER) を用いる. これらは工数予測手法の精度評価で良く用いられ
る指標であり [25], それぞれ式 (23)~(26) で定義される.

$$MMRE = \frac{1}{n} \sum_{i=1}^n MRE_i \quad (23)$$

$$MdMRE = \text{median of } MREs \quad (24)$$

$$Pred(0.25) = \frac{\# \text{ projects with } MRE \leq 0.25}{\# \text{ projects}} \quad (25)$$

$$MMER = \frac{1}{n} \sum_{i=1}^n MER_i \quad (26)$$

4.3 工数予測手法の選択

複数の工数予測手法について, それぞれ現行プロジェクトの工数の予測値と, その予測値に対する
期待される精度を求める. 本研究で提案する期待される精度の評価指標のうち, MMRE, MdMRE,

MMER は値が小さいほど、Pred(0.25) は値が大きいほど期待される精度は高いといえる。したがって、使用した工数予測手法のうち、MMRE, MdMRE, MMER がより小さく、Pred(0.25) がより大きい工数予測手法を選択し、その工数予測手法によって得られた工数の予測値を採用することで、現行プロジェクトの工数の実測値との誤差の小さい予測が可能であると考えられる。

5 工数予測ツール「 e^3 」

5.1 概要

工数予測ツール「 e^3 」は、本研究で提案する期待される精度による工数予測手法の選択を支援するツールである。1件の現行プロジェクトに対して複数の工数予測手法を適用し、工数予測手法ごとに工数の予測値と期待される精度を表示する。1度の実行で複数の工数予測手法の予測結果を確認ことができ、期待される精度の比較によって、特定の工数予測手法に依存しない柔軟な工数予測が可能となる。なお、使用する工数予測手法は、既にツールに実装されているもののほか、ユーザがプログラミングにより比較的簡単に追加実装できるよう設計されている。

本ツールの構成図を図2に示す。ユーザが現行プロジェクトと過去プロジェクトの実績データ（プロジェクトデータ）を与えると、まず内部でプロジェクトデータの編集が行われる。その後、編集後のプロジェクトデータと実装されている工数予測手法から現行プロジェクトの工数予測と、期待される精度の計算が行われる。ユーザは、本ツールから得られた工数の予測値と期待される精度をもとに、採用する予測値を選択できる。

なお、汎用性を重視し、本ツールは Ruby on Rails を用いて Web アプリケーションとして実装している。

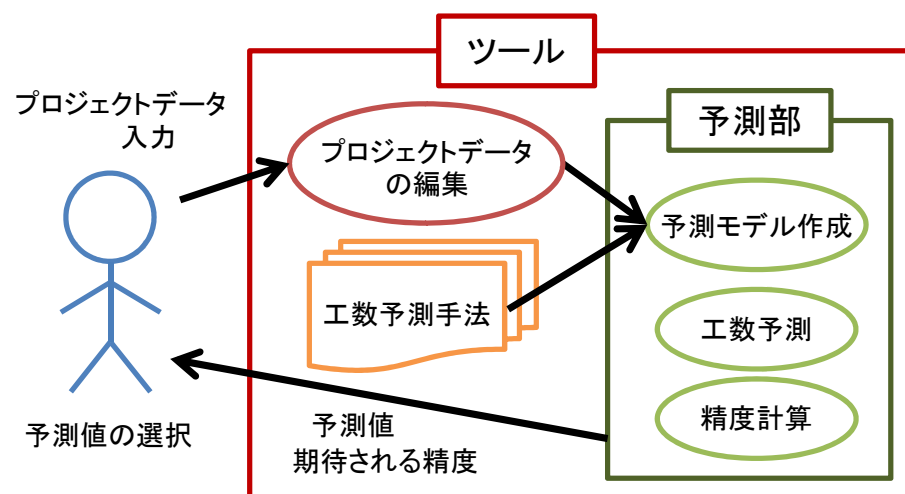


図 2: ツールの構成図

5.2 入出力

5.2.1 入力画面

本ツールの入力画面を図 3 に示す。以降、各入力項目について説明する。

- プロジェクトデータ

現行プロジェクトと過去プロジェクトのデータを記述した CSV ファイルを指定する。

- 工数予測手法

適用する工数予測手法をチェックボックスにより選択する。ユーザが追加実装した工数予測手法も含め、本ツールに実装されているすべての工数予測手法が一覧で表示される。

- 精度の計算頻度

精度の計算頻度を指定する。本来は「1 プロジェクト毎」が望ましいが、数字を大きくすることで計算を省略し、実行時間を短縮することが可能である。

単純な LOO を行う場合、過去プロジェクトが n 件あったとき、本来の現行プロジェクトの工数予測とは別に n 回の工数予測を実行しなければならないが、実行時間が長くなってしまうことが考えられる。そこで、本ツール導入の初期段階で使用する過去プロジェクトやメトリクスを変更して複数回試験的にツールを実行する状況を考慮し、時間短縮のオプションとして入力項目

Estimator 3 - Mozilla Firefox

ファイル(E) 編集(E) 表示(V) 履歴(S) ブックマーク(B) ツール(I) ヘルプ(H)

http://127.0.0.1:3000/estimator/start

工数予測ツール e³

プロジェクト情報

参照

プロジェクトデータ

予測手法

- ☒ Average: 平均値
- ☒ Average: 中央値
- ☒ Eba: 類似性に基づく手法 (k=1)
- ☒ Eba: 類似性に基づく手法 (k=2)
- ☒ Eba: 類似性に基づく手法 (k=3)
- ☒ Eba: 類似性に基づく手法 (k=4)
- ☒ Eba: 類似性に基づく手法 (k=5)
- ☒ Regression: 直線単回帰分析
- ☒ Regression: 累乗単回帰分析
- ☒ Regression: 重回帰分析 (全メトリクス使用)
- ☒ Regression: 重回帰分析 (変数増加法)

工数予測手法

精度の計算頻度

1 プロジェクト毎

精度の計算頻度

実行

図 3: 入力画面

「精度の計算頻度」を設けた。これは、「 x プロジェクト毎」の x の値を指定するもので、4.2 の LOO の手順 4 において全ての過去プロジェクトを 1 件ずつ抽出するのではなく、何件に 1 件の過去プロジェクトを抽出するかを指定する項目である。例えば、 x が 10 であれば過去プロジェクト 10 件に 1 件を、100 であれば過去プロジェクト 100 件に 1 件を予測対象プロジェクトとして抽出し、精度計算のための予測を行う。これにより、単純な LOO より高速に結果を表示することが可能となる。ただし、 x を大きくするほど出力される期待される精度は大まかなものとなる。精度の計算頻度の違いによる実行時間の変化については、7.2 で簡単な実験を行う。

5.2.2 プロジェクトデータファイル

プロジェクトデータを記述した CSV ファイルは図 4 のように書式が決まっている。各行には、1 行目にメトリクスの名前、2 行目にメトリクスの種類、3 行目に現行プロジェクトのデータ、4 行目以降に過去プロジェクトのデータを記述する。メトリクスの種類としては、

- 比例尺度 (1)
- 名義尺度 (2)
- その他 (0)

に対応しており、2 行目にはそれぞれ対応する数字を記述する。「その他」が指定されたメトリクスは無視され、工数予測には使用されない。各列には、1 列目にプロジェクト名、2 列目に工数、3 列目に規模（ファンクションポイント (FP) やコード行数 (LOC)）、4 列目以降にその他のメトリクスを記述する。なお、工数と規模は比例尺度のメトリクスであることから、2 列目と 3 列目のメトリクスの種類としては「1」を指定しなければならない。

プロジェクト名	工数	規模	その他のメトリクス	
メトリクス名	メトリクスの種類			
Sample1.csv	Effort	FP	Language	Team Size
ver.	1	1	2	1
Current Project		220	Java	10
Sample Project A	1059	133	Java	8
Sample Project B	3700	791	C	
Sample Project C	1950	360	Java	21
Sample Project D	883	191		16

図 4: CSV ファイルを Excel 等で開いたイメージ

5.2.3 プロジェクトデータの編集

CSV ファイルによって入力されたプロジェクトデータは、予測の実行の前に一部編集が行われる。なお、これはツール内部の処理であり、実際の CSV ファイルは変更されない。

- 不必要なメトリクスの削除

メトリクスの種類が「その他」のメトリクス、および全過去プロジェクトにおいて値が欠損しているメトリクスは、予測の前に削除される。

- 欠損値の自動補完

プロジェクトデータに欠損値が含まれている場合、自動補完が行われる。補完方法はメトリクスの種類によって異なる。

比例尺度のメトリクス

過去プロジェクトの当該メトリクスの平均値を自動補完。

名義尺度のメトリクス

過去プロジェクトの当該メトリクスの最頻値を自動補完。

自動補完を望まない場合は、欠損値を含む過去プロジェクトやメトリクスを事前に CSV ファイルから削除しておくなどの対応が必要である。なお、工数以外の現行プロジェクトの欠損値も同様に自動補完される。

- 名義尺度メトリクスのダミー変数化

多くの工数予測手法において、名義尺度のメトリクスはそのままでは処理ができない。多様な工数予測手法の実装を想定している本ツールにおいては、名義尺度のメトリクスも比例尺度のメトリクスと同様に処理が行えるよう、予測前にダミー変数化が行われる。3.1 の EbA の手順 1 に相当する。

例として、Language という名義尺度のメトリクスがあり、その値の候補として Java, C, Ruby があつたとすると、メトリクス Language は削除され、新たに Language が Java であるかどうかを表す Language='Java' と、Language が C であるかどうかを表す Language='C' という 2 つの比例尺度メトリクスが生成される。説明変数間の相関の強さに由来する多重共線性を防ぐため、Language='Ruby' は生成されない。新たな各メトリクスの値としては、もとの Language が Java であつたプロジェクトについては、Language='Java' は 1, Language='C' は 0 となり、もとの Language が Ruby であつたプロジェクトについては、Language='Java', Language='C' ともに 0 となる。

5.2.4 出力画面

出力画面を図 5 に示す。表の 1 つの行が 1 つの工数予測手法による予測結果を表す。1 列目に工数予

予測手法	Effort	MMRE	MdMMRE	pred(0.25)	MMER	順位
平均値	4023.273	4.125	3.786	0.000 %	0.662	5 (4.000)
類似性に基づく手法 (k=3)	344.763	0.410	0.460	50.000 %	0.332	3 (1.223)
類似性に基づく手法 (k=5)	376.465	0.472	0.290	50.000 %	0.314	1 (1.185)
累乗単回帰分析	1324.431	0.480	0.101	50.000 %	0.359	2 (1.185)
重回帰分析 (全モデル使用)	994.133	1.273	0.488	20.000 %	0.623	4 (2.180)

[戻る](#)

図 5: 出力画面

測手法の名前，2 列目に現行プロジェクトの工数の予測値，3 列目以降に期待される精度とそれをもとに算出した総合順位が表示される．期待される精度とは 4.2 で述べた MMRE，MdMMRE，Pred(0.25)，MMER であり，総合順位とは 4 種類の期待される精度を 0 から 1 で正規化（数字が小さいほど精度が高い）し，その和の小さい順に順位付けをしたものである．各評価指標において，最も期待される精度が高いものは赤く表示される．また，工数予測手法の名前をクリックすることで，回帰式や決定係数など，工数予測手法ごとの個別の情報を表示することも可能である．

なお，現行プロジェクトの工数の予測結果として「予測不能」と表示されることがあるが，これは当該手法の予測結果として 0 以下の数が算出されたためにツールが「予測不能」と表示している場合と，手法自体から「予測不能である」と判断された場合がある．前者の場合，手法の名前をクリックすると「(予測値が 0 以下の数であったため予測不能としました)」と表示される．また，精度計算のための予測を実行中にも一部の予測で「予測不能」と判断される場合もあるが，その場合，本ツールでは当該予測をとばし，その他の予測のみを用いて MMRE 等を計算する．

5.3 実装済みの工数予測手法

本ツールに実装されている工数予測手法について説明する．

- 平均値

過去プロジェクトの工数の平均値を現行プロジェクトの工数の予測値とする．

- 中央値

過去プロジェクトの工数の中央値を現行プロジェクトの工数の予測値とする．

- 類似性に基づく手法 ($k=1\sim5$)

類似性に基づく手法による工数予測を行う。出力画面において手法名がクリックされたときは、上位 $k+2$ 件の類似プロジェクトの情報が表示される。

手法の性質上、 $k=n$ のこの手法を選択する場合、入力するプロジェクトデータとして $n+4$ 件（現行プロジェクト含む）のプロジェクトが必要である。

- 直線単回帰分析

直線単回帰分析による工数予測を行う。出力画面において手法名がクリックされたときは、導出された回帰式と決定係数が表示される。

手法の性質上、全過去プロジェクトを通して工数または規模の値が同一であってはならない。

- 累乗単回帰分析

累乗単回帰分析による工数予測を行う。出力画面において手法名がクリックされたときは、導出された回帰式と決定係数が表示される。

手法の性質上、全過去プロジェクトを通して工数または規模の値が同一であってはならない。また、工数または規模の値が 0 以下のプロジェクトが含まれていてはならない。

- 重回帰分析（全メトリクス使用）

CSV ファイル中の全メトリクスを使用し、重回帰分析による工数予測を行う。出力画面において手法名がクリックされたときは、導出された回帰式と決定係数、自由度調整済み決定係数が表示される。

使用する過去プロジェクトによっては、式 (19) の行列 S に逆行列が存在せず、「予測不能」となることがある。

- 重回帰分析（変数増加法）

変数増加法によって自動で使用するメトリクスを選択し、重回帰分析による工数予測を行う。出力画面において手法名がクリックされたときは、導出された回帰式と決定係数、自由度調整済み決定係数が表示される。

使用する過去プロジェクトによっては、式 (19) の行列 S に逆行列が存在せず、「予測不能」となることがある。

6 実験

6.1 概要

本研究で提案する期待される精度の比較による工数予測手法の選択の有効性を確認するために実験を行った。実験には工数予測ツール「e³」を用い、複数回工数予測を行う際に、

- 常に同じ工数予測手法を使用し続けた場合
- ツールの表示する期待される精度の最も高い工数予測手法をその都度選択して使用した場合

とで現行プロジェクトの工数の実測値と予測値の誤差を比較する。実験に使用する工数予測手法は、ツールに実装されている 11 の工数予測手法である。

6.2 使用したデータセット

実験には、以下の 3 種類のデータセット¹を使用した。

- ISBSG DATA Release 11[26]
- A 社のデータセット
- B 社のデータセット

ISBSG DATA Release 11 は、The International Software Benchmarking Standards Group (ISBSG) によって世界 24 国から収集された 5052 件の過去プロジェクトによって構成されるデータセットである。本実験では、その中から、同じく ISBSG のデータセットを用いて実験を行っている C. Lokan らの研究 [27] を参考に、以下の条件に当てはまるプロジェクトのみを抽出し、使用した。

- データの質ランクが A である。
- FP の計測方法が IFPUG 法 [28] である。
- 正規化済み工数と総工数が等しい。
- Web 開発ではない。

また、開発形態は最も多いエンハンスメントのみを抽出した。その結果、122 件のプロジェクトが残った。メトリクスとしては、欠損値の多いものを除外し、予測対象である総工数のほか、調整済み FP、最大要員数、無調整 FP の外部入力・外部出力・外部照合・内部論理ファイル・外部インターフェースの内訳、無調整 FP の追加分・変更分・削除分の内訳、の 10 の比例尺度メトリクスと、言語タイプ、アーキテクチャ、クライアントサーバ方式であるかどうか、の 3 つの名義尺度メトリクスを使用した。表 1 に、使用したメトリクスと 122 件のプロジェクト内におけるメトリクスごとの欠損値の割合を示す。

¹一部のデータセットでは、本研究の趣旨に影響しない範囲でメトリクスの数値に加工がなされている。

表 1: ISBSG : 使用メトリクスと欠損値の割合

メトリクス名	メトリクスの種類	欠損値の割合
Summary Work Effort	比例尺度	0%
Adjusted Function Points	比例尺度	0%
Max Team Size	比例尺度	26.2%
Input count	比例尺度	9.0%
Output count	比例尺度	9.0%
Enquiry count	比例尺度	9.0%
File count	比例尺度	9.0%
Interface count	比例尺度	9.0%
Added count	比例尺度	9.0%
Changed count	比例尺度	9.0%
Deleted count	比例尺度	9.0%
Language Type	名義尺度	9.0%
Architecture	名義尺度	18.0%
Client Server?	名義尺度	18.9%

A 社のデータセットは、ある業界へ向けた A 社のプロジェクトを収集したデータセットである。53 件のプロジェクトデータで構成され、メトリクスとしては、予測対象である工数のほか、1 つの比例尺度メトリクスと、11 の名義尺度メトリクスがある。データセット中に欠損値は含まれていない。

B 社のデータセットは、あるアプリケーションに関する B 社のプロジェクトを収集したデータセットである。37 件のプロジェクトデータで構成され、メトリクスとしては、予測対象である工数のほか、1 つの比例尺度メトリクスと、2 つの名義尺度メトリクスがある。データセット中に欠損値は含まれていない。

6.3 手順

工数の実測値が既知であるプロジェクトを現行プロジェクトとして本ツールを実行することにより、現行プロジェクトの工数の実測値と予測値の誤差を求めることが可能である。実験には LOO を用い、3 つのデータセット別に、1 件ずつ現行プロジェクトとみなすプロジェクトを抽出してツールを実行することをデータセットに含まれる全プロジェクト分繰り返した。以降、実験における LOO の手順中の 1 回のツールの実行を 1 回の試行と呼ぶこととする。各試行では、表 2 に示すように、現行プロジェクトの工数の実測値と各工数予測手法による工数の予測値との誤差である MRE と MER を計算し、同時にツールの表示する期待される精度の最も高い工数予測手法を記録する。表中の $M_1 \sim$

表 2: 実験で収集されるデータ

試行	誤差の種類	M_1	M_2	$\cdots$	M_{11}	best MMRE	best MMER	best Rank
1	MRE	$mre_{1,1}$	$mre_{1,2}$	$\cdots$	$mre_{1,11}$	X_1	Y_1	Z_1
	MER	$mer_{1,1}$	$mer_{1,2}$	$\cdots$	$mer_{1,11}$			
2	MRE	$mre_{2,1}$	$mre_{2,2}$	$\cdots$	$mre_{2,11}$	X_2	Y_2	Z_2
	MER	$mer_{2,1}$	$mer_{2,2}$	$\cdots$	$mer_{2,11}$			
$\vdots$	$\vdots$			$\cdots$		$\vdots$	$\vdots$	$\vdots$
n	MRE	$mre_{n,1}$	$mre_{n,2}$	$\cdots$	$mre_{n,11}$	X_n	Y_n	Z_n
	MER	$mer_{n,1}$	$mer_{n,2}$	$\cdots$	$mer_{n,11}$			

M_{11} は工数予測手法を, $mre_{a,b}$ と $mer_{a,b}$ は, a 回目の試行における, 現行プロジェクトの工数の実測値と工数予測手法 M_b による工数の予測値との誤差 **MRE** と **MER** を示す. また, X_a , Y_a , Z_a は, それぞれ a 回目の試行においてツールの表示する **MMRE**, **MMER**, 総合順位が最も良かった工数予測手法の番号を示す. すると, 工数予測手法 M_b を使用し続けた場合の **MRE** の平均値 $MMRE_{Mb}$ と **MER** の平均値は $MMER_{Mb}$ は, それぞれ式 (27), (28) で表される.

$$MMRE_{Mb} = \frac{1}{n} \sum_{i=1}^n mre_{i,b} \quad (27)$$

$$MMER_{Mb} = \frac{1}{n} \sum_{i=1}^n mer_{i,b} \quad (28)$$

また, 各試行でツールの表示する **MMRE** が最も良かった工数予測手法を使用した場合の **MRE** の平均値 $MMRE_{mmre}$, **MMER** が最も良かった工数予測手法を使用した場合の **MER** の平均値 $MMER_{mmer}$, 総合順位が最も良かった工数予測手法を使用した場合の **MRE** の平均値 $MMRE_{rank}$, **MER** の平均値 $MMER_{rank}$ はそれぞれ式 (29)~(32) で表される.

$$MMRE_{mmre} = \frac{1}{n} \sum_{i=1}^n mre_{i,X_i} \quad (29)$$

$$MMER_{mmer} = \frac{1}{n} \sum_{i=1}^n mer_{i,Y_i} \quad (30)$$

$$MMRE_{rank} = \frac{1}{n} \sum_{i=1}^n mre_{i,Z_i} \quad (31)$$

$$MMER_{rank} = \frac{1}{n} \sum_{i=1}^n mer_{i,Z_i} \quad (32)$$

実験は, 式 (27)~(32) の比較によって評価を行う. それぞれの値は誤差であるので, 値は小さいほど良いといえる.

6.4 結果

実験結果を図 6～図 11 に示す。図 6 と図 7 は ISBSG のデータセットを使用した実験結果を、図 8 と図 9 は A 社のデータセットを使用した実験結果を、図 10 と図 11 は B 社のデータセットを使用した実験結果をそれぞれ示している。グラフの縦軸は n 回の試行における MRE の平均値である MMRE と MER の平均値である MMER を表し、横軸において mean, eba1 など工数予測手法の名前の部分は式 (27), (28) の $MMRE_{Mb}$ や $MMER_{Mb}$ を, best MMRE, best MMER, MMRE を比較するグラフの best Rank, MMER を比較するグラフの best Rank はそれぞれ式 (29)～(32) の $MMRE_{mmre}$, $MMER_{mmer}$, $MMRE_{rank}$, $MMER_{rank}$ を示している。MMRE や MMER の値の小さい工数予測手法ほど、小さい誤差で工数予測が行えたことを表す。なお, mean は平均値, median は中央値, eban は $k=n$ の EbA, linear は直線単回帰分析, power は累乗単回帰分析, multi は重回帰分析 (全メトリクス使用), forward は重回帰分析 (変数増加法) をそれぞれ表す。

図 6, 図 7 から, ISBSG のデータセットにおいて特定の工数予測手法を使用し続けた場合, 良い結果を出したのは k の値の大きい EbA や累乗単回帰分析であった。MMRE の小さい工数予測手法は, 順に EbA($k=3$) (0.748), EbA($k=5$) (0.751), 累乗単回帰分析 (0.762) であり, MMER の小さい工数予測手法は, 順に EbA($k=5$) (0.633), 直線単回帰分析 (0.645), EbA($k=4$) (0.662) である。一方, ツールの表示する期待される精度の最も高い工数予測手法をその都度選択して使用した場合も, best MMRE で MMRE が 0.849, best MMER で MMER が 0.651, best rank で MMRE が 0.810, MMER が 0.781 と, 比較的良好な結果が得られている。

図 8, 図 9 から, A 社のデータセットにおいて特定の工数予測手法を使用し続けた場合, 良い結果を出したのは累乗単回帰分析や直線単回帰分析であった。MMRE の小さい工数予測手法は, 順に累乗単回帰分析 (0.402), 重回帰分析 (変数増加法) (0.567), 直線単回帰分析 (0.574) であり, MMER の小さい工数予測手法は, 順に直線単回帰分析 (0.367), 累乗単回帰分析 (0.405), EbA($k=2$) (0.465) である。一方, ツールの表示する期待される精度の最も高い工数予測手法をその都度選択して使用した場合も, best MMRE で MMRE が 0.402, best MMER で MMER が 0.436, best rank で MMRE が 0.434, MMER が 0.489 と, 使用し続けた場合に良い結果を出した累乗単回帰分析や直線単回帰分析と近い精度で予測が行えている。

図 10, 図 11 から, B 社のデータセットにおいて特定の工数予測手法を使用し続けた場合, 良い結果を出したのは累乗単回帰分析であった。MMRE の小さい工数予測手法は, 順に累乗単回帰分析 (0.511), 直線単回帰分析 (0.614), 重回帰分析 (全メトリクス使用) (0.691) であり, MMER の小さい工数予測手法は, 順に累乗単回帰分析 (0.483), 重回帰分析 (変数増加法) (0.712), 重回帰分析 (全メトリクス使用) (0.727) である。一方, ツールの表示する期待される精度の最も高い工数予測手法をその都度選択して使用した場合も, best MMRE で MMRE が 0.511, best MMER で MMER が 0.483, best rank で MMRE が 0.534, MMER が 0.460 と, 使用し続けた場合に良い結果を出した累乗単回帰分析と近い精度で予測が行えている。

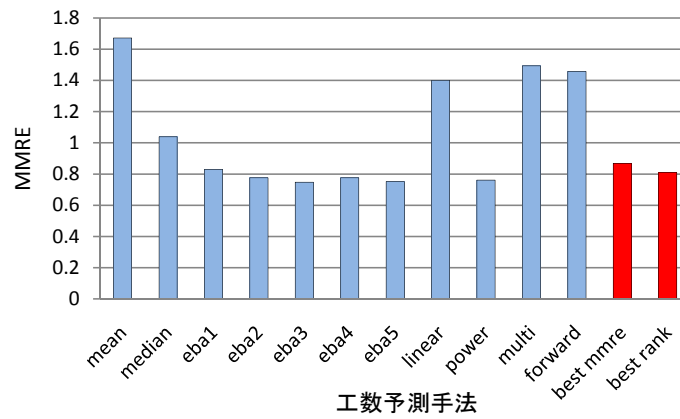


図 6: ISBSG : MMRE の比較

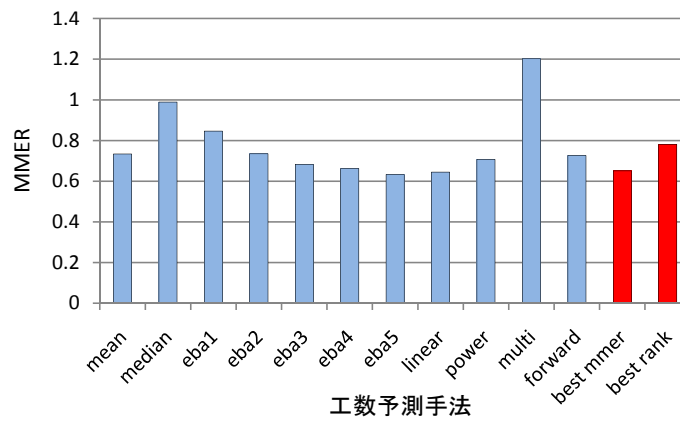


図 7: ISBSG : MMER の比較

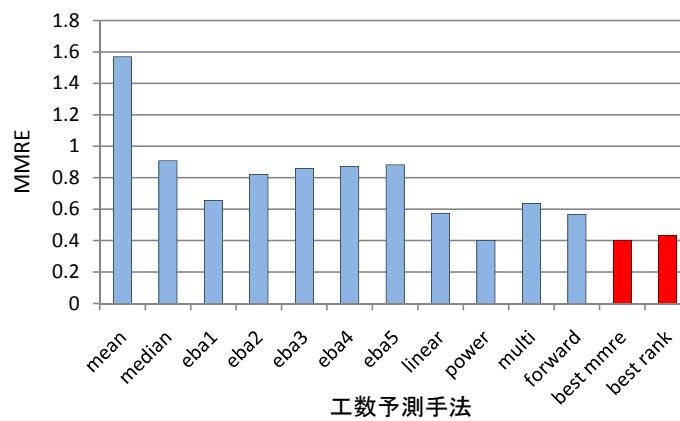


図 8: A 社 : MMRE の比較

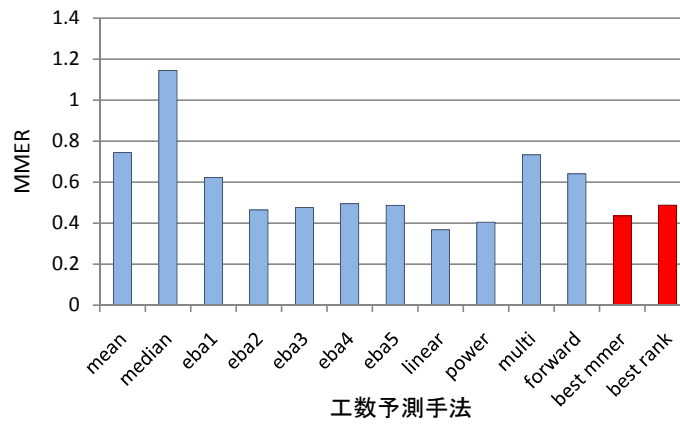


図 9: A 社 : MMER の比較

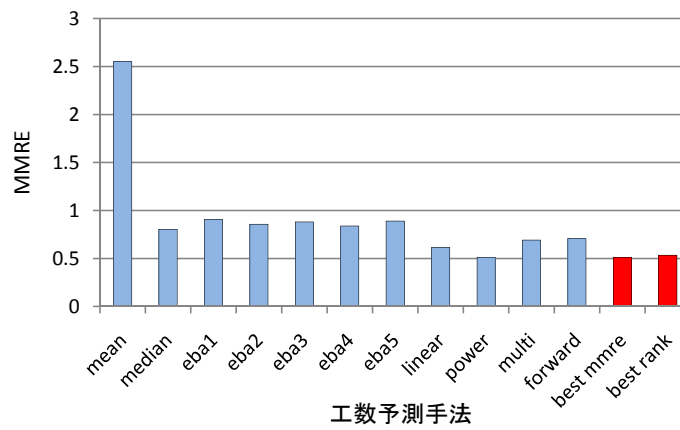


図 10: B 社 : MMRE の比較

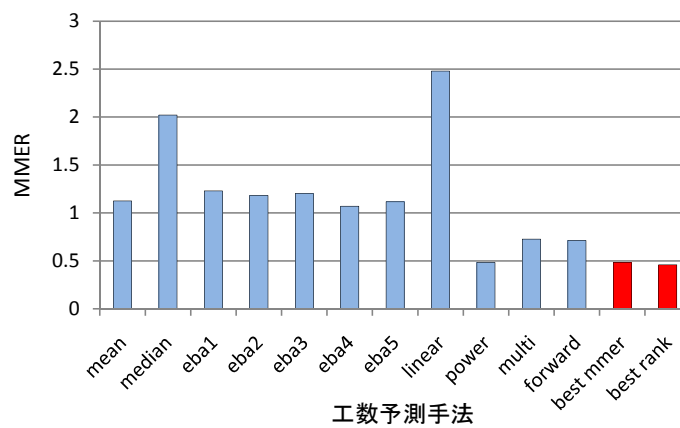


図 11: B 社 : MMER の比較

7 考察

7.1 実験結果の考察

実験の結果、データセットごとに有効な工数予測手法は異なり、本研究の提案する期待される精度をもとに工数予測手法を選択する方法で、それぞれのデータセットで有効な工数予測手法と同程度の誤差で予測が行えることがわかった。

ISBSG のデータセットに対しては、 k の値の大きい EbA が有効であった。ISBSG のデータセットは、世界中の様々な条件のもと実施されたプロジェクトのデータから構成されている。そのため、単純な関数で工数を表現することが難しく、類似しているプロジェクトのみを用いて工数を予測する EbA が有効であったと考えられる。表 3 は、best MMRE や best MMER で実際に選択された工数予測手法の割合（小数第 2 位で四捨五入）を示している。ISBSG のデータセットでは、best MMRE では MMRE の比較において最も良い結果を出した EbA($k=3$) が、best MMER では MMER の比較において最も良い結果を出した EbA($k=5$) がそれぞれ最も多く選択されており、期待される精度による工数予測手法の選択は有効であったといえる。

A 社と B 社のデータセットに対しては、累乗単回帰分析などの回帰分析が有効であった。特定の業界へ向けたプロジェクトや特定のアプリケーションに対するプロジェクトなど、1 つの企業の中でも何かしらの条件でプロジェクトを層別することで、類似したプロジェクトが集まり、工数を関数で表現しやすかったためと考えられる。また、回帰分析の中でも、プロジェクトの規模を説明変数とした単回帰分析が良い結果を出しており、工数予測を行う前に正確に規模を測定することの重要性が明らかになった。B 社のデータセットに含まれる規模を表すメトリクスは、簡易測定法によって測定された FP である。その FP を用いた累乗単回帰分析の結果、EbA などの他の工数予測手法より小さい誤差で予測が行えていることから、NESMA[29] などの簡易測定法を用いてプロジェクトの初期段階で規模を測定することは重要であるといえる。なお、期待される精度をもとに工数予測手法を選択した場合、表 3 より、それぞれのデータセットにおいて最も良い結果を出した工数予測手法をほぼ 100% の割合で選択できており、データセットに関わらず、期待される精度による工数予測手法の選択の有効性が確認された。

しかしながら、期待される精度によって複数の工数予測手法の中からその都度適切と思われるものを選択することで特定の工数予測手法より高い予測精度が得られたわけではなかった。図 12 は、ISBSG のデータセットに対する実験中の各試行において、最小の誤差で予測が行えた工数予測手法の割合（小数第 2 位で四捨五入）を示した円グラフである。グラフからわかるように、同じデータセット中でも、試行ごとに最善の工数予測手法は異なり、各試行でグラフ中の工数予測手法を選択できることが最も望ましい。しかし工数の実測値が未知である工数予測の段階で最善の工数予測手法を見極めることは非常に困難であり、実験結果のように、当該データセットに対して平均的に小さい誤差で予測が行える工数予測手法を判別できたことには意義があるといえる。なお、図 12 でそれぞれ 10% 程度の割合を占めている平均値、中央値、重回帰分析が図 6 や図 7 で悪い結果となっている

表 3: 選択された工数予測手法の割合

データセット	手法	1 位	2 位	3 位
ISBSG	best MMRE	EbA($k=3$) (77.9%)	EbA($k=5$) (15.6%)	power (4.1%)
	best MMER	EbA($k=5$) (90.2%)	power (7.4%)	EbA($k=4$) (2.6%)
A 社	best MMRE	power (100%)	-	-
	best MMER	linear (98.1%)	forward (1.9%)	-
B 社	best MMRE	power (100%)	-	-
	best MMER	power (100%)	-	-

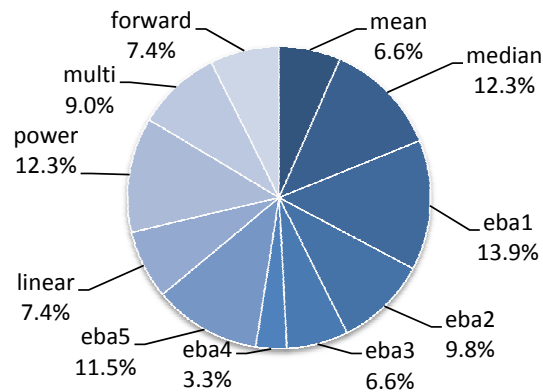


図 12: ISBSG : 各試行で最小の誤差で予測が行えた工数予測手法の割合

のは、予測誤差のばらつきが大きいためである。

以上より、本研究で提案する期待される精度に基づく工数予測手法の選択では、当該データセットにおいて平均的に小さい誤差で予測が可能な工数予測手法を適切に判定することが可能であるといえる。提案手法は、新しい環境で工数予測を行う場合など、どの工数予測手法を用いるべきかわからない場合に有用であると考えられる。

7.2 ツールの評価

ツールの実用性に関する評価として、ツールを試用して頂いた 3 社の企業の方へアンケートを実施した。アンケートは表 4 に示す操作性と機能性に関する項目のほか、自由記述としてツールの良い点と改善点を伺った。

表 4 より、操作性に関しては概ね良い評価を頂いている。ただし、用語のわかりにくさや CSV ファイルの作成方法の難しさに対する指摘もあった。CSV ファイルについては、メトリクスとして比例尺度と名義尺度にしか対応していない点、その上で任意のメトリクスを使用可能である点、欠損値補

完の方法が固定である点などを考慮して作成しなければならないが、これらの制約や自由を設けることで様々な工数予測手法に対応可能であるというメリットが生まれている。機能性に関しても概ね良い評価を頂いた。図 5 は、精度の計算頻度に応じたツールの実行時間に関する実験結果である。実験は過去プロジェクト数 121、工数の他のメトリクス数は比例尺度 10、名義尺度 3、ツールに実装済みの 11 の工数予測手法を用いて行った。実験環境は CPU が Intel Core 2 Duo 1.20GHz、メモリは 2GB、OS は Windows Vista Business SP2 である。Ruby という言語の特性上、実行時間は短いとはいえないが、正確な工数予測の重要性と精度の計算頻度の指定による時間短縮も可能である点も考慮に入れると、許容範囲であるとの評価を頂いたようである。

自由記述でのツールの良い点としては、主に以下の 3 点を挙げて頂いた。

- 導入が容易
- 任意のメトリクスを使用可能
- 独自の工数予測手法を追加可能

まず、本ツールは Ruby on Rails を用いて開発を行っているため、Instant Rails[30] を用いることで簡単に実行環境を構築することが可能である。ツールのマニュアルにはインストール方法の詳細が記載されており、通常の PC ですぐにスタンドアロンのサーバを構築して実行可能である点を評価して頂いた。また、他の工数予測ツールは特定の工数予測手法に特化した入出力を特徴としているため、使用するメトリクスも指定されていることが多い。それに対して本ツールは、組織ごとに独自に収集したメトリクスを使用可能であり、工数予測手法の追加機能と併せてより独自性の高い工数予測が可能である。

表 4: ツールに関するアンケート結果

アンケート項目	A 社	B 社	C 社
入力画面の操作性 (4 点満点)	2	4	4
出力画面の操作性 (4 点満点)	3	4	4
出力結果の満足度 (4 点満点)	4	3	3
実行時間の満足度 (4 点満点)	4	3	4

表 5: ツールの実行時間

精度の計算頻度	実行時間 [s]
1 プロジェクト毎	389.0
10 プロジェクト毎	42.1

反対に、改善点としては主に以下の 2 点を挙げて頂いた。

- CSV ファイルの作成にコツが必要
- 出力結果をどのように解釈して工数予測手法を選択すべきかわかりにくい

出力結果については、様々な情報をもとに最終的にユーザが工数予測手法を選択する自由度を持たせるため、本ツールではあえて 1 つの工数予測手法を推薦することは行っていない。4 つの精度の評価指標や、EbA で選択された類似プロジェクト、回帰分析の決定係数などを参考に、ユーザが総合的に判断する必要がある。3 つのデータセットを用いた提案手法の実験の結果、MMRE や MMER の最も良い工数予測手法を機械的に選択することで平均して小さい誤差で予測が行えたが、さらに多くの情報を出力する本ツールを利用する場合に具体的にどの工数予測手法を選択すべきかについては、今後の検討が必要である。

8 あとがき

本研究では、1 件の現行プロジェクトに対して複数の工数予測手法を同時に適用し、期待される精度に基づく工数予測手法の選択支援について提案した。期待される精度とは、LOO を用いて過去プロジェクトの工数を予測した際の誤差によって算出される値であり、予測の都度期待される精度の高い工数予測手法を選択することで、データセットの特徴に応じた工数予測が可能となる。本研究で開発した工数予測ツール「 e^3 」を用いると、簡単に複数の工数予測手法による期待される精度の比較が可能である。3 種類のデータセットを用いた実験の結果、本研究で提案する期待される精度をもとに使用する工数予測手法を選択することで、データセットごとに最も有効な工数予測手法を使用し続けた場合と同程度の精度で予測が行えた。したがって、期待される精度に基づく工数予測手法の選択は、初めての環境で工数予測を行う場合など、どの工数予測手法を使用したら良いかわからない場合に特に有用であるといえる。また、ツールの実用性に関するアンケートでは、試用して頂いた企業の方から良い評価を頂いた。

今後の課題としては、4 種類の期待される精度ごとの解釈方法の検討が挙げられる。ユーザが自身のニーズに応じて適切な工数予測手法を選択するためには、それぞれの期待される精度の指標の特徴を明らかにする必要がある。例えば、MMRE は実測値より予測値の方が大きくなる過大予測時に特に値が大きくなり、MMER は実測値より予測値の方が小さくなる過小予測時に特に値が大きくなる。したがって、過小予測を避けたいときは MMER の大きい工数予測手法を避けるべきと考えられるが、実際の効果を確認するためには新たな実験が必要である。また、期待される精度の他にも、EbA の類似プロジェクトや回帰分析の決定係数など、ツールの表示する様々な情報を総合して工数予測手法を選択するノウハウの確立が求められる。そのためには、今後も多くの実証的研究が必要である。

謝辞

本研究を行うにあたり、理解ある御指導及び的確な御助言を賜りました大阪大学大学院情報科学研究科コンピュータサイエンス専攻 楠本真二教授に深く感謝申し上げます。

本研究を行うにあたり、適切かつ丁寧な御指導を賜りました大阪大学大学院情報科学研究科コンピュータサイエンス専攻 岡野浩三准教授に深く感謝申し上げます。

本研究を行うにあたり、多大なる御助言を頂きました大阪大学大学院情報科学研究科コンピュータサイエンス専攻 肥後芳樹助教に深く感謝致します。

本研究の全過程を通し、丁寧な御指導を頂きました香川高等専門学校電気情報工学科 柿元健講師に深く感謝致します。

本研究にご協力頂きました、日本ファンクションポイントユーザ会 FP 法利用検討委員会主催のワークショップ参加者の皆様に感謝致します。

ツールの試用にご協力頂きました、株式会社エヌ・ティ・ティ・データ 鶴澤仁氏、クボタシステム開発株式会社 板垣庄治氏、日本アイ・ビー・エム株式会社 梶山昌之氏に感謝致します。

最後に、本研究の全過程を通してお世話になりました大阪大学大学院情報科学研究科コンピュータサイエンス専攻楠本研究室の皆様に感謝致します。

参考文献

- [1] M. Shepperd, and C. Schofield, “Estimating software project effort using analogies,” IEEE Transactions on Software Engineering, vol.23, no.12, pp.736–743, Nov. 1997.
- [2] K.D. Maxwell, Software Quality Institute Series: Applied statistics for software managers, Prentice Hall, NY, 2002.
- [3] S.J. Huang, and N.H. Chiu, “Optimization of analogy weights by genetic algorithm for Macro 1software effort estimation,” Information and Software Technology, vol.48 no.11, pp.1034–1045, Nov. 2006.
- [4] Y. Kultur, B. Turhan, and A.B. Bener, “ENNA: Software effort estimation using ensemble of neural networks with associative memory,” Proceedings of the 16th ACM SIGSOFT International Symposium on Foundations of Software Engineering, pp.330–338, Atlanta, GA, USA, Nov. 2008.
- [5] 情報処理推進機構ソフトウェア・エンジニアリング・センター, ソフトウェア開発見積りガイドブック～IT ユーザとベンダにおける定量的見積りの実現～, オーム社, 2006.
- [6] B. Boehm, Software Engineering Economics, Prentice Hall, 1981.
- [7] M. Shepperd, and G. Kadoda, “Using simulation to evaluate prediction techniques,” Proceedings of the IEEE 7th International Software Metrics Symposium, pp.349–358, London, UK, Apr. 2001.
- [8] E. Mendes, S.D. Martino, F. Ferrucci, and C. Gravino, “Effort estimation: How valuable is it for a web company to use a cross-company data set compared to using its own single-company data set?,” Proceedings of the 16th International World Wide Web, pp.963–972, Banff, Alberta, Canada, May 2007.
- [9] A. De Lucia, E. Pompella, and S. Stefanucci, “Effort estimation for corrective software maintenance,” Proceedings of the 14th International Conference on Software Engineering and Knowledge Engineering, pp.15–19, Ischia, Italy, Jul. 2002.
- [10] 生方克馬, 柿元健, 楠本真二, “複数の手法による予測結果が比較可能な工数予測ツールの開発,” 2010 年電子情報通信学会総合大会講演論文集, 情報・システム講演論文集 1, no.D-3-2, p.19, Mar. 2010.
- [11] 生方克馬, 柿元健, 楠本真二, “複数の手法による予測結果が比較可能な工数予測ツールの開発と評価,” 電子情報通信学会技術研究報告, vol.110, no.336, pp.7–12, Dec. 2010.

- [12] R. Jeffery, M. Ruhe, and I. Wiecek, “A comparative study of two software development cost modeling techniques using multi-organizational and company-specific data,” *Information and Software Technology*, vol.42, pp.1009-1016, 2000.
- [13] B.A. Kitchenham, and E. Mendes, “A comparison of cross-company and single-company effort estimation models for web applications,” *Proceedings of 8th International Conference on Empirical Assessment in Software Engineering*, pp.47–55, 2004.
- [14] E. Mendes, and B. Kitchenham, “Further comparison of cross-company and within-company effort estimation models for web applications,” *Proceedings of 10th IEEE International Software Metrics Symposium*, pp.348–357, Chicago, IL, USA, Sept. 2004.
- [15] I. Wiecek, and M. Ruhe, “How valuable is company-specific data compared to multi-company data for software cost estimation?,” *Proceedings of 8th IEEE International Software Metrics Symposium*, pp.237–246, Ottawa, Canada, Jun. 2002.
- [16] 角田雅照, 大杉直樹, 門田暁人, 松本健一, 佐藤慎一, “協調フィルタリングを用いたソフトウェア開発工数予測方法,” *情報処理学会論文誌*, vol.46, no.5, pp.1155–1164, May 2005.
- [17] 大杉直樹, 角田雅照, 柿元健, “ワンクリック見積&データ品質診断ツール Magi の紹介,” 奈良先端科学技術大学院大学, <http://se.naist.jp/magi/>.
- [18] 中村宏美, “CoBRA 法に基づく見積り支援ツール —プロジェクトの定量的見積り評価を易しく支援する Web ツールの提供—,” *SEC Journal*, vol.5, no.6, pp.377–379, Dec. 2009.
- [19] ソフト工学センター, “ソフトウェア見積ツール KnowledgePLAN,” 構造計画研究所, <http://www4.kke.co.jp/sec/service/01.html>.
- [20] L.C. Briand, K.E. Emam, and F. Bomarius, “COBRA: A hybrid method for software cost estimation, benchmarking, and risk assessment,” *Proceedings of the 20th International Conference on Software Engineering*, pp.390–399, Kyoto, Japan, Apr. 1998.
- [21] 中村哲彬, 柿元健, 楠本真二, “類似性に基づく工数予測における予測回避の効果,” *ソフトウェア信頼性研究会第6回ワークショップ論文集*, pp.37–44, Mar. 2010.
- [22] B. Sarwar, G. Karypis, J. Konstan, and J. Riedl, “Item-based collaborative filtering recommendation algorithms,” *Proceedings of the 10th International Conference on World Wide Web*, pp.285–295, Hong Kong, May 2001.
- [23] 田中豊, 垂水共之, 多変量解析, 共立出版, 東京, 1995.

- [24] T. Foss, E. Stensrud, B. Kitchenham, and I. Myrtveit, "A simulation study of the model evaluation criterion MMRE," *IEEE Transactions on Software Engineering*, vol.29, no.11, pp.985–995, Nov. 2003.
- [25] M. Azzeh, D. Neagu, and P. Cowling, "Improving analogy software effort estimation using fuzzy feature subset selection algorithm," *Proceedings of the 4th International Workshop on Predictor Models in Software Engineering*, pp.71–78, Leipzig, Germany, May 2008.
- [26] ISBSG Estimating, Benchmarking and research suite release 11, International Software Benchmarking Standards Group (ISBSG), 2009.
- [27] C. Lokan, and E. Mendes, "Applying moving windows to software effort estimation," *Proceedings of the 3rd International Symposium on Empirical Software Engineering and Measurement*, pp.111–122, Lake Buena Vista, FL, USA, Oct. 2009.
- [28] "Function point counting practices manual release 4.3.1," International Function Point Users' Group (IFPUG), Jan. 2010.
- [29] "Early function point counting," Netherlands Software Metrics Association (NESMA), <http://www.nesma.nl/section/home/>.
- [30] Instant Rails, <http://rubyforge.org/projects/instantrails/>.

A 工数予測手法の追加

工数予測ツール「e³」は、ユーザが比較的簡単に独自の工数予測手法を追加登録できるよう設計されている。ここでは、工数予測手法の追加方法の詳細について述べる。

A.1 概要

工数予測手法は、モジュールとその中のメソッドとして実装されている。工数予測手法を追加する場合、指定のディレクトリに指定の書式でメソッドを記述すれば良い。記述はプログラミング言語 Ruby を用いて行う。モジュールは工数予測手法の分類、メソッドは工数予測手法そのものに対応する。

メソッドのインターフェースは以下のとおり。

- メソッドの引数
メトリクス情報、プロジェクト名情報、プロジェクト情報（メトリクスの値）を格納した3つの配列
- メソッドの返り値
要素として「現行プロジェクトの工数の予測値」と「メモ（個別表示の内容）」を持つ大きさ2の配列

A.2 命名規則

ファイル名、モジュール名、メソッド名には命名規則があり、それに従わなければならない。ファイル名とメソッド名はスネークケース（全て小文字で、複合語をアンダースコアで区切った書き方）で、モジュール名はファイル名と同じ名前をアップーキャメルケース（全ての複合語の先頭のみ大文字で書き、区切り文字を使わない書き方）で記述する。例としては、ファイル名「sample_module.rb」、モジュール名「SampleModule」、メソッド名「sample_method」など。

A.3 ファイルの配置

1つのファイルは1つのモジュールであり、1つのモジュールの中に複数のメソッドとプライベートメソッドを記述することができる。本ツールでは、モジュールは工数予測手法の分類に対応し、その中のメソッドは工数予測手法そのものに対応する。モジュール中に記述したプライベートメソッドは、メソッド記述時のサブルーチンとして使用可能である。

手法記述ファイルは、以下の場所に配置する。

`estimator3/lib/my_methods/`

図 13 にファイルの配置場所と記述例を示す。

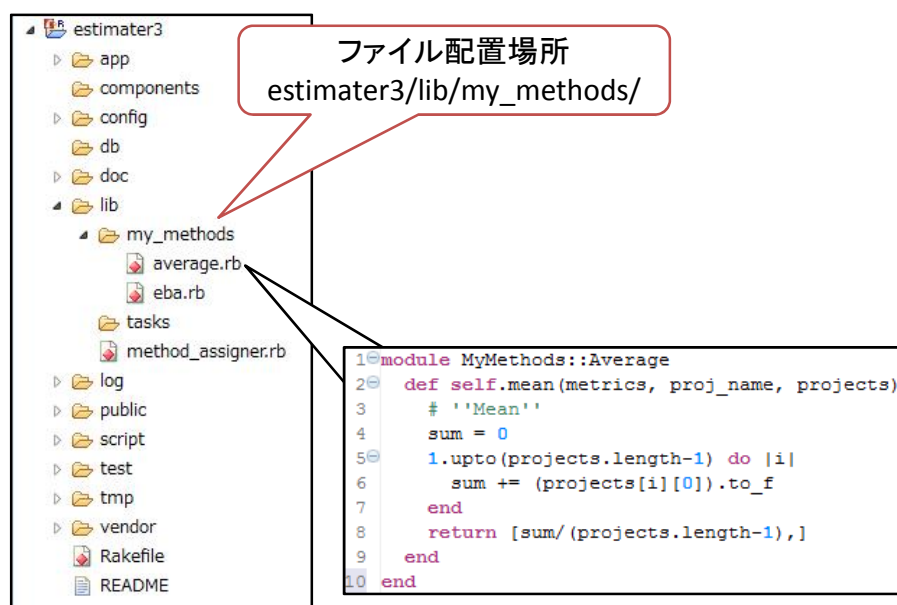


図 13: ファイルの配置場所と記述例

A.4 メソッドの引数

メソッドの引数として与えられる配列は，入力された CSV ファイルの内容を整形したものである．以降，図 14 の CSV ファイルが入力として与えられたときの各配列の例を示す．

- メトリクス情報の配列

メトリクス名とメトリクスの種類を格納した 2 次元配列である（図 15）．名義尺度のメトリクスの場合，ダミー変数化後の変数が格納される．メトリクスごとの候補となるカテゴリは，最初のダミー変数の後ろに保持されている．

- プロジェクト名情報の配列

プロジェクト名を格納した配列である（図 16）．1 つ目の要素に現行プロジェクトの名前，2 つ目以降に順に過去プロジェクトの名前が格納される．

- プロジェクト情報の配列

各プロジェクトのメトリクスの値を格納した 2 次元配列である（図 17）．現行プロジェクトの工数の値は `nil` または CSV ファイルに記載されていた値であるが，その他の要素は欠損値補完がされているので，欠損値はない．

sample2.csv	Effort	FP	Language	Team Size
ver.1	1	1	2	1
Current Project		220	Java	10
Sample Project A	1120	90	C	18
Sample Project B	2607	282	Ruby	11
Sample Project C	996	138	Java	9

図 14: CSV ファイルの例

Effort	1			
FP	1			
Language='Java'	2	Java	C	Ruby
Language='C'	2			
Team Size	1			

図 15: メトリクス情報の配列

Current Project
Sample Project A
Sample Project B
Sample Project C

図 16: プロジェクト名情報の配列

	Effort	FP	Language ='Java'	Language ='C'	Team Size
Current Project		220	1	0	10
Sample Project A	1120	90	0	1	18
Sample Project B	2607	282	0	0	11
Sample Project C	996	138	1	0	9

図 17: プロジェクト情報の配列

A.5 記述例

工数予測手法記述ファイルの記述例と各行の説明を以下に示す。ファイル名は `sample_module.rb` とする。

```
1: module MyMethods::SampleModule
2:   def self.sample_method1(metrics, proj_names, projects)
3:     # '' 手法その 1''
4:     ...
5:     return [result, memo]
6:   end
7:
8:   def self.sample_method2(metrics, proj_names, projects)
9:     # '' 手法その 2''
10:    ...
11:    if (予測不能)
12:      return [nil, memo]
13:    end
14:    ...
15:    return [result, memo]
16:  end
17:
18:  private
19:
20:  def self.private_method1(a, b)
21:    ...
22:    return result
23:  end
24: end
```

1 行目 1 行目にモジュール名を記述する。SampleModule 以外の部分は定型である。

2 行目 2 行目以降にメソッド（工数予測手法）を記述することができる。

3 行目 メソッドの記述を開始した次の行に、ブラウザに表示させる名前（日本語も可能）をコメント文中に記述する。シングルクォート（'）を 2 つ続けた間に書く。

5 行目 返り値は、大きさ 2 の配列である。配列の最初の要素は数字、2 つ目の要素は文字列でなければならない。

12 行目 予測不能と判断した場合、返り値の配列の最初の要素を `nil` として返せば良い。このとき、2 つ目の要素としてエラーメッセージを持たせると、ブラウザで表示することができる。ただし、精度計算のための予測時のエラーの場合は、この個別エラーメッセージは表示されない。

18 行目 `private` と記述した行以降のメソッドはプライベートメソッドであるため、工数予測手法とはみなされない。プライベートメソッドは、同一モジュール内のメソッドから呼び出すことができる。