

特許間の技術的距離と M&A 後におけるイノベーションに関する分析

玉川 希[†] 高橋 大志[†]慶應義塾大学大学院経営管理研究科[†]

1. 背景

近年、多くの企業戦略で行われている M&A においては、主要な目的の一つとして技術的な成果が挙げられる。そして、その中でもイノベーションに着目した、M&A がイノベーションに与える影響については様々な議論が行われており、その重要性が窺える[1][2]。しかし、先行研究で用いられている特許の出願数や、国際特許分類、産業分類などの一定の基準に基づいて計測された指標のみでは、企業が保有している技術の詳細を加味できていない可能性があるという課題が存在している。そのため、本分析では非構造化データである特許文書を解析することで技術の詳細を加味した分析を試みる。具体的には、先行研究を参考に、M&A 前のイノベーションの規模を示すナレッジベースに焦点を当てて、以下の仮説について検討を行う[1]。

仮説 1：案件公表前において、買収企業に対する被買収企業の相対的なナレッジベースが大きい場合には、M&A の公表後のイノベーションアウトプットが減少する。

仮説 2：案件公表後においてイノベーションアウトプットが減少する M&A は、買収企業と被買収企業の技術的な類似度が高い傾向にある。

2. データ

本章では、使用するデータについて示す。本稿では、特許データと M&A の案件データを使用する。まず、特許データは Derwent World Patent Index (以下、DWPI) における 1970 年から 2015 年までに公報発行が行われている、日本、アメリカ、ドイツの特許データを活用する。次に、M&A 案件のデータについては、Refinitiv Eikon から取得する。以下に、分析対象とする M&A 案件の抽出方法について、3 つのステップに分けて示す。まず、ステップ 1 では、1990 年から 2012 年までに案件公表が行われており、日本市場で上場企業の間で行われた M&A 案件を抽出する。加えて、その中でも Refinitiv Eikon 内にて取引形態が「Merger」、「Acquisition of Majority Assets」、案件ステータスが「Completed」に分類されている案件を抽出する。次に、ステップ 2 では買収企業と被買収企業が

どちらも DWPI にて標準コードを付与されている案件を抽出する。それに加えて、案件を公表した年からその 3 年前までに買収企業と被買収企業において少なくとも 1 件以上の特許が公報発行されている企業のみを抽出する。そして最後に、ステップ 3 では M&A 完了後の買収企業の持分が 50%以下となっていた案件と、M&A にて取得した被買収企業の株式比率が 10%以下の案件を排除し、103 件の M&A を最終的な分析対象とする。

3. 分析手法

本章では、イノベーションアウトプットの測定、技術的距離の測定、ナレッジベースの規模の測定について示す。

3.1. イノベーションアウトプットの測定

本稿では、BENA and LI(2014)を参考に、日本の特許データを用いて M&A 案件ごとの *Patent Index* を算出し、これをイノベーションアウトプットの指標として用いる[2]。計測する期間は、案件の公表が行われた年(以下、*ayr*)を基準に、3 年前である *ayr-3* から 3 年後の *ayr+3* までの 7 年間における各年の *Patent Index* を計測する。そして、最後にそれぞれの M&A の各年における *Patent Index* の値を、7 年間の *Patent Index* の合計値で除することによって、7 年間の合計に対する各年(*ayr-3* から *ayr+3* まで)の *Patent Index* の割合を表す値を算出する。

3.2. 技術的距離の測定

本分析では、特許文書ベクトルを用いて測定した、買収企業と被買収企業の技術的距離を、技術的な類似性の指標として用いる。

まず、特許文書データに対して自然言語処理技術を用いることで、特許文書ベクトルを算出する。その際には、DWPI における抄録の内容を、Mekala et al. (2017)を参考に Sparse Composite Document Vector を用いて、ベクトル化を行った[3]。そして、M&A 取引前に買収企業と被買収企業が出願し、公報発行が行われた特許の文書ベクトルの重心間距離をそれぞれ算出した。そして、この値を各 M&A 案件の技術的距離として用いた。また、それらの特許が引用している、日本、アメリカ、ドイツにおける過去の特許群を抽出し、それらの重心間の距離を計測した値

を、引用特許間の技術的距離として用いる。尚、全分析対象である 103 件の中で、6 件は M&A の公表前の引用特許を検出できなかったため、分析対象から除外した。

3.3. ナレッジベースの測定

本分析においては、企業が M&A の公表前に広報発行を行った特許数を、M&A におけるナレッジベースの大きさの指標として用いる。そして、相対的なナレッジベースの大きさに加えて、取引形態を用いて M&A を分類する。具体的には、買収企業に対して被買収企業の相対的なナレッジベースが小さい買収案件の Acquisition Group 1 と、合併案件の Merger Group 1 に分類する。また、被買収企業の相対的なナレッジベースが大きい買収案件の Acquisition Group 2 と、合併案件の Merger Group 2 に分類する。分類の結果、Acquisition Group 1、Acquisition Group 2、Merger Group 1、Merger Group 2 に分類された案件数はそれぞれ、36 件、8 件、40 件、13 件となった。

4. 分析結果

図 1 は、グループごとに分類されている全ての M&A と、分析対象全体の *Patent Index* の割合の平均値を示している。図の横軸は、*ayr* を基準として前後 7 年間の時間を示している。図における縦軸は、7 年間の *Patent Index* の合計値に対する各年の *Patent Index* の割合を示している。

図 1 より、被買収企業の相対的なナレッジベースが大きい合併案件のグループ Merger Group 2 は、案件公表前後において *Patent Index* が低下する傾向にあることが確認できる。また、その他のグループについては、概ね分析対象全体の平均と同様に推移していることが確認できる。この点について、各グループ間の買収後のイノベーションアウトプットの割合について平均値の差の検定を行うと、Merger Group 2 とその他のグループでは統計的に有意な結果が示された。

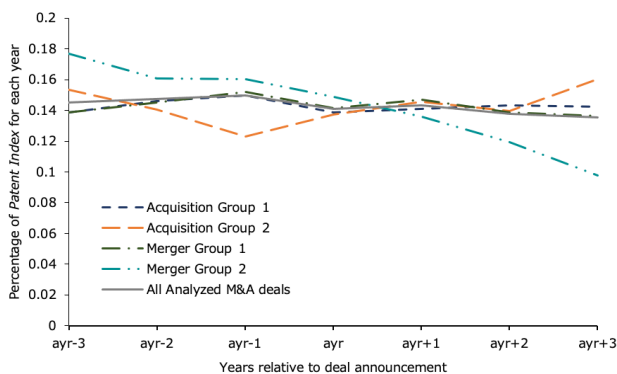


図 1. M&A 前後のイノベーションアウトプット

次に、各グループの技術的な類似性を比較すると、Merger Group 2 が技術的距離と引用特許間の技術的距離において最も小さい値を取ることが確認できた。また、差の検定を行ったところ、特に引用特許間の技術的距離については他のグループとの統計的に有意な差がみられた。

5. 考察

本章では、上記の分析結果についての考察を行う。分析結果から、被買収企業の相対的なナレッジベースが大きい合併案件である Merger Group2 のイノベーションアウトプットが案件公表後に低下している点については、仮説 1 と整合的な結果といえる。また、Merger Group2 は、特に引用特許間の技術的距離において他のグループに比べて最も小さい値であることが確認できた点は、仮説 2 と整合的な結果であると言える。これらの結果から、Merger Group2 の技術的距離が近いことから、M&A 後に重複する技術を削除したことによってイノベーションアウトプットが低下した可能性などが考えられる。

6. まとめ

本稿では、大規模な特許データを活用し、計測した M&A の公表前後のイノベーションアウトプットを用いて、買収企業と被買収企業の相対的なナレッジベースと技術的な類似性に焦点を当て、M&A の公表後のイノベーションアウトプットに与える影響について分析を行った。分析の結果、被買収企業の相対的なナレッジベースが大きい合併案件は、案件公表後にイノベーションアウトプットが低下する傾向があり、買収企業と被買収企業の技術的類似性が高いことが確認できた。より詳細な分析については、今後の課題である。

参考文献

- [1] Ahuja, G., Katila, R: Technological acquisitions and the innovations performance of acquiring firms: a longitudinal study. *Strategic Management Journal* 22(3), 197-220, 2001.
- [2] Bena, J., Li, K: Corporate innovations and mergers and acquisitions. *The Journal of Finance*, 69(5), 1923-1960, 2014.
- [3] Mekala, D., Gupta, V., Paranjape, B., Karnick, H: SCDV: Sparse Composite Document Vector using soft clustering over distributional representations, Association for Computational Linguistics. *Proceedings of the 2017 Conference on Empirical Methods in Natural Language Processing*, 659-669, 2017.