

**The innovative characteristics of
target firms:
An analysis of Swedish data from
1998 to 2014**

Patrik Gunnvald ¹ & Victor Sylwander ²

Bachelor Thesis in Finance

Stockholm School of Economics

Tutor: Cristina Cella

May 18, 2015

¹22625@student.hhs.se

²22626@student.hhs.se

Abstract

This paper investigates firm-specific innovative characteristics of acquired companies using descriptive statistics and regression analysis. This is done on a data sample on Swedish acquisitions among listed firms from the time period 1998 to 2014. The sample was narrowed down to include only transactions where the acquirer changed from holding a minor stake to a major stake and thus having controlling interest. Descriptive statistics is used to find differences in innovation output and property between the target firms and firms not participating in M&As. The probit regression is used to find the probability of becoming a target firm given certain innovative characteristics. A discussion of assumptions and sample bias is made. The analysed characteristics are R&D expenses over total sales and intangible assets (adjusted by extracting goodwill) over total assets. The quotas are used in order for us to compare firms with different size. Results are robustness checked and the model's fit is assessed. First, we show that R&D expenses lowers the probability of a firm becoming a target. Second, we show that the target firms within health care, technology and consumer discretionary have higher quotas of intangible assets compared to the non-target benchmark within the same industry. The targets' innovative property was however lower than the benchmark for the remaining industries. Industry- and time effects are controlled for in the regression but also discussed separately.

Contents

1	Introduction	1
1.1	Background	2
1.2	Aim of the paper	5
1.3	Limitations	5
2	Theoretical framework	7
2.1	Previous research	7
2.2	Mathematical background and application	10
2.2.1	Briefly about hypothesis testing	10
2.3	Regression analysis	12
2.3.1	Linear regression	12
2.3.2	Probit regression	18
3	Data	24
3.1	Acquisition of data	24
3.2	Processing of data	25
3.3	Final sample	26
4	Methodology	27
4.1	Choosing relevant target data	27
4.2	Characteristics	27
4.2.1	Research & Development	28
4.2.2	Intangible assets	29
4.3	Descriptive statistics	31
4.4	Regression model selection	31
4.4.1	Controls	33
4.4.2	Possible sample selection biases	34
4.4.3	Final probit model	35
4.4.4	Measurements of fit	35
5	Results	36
5.1	Characteristic results	36
5.1.1	R&D Expenses	36
5.1.2	Intangible assets	39
5.2	Industry and time effects	42
5.2.1	Time effects	42

5.2.2	Industry effects	45
5.3	Robustness	48
6	Conclusions & discussion	50
6.1	R&D Expenses	50
6.2	Intangible assets	52
6.3	Industry, size and time effects	54
7	Suggestions for further studies	56
8	Appendix A, Tables	58
9	Appendix B, Correlation	62
10	Appendix C, Code	64
11	References	66

1 Introduction

There are many reasons for companies to conduct M&As. It could be to smother competition, to share costs or to get hands on some recent development in technology. However, in a market there are always trends. To be able to chart these trends could prove valuable for shareholders, investors or other market participants. While studying M&As, there are many different patterns one could examine. Macroeconomic factors and post merger firm performance is something that researchers constantly investigate. In this study however, the aim is to map the intangible property and research & development expenses of Swedish firms in order to more accurately distinguish the target firms. As technology, process improvements and innovation becomes more of a differentiating factor with modern firms, we found it interesting to investigate if this has an affect on the M&A landscape. The fact that Sweden is often placed at the very top when it comes to R&D and innovation was another argument for us choosing the topic (Bloomberg, 2015). According to a recent article in SvD, the Swedish government has increased the research- and innovation budget by almost 10 billion per year during the last decade (SvD, 2015). On the other hand, statistics from OECD show that the gross domestic spending on R&D has decreased slightly (0.4 percentage points) during the last 15 years (OECD Data).

With the aim to study innovative property and its affect on Swedish M&A's, this study is divided into two main areas; the targets' R&D

expenses and intangible assets. To analyse the target firms and the potential innovative characteristics within these firms, descriptive statistics and probit regression models are used. The probit regression is useful when the outcome of our interest is binary; the innovation expenses and properties potential influence on the firm becoming a target or not. Control variables for years and industries are used to remove potential biases in the regression.

The paper is structured in the following manners: To begin with, a general background on the Swedish M&A landscape and the innovation development will be presented. Secondly, previous research on the topic and a mathematical background of hypothesis testing, ANOVA test and regression analysis is provided. A description of the data is then followed by the methods used in the study. Assumptions of the characteristics and how the descriptive data and regression analyses are applicable on these items are presented in this section. The results section contains both descriptive results as well as regression results. Finally, the conclusion section summarizes the results, discusses implications and potential further research on the topic.

1.1 Background

In Corrado and Hulten's article in the Conference Board, reports from US-based firms show that investments in intangible assets and R&D activity have almost doubled during the last two decades. The rapid expansion is described as a key feature of the recent US economic growth.

In an article in Bloomberg, Peter Coy describes the new look on research & development expenses and intangible assets, as they nowadays are included in the gross domestic product of the US. As firms like Facebook, Twitter and Google are the key players of today's economy; the promotion of innovative activity becomes more important. In our research, Swedish M&As between 1998-2014 will be analyzed, and the characteristics of target firms will be explained. Numerous studies of M&A characteristics and drivers have been made before but our approach will be concentrated to the innovation climate in Sweden.

In the Journal of Finance, Bena and Li finds that, in the US, there are some technological and innovative overlaps between acquiring and target firms. Their findings are focused on number of patents and R&D costs. Further, they evaluate how these asset complementarities affect the outcome of the merger. Inspired by this article, the focus of this paper will be characteristics that distinguishes target firms in Swedish M&As. Since the US-study, by Bena and Li, showed significant results (e.g. targets having synergies with the acquirer) we found it interesting to chart and extend their study, but focusing on the Swedish market and its innovation landscape which was not touched upon in their study.

In the last 17 years, 1257 mergers, acquisitions and changes in minority interests has been conducted between Swedish firms. Of these 1257 observations, we only kept deals where the acquirer owned less than 50% prior to the deal and more than 50% after the deal. This makes

sense when looking thoroughly into the data as many deals are share buybacks, minority stakes or capital increases. Including smaller capital increases would probably skew our results as firms perform these on daily basis. Further, the M&As studied are the ones where both actors are Swedish firms, mainly because trends differ between different markets. Only listed and delisted target firms were kept as retrieving detailed financial information of private firms are sometimes very difficult. We also want to reduce the noise due to difference in regulation, corporate environment etc. that may arise when including acquirers acting in foreign markets. Finally, compared to the rest of Europe, the Swedish M&A market has remained fairly stable the last 10 years (Vinge, 2013).

The transactions in the Eurozone has remained fairly stable during the last 10 years, except some years of peaking numbers. Most deals in Sweden were made during 1999 and this could perhaps be explained by the Information Technology bubble during the turn of the century but the Eurozone debt crisis during 2009-2011 caused a slight drop in transactions. The sectors with most transaction activity are healthcare, energy, telecoms and media. All these sectors are driven by both private equity firms and corporates. With the financial market changing, the bids from the private equity firms dropped, facilitating the corporate acquisitions. However, the private equity firms dominated the market in the years of 2005-2008, but changing availability on price and financing as well as lack of experience in volatile times is a chal-

lenging factor for the private equity firms. Macroeconomic factors are probably affecting the M&A's in Sweden and we will therefore control for these variables when performing our tests.

The goal is to find intangible property and innovation expenses within the target firms and analyse if these could affect the probability of the firm being acquired. These characteristics would, as previously stated, be essential to keep track of the actors on the Swedish market. Our findings could help these actors to foresee upcoming as well as charting the Swedish innovation-driven M&A's (KPMG, 2014).

1.2 Aim of the paper

The aim of this paper is to study firms' innovative characteristics and whether that affect firms prospects of becoming a target in a M&A transaction. The focus is on the target of the transaction as this is relevant from a shareholder perspective since targets in M&A can be bought at a substantial premium (Evans and Mellen, 2010). These innovative characteristics are divided into two main areas: R&D expenses and intangible assets.

1.3 Limitations

The study is conducted on Swedish firms, where both the target and the acquirer are Swedish. This selection is made for several reasons; Swe-

den prominent role as a innovative country, the greater acknowledge of firms intangible assets during the last decade, as well as mitigating the risk that the firm was bought by a non Swedish firm for other reasons such as tax rules and regulations.

Furthermore this paper only deal with transactions where an acquirer went from a minority to a majority stake, as this makes it more plausible that the acquirer via controlling interest of the target firm is interested in its innovations such as R&D, patents software or similar.

2 Theoretical framework

2.1 Previous research

The previous literature and research within M&As have been important in the shaping of our work. The main topic in earlier work (Harford, 2005) has been macro-level changes and characteristics, as high industry liquidity or technological shocks, that affects the amount of mergers and acquisitions. They conclude that the liquidity component causes M&A activity to cluster, even though industry shocks do not occur. Further on Maksimovic, Phillips and Yang (2012) compare public and private firms and their activity during and after M&A waves. For instance, the public firms acquisitions realize higher productivity gains when their stock is liquid and highly valued.

Many research papers are examining the drivers of M&As on a macro-level and also the performance and results of firms after certain industry shocks and waves of M&A activity. However, few papers examine the innovative firm characteristics and their impact on the possibility of upcoming M&As. Bena and Li studied the innovation and technology overlap between firms by using a patent-merger data set in their article *Corporate Innovations and Merger and Acquisitions*, written 2011. The paper is confined to the US market and includes data until 2006 and they examine the relationship between the acquirer and the target and what combination of technological overlap that creates M&As. The focus in their paper lays within the information asymmetry between

participating firms and they test this by using a dataset of patents and R&D expenses.

Asset complementarities between M&A participants are discussed in (Rhodes-Kropf and Robinson (2008)) where they find a relationship between acquirer and target market-to-book ratios. (Hoberg and Philips (2010)) examines asset similarities and find that post-merger stock returns are higher when the target is less similar to the acquirer's closest rivals. (Phillips and Zhdanov (2013)) uses R&D, acquisitions and firm size to determine if firms are increasing their R&D if they are considered becoming a target.

In Webb's *How to acquire a company*, a number of reasons for acquisition are brought up. Management, diversification and corporate psychosis are brought up but few traits are measurable. Organic growth is stated as an unquestioned trait of acquisition but it is more difficult to measure. In our study we focus on innovative property and if this can lead to a firm becoming a target.

Corporate governance changes and relative valuations are studied in *Shleifer and Vishny: Stock market driven acquisitions, 2001* as well as in *Holmstrom and Kaplan: Corporate governance and merger activity in the US; making sense of the 1980s and 1990s, 2001*. They find that the changes in the managerial climate affected the M&A activity in the US as well as that higher stock market valuation periods are followed

by merger waves.

In conclusion, one can say that previous literature aims their research on industry and macro-level impact on M&As and the performance and results that follows after the deal. Our study includes past innovation output as we also analyse the intangible assets. Using a descriptive statistics and a probit regression on solely Swedish firms, we provide new information applicable to the Swedish innovation development and M&A climate.

Our paper is closely related to Bena and Li's study but instead of analysing the information asymmetry and overlap between M&A participants, we investigate the innovative characteristics within Swedish target firms. We are also using data until 2014, which means we will capture the period when firms were highly valued based on their intangible assets. Sweden is also argued to be one of the most innovative countries in the world according to the Global Innovation Index ranking, which increases the relevance of our study. Hopefully, our findings will contribute and add sense to the research of the Swedish M&A landscape and also show if actors could foresee upcoming M&A's in the future.

2.2 Mathematical background and application

2.2.1 Briefly about hypothesis testing

When one shall test a hypothesis one usually have a so called *null hypothesis*, H_0 and an alternative hypothesis, H_1 . Then one decides upon a *significance level* usually denoted α . Common significance levels are 0.10, 0.05 and 0.01 respectively.

We have chosen $\alpha = 0.05$ in this study for our tests, this means that we at most can accept a risk of 5% of committing a so called *type I error* - rejecting the null hypothesis when it is in fact true (Stock and Watson, 2012). Thus any outcomes from tests where p-values are above 0.05 will make us keep the null hypothesis.

The null hypothesis is usually stated such that "there is no difference", "the variable does not have any explanation value" or "the slope parameter is zero" in hypothesis testing or linear regressions. Note that if one wants to test the null hypothesis that two mean values of firm A and firm B are equal, i.e. $\mu_A = \mu_B$, this can be reformulated as $\mu_A - \mu_B = 0$ ("there is no difference").

Depending on what the aim is to test, you can perform a *one-tailed (one-sided) test* or a *two-tailed (two-sided) test* and the null hypothesis has to be set up accordingly. Again considering two samples from our firms A and B, a typical one-tailed test will be "mean of firm A is

less than mean of firm B”, with the following hypotheses:

$$H_0 : \mu_A < \mu_B \quad \text{against}$$

$$H_1 : \mu_A > \mu_B$$

A two-tailed test is generally ”something is *equal* to...” instead of ”less/-more than”. Again with our two firms as an example you would want to test ”mean of firm A is equal to mean of firm B”, with the following hypotheses:

$$H_0 : \mu_A = \mu_B \quad \text{against}$$

$$H_1 : \mu_A \leq \mu_B$$

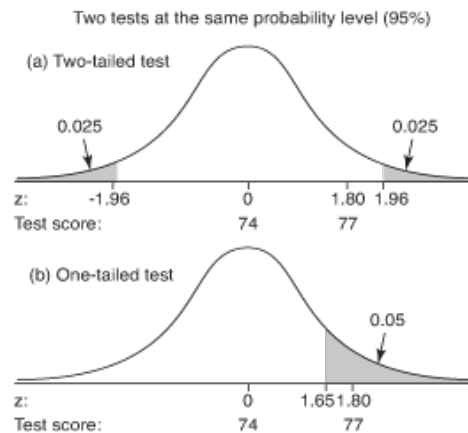


Figure 1: An example showing the difference between one and two-sided tests

2.3 Regression analysis

This section will briefly discuss some mathematics and terminology of the linear and probit regression models. The linear regression is a good starting point to explain and understand as much of the terminology and characteristics are transferable to the probit regression model.

2.3.1 Linear regression

A linear regression model is:

$$Y_i = \beta_0 + \beta_1 X_{1i} + \beta_2 X_{2i} + \dots + \beta_k X_{ki} + e_i, \quad (1)$$

where

i is the number of observations, $i = 1, \dots, n$

Y_i is the *dependent variable*, or *explained variable*;

X_{ji} is the *independent variable*, *explanatory variable*, *covariate* or *control variable*,

j is the number of covariates, $j = 1, \dots, k$

β_0 is the *intercept*;

β_j are the *slope parameters*, *regression coefficients* or *beta coefficients*, note that some simply use the word *coefficients* but in this paper the term *beta coefficient* will be used.

e_i is the error term related to observation i .

Shortly about the interpretation of this equation:

- Y_i is the observation of something e.g. wage, the price of a car,

the score on a test etc.

- X_{ji} are the variables that should explain the changes in Y_i .
- If Y_i is the observed (*wage*) in a sample, then one would probably like to include explanatory variables, or control variables, (*age*), (*years of education*) and more. Note that including other (uncorrelated) control variables "purges" the effect of e.g. (*age*) on (*wage*). If one would not include (*years of education*) as a control variable, the beta coefficient for (*age*) might be unexpectedly high and have a bigger impact on (*wage*) when most of the effect is due to education rather than sheer age. This is a form of omitted variable bias.
- The beta coefficient β_m represents the expected change in Y_i for one unit change in X_m given all other control variables $X_j, j \neq m$ are held constant.
- As is evident from Equation 1, the intercept β_0 represents the expected value of Y_i given $X_j = 0$ for all j
- e_i is the error term for observation i , or the part of Y_i that is still *unexplained* in the equation given the values of all $\beta_j X_{ji}$ and β_0 .

Note that sometimes the true (unobserved) model is denoted with β as coefficients while the *estimated* beta coefficients are denoted $\hat{\beta}$. Also the error e_i refers to the difference between the observed Y_i and the (unobservable) true model while the difference between Y_i and the *es-*

estimated model is called *residual* and usually denoted $\hat{u}_i$.

The model is estimated using the Ordinary Least Squared (OLS) Estimator which minimizes the sum of squared residuals (Stock and Watson, 2012). More formally we have that the OLS estimators $\hat{\beta}_0, \hat{\beta}_1, \dots, \hat{\beta}_k$ are the values $b_0, b_1, \dots, b_k$ that minimizes

$$\sum_{i=1}^n (Y_i - b_0 - b_1 X_{1i} - \dots - b_k X_{ki})^2, \quad (2)$$

and the residual, $\hat{u}_i$ is as earlier mentioned defined as

$$\hat{u}_i = Y_i - \hat{Y}_i, \quad \text{where} \quad (3)$$

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 X_{1i} + \dots + \hat{\beta}_k X_{ki} \quad (4)$$

For linear regression analysis and specifically the OLS estimation to be valid, some assumptions must hold (these are formulated a bit differently in different textbooks, these are mainly from Kennedy and Lang).

Underlying assumptions

1. Y_i and X_{ji} are independent and identically distributed across observations and Y_i have a linear relationship with X_{ji} . As formulated in Kennedy p. 93 "the conditional expectation of the dependent variable is an unchanging function of known independent variables. It is usually referred to as the model specification".
2. The expected value of the error term is zero, given the covariates,

$$\text{i.e. } E[e_i | X_{1i}, X_{2i}, \dots, X_{ki}] = 0$$

3. The error terms are independent and homoskedastic,
 $Var[e_i | X_{1i}, X_{2i}, \dots, X_{ki}] = \sigma^2$. Note that given assumption 2
 this could be rewritten as $E[e_i^2 | X_{1i}, X_{2i}, \dots, X_{ki}] = \sigma^2$
4. The error terms are identically distributed, usually an assumption
 of normally distributed error terms is made.
5. No perfect multicollinearity

Common reasons for violating the assumptions

Kennedy elaborates further on common reasons to when these assumptions are violated and suggests practical remedies, for the interested reader we refer to this textbook for further details. However it worth mentioning some common setups where violations can be made. This is something that both we as modelers have to keep in mind when constructing the model but it is also beneficial for the reader to think about when reading this or other reports, as well as constructing own models.

Violation of assumption 2: is called *endogeneity* (Lang, 2013) and technically means that the error term is correlated with at least one of the control variables. This might be due to several reasons such as

- *Sample selection bias*, where the selection of data incurs a bias,
 e.g. if I want to examine the average income in Stockholm and
 choose "random" people of the street, but the street itself might

be in an upper class area. Another common bias is the *self selection bias*, which is reasonable to occur in e.g. sales of third party insurances. People buying those insurances are probably more prone to accidents and mishaps than a sample of the entire population.

- *Simultaneity*, when not only X affects Y (which is desirable), but also Y affects X , an example would be if Y is supply and X is demand.
- *Omitted variable bias*, explained earlier.
- *Measurement errors*.

Remedies include; reformulating the model, finding the omitted variable, adjusting for measurement errors (easy if it is a constant error) or employing an *Instrumental Variable* (IV) which is highly correlated with the endogenous variable (preferably perfect correlation, $\text{corr} = 1$) and almost uncorrelated (preferably uncorrelated, $\text{corr} = 0$) with the error term.

Violation of assumption 3: is called heteroskedasticity, which means that the variance is dependent on X_i , i.e. the conditional distribution of u_i given X_i is *not* constant. Heteroskedasticity is actually reasonable to assume in real world data. A remedy is to apply (White's) robust standard errors.

Violation of assumption 4: can be checked by calculating the standard error for each observation and then graph the distribution of the standard errors and compare it to a normal distribution.

Violation of assumption 5: when the intercept and covariates or control variables are linearly dependent. It can appear if one is not careful when implementing dummy variables, i.e. variables that are only assigned a 1 or 0. Consider again the example where Y_i is the wage. A modeler might expect the wage to be determined by gender and so adds (*male*) and (*female*) as control variables. The problem is that if we know that the person is a male so that (*male*) = 1, then we can directly conclude that (*female*) = 0. Then beta coefficients cannot be determined uniquely since adding an arbitrary number c to the beta coefficients of both (*male*) and (*female*) and subtract it from the intercept β_0 we will get the same residual. Thus OLS estimation will not work. Multicollinearity is often spotted by very high standard errors or beta coefficients (Lang, 2013).

As a final note there is a measurement of goodness of fit, R^2 , which tells how well the structural equation (right hand side of Equation 1) explains the variation in Y , $0 \leq R^2 \leq 1$ where 0 is no explanation and 1 is complete explanation and the underlying true model is found (unrealistic in real world). However adding more control variables *always* increases R^2 which might lead to overly complex models (Stock and Watson, 2012). Instead it is desirable to use an adjusted R^2 , or

R_{adj}^2 . This measurement of fit penalizes addition of new covariates so that the part of increased explained variance from the covariate must be higher than the penalty for adding it. Mathematically we have:

$$R^2 = \frac{ESS}{TSS} = 1 - \frac{SSR}{TSS}, \quad \text{where} \quad (5)$$

$$ESS = \sum_{i=1}^n (\hat{Y} - \bar{Y})^2, \quad TSS = \sum_{i=1}^n (Y_i - \bar{Y})^2 \quad (6)$$

and $\bar{Y}$ the sample mean.

$$R_{adj}^2 = 1 - \frac{n-1}{n-k-1} \frac{SSR}{TSS} \quad (7)$$

where n is the sample size and k is the total number of control variables.

2.3.2 Probit regression

There are several real world scenarios where the outcome of interest is binary. This could be for example if a person becomes ill or not, if a person buys a product or not, or if a person passes an exam or not (where the score is "irrelevant" as long as they pass, e.g. Swedish drivers license exam). It is natural to assign the dependent variable a binary outcome, such as 1 for "buy", "pass" or 0 for "not buy", "fail" and similarly. The name "Probit" stems from the words *Probability* and *unit*. Most of the terminology, notation and assumptions in the linear regression also applies to the probit regression which is one of the reasons it is covered to the extent it is. However there are some

important differences which will be covered later in this section.

Linear probability model

The linear (probability) regression to a binary event is still possible. Using the same terminology as in that of the linear regression, consider Equation 1. As shown in Stock and Watson (2012) in a binary setting the expected value of Y_i conditional on X_{ji} becomes:

$$Pr(Y = 1 | X_1, X_2, \dots, X_k) = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_k X_k \quad (8)$$

Note that this implies a change in interpretation of the beta coefficient, namely that the beta coefficient represents the *change in probability* for one unit change of the control variable, holding the others constant.

However applying a linear probability regression brings forth questions. First and foremost that a probability outcome should be confined in the interval $I = [0, 1]$. A linear probability model will allow for Y -vales above 1 or below 0, which is not consistent with a probability measure. This becomes evident when the regression line is fitted to the data. As an example Stock and Watson (2012) have run a linear probability regression on applications for mortgages and whether they are denied a mortgage. When the mortgage is denied an observation $Y = 1$ is made and thus when it is approved the observation $Y = 0$ is made. The covariate used is the payment-to-income, or (P/I) , ratio.

As is evident from Figure 2 even though a vast majority of the obser-

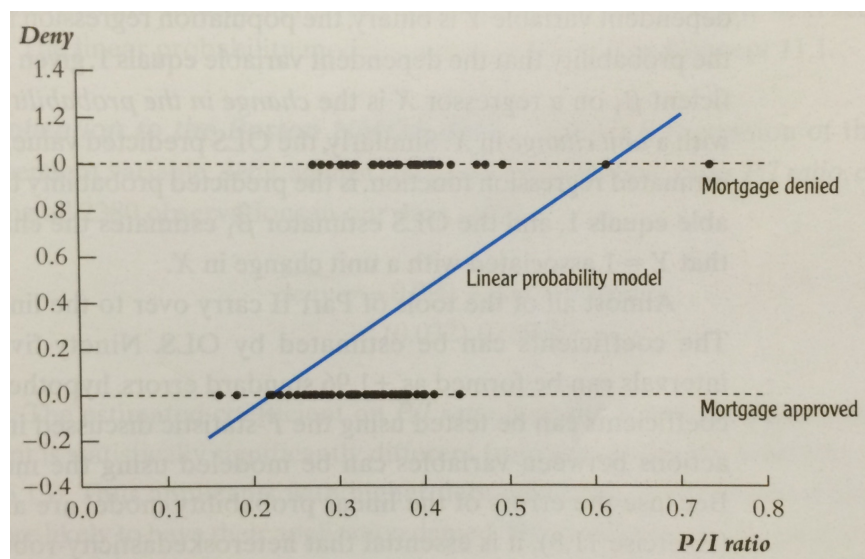


Figure 2: The linear probit model, Figure source: Stock, Watson p. 425

variations will render a probability value in the interval I , a P/I ratio below approximately 0.2 gives a negative probability (!) value and a P/I ratio above approximately 0.6 will yield a probability value above 1. The fact that this is possible is a troublesome flaw of the model.

Probit model

The probit model deals with this problem of probabilities being outside of the allowed interval for a probability measure. The probit model uses a *link* that maps all outcomes or observations into the interval $I = [0, 1]$, for the probit model this link is the cumulative standard normal function (there are other models such as the logistic, or logit, model that uses other links). The probit model does not use the OLS Estimator but a Maximum-Likelihood Estimator (MLE) to find the beta coefficients. It does require the same assumptions as stated for

the linear regression (Garson, 2012).

Difference in assumptions:

The difference between the OLS linear regression model and the probit model is that it does *not* assume a linear relationship between Y and X , homoskedasticity or normally distributed variables. The MLE is consistent and normally distributed in large samples, this means that t-statistics and confidence intervals can be computed as in the linear regression case.

Using the same notation we have that the Probit model is:

$$Pr(Y = 1 \mid X_1, X_2, \dots, X_k) = \Phi(\beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_k X_k) \quad (9)$$

where $\Phi(z)$ is the standard normal cumulative function, $z \in [-\infty, +\infty]$.

This is shown in Figure 3.

Shortly on the interpretation of the probit model:

- The function $\Phi(z)$ maps any real value into the interval $I = [0, 1]$.
- The left hand side is the probability of the event $Y = 1$ *in a sample*. Consider the event $Y = 1$ representing a person getting a job. So an estimated probability $Pr(Y = 1 \mid X_1, X_2, \dots, X_k) = 0.6$ means that in a large sample, 60% of the people with the characteristics that generate a z-value that gives $\Phi(z) = 0.6$ is expected to get a job.

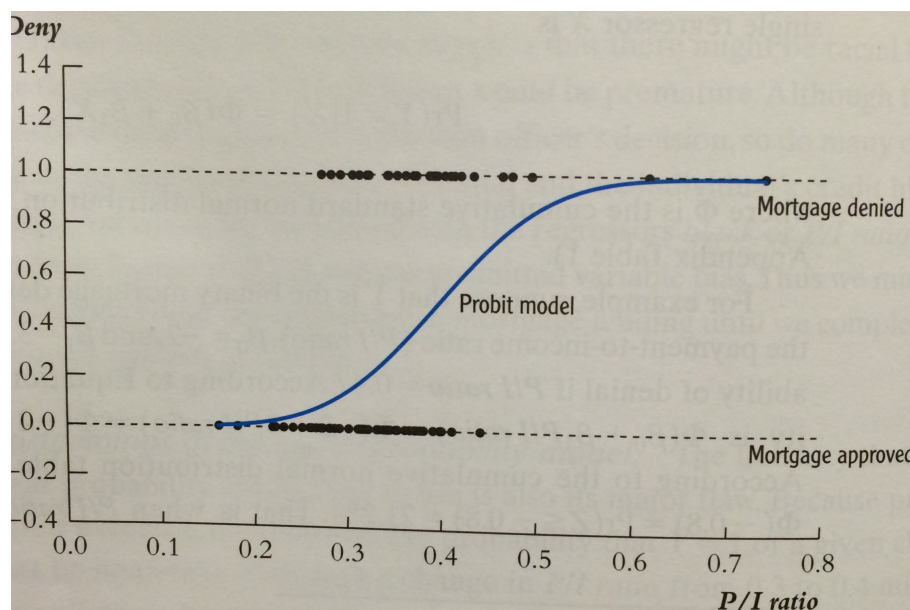


Figure 3: The probit model, Figure source: Stock, Watson p. 430

- Implies from the above statement, a probability of 0.6 means that 6 out of 10 observations is expected to be an event, i.e. $Y = 1$, the remaining 4 are non-events i.e. $Y = 0$.
- The beta coefficient represent *a change in z-value* for a unit change of the corresponding covariate, all other covariates held constant.
- It is not as straight-forward to interpret the beta coefficients effect on the *probability* as it depends on the values of the other covariates. This can quickly be realized by studying the standard normal cumulative function where a $\Delta z = 0.1$ has a bigger impact around $z = 0$ than for very large positive or negative values of z . Nonetheless a negative coefficient will lower the probability

and a positive one increases it.

Finally, without being too technical about the Maximum-Likelihood Estimator, it maximizes the likelihood function. The likelihood function is a joint probability distribution of the data observed. The MLE then "wants" the observed data to reflect a sample of the underlying distribution. The distribution changes as the unknown coefficient changes, so the MLE simply determines the coefficients by *maximizing the likelihood* that the sample was drawn from the underlying distribution, or likelihood function. In other words this distribution function determines the coefficients that are most likely to have produced the observed data.

3 Data

3.1 Acquisition of data

The data used in the study is collected from Zephyr and Bloomberg. The Zephyr database provided us with detailed information about Swedish mergers and acquisitions over the period 1998 - 2014. This gave us a list of 1257 mergers, acquisitions and changes in minority interest. The list includes from small stake increases to major stake and acquisitions of complete companies. In this thesis we focus on the innovative properties of R&D expenses and other intangibles such as patents and software and whether those properties are significant for targets in M&A transactions. It makes sense only to examine transactions where the acquirer moves from a non-controlling to a controlling interest. We thus narrowed down the results by only keeping acquisitions where the target firm was listed and where the acquirer owns less than 50% of the target firm before the deal and more than 50% when the deal is closed, or in any other way acquiring a major stake in the company. This gave us a list of 215 observations of target firms.

From Bloomberg we obtained firm specific information about R&D as reported on the income statement, total sales, total assets, disclosed intangibles and goodwill. Disclosed intangibles includes patents, copyrights and trade names but also goodwill, on which we decided to collect the goodwill amount separately and then subtract it from the disclosed intangible assets. Bloomberg provided us with data of both target and

benchmark firms.

Finally, we obtained the BICS classification industry for each firm by sector (level 1). Where industries could not be found from Bloomberg, Avanza was used.

3.2 Processing of data

The collected data from Bloomberg was exported into excel and structured. The dataset was checked for any obvious errors. Some observations had R&D expenses over total sales of less than zero, which is not reasonable. For these datapoints efforts were made to backtrack the value by using historical financial reports for the firms or retrieving the information via alternative sources (homepages, news releases, press releases etc.). If there were any uncertainties whether this would be correct, the value was dropped (however there could still be observations for other characteristics which of course were kept). The same procedure was done in cases where Bloomberg reported missing values for the target firms, in order to keep as much information as possible in our target sample. The benchmark consists of yearly information (1998 - 2014) from the listed firms at OMXS; a total of 314 firms per year. Once the data sample was cleaned and as much information as possible collected and maintained the calculation of quotas was made.

Furthermore, as controls and dummy variables will be used in the regression analysis, the final sample was extended in excel before import-

ing it into the statistical software. A dummy variable for each industry and each year was added.

3.3 Final sample

After checking for errors and completing missing info as much as possible via annual reports etc. we arrived at a final sample of 151 *R&D* observations and 116 *Intangible assets* observations. As far as controls are considered, the entire target sample and benchmark sample have information about industry and year.

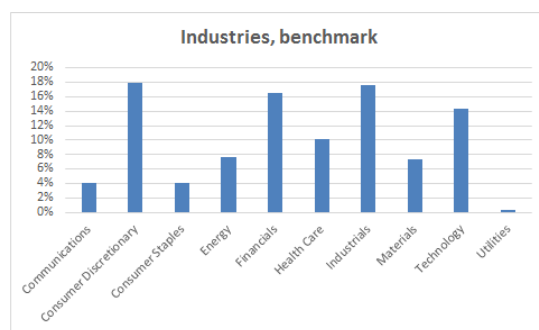


Figure 4: Showing the benchmark split by industries.

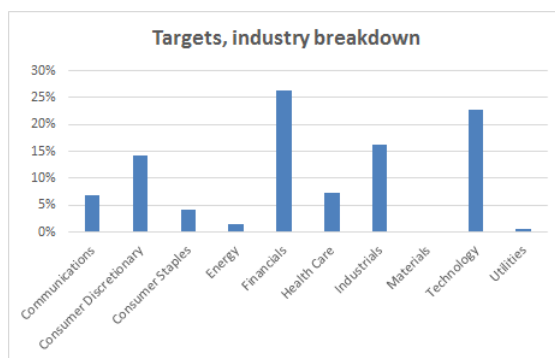


Figure 5: Showing the targets split by industries.

4 Methodology

4.1 Choosing relevant target data

For our target sample we assume that the data the year prior to the acquisition is the relevant data to examine when acquirers shall find their target. This is because of several reasons. First, it is the most recent, complete data over the firm's performance. Second it is reasonable to believe that the decision to acquire another firm isn't made overnight. Which means that the acquirer probably have been following the target over some time, and thus based their decision largely on that data. Finally, the alternatives are much worse. Either one would look two years prior to the acquisition announcement which is both illogical and the data might be outdated, or one would look at the data for the same year the acquisition is was made. This however would theoretically be flawed if the acquisition was made earlier than last December, but flawed enough for all acquisitions made when a relatively large part of the year still remains. The above mentioned statements is the rationale why, if an acquisition was made year t , we look at the data of the target for year $t - 1$.

4.2 Characteristics

The characteristics examined were constructed such that they were independent of company size, this means that they are percentages or other quotas in order to make the study apply in a broader context. As an example, R&D Expenses were measured as total R&D expenses divided

by total sales. This makes it easier to compare the observations across different sized companies as the absolute value of total R&D would be limited to comparisons of R&D numbers of equally sized companies, or in other ways controlled for. Below follows a short description of our characteristics

4.2.1 Research & Development

The total research & development expenditure is a measurement, which includes R&D in profit and loss account and capitalized R&D during the period. As previously discussed, our R&D measurement is calculated as total R&D expenses over total sales. Depending on accounting standards and the characteristics of the R&D expense, one can capitalize the expenses on the balance sheet and then amortize the costs on the income statement, or if treated as a cost instead of an asset; the R&D appears directly on the income statement.

The term R&D covers three activities: basic research, applied research and experimental development. Basic research could be experimental or theoretical work undertaken primarily to acquire new knowledge. Applied research is also original investigation undertaken in order to acquire new knowledge. It is directed primarily towards a specific practical aim or objective. Experimental development however, is systematic work, drawing on existing knowledge gained from research. Production of new materials, devices, installation of processes, systems and services are examples of typical experimental development. According to Swedish accounting legislation, research expenses must always ap-

pear as an expense on the income statement. Development expenses, on the other hand, will be accounted for as an asset if the following criteria are fulfilled:

1. The development expenses can be finished and the asset can be used or sold.
2. The firm's intentions is to finish the development expenses and thereafter sell or use the asset.
3. The firm has qualifications to sell or use the final asset.
4. It is likely that the development expense will create revenues in the future.
5. There are resources available for finishing the development and thereafter sell or use the asset.
6. The firm can reliably measure the costs that stems from the asset.

4.2.2 Intangible assets

The intangible assets includes goodwill, patents, copyrights, trade names, property rights and trademarks in the balance sheet. This item shows innovation investments performed in the past that is still valuable to the firm, and this is the main reason why we also want to include this in our analysis. In contrast to R&D expenses, which only refers to the period prior to the deal, the intangible assets will give another important angle of the innovation analysis of the target firms.

Goodwill was the only item we did not want to include in the intangible assets. The amount we extract is acquired goodwill, i.e the excess price paid over the fair market value of assets in an acquisition accounted for by the purchase method. Since this item has nothing to do with the target's innovation property, we decided not to include it in our analysis. Internally developed goodwill is not reported in the balance sheet since the company can not measure this reliably, therefore this item will not be extracted from the intangible assets (Redovisningsradet, 2000).

When it comes to measuring the intangible assets, Swedish accounting legislation says the firms must be able to measure the asset in reliable manners. Otherwise, the item can not be classified as an assets in the balance sheet. Some intangible assets have active markets that allows the firm to easily determine the value, but many times, this is not the case. Guidelines are provided from the accounting council but one could argue that the valuation is still more or less an approximation (Redovisningsradet, 2000). Potential errors in the valuation could be exploited by the acquiring firm but if this is a driver of the deal is difficult to determine. What values would imply that the asset is undervalued? If the acquirer has information of the assets true value, the firm could indeed use this information to buy undervalued assets. This is something our data will not reveal and no further analysis will therefore be made on this topic.

4.3 Descriptive statistics

The data sample is examined thoroughly and descriptive statistics is produced and graphed. Standard statistics such as mean, median, minimum and maximum values are available for each variable, year and industry for the data sample. These are complemented with industry distributions in both the target and benchmark sample, M&A activity by year, M&A activity within industries over time, mean values of the characteristics both over time and by industry. Tables and graphs are mainly presented in the Results section as well as in Appendix. These results are analyzed to see the characteristics of our samples and sub samples, which in turn is used to draw conclusions on whether targets and benchmark firms differ and about the significance of innovative properties for becoming a target.

4.4 Regression model selection

As described in Section 2.3 there are different types of regression models and it is important to acknowledge the differences and interpretations of any chosen model. In our data sample we want to examine whether innovative characteristics affect a firm of becoming a target. In our sample we have both targets and benchmark firms and so the observations of Y are binary, $Y = 1$ if the firm is (was) a target and 0 otherwise. The, for many well-known, OLS regression has shortcoming when it comes to modelling binary outcomes. The OLS regression however serves as a good starting point of describing the terminol-

ogy, function and interpretation in a regression analysis. Thereafter by showing differences to other models, these new models are easier explained and understood and also why it has a section in the mathematical background. However a linear probability model regression on a binary variable Y does not map all outcomes into the interval $I = [0, 1]$ where probability measures are defined.

We therefore choose the probit model which is more suitable in our case. The probit model uses a link function that maps any real number into the desired interval I . We are aware of the fact that the probit model is only one model in a larger family of generalized linear models which will be used to check robustness. As mentioned in Section 2.3, the probit model has less stringent assumptions. However we still have to have independent observations which we argue is fulfilled. One merger in one firm has little to do with a merger of another firm and Figure 10 does not show any strong dependence. If there are heightened activity in certain industries or years, this will be remedied by imposing controls as will be described in more detail in the next section. Furthermore the assumption of no perfect multicollinearity must be withheld. This assumption is checked via a different methods: first, one dummy variable for the yearly controls and one dummy variable for industry is kept in the benchmark, second the standard errors are checked for any anomalies, the sign and size of the coefficients are compared with the descriptive statistics of the data sample to see that they are reasonable and finally a correlation matrix is presented in appendix.

The assumption of linear dependence between Y and X does not need to be fulfilled, but there has to be linearity between the independent variables and the logit of the dependent. This assumption is fulfilled by the model construction, please see Section 2.3.2 and Section 4.4.3.

4.4.1 Controls

The focus of this study is to see if it is possible to find innovative characteristics among target firms. However we are aware of that there are e.g. research heavy industries. If they are overrepresented in the target sample compared to the benchmark sample, results might show that R&D affects the probability of becoming a target to a larger extent than it really does where the main reason might have been a consolidation in the industry. Industry controls are therefore added.

It is also reasonable to assume that M&A activity can vary over time due to swings in macroeconomic or other factors. We investigate how M&A activity vary over time, both for the sample as a whole and within industry. In the regression it is therefore warranted to add dummy variables to control for time. We leave one industry dummy and one year dummy out, this means that that effect will be in the intercept. Furthermore it does also mean that a beta coefficient for e.g. a year dummy represents the probability (or z-value) change with respect to the year dummy in the intercept. The benchmark year in our model is chosen to be close to the average yearly M&A activity, the benchmark industry was chosen to be as close to 0 as possible in Figure 11. The

benchmark year is 2013 and the benchmark industry is industrials (due to few observations on utilities). Finally, as the number of targets in our sample is quite small we want to avoid adding more control variables if possible. This lead to the decision of constructing our covariates, R&D expenses and intangible assets, as quotas instead of absolute values to at least to some extent control for size of the firm. We are aware of the fact that there might be a difference in these quotas that depend on size that is not controlled for.

4.4.2 Possible sample selection biases

The data sample consists of only listed, or previously listed companies for the reason of better accessible data. It is possible that certain tech start-ups or other companies whose research or software starts to get traction, get's bought. It is hard to approximate in which way the pendulum swings when not including private firms, but this should be kept in mind.

We have also limited the study to include transactions of a Swedish firm, by a Swedish firm. This might bias the sample in different ways; in a bidding war between a Swedish firm and a large multinational over a tech company, the large international firm might have an advantage. We do not have data on this but an article in SvD that was published a week prior to this paper, suggests that many innovative, Swedish firms ends up in the hands of non-Swedish companies (SvD, 2015)

4.4.3 Final probit model

Summing this section up, our probit model is:

$$\begin{aligned}
 Pr(Y = 1 \mid \mathbf{X}) = & \Phi(\beta_0 + \beta_1(R\&D) + \beta_2(Intangibles) \\
 & + \beta_3(1997) + \beta_4(1998) + \dots + \beta_{18}(2012) \\
 & + \beta_{19}(Communications) + \beta_{20}(Cons.Discr.) + \beta_{21}(Cons.Staples) \\
 & + \beta_{22}(Energy) + \beta_{23}(Financials) + \beta_{24}(HealthCare) + \beta_{25}(Materials) \\
 & + \beta_{26}(Technology) + \beta_{27}(Utilities)) \quad (10)
 \end{aligned}$$

where $\mathbf{X}$ denotes the vector of covariates and control variables.

4.4.4 Measurements of fit

The R^2 as defined in Equation 5 does not translate into the probit model, thus alternative ways of evaluating the fit of the model is needed. We implement two such measurements, one is the amount of correctly predicted variables, the other one is McFadden's Pseudo R^2

5 Results

5.1 Characteristic results

5.1.1 R&D Expenses

OECD reports a decline in gross domestic spending on R&D and we notice a similar decline in Figure 9, describing the full sample R&D quota over the period. During 1999-2001 we can see a peak in the R&D expenses on the total sample and worth noticing was the same peak in M&A activity during approximately the same years. Comparing this to the average intangible asset quota on Figure 10, we can see that the development over the 15 year period is almost the opposite. When R&D spending peaked, the intangible assets declined. The industry with the highest R&D over sales quota is not surprisingly the health care industry with 0.538, followed by an average of 0.2000 within the technology industry. The health care industry is known for heavy R&D spending since new pharmaceuticals and medications are constantly developed and invented. The second highest number in the technological industry is also not that surprising since our time period includes the IT-bubble where large spendings on R&D probably were made.

To complement our previous analysis a probit regression was run. As can be seen from Table 5, R&D quota is statistically significant on a 1% level and almost on a 0.1% level too. The marginal effect is -0.173 which means that increases in R&D quota, *ceteris paribus*, lowers the

probability of becoming a target by 17.3%. As explained in Section 2.3 it is not as straight forward of interpreting the actual change in probability from the beta coefficient alone as it depends on other coefficients too. So this marginal is calculated by entering the average value to all other covariates, then by changing R&D quota by one unit, how much does this affect the probability. Note that the marginal effect is *not* the coefficient, but rather the *equivalent* of a beta coefficient in a linear regression. The coefficients themselves are not reported in the table as they are not easily interpreted on a stand alone basis (see Section 1, probit regression for details). Recall that a probability of 0.6 should be interpreted as "in a large sample, 60% of the observations are expected to be targets, i.e. observed ones, and the remaining 40% observed zeroes", however we will loosely speak of it as "probability of becoming a target".

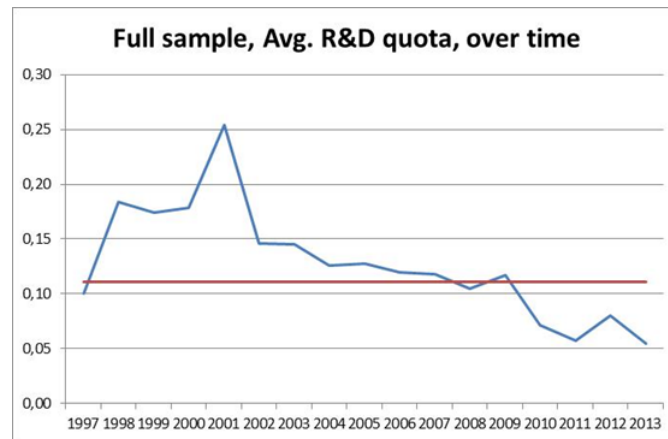


Figure 6: The full sample's (both targets and benchmark) R&D Expenses quota over time in blue and the average in red.

Table 1: R&D Expenses quota by industry for the full sample.

Industry	Average R&D quota
Communications	0.034
Consumer Discretionary	0.011
Consumer Staples	0.005
Energy	0.014
Financials	0.025
Health Care	0.538
Industrials	0.058
Materials	0.035
Technology	0.200
Utilities	0.004
Grand Total	0.111

5.1.2 Intangible assets

The intangible assets quota has varied over time and fluctuates around our sample's long term average of approximately 0.08 as shown in Figure 7. The quota seem to have followed the boom and bust of the IT-era in the late 90's and early 2000, thereafter it has experienced a steady increase.

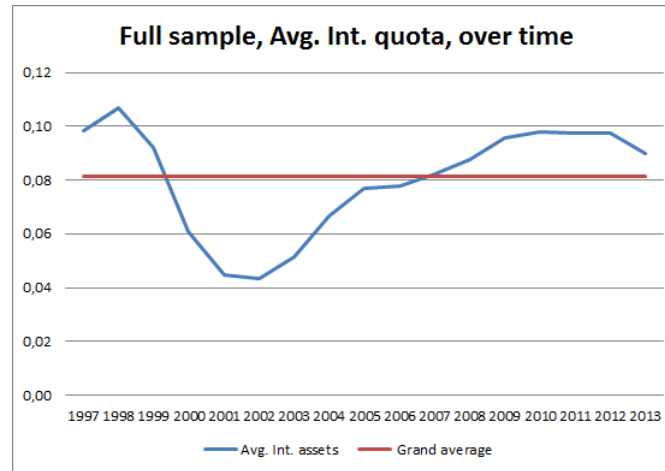


Figure 7: The full sample's (both targets and benchmark) intangible assets quota over time in blue and the average in red.

When split by industry the highest quotas belong to the health care industry, followed by communications and technology with quotas of 0.17, 0.15 and 0.12 respectively. Notably the top duo in the R&D quota, health care and technology, is again represented here in the top 3. The industry is the second highest when it comes to intangible assets which could be explained by the assets being concentrated to large Swedish communication firms, according to OECD:s technology and industry outlook (OECD Science, 2014). All average quotas can be found in

Table 2. It should also be mentioned that very few observations are made in the utilities sector.

Table 2: Intangible assets quota by industry for the full sample.

Industry	Average int. assets quota
Communications	0.15
Consumer Discretionary	0.07
Consumer Staples	0.08
Energy	0.09
Financials	0.03
Health Care	0.17
Industrials	0.05
Materials	0.03
Technology	0.12
Utilities	0.00
Grand Total	0.08

When comparing Figure 6 and Figure 7 there seems to be an inverse relation. The average intangible asset quota of the total sample has developed in the opposite way to R&D expenses over time. Statistically the correlation between the two is approximately -0.555 which will be touched upon later in the discussion even though correlation is a statistical measure and does not in itself imply causality.

The result from the regression showed a marginal effect of approximately 0.003. Since this is very small it suggests that intangible assets does not have a large effect on the probability of becoming a target. Note however that the p-value is high, 0.946, which means that it is far from being statistically significant since this reflects a high risk of com-

mitting a type I error. For further details about hypothesis testing we refer to Section 2.2.1. The intangible assets were investigated further by graphing the average intangible assets by industry, split into target and benchmark which can be seen in Figure 8. The mean values for each industry is similar for all industries except Energy, showing that intangible assets does not differ between benchmark and target, which would indicate that this is not a reason for becoming a target in an M&A transaction.

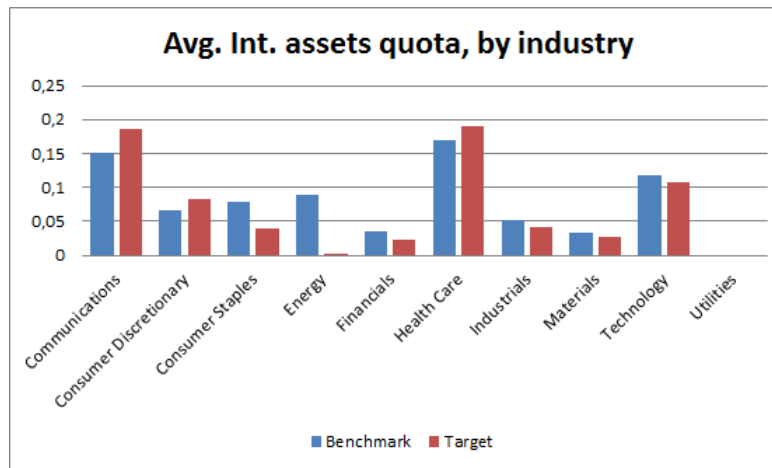


Figure 8: Benchmark and target intangible assets qouta, by industries

5.2 Industry and time effects

5.2.1 Time effects

The data sample contains the targets of M&A transactions in the period 1998 to 2014. In order to get more unbiased estimates of the effect of R&D expenses as well as intangible assets in a regression, it is important to impose control variables. It is reasonable to believe that the frequency of M&A transactions can vary over time and if so, it has to be controlled for.

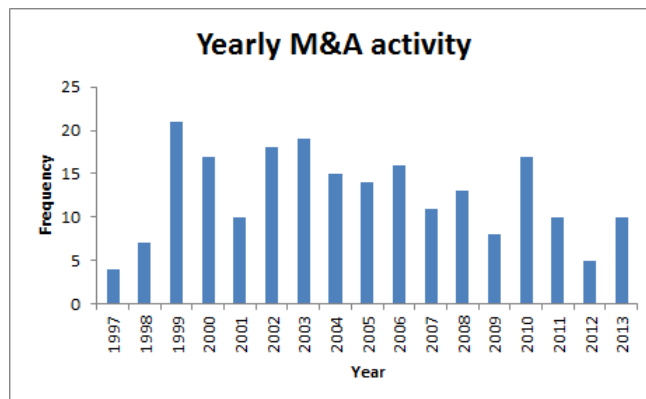


Figure 9: Target observations over time.

As shown in Figure 9, the Swedish M&A activity has remained fairly stable during the last 15 years. M&A reports also describes the past years as stable with exception of the booming years of 2006-2008 (Vinge, 2013). Since our data only includes Swedish firms acquiring other domestic firms, the numbers will be slightly different. The years 2006-2008 are not showing particularly high frequencies. One explanation, which will be discussed later on, could be the fact that most Swedish

firms were being targeted by foreign acquirers. The amount of transactions in Sweden were highest during 1999. Figure 10 shows that during this year, most deals were made within the financial sector, with the technology sector also being partially affected. For other industries and years the activity was relatively evenly spread out over time.

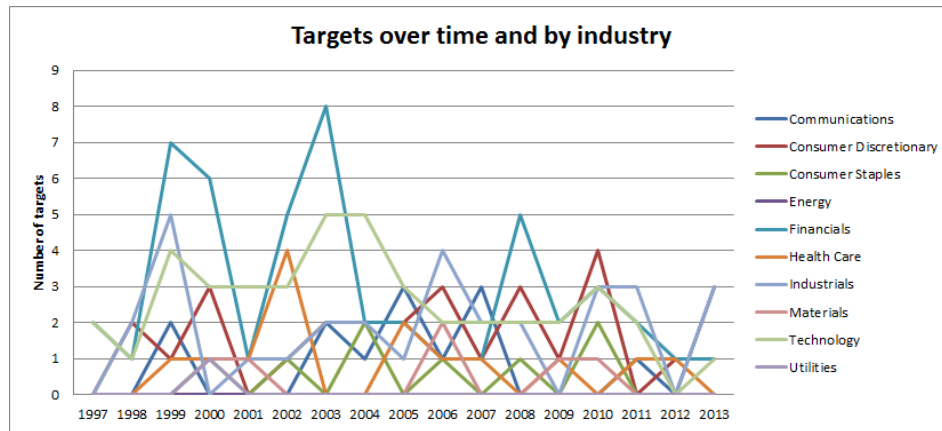


Figure 10: Target observations over time, split by industry. The overall activity is quite stable over time with exception of the years 1999-2000 and 2002-2003. During these years activity was especially strong in the financials and technology industry.

The regression results show significance on the 0.1% level for the years 1999, 2000 and on the 1% level for 2002, 2003 and 2006. In addition there were three more years significant on the 5% level and yet another three on the 10% level, details can be found in Table 5. Recall that 2013 serves as a benchmark year. The years around the IT-era as well as 2002-2003 constitute a solid period where the year had a significant effect on the probability of becoming a target, compared to 2013. Note that 2013 was by no means a low activity year in our sample, still during

the mentioned years the probability of becoming a target compared to 2013 was approximately 7.9% to 10.7%. Of course these yearly covariate results cannot be used to find other targets since obviously those days have passed. However it purges the other covariates and gives a hint on that timing is not irrelevant.

5.2.2 Industry effects

The industry variation is investigated firstly by graphing the distributions of the industries but also over industries. This is As shown in Figure 4 and Figure 5, the sample is spread out across different industries. If it was the case that industry did not matter at all one would expect the targets to follow a similar distribution to that of the benchmark, given a large number of observations. Figure 11 shows that over the entire period, the targets differ from the benchmark when it comes to industry distribution. Some industries are approximately represented in the same frequency, but Energy, Financials, Materials and Technology stands out.

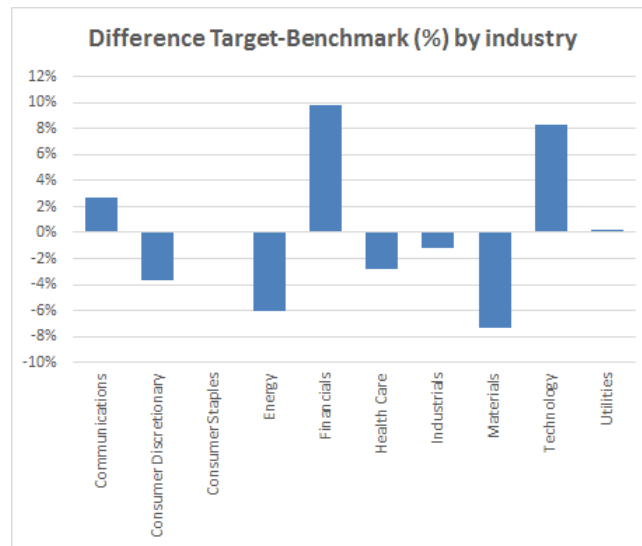


Figure 11: Target vs benchmark by industries.

Most deals were made within the financial sector and as mentioned

above, with a spike in transactions in 1999. In the technology sector, many deals were made during 2002-2003, and this is the second industry that was overrepresented among the target firms. There is a possibility to include combined dummy variables for these occasions. That is if one believes that there is an extra effect of the *combination* of being a financial firm that specific year. However this is not included in our model with regard to the controls that already exists and bearing in mind that the target observations in our sample are limited.

The regression analysis further confirmed this picture as Technology was statistically significant on a 0.1% level showing a average marginal effect of approximately 0.050 meaning that being a technology firm increased probability of becoming a target by 5% compared to Industrials. The analysis also strengthened the earlier results of Financials becoming targets, being significant on a 10% level with a marginal effect of of 2.8%. Finally the health care sector was significant on a 5% level with an increase in probability of roughly 4.5% of becoming a target. It should be mentioned that there were very few observations of utilities during the target period which might explain the large marginal effect. However the p-value is very high, almost 1, so there is still no point in analyzing the marginal effect. We emphasize to always bear the significance from the regression results in mind.

Table 3: Target vs benchmark mean

Industry	R&D	Other Intangibles
Communication	13%	38%
Consumer Discretionary	23%	28%
Consumer Staples	14%	17%
Energy	0%	0%
Financials	0%	8%
Health Care	11%	44%
Industrials	5%	33%
Materials	0%	50%
Technology	3%	32%
Utilities	0%	0%

Table showing number of target observations *above* benchmark means. I.e. a number of 23% means that 23% of target observations was found above the benchmark mean. The system used to partition the firms into industries is BICS.

Notice however that quite a substantial amount of target firms reported 0 in R&D, while some relatively large quotas were found in the benchmark sample. This increases mean value for the benchmark which has an impact on these numbers. The R&D measure is especially exposed to this as the R&D quota is defined as R&D expenditure over total sales. For firms being research intense but might rely on venture capital rather than sales, or in any other way having a small sales figure, this quota can get large.

5.3 Robustness

In order to check the robustness of our results we employ a logistic regression to our data sample. It differs from the probit regression in the sense that the link that maps any value on the real axis into the interval $I = [0, 1]$ differs. Instead of the standard normal cumulative function Φ in the probit model, the logistic model has the function

$$Pr(Y = 1 \mid X_1, X_2, \dots, X_k) = \frac{1}{1 - e^{-(\beta_0 + \beta_1 X_1 + \dots + \beta_k X_k)}}. \quad (11)$$

The cumulative function in RHS has slightly flatter tails than that of the standard normal cumulative function. The coefficients are found using MLE, as they are in the probit model. The results are similar to those of the probit regression, with the same covariates being similar. The year dummies 1998 and 2003 both increased in significance level from 10% and 1% to 5% and 0.1% respectively. Technology also increased from 1% to 0.1% in the logit regression. The p-values and marginal effects also changed but the conclusions from the probit regression still remain. Notably R&D expenses marginal effect decreased from -17.3% to 26.2%. Overall, many of the marginal effects increased in absolute value compared to the probit regression but the sign and significance strengthened the earlier results. For details please see Table 6.

With the estimated coefficients, the probit model's ability to predict observations correctly can be assessed. Taking one observation, calculate the RHS of Equation 10 with the predicted value, and compare

it to the LHS. If the right hand side of the equation is 0.5 or greater when LHS is observed as a 1, then it is correctly predicted. Similarly if RHS is calculated as less than 0.5 and LHS is 0, then it's also correctly predicted.

Table 4: Summary statistics: Predicted probit model.

Measure	Value
Min.	0.000
1st Quantile	0.015
Median	0.036
Mean	0.050
3rd Quantile	0.070
Max.	0.319

Statistics of the data sample using the predicted probit model to calculate the observations.

As Table 4 shows, the median and mean reflects quite well the original sample, but the maximum is 0.319. This means that the model does not predict any targets in our data sample, rendering it will always predict non-event observations correctly and event observations, i.e. targets incorrectly. As our sample has a low portion of targets, this will probably not be a fair measure. Instead we employ McFadden's Pseudo R^2 to determine the fit of the model.

$$R_{McFadden}^2 = 1 - \frac{\ln(L_{model})}{\ln(L_0)} \quad (12)$$

where L_{model} is the likelihood of the model and L_0 the likelihood without any covariates, i.e. using only the intercept, which gives a measure of error reduction. Our reported $R_{McFadden}^2$ is approximately 0.556.

6 Conclusions & discussion

6.1 R&D Expenses

As earlier mentioned, the total sample R&D expenditures and intangible assets during our time frame was almost the opposite to each other. A possible interpretation of the relationship between these graphs (Figure 6 & 7) could be explained by the IT-bubble during the turn of the century. Large amounts were probably spent on research and development since many firms saw their value rising quickly from rapid growth within the IT-sector. When the bubble collapsed during 2002, the R&D spendings declined substantially and the intangibles started to increase. Perhaps the large spendings finally paid off to the firms with high R&D expenses? One must have in mind that other R&D intensive industries might push the overall spendings up and down. Health care has the highest quota of all industries and it was probably not affected the same way as technology firms during the IT-bubble. Something is however increasing the spendings again as they are now back at the same levels as before the bubble. Perhaps we are entering a even more R&D intensive period within the health care industry?

The results are showing that the characteristic is actually decreasing the possibility of becoming a target and strengthening the conclusion Bena and Li had in their paper, which suggests that even though Bena and Li studied the US, the pattern persist in Sweden too and might be more general. Another important aspect is that Bena and Li's results

does not include observations from 2006- and onwards. Potential differences in the results would have been expected since we include a period in Sweden where many hyped firms have been acquired. However, the results are similar and the conclusion from this could be that acquiring companies are not targeting R&D intensive firms. Once again does this measure only capture the selected period's innovative input and not past investments.

Before comparing the results of our study with previous literature we should have in mind the following:

- The most R&D intensive industries are technology and health care. Sweden is one of the most R&D intensive countries since there are many firms within the high-technological industry, as defined by OECD. In R&D expenses compared to GDP, Sweden is one of the top-five countries (SCB, 2011). Sweden's R&D intensive industries and firms should however give us higher values in both target and benchmark firms, not only in the benchmark. Another explanation to differences between the studies could be that the R&D intensive acquisitions have been Swedish firms acquiring foreign companies like AstraZeneca and Pharmacia&Upjohn's merger with Monsanto. Our data only consists of listed acquirers and targets in Sweden; therefore these acquisitions are not included in the sample.
- Different accounting regulations between Sweden and the U.S. The US GAAP says that internal research and development ex-

penditure is expensed as it is incurred, unlike IFRS where internal development expenditure is capitalized if specific criteria are met. This could be a potential explanation of the larger amount of R&D expenses within U.S target firms (KPMG, 2013).

6.2 Intangible assets

When looking at the descriptive statistics of our second quota, intangible assets, some numbers are similar to the R&D data. The industries with the highest quotas are about the same, when looking at the entire sample. However, Figure 8 shows that intangible assets in target firms, within the health care industry, have even higher quotas than the benchmark. Past research (that has been capitalized) and innovative property seems to be valuable to acquirers as the target firms have higher values. Comparing this to our findings within the R&D quota, where target firms actually had lower R&D expenses than the benchmark, raises many questions. Are the some acquirers preferring past innovation output before recent R&D activity or is it some other driver that causes these numbers to differ? Or is it obvious that verified innovation output, reported as assets, are more attractive than recent R&D activities that perhaps could result in nothing but expenses?

The industries with higher values in targets compared to benchmark were health care, communications and consumer discretionary. Why are these industries containing target firms with higher intangible quotas? Health care and communications are indeed R&D intensive industries but something is pushing the target firms above the benchmark.

One explanation could be that the intangible property that shows up in the balance sheet of the target firms, within these industries, are easier to measure and value properly. Perhaps a verified patent within the health care industry is more likely to have a correct value than a patent within the energy industry? If the acquirer is aware of the true value of the acquired assets, the overall M&A driver might be pushed towards the intangible assets and their value, making highly innovative firms becoming more attractive. The same intangible asset in an energy-classified firm might not be as attractive since the acquirer is not sure of the true value of the asset. As discussed earlier, intangible assets are indeed difficult to measure and firms might not always have guidelines of how to measure the asset properly. Since intangible assets are a huge part of the health care industry, evaluating these might be done in a more correct way than in any other industry.

To summarize the descriptive output of the intangible quota one can say that some industries stood out, but overall does it seem that intangible property is not a key driver of M&As. Some industries have higher quotas within the target firms and a plausible explanation of this could be the difficulties of measuring intangible property correctly. The regression analysis showed very little marginal effect of the intangible asset quota, but though the p-value was too high to be significant. The sample was split into a benchmark and a target sample, thereafter split by industry to investigate this further. This showed that split by industry, the target and benchmark had similar mean values of intangi-

ble assets which is in line with the low marginal effect of our regression analysis.

6.3 Industry, size and time effects

In the probit model, time and industry effects were controlled for. As shown in our results quite a lot of the years were significant. Even though the main reason for including them is to purge the R&D and intangible covariates, additional conclusions can be drawn. The period years 1999-2000 and 2002-2003 all showed an increase in probability of 7.9%-10.7% to become a target compared to 2013. These years have of course passed so this won't explicitly help detection of new targets, it does however show that the period and timing does matter. All these years were statistically significant on the 0.1-1% level. This suggests that if investors are looking to identify possible targets, certain time periods are extra favourable.

The peaking probability of becoming a target during these years could be affected by the IT-bubble. The bubble only lasted until 2001 and hyped firms with high expectations and returns were probably popular during this time. When the bubble popped and firms value were dropping, some firms could still be exploited but now for another reason: after the bust acquirers now had a better chance of buying firms at discount. As an example Ericsson B shares was traded for 830 SEK at its peak in early 2000 but could be bought for only 3.8 SEK at a shares issue two years later.

Our sample did not show the same M&A activity as was described in Vinge's M&A report where Sweden had booming years during 2006-2008. Again, our sample only includes Swedish acquirers buying Swedish targets which might affect the sample.

As for industries most notably being a technology, health care or financial firm had a statistically significant effect of increasing the probability of becoming a target compared to the benchmark industrials. They are also among those who were most frequently bought and had high intangible assets quota. Whether this is a result of the innovative properties themselves, consolidation in the market or general hype is harder to tell. Figure 10 does indeed show some spikes around the turn of the millenium, but the interest for firms in these industries still persisted. The answer to the increased probability of becoming a target for those sectors might be a combination of the above stated factors.

7 Suggestions for further studies

Additional research of the characteristics of M&As would be of high value to investors, shareholders and for other academic research fields. Deeper knowledge into the target traits, for example organic growth, would be important to many M&A participants. The increase in volume of a relatively unchanged service or product is a widely known source of profit. Measuring and collecting data regarding the organic growth might shed further light on which firms that might be acquired. Another suggestion for future research is the Swedish acquiring firms' characteristics, especially one theory that Michael Webb brought up in *How to acquire a company*. He claims that firms with over-committed in areas or industries that are prone to cyclical variation might want to diversify or by other means mitigate their cyclical pattern by acquiring other firms. Finally, one would like to examine not only whether the firm is going to become a target, but also if it will be a target in a bidding war as these wars might push the market price of the target higher.

8 Appendix A, Tables

Table 5: Probit regression results.

Independent variable	p-value	Statistical significance	Avg. marginal effect
Intercept	0.000	***	-0.204
R&D expenses quota	0.003	**	-0.173
Intangible assets quota	0.946	N	0.003
1997	0.097	.	0.064
1998	0.051	.	0.063
1999	0.000	***	0.102
2000	0.000	***	0.088
2001	0.624	N	0.017
2002	0.003	**	0.071
2003	0.001	**	0.075
2004	0.093	.	0.042
2005	0.103		0.041
2006	0.001	**	0.059
2007	0.021	*	0.052
2008	0.061	.	0.043
2009	0.375	N	0.021
2010	0.124	N	0.032
2011	0.813	N	-0.006
2012	0.307	N	-0.027
2013	BM	BM	BM
Communication	0.405	N	0.017
Consumer Discretionary	0.181	N	0.020
Consumer Staples	0.155	N	0.033
Energy	0.213	N	-0.044
Financials	0.057	.	0.028
Health Care	0.042	*	0.030
Industrials	BM	BM	BM
Materials	0.814	N	-0.006
Technology	0.002	**	0.046
Utilities	0.986	N	-0.249

Marginal effects, p-values and significance. "BM" is for Benchmark among dummy control variables. Significance levels: "****" 0.1%, "***" 1%, "**" 5%, "." 10%, "N" no significance (higher than 10%).

Table 6: Logit regression results.

Independent variable	p-value	Statistical significance	Avg. marginal effect
Intercept	0.000	***	-0.190
R&D expenses quota	0.001	**	-0.262
Intangible assets quota	0.719	N	0.014
1997	0.065	.	0.072
1998	0.034	*	0.071
1999	0.000	***	0.107
2000	0.000	***	0.090
2001	0.533	N	0.023
2002	0.001	**	0.079
2003	0.001	***	0.082
2004	0.057	.	0.051
2005	0.070	.	0.049
2006	0.001	**	0.066
2007	0.015	*	0.060
2008	0.042	*	0.051
2009	0.283	N	0.029
2010	0.103	N	0.038
2011	0.929	N	-0.002
2012	0.353	N	-0.028
2013	BM	BM	BM
Communication	0.403	N	0.017
Consumer Discretionary	0.186	N	0.021
Consumer Staples	0.139	N	0.033
Energy	0.207	N	-0.059
Financials	0.062	.	0.028
Health Care	0.030	*	0.045
Industrials	BM	BM	BM
Materials	0.743	N	-0.008
Technology	0.001	***	0.050
Utilities	0.986	N	-0.423

Beta coefficients, p-values and significance. "BM" is for Benchmark among dummy control variables. Significance levels: "****" 0.1%, "****" 1%, "**" 5%, "." 10%, "N" no significance (higher than 10%).

Target
 Min. :0.00000
 1st Qu.:0.00000
 Median :0.00000
 Mean :0.03872
 3rd Qu.:0.00000
 Max. :1.00000

Figure 12: Summary statistics

RnDqouta	OtherIntqouta	X1997	X1998
Min. :0.00000	Min. :0.00000	Min. :0.00000	Min. :0.00000
1st Qu.:0.00000	1st Qu.:0.0023	1st Qu.:0.00000	1st Qu.:0.00000
Median :0.0002	Median :0.0232	Median :0.00000	Median :0.00000
Mean :0.0999	Mean :0.0814	Mean :0.05727	Mean :0.05781
3rd Qu.:0.0344	3rd Qu.:0.1051	3rd Qu.:0.00000	3rd Qu.:0.00000
Max. :3.8507	Max. :0.9559	Max. :1.00000	Max. :1.00000
NA's :2577	NA's :1843		
X1999	X2000	X2001	X2002
Min. :0.00000	Min. :0.00000	Min. :0.00000	Min. :0.00000
1st Qu.:0.00000	1st Qu.:0.00000	1st Qu.:0.00000	1st Qu.:0.00000
Median :0.00000	Median :0.00000	Median :0.00000	Median :0.00000
Mean :0.06033	Mean :0.05961	Mean :0.05835	Mean :0.05979
3rd Qu.:0.00000	3rd Qu.:0.00000	3rd Qu.:0.00000	3rd Qu.:0.00000
Max. :1.00000	Max. :1.00000	Max. :1.00000	Max. :1.00000
X2003	X2004	X2005	X2006
Min. :0.00000	Min. :0.00000	Min. :0.00000	Min. :0.00000
1st Qu.:0.00000	1st Qu.:0.00000	1st Qu.:0.00000	1st Qu.:0.00000
Median :0.00000	Median :0.00000	Median :0.00000	Median :0.00000
Mean :0.05997	Mean :0.05925	Mean :0.05907	Mean :0.05943
3rd Qu.:0.00000	3rd Qu.:0.00000	3rd Qu.:0.00000	3rd Qu.:0.00000
Max. :1.00000	Max. :1.00000	Max. :1.00000	Max. :1.00000

Figure 13: Summary statistics

X2007	X2008	X2009	X2010
Min. :0.00000	Min. :0.00000	Min. :0.00000	Min. :0.00000
1st Qu.:0.00000	1st Qu.:0.00000	1st Qu.:0.00000	1st Qu.:0.00000
Median :0.00000	Median :0.00000	Median :0.00000	Median :0.00000
Mean :0.05853	Mean :0.05889	Mean :0.05799	Mean :0.05961
3rd Qu.:0.00000	3rd Qu.:0.00000	3rd Qu.:0.00000	3rd Qu.:0.00000
Max. :1.00000	Max. :1.00000	Max. :1.00000	Max. :1.00000

X2011	X2012	X2013	Communication
Min. :0.00000	Min. :0.00000	Min. :0.00000	Min. :0.00000
1st Qu.:0.00000	1st Qu.:0.00000	1st Qu.:0.00000	1st Qu.:0.00000
Median :0.00000	Median :0.00000	Median :0.00000	Median :0.00000
Mean :0.05835	Mean :0.05745	Mean :0.05835	Mean :0.04556
3rd Qu.:0.00000	3rd Qu.:0.00000	3rd Qu.:0.00000	3rd Qu.:0.00000
Max. :1.00000	Max. :1.00000	Max. :1.00000	Max. :1.00000

Consumer.Discretionary	Consumer.Staples	Energy	Financials
Min. :0.0000	Min. :0.00000	Min. :0.00000	Min. :0.0000
1st Qu.:0.0000	1st Qu.:0.00000	1st Qu.:0.00000	1st Qu.:0.0000
Median :0.0000	Median :0.00000	Median :0.00000	Median :0.0000
Mean :0.1732	Mean :0.04124	Mean :0.07401	Mean :0.1774
3rd Qu.:0.0000	3rd Qu.:0.00000	3rd Qu.:0.00000	3rd Qu.:0.0000
Max. :1.0000	Max. :1.00000	Max. :1.00000	Max. :1.0000

Health.Care	Industrials	Materials	Technology
Min. :0.0000	Min. :0.0000	Min. :0.00000	Min. :0.0000
1st Qu.:0.0000	1st Qu.:0.0000	1st Qu.:0.00000	1st Qu.:0.0000
Median :0.0000	Median :0.0000	Median :0.00000	Median :0.0000
Mean :0.1005	Mean :0.1709	Mean :0.07149	Mean :0.1424
3rd Qu.:0.0000	3rd Qu.:0.0000	3rd Qu.:0.00000	3rd Qu.:0.0000
Max. :1.0000	Max. :1.0000	Max. :1.00000	Max. :1.0000

Figure 14: Summary statistics

Utilities
Min. :0.000000
1st Qu.:0.000000
Median :0.000000
Mean :0.003242
3rd Qu.:0.000000
Max. :1.000000

Figure 15: Summary statistics

9 Appendix B, Correlation

	Rndqouta	OtherIntq	1997	1998	1999	2000	2001	2002	2003	2004	2005	2006	2007
Rndqouta	1.00												
OtherIntq	0.26	1.00											
1997	0.00	0.02	1.00										
1998	0.03	0.04	-0.06	1.00									
1999	0.03	0.02	-0.06	-0.06	1.00								
2000	0.03	-0.03	-0.06	-0.06	-0.06	1.00							
2001	0.07	-0.06	-0.06	-0.06	-0.06	-0.06	1.00						
2002	0.02	-0.07	-0.06	-0.06	-0.06	-0.06	-0.06	1.00					
2003	0.02	-0.05	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	1.00				
2004	0.01	-0.03	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	1.00			
2005	0.01	-0.01	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	1.00		
2006	0.01	-0.01	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	1.00	
2007	0.00	0.00	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	1.00
2008	0.00	0.01	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06
2009	0.00	0.03	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06
2010	-0.03	0.03	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06
2011	-0.04	0.03	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06
2012	-0.03	0.03	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06
2013	-0.05	0.02	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06	-0.06
Communications	-0.04	0.12	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.01	0.00	0.01
Consumer Discretionary	-0.10	-0.05	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00
Consumer Staples	-0.05	-0.01	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.01	0.00	0.00	0.00
Energy	-0.06	0.01	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00
Financials	-0.11	-0.13	0.00	0.00	0.01	0.00	0.00	0.00	0.01	0.00	0.00	-0.01	0.00
Health Care	0.38	0.22	0.00	0.00	0.00	0.00	0.00	0.01	0.00	0.00	0.00	0.00	0.00
Industrials	-0.07	-0.12	0.00	0.00	0.00	-0.01	0.00	0.00	0.00	0.00	0.00	0.00	0.00
Materials	-0.05	-0.11	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00
Technology	0.09	0.13	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00
Utilities	-0.01	-0.01	0.00	0.00	0.00	0.01	0.00	0.00	0.00	0.00	0.00	0.00	0.00

Figure 16: Correlation between covariates, part 1 of 2.

	2008	2009	2010	2011	2012	2013	Communications	Consumer Discretionary	Consumer Staples	Energy	Financials	Health Care	Industrials	Materials	Technology	Utilities
Rndqouta																
Otherintqouta																
1997																
1998																
1999																
2000																
2001																
2002																
2003																
2004																
2005																
2006																
2007																
2008	1.00															
2009	-0.06	1.00														
2010	-0.06	-0.06	1.00													
2011	-0.06	-0.06	-0.06	1.00												
2012	-0.06	-0.06	-0.06	-0.06	1.00											
2013	-0.06	-0.06	-0.06	-0.06	-0.06	1.00										
Communications	0.00	0.00	0.00	0.00	0.00	0.00	1.00									
Consumer Discretionary	0.00	0.00	0.00	0.00	0.00	0.00	-0.10	1.00								
Consumer Staples	0.00	0.00	0.01	0.00	0.00	0.00	-0.04	-0.09	1.00							
Energy	0.00	0.00	0.00	0.00	0.00	0.01	-0.06	-0.13	-0.06	1.00						
Financials	0.00	0.00	0.00	0.00	0.00	0.00	-0.10	-0.21	-0.10	-0.13	1.00					
Health Care	0.00	0.00	0.00	0.00	0.00	0.00	-0.07	-0.15	-0.07	-0.09	-0.16	1.00				
Industrials	0.00	0.00	0.00	0.00	0.00	0.00	-0.10	-0.21	-0.09	-0.13	-0.21	-0.15	1.00			
Materials	0.00	0.00	0.00	0.00	0.00	0.00	-0.06	-0.13	-0.06	-0.08	-0.13	-0.09	-0.13	1.00		
Technology	0.00	0.00	0.00	0.00	0.00	0.00	-0.09	-0.19	-0.08	-0.12	-0.19	-0.14	-0.19	-0.11	1.00	
Utilities	0.00	0.00	0.00	0.00	0.00	0.00	-0.01	-0.03	-0.01	-0.02	-0.03	-0.02	-0.03	-0.02	-0.02	1.00

Figure 17: Correlation between covariates, part 2 of 2.

10 Appendix C, Code

```
blu <- attach(Probit_data2_controls_v9)
Y <- cbind(Target)
X <- cbind(RnDqouta, OtherIntqouta, X1997, X1998, X1999, X2000,
          X2001, X2002, X2003, X2004, X2005, X2006, X2007, X2008,
          X2009, X2010, X2011, X2012,
          Communication, Consumer.Discretionary, Consumer.Staples,
          Energy, Financials, Health.Care, Materials, Technology, Utilities)

#Base case: Industry: Industrials   Year: 2013
summary(Y)
summary(X)

#Probit model coefficients
probit <- glm(Y ~ X , family = binomial (link="probit"))
summary(probit)

#Logit model coefficients
logit <- glm(Y ~ X , family = binomial (link="logit"))
summary(logit)

#Odds ratio
exp(logit$coefficients)

#Probit model average marginal effects
```

```
ProbitScalar <- mean(dnorm(predict(probit, type= "link")))
ProbitScalar * coef(probit)

#Logit model average marginal effects
LogitScalar <- mean(dlogis(predict(logit, type= "link")))
LogitScalar * coef(logit)

#Probit model predicted probabilities
pprobit <- predict(probit, type="response")
summary(pprobit)

#McFadden's Pseudo R-squared
probit0 <- update(probit, formula = Y ~ 1)
McFadden <- 1-as.vector(logLik(probit)/logLik(probit0))
McFadden
```

11 References

- [1] Andrade, G. Mitchell, M. and Stafford, E. (2001) *New Evidence and Perspectives on Mergers*, Journal of Economic Perspectives, Vol. 15, p. 103-120

- [2] Bena, J. and Li, K (2011) *Corporate Innovations and Mergers and Acquisitions*

- [3] Bloomberg *Index Methodology, Global Fixed Income Family*, available at <http://www.bloombergindeces.com/content/uploads/sites/3/2013/05/Index-Family-Methodology.pdf>, accessed at 16 April, 2015

- [3] Bloomberg, 2015 *The Bloomberg Innovation Index*, available at <http://www.bloomberg.com/graphics/2015-innovative-countries/>, accessed at 11 May, 2015

- [3] Curtis, G. (2009) *Trademarks Of A Takeover Target*, available at http://www.investopedia.com/articles/stocks/07/takeover_target.asp? accessed at 9 May, 2014

- [4] Dagens Industri, *Tio heta uppköpskandidater*, available at <http://www.di.se/artiklar/2014/4/8/tio-heta-uppkopskandidater/>, accessed 27 April, 2014

- [5] Fama, E. and French, K (1992) *The Cross-Section of Expected Stock Returns*, The Journal of Finance, Vol. XLVII No.2, p. 427-465

[5] Evans, F. and Mellen, C (2010) *Valuation for M&A: Building Value in Private Companies*, Second edition, Wiley and Sons Inc. ISBN: 978-0-470-60441-0

[5] Garson, G. D. (2012). *Probit Regression and Response Models*. Asheboro, NC: Statistical Associates Publishers.

[6] Harford, J (2005). *What drives merger waves?*, Journal of Financial Economics, Vol. 77, p 529-560

[10] Kennedy, P. (2003) *A Guide to Econometrics*, Sixth edition, Blackwell Publishing Ltd.

[11] KPMG (2013) *IFRS Compare to US GAAP: An overview*, available at <http://www.kpmg.com/CN/en/IssuesAndInsights/ArticlesPublications/Documents/IFRS-compared-to-US-GAAP-An-overview-O-201311.pdf>, accessed 11 May, 2014

[12] KPMG (2014) *2014, M&A Outlook Survey Report*, available at <http://www.kpmgsurvey-ma.com/pdf/2014%20MA%20Outlook%20Survey.pdf>, accessed 8 May, 2014

[13] KPMG (2010) *The Determinants of M&A Success*, available at <http://www.kpmg.com/NZ/en/IssuesAndInsights/>

ArticlesPublications/ Documents/Determinants-of-MandA-Success- report.pdf,
accessed at 6 May, 2014

[14] Lang, H. (2013) *A Brief Introduction To Regression Analysis*, KTH
Mathematics, div. of Mathematical Statistics.

[15] Maksimovic, V. Phillips, G. and Yang, L. (2013) *Private and
public merger waves*, The Journal of Finance, Vol. LXVIII No.5, p.
2177-2217

[16] OECD Data. *Research and Development (R&D)*,
available at <https://data.oecd.org/rd/gross-domestic-spending-on-r-d.htm>,
accessed at 8 May, 2015

[16] OECD Science. (2014) *Technology and Industry Outlook 2014*,
available at [http://www.keepeek.com/Digital-Asset-Management/oecd/science-
and-technology/oecd-science-technology-and-industry-outlook-2014/sweden_sti_outlook-
2014-76-en#page1](http://www.keepeek.com/Digital-Asset-Management/oecd/science-and-technology/oecd-science-technology-and-industry-outlook-2014/sweden_sti_outlook-2014-76-en#page1) , accessed at 7 May, 2015

[16] Redovisningsradet. (2000) *RR 15, Immateriella Tillgångar*,
available at <http://www.bfn.se/redovisning/radet/RR/RR15.pdf> accessed
at 4 May, 2015

[17] SCB (2011) *Appendix till: Forskning och utveckling inom
företagssektorn 2011*, available at <http://www.scb.se/Statistik/UF/UF0302/2011A01F/>

UF0302.2011A01F_SM_UF14SM1301.pdf, accessed at 11 May, 2014

[18] Stock, H. and Watson, J. (2012) *Introduction to Econometrics*, Third Edition. Harlow, Pearson Education Limited.

[19] SvD (2015) *Varfor stannar inte tillvaxtforetagen?*, available at http://www.svd.se/opinion/brannpunkt/varfor-stannar-inte-tillvaxtforetagen_4553560.svd, accessed 11 May, 2014

[19] Vinge (2013) *Vinge Q&A: Private equity and the Swedish M&A Market*, available at <http://www.vinge.com/Global/Bilagor%20till%20nyheter%20etc/Vinge-QandA.pdf>, accessed 9 May, 2014

[20] Webb, M (1974) *How to Acquire a Company*, Gower Press, ISBN 0-7161-0247-1