



MASTERARBEIT | MASTER'S THESIS

Titel | Title

Relative Performance Evaluation and Managerial Power

verfasst von | submitted by

Dott. Mara Mück

angestrebter akademischer Grad | in partial fulfilment of the requirements for the degree of
Master of Science (MSc)

Wien | Vienna, 2026

Studienkennzahl lt. Studienblatt | Degree
programme code as it appears on the
student record sheet:

UA 066 915

Studienrichtung lt. Studienblatt | Degree
programme as it appears on the student
record sheet:

Masterstudium Betriebswirtschaft

Betreut von | Supervisor:

Ass.-Prof. Dipl.-Math. oec. Dr. Ulrich Schäfer

Abstract

This thesis examines the impact of managerial power on the design and effectiveness of Relative Performance Evaluation (RPE) in executive compensation. While RPE is theoretically efficient in filtering out systematic risk from managerial pay, empirical evidence suggests its limited use in practice, giving rise to the so-called RPE puzzle. By conducting a theoretical analysis of the model “*CEO Power and RPE*” by Dikolli et al. (2018) this thesis investigates how CEO power affects the implementation of RPE. The model demonstrates that powerful CEOs can distort peer-group weights or influence peer selection to their advantage, leading to inefficient risk filtering and reduced firm value. In some cases, boards may optimally refrain from implementing RPE altogether when expected peer performance is high. The results of the model are critically examined in relation to the relevant literature, illustrated through a numerical example, and all mathematical proofs are independently derived. Overall, the findings contribute to a better understanding of how managerial power shapes compensation design and help explain the mixed empirical evidence on RPE.

Zusammenfassung

Diese Masterarbeit untersucht den Einfluss von Managementmacht auf die Gestaltung und Wirksamkeit der relativen Leistungsbewertung (Relative Performance Evaluation, RPE) in der Vorstandsvergütung. Während RPE theoretisch effizient ist, um systematische Risiken aus der Managementvergütung herauszufiltern, deutet die empirische Evidenz auf eine begrenzte Anwendung in der Praxis hin. Dieses Phänomen wird in der Literatur als „RPE-Puzzle“ bezeichnet. Durch eine theoretische Analyse des Modells „CEO Power and RPE“ von Dikolli et al. (2018) untersucht diese Arbeit, inwiefern die Macht des CEOs die Implementierung von RPE beeinflusst. Das Modell verdeutlicht, dass einflussreiche CEOs die Gewichtung und Auswahl der Peer Group zu ihrem Vorteil manipulieren können. Dies resultiert in einer ineffizienten Risikofilterung und mindert letztlich den Unternehmenswert. Darüber hinaus zeigt sich, dass es für den Aufsichtsrat optimal sein kann, vollständig auf die RPE zu verzichten, wenn die erwartete Performance der Peer Group hinreichend hoch ist. Die Modellergebnisse werden im Kontext der relevanten Literatur kritisch eingeordnet, anhand eines numerischen Beispiels veranschaulicht und sämtliche mathematischen Beweise eigenständig hergeleitet. Insgesamt tragen die Ergebnisse zu einem besseren Verständnis darüber bei, wie Managementmacht die Gestaltung von Vergütungssystemen beeinflusst, und liefern eine mögliche Erklärung für das RPE-Puzzle.

Table of Contents

1. Introduction	1
2. Theoretical Foundations	3
2.1 Fundamentals of CEO Compensation	3
2.1.1 CEO Compensation in Practice	3
2.1.2 Relative Performance Evaluation (RPE)	5
2.1.3 Determining Compensation and Managerial Power Theory	6
2.2 Fundamentals of Principal-Agent Theory	8
2.2.1 Introduction to Principal-Agent Theory	8
2.2.2 The Agency Problem in Public companies	9
2.2.3 The LEN Model	11
2.2.3.1 A General Agency Model	11
2.2.3.2 The Assumptions of the LEN Model	13
2.3 Literature Review	16
2.3.1 The Rise of CEO Compensation and Managerial Power	16
2.3.2 RPE Puzzle and Managerial Power Hypothesis	17
2.3.3 Alternative Explanations for the RPE Puzzle	18
3. Model Analysis	23
3.1 Model Assumptions	23
3.2 Relative Performance Evaluation and Efficient Risk Filtering	25
3.3 CEO Power and Relative Performance Evaluation	26
3.3.1 The Case of a Powerful Board of Directors	26
3.3.2 The Case of a Powerful CEO	28
3.3.2.1 CEO Choice of Peer-Group Incentive Weight	28
3.3.2.2 CEO Choice in Peer Selection Process	33
3.4 Implications for Empirical Analysis	36
3.4.1 Theoretical Derivation of Expected Regression Coefficients	36
3.4.2 Empirical Implications	39
3.4.2.1 Hypotheses on the Role of Powerful CEOs in RPE Design	39

3.4.2.2	Implications for Interpreting Conventional RPE Tests	40
3.4.2.3	RPE-based Indicator Variable of CEO Power	43
3.5	Illustration with a Numerical Example	44
3.5.1	CEO Choice of Peer-Group Incentive Weight	44
3.5.1.1	Baseline Scenario	44
3.5.1.2	Comparative Statics: The Effect of CEO Risk Aversion	48
3.5.1.3	Comparative Statics: The Effect of Market Risk	51
3.5.1.4	Comparative Statics: The Effect of Market Risk Exposure	54
3.5.1.5	Comparative Statics: Effect of Firm-Specific Risk	57
3.5.2	CEO Choice in Peer Selection Process	60
3.6	Discussion of Model Results	63
4.	Summary and Conclusion	69
5.	List of References	71
6.	Appendix – Proof of Results	76

List of Tables

Table 1 - Results for Baseline Scenario _____	45
Table 2 - Sensitivity Analysis for Low and High Risk Aversion _____	48
Table 3 - Sensitivity Analysis for Low and High Market Risk _____	51
Table 4 - Sensitivity Analysis for Low and High Beta _____	54
Table 5 - Sensitivity Analysis for Low and High Firm Risk _____	57
Table 6 – Numerical Illustration of Strategic Peer Selection by a Powerful CEO _____	60

List of Figures

Figure 1 - Decision Timeline in the Case of a Powerful BoD _____	23
Figure 2 - Decision Timeline in the Case of a Powerful CEO (Choice of Peer-Group Incentive Weight) _____	28
Figure 3 - Decision Timeline in the Case of a Powerful CEO (Choice in Peer Selection) ____	33
Figure 4 - Impact of Expected Peer Performance on the Decision to Implement RPE ____	45
Figure 5 - Effect of Risk Aversion on Expected Firm Value _____	49
Figure 6 - Effect of Risk Aversion on Firm Incentive Weights _____	49
Figure 7 - Effect of Risk Aversion on Peer Incentive Weights _____	50
Figure 8 - Effect of Market Risk on Peer Incentive Weights _____	52
Figure 9 - Effect of Market Risk on Firm Incentive Weights _____	52
Figure 10 - Effect of Market Risk on Expected Firm Values _____	53
Figure 11 - Effect of Market Risk Exposure on Peer Incentive Weights _____	55
Figure 12 - Effect of Market Risk Exposure on Firm Incentive Weights _____	55
Figure 13 - Effect of Market Risk Exposure on Expected Firm Values _____	56
Figure 14 - Effect of Firm-Specific Risk on Firm Incentive Weights _____	58
Figure 15 - Effect of Firm-Specific Risk on Peer Incentive Weights _____	59
Figure 16 - Effect of Firm-Specific Risk on Expected Firm Values _____	59
Figure 17 - Effect of Expected Peer Performance on Peer Incentive Weights _____	61
Figure 18 - Effect of Expected Peer Performance on Expected Firm Value and CE _____	61

1. Introduction

The design of executive compensation is an important instrument to align the interests of a firm's shareholders and management. Empirical studies document a significant increase in compensation payments over the past decades. There are various justifications and reasons for this trend. While some link it to the firms' economic growth others point to inefficiencies in the design of compensation contracts. A common theoretical justification lies in the inherent moral hazard problem that shareholders face when hiring top management. As a consequence, firms must implement incentives to guide managerial behaviour (Göx, 2016). A specific form of such incentives aimed at aligning executive behaviour with shareholder interests is Relative Performance Evaluation (RPE). This mechanism is based on the theoretical idea that Chief Executive Officers (CEOs) should only be evaluated and paid according to their individual performance. In fact, a theoretical advantage of RPE is that it can be designed to effectively filter out risk. Effective performance measurement systems exclude all factors that are beyond a CEO's control which should enable a more precise measure of managerial effort (Bolton & Dewatripont, 2005; Gibbons & Murphy, 1990).

Even though the theoretical models in executive compensation largely support the use of RPE, most empirical studies do not find any or only very little evidence for its use in practice. The literature describes this mismatch between theory and reality as RPE puzzle and various studies provide potential explanations for this observation (Core et al., 2003; Kabitz, 2017; Maug, 2000). One possible explanation is to consider the puzzle in the light of the managerial power hypothesis advocated most prominently by Bebchuk and Fried (2004). Their analysis suggests that managers can influence their own compensation agreements. Therefore, they will use this power to make their compensation as favourable as possible, by weakening the link between pay and performance.

The practical problem for firms that arises from managerial power in connection to RPE is that it can lead to excess compensation for the CEOs and therefore reduce the effectiveness of the incentive mechanism. As a consequence, the dilemma arises whether RPE should be implemented at all, considering that the managers' influence is likely to distort its intended purpose. To investigate this question, it is important to understand how and under which conditions CEOs can change the RPE mechanism to their advantage and the way in which the Board of Directors (BoD) will respond to this (Dikolli et al., 2018). This discussion on CEO power and its implications for compensation design is crucial, as it is connected to broader issues of corporate accountability, firm performance, and governance efficiency.

The objective of this thesis is to gain insights in the impact that managerial power can have on the optimal use and efficacy of RPE in executive compensation. If firms implement RPE, CEOs might use their power to influence peer-group selection or the incentive weight assigned to the peer index. I study the research question in a theoretical agency model adapted from the framework in Dikolli et al. (2018). My analysis sheds light on how managerial power affects both peer-group design and RPE implementation. Based on the theoretical foundations and existing literature in this field, I discuss the implications of managerial power on incentive design and the efficiency of risk filtering, before developing a set of empirical predictions.

The remainder of this thesis is structured as follows. Section **2** presents the theoretical foundations of the analysis. It introduces the fundamentals of CEO compensation and RPE, as well as the basic concepts of principal-agent theory. In addition, the LEN model is presented to establish the theoretical framework for the agency model. The section concludes with a literature review that discusses the most relevant contributions on RPE and managerial power. Section **3** contains the main analysis of the thesis. It presents the theoretical model and derives its key results. The section also discusses the empirical implications of the model and provides a numerical example to illustrate the comparative statics. The results are then discussed in relation to the existing literature, and the limitations of the model are highlighted. Section **4** concludes the thesis with a summary of the main findings. Detailed mathematical proofs are provided in the appendix.

2. Theoretical Foundations

2.1 Fundamentals of CEO Compensation

This section establishes the foundations of executive compensation by describing pay structures and the pay-setting process. It introduces RPE as a tool to filter market risk and contrasts the optimal contracting theory with the rent-extraction perspective of managerial power theory.

2.1.1 CEO Compensation in Practice

To understand the incentive effects of CEO pay, it is essential to first analyse how compensation packages are structured. In general, the design of such incentive contracts comprises three core components (Ewert et al., 2023):

- 1) **Types of Compensation:** This refers to the form in which pay is received. A primary distinction is drawn between tangible and intangible compensation. Intangible compensation includes non-financial rewards, such as titles or public recognition. Tangible compensation is further categorized into monetary (cash, equity) and non-monetary elements (tangible benefits like company cars or childcare services).
- 2) **Performance Measures:** These constitute the basis for evaluating managerial output. They can either be quantitative or qualitative. While the former can be assessed objectively through some key performance indicators, the latter depend on subjective judgement. Typically, firms use a mix of performance measures to incorporate multiple evaluation criteria.
- 3) **Compensation Function:** This defines the relationship between the performance measure(s) and the payout. This link can be linear, non-linear, or take the form of a bonus that is only granted if a specific performance threshold is reached.

In practice, these theoretical concepts translate into complex pay packages. Most executive compensation consists of five basic components: fixed salary, annual bonuses, payouts from Long-Term Incentive Plans (LTIPs), stock options and stock grants. In addition, executives often receive perks, defined pension plans and severance payments (Coenenberg et al., 2015; Conyon & Peck, 2012; Edmans et al., 2017; Gaßner, 2012).

The **base salary** is usually paid in cash and is fixed, as it does not depend on firm or managerial performance. Consequently, it carries no risk for the executive. This cash element is typically benchmarked against the market to ensure competitiveness (Conyon & Peck, 2012; Gaßner, 2012).

Over the last few decades, the use of **variable compensation** has increased significantly. This form of compensation is tied to performance measures and serves as the primary incentive mechanism for aligning managerial incentives with shareholder interests. Due to its variable nature, it incorporates risk for the CEO and is therefore also often referred to as “pay at risk”. Currently, it constitutes the main component of executive pay and typically includes annual bonus plans, which are usually tied to accounting metrics as well as stock options and stock grants, which are directly linked to the firm's share price (Conyon & Peck, 2012).

LTIPs are bonus plans based on multi-year performance, which are often paid out over several years, either in cash or stock. They are usually based on a combination of accounting metrics, where performance is measured over one or across multiple years. The payoff structure is often non-linear, where no bonus is paid until a minimum performance threshold is met, at which point the payout jumps discretely. Then, the bonus increases with performance within the so-called “incentive zone” until it hits a cap, after which it does not increase further. Critics argue that this non-linear structure creates incentives for executives to manipulate earnings to ensure they fall within the profitable incentive zone (Edmans et al., 2017).

The primary aim of granting **equity-based compensation** is to create a stronger link between shareholder value and CEO pay. **Stock grants** achieve this by directly tying managerial wealth to the company's share price. However, a significant drawback of this approach is that the CEO is exposed to external factors such as market fluctuations. Consequently, executives may be rewarded or penalized for general market movements rather than firm-specific performance (Coenenberg et al., 2015; Core et al., 2003). Alternatively, **stock options** give their owner the right, but not the obligation to purchase a share of the firm at a prespecified price at some date in the future. Since the value of an option increases as the company's value rises above the exercise price, it functions as a strong long-term incentive. However, unlike stock grants, the value of an option also depends positively on stock volatility. This can encourage excessive risk-taking rather than sustainable value creation (Coenenberg et al., 2015; Conyon & Peck, 2012).

A critical aspect of performance-based pay is that it is riskier and, therefore, often valued lower by the CEO than its market value. This is particularly relevant when CEOs are undiversified, having invested both their human capital and personal wealth in the same company. Consequently, a higher level of absolute compensation may be required to compensate for the additional risk stemming from performance-based pay (Core et al., 2003; Core & Guay, 2010; Edmans et al., 2017).

Finally, top executives often receive perks, which encompass a variety of goods and services, such as company cars, corporate jets, club memberships and housing. Along with pension plans and severance pay, these components are often less transparently disclosed than salary

or bonuses, yet they can account for a significant portion of total compensation (Edmans et al., 2017; Frydman & Jenter, 2010; Gaßner, 2012).

Recent empirical analyses by Clementi and Cooley (2023) document several critical characteristics of CEO compensation. First, they find that the distribution of compensation is extremely right skewed. This means that only a small number of CEOs earn incredibly high amounts of compensation, mostly driven by an increase in the valuation of their stock or option holdings. Contrary to the perception that executives always gain, the authors note that the significant exposure to stock prices through equity holdings implies that a sizeable fraction of CEOs suffer financial losses in any given year. Furthermore, they find that the variation in CEO pay and therefore also the incentive effect is driven by the deferred component of compensation. This component is defined as the yearly changes in the value of stock and options held in the portfolio, retirement benefits, and expected future salaries. In contrast, current cash payments display limited variation. The sensitivity of pay to performance is mostly driven by the revaluation of the CEO's existing portfolio of stock and options, rather than by adjustments to annual salary or bonuses. However, the authors find that this sensitivity has declined in the post-2008 period. They attribute this weakening of incentives to a general decline in the use of equity compensation as well as a shift in compensation practices away from stock options.

2.1.2 Relative Performance Evaluation (RPE)

As outlined in the previous subsection, variable compensation is often used to link managers' pay to their performance. An important issue in implementing such compensation is the choice of performance measures, as not all measures reflect the managers' level of effort equally well. Among these, RPE has been established as a tool to design executive compensation in a way that better reflects managerial actions.

A key critique of absolute performance measures is that they are affected by a series of random factors beyond the executives' control. This is inconsistent with the **controllability principle**, which states that fair compensation should depend solely on factors managers can influence (Milgrom & Roberts, 1992). Including uncontrollable factors in pay increases variability in compensation and exposes CEOs to unnecessary risk. Since this variability is random, it does not contribute to motivating or incentivizing managers (Coenenberg et al., 2015; Ewert et al., 2023). As an illustration, one can consider remuneration based on the share price. This performance measure depends on a variety of factors, which are not entirely under managerial control. As a result, CEOs may be penalized or rewarded for market fluctuations unrelated to their actions. One way to address this limitation is to extend the performance

measure to similar firms, such as competitors in the same industry. These firms are likely to be exposed to the same industry or market movements and by comparing them, common effects can be excluded. Accordingly, compensation can more accurately reflect managerial actions (Gibbons & Murphy, 1990).

RPE provides three primary advantages. First, it reduces managers' risk exposure and uncertainty by filtering out common shocks. This risk reduction may result in lower contracting costs, as a lower risk premium is demanded by managers. Second, it strengthens incentives and induces more effort provision, as pay is more closely tied to performance and therefore more directly linked to managerial decisions. Third, RPE can be perceived as a fairer system of pay, since it excludes factors beyond managerial control (Bolton & Dewatripont, 2005; Gibbons & Murphy, 1990). While RPE has clear advantages in theory, there are some limitations to its implementation in practice, which are discussed in the following sections.

2.1.3 Determining Compensation and Managerial Power Theory

To better understand the challenges associated with managerial pay, it is useful to examine the compensation setting process within firms. As shareholders are too numerous and diverse, they cannot set executive pay directly. Instead, this task is delegated to a set of institutions designed to represent their interests. The BoD is responsible for overseeing management on behalf of shareholders. In practice, the board delegates the detailed work of designing executive pay packages to a specialized compensation committee. Additionally, compensation consultants, which are external professionals with expertise in executive remuneration, are often employed to advise the BoD on market practices and pay structures. The literature suggests that the effectiveness of these institutions largely depends on their independence, as this reduces the risk of biased or self-serving decisions (Conyon & Peck, 2012; Florin et al., 2010). More recently, Say on Pay regulations have emerged, giving shareholders the right to a binding or advisory vote (depending on jurisdiction) on the approval of executive compensation during annual shareholder meetings. While not always decisive, Say on Pay provides an additional check on managerial remuneration and signals shareholder dissatisfaction when pay levels are perceived as excessive (Conyon & Peck, 2012; Edmans et al., 2017).

Within this institutional framework, two main theories attempt to explain the level of CEO pay. Under the **optimal contracting view**, executive compensation is seen as an optimal incentive system through which shareholders motivate executives to maximize firm value. Moreover, pay setting is believed to be the efficient outcome of a market mechanism where firms design compensation packages to attract and retain managerial talent in the labour market. In this

view, well-functioning boards and compensation committees negotiate pay levels that reflect the executive's market value. Consequently, high levels of pay are interpreted as an efficient contracting outcome resulting from the competition for managerial skills, the necessary compensation for their effort and the risk premiums associated with performance-based compensation (Core & Guay, 2010; Frydman & Jenter, 2010; Murphy & Zábojník, 2004).

However, when these institutions fail to operate effectively and independently, executives may be able to influence their own pay, thereby achieving higher compensation compared to what would be optimal under efficient contracting. This observation is central to the **managerial power theory** (Bebchuk & Fried, 2004), which argues that CEOs can “extract rents” by leveraging weaknesses in corporate governance, specifically the lack of independence of directors. Bebchuk and Fried (2004) argue that BoDs often possess economic and non-financial incentives to favour or at least get along with executives. This dynamic can weaken their commitment to maximizing shareholder value. For instance, directors might favour executives to maintain personal relationships, secure their re-election to the board, or defer to authority, thereby allowing compensation to exceed the level justified by performance and structure pay in a way that is less sensitive to performance. Consequently, managerial power and excessive pay can be viewed as the result of failures in the very institutional structures intended to protect shareholder interests. Nevertheless, the authors note that rent extraction is constrained by “outrage costs”, defined as the potential negative reaction from shareholders and the public. To minimize these costs, executives may engage in practices to obscure the true value of their pay, thereby ensuring that rent extraction remains less visible to outsiders.

2.2 Fundamentals of Principal-Agent Theory

The relationship between shareholders and management in a company is best described through principal-agent theory. This section explains the theoretical basics and introduces the LEN model, which serves as the foundation for the subsequent analysis.

2.2.1 Introduction to Principal-Agent Theory

Principal-agent theory provides a framework for understanding relationships where one party (the principal) delegates decision-making authority or tasks to another party (the agent), who is responsible for acting on the principal's behalf. In this setting, the outcome and thus the welfare of the principal depends on the agent's decisions. In return for performing these tasks, the agent receives compensation from the principal. Delegation is often necessary because the principal may lack the specific expertise, time or information possessed by the agent. However, **conflicts of interest** arise when the objectives of the two parties are misaligned. Since both seek to maximize their own utility, their preferred course of actions may diverge. Commonly, these conflicts of interest stem from the fact that higher levels of effort by the agent tend to result in better outcomes for the principal, yet exerting greater effort creates disutility for the agent. Therefore, the agent requires compensation to offset this cost, which in turn reduces the principal's net payoff. Fundamentally, the problem is that the agent does not bear the full financial consequences of their actions as they bear the full cost of their effort but capture only a portion of the resulting gain (Ewert et al., 2023; Küpper et al., 2024; Laffont & Martimort, 2002).

In a theoretical setting of perfect information, this would not pose a problem as the principal could directly observe the agent's effort and enforce a contract based on that specific effort level. However, information asymmetry impedes this direct monitoring, making it difficult to achieve the optimal outcome. The following paragraphs provide an overview of the three main forms of information asymmetry and their implications.

Hidden characteristics refer to a situation in which the principal lacks information about the agent's qualifications or properties prior to entering a contractual agreement. This can lead to a problem known as adverse selection, where contracts offered may fail to attract the most suitable agents or lead to market failure. To mitigate adverse selection, several mechanisms can be employed. On the side of the principal, screening methods such as interviews or tests can be used to gather additional information. On the agent's side, self-selection occurs when agents choose from a set of contract options, ideally selecting the one that best aligns with their private characteristics. Additionally, agents may engage in signalling by voluntarily

providing credible information (e.g., certifications, prior achievements) to convey their suitability to the principal (Küpper et al., 2024).

In the situation of **hidden information**, the agent possesses an informational advantage over the principal after the contract is established but before the decision is made. For instance, this advantage may stem from the agent's greater experience or expertise gained through performing the assigned tasks. Consequently, the agent may selectively disclose only the information that serves their own interest. This can give rise to moral hazard, where the agent chooses alternatives that maximize their own utility rather than the principal's. However, because of the agent's informational advantage, the principal cannot evaluate the decision and the resulting outcome. As a solution, the principal could design incentive systems that induce truthful reporting or implement control systems to monitor the agent's actions (Küpper et al., 2024).

Hidden action refers to a situation where the principal can observe the outcomes of the agent's actions, but not the specific activities or effort level. This creates a challenge, as the principal cannot distinguish whether results are due to the agent's efforts or external random factors. As effort is costly for the agent, this situation can lead to the problem of shirking, where the agent exerts minimal effort knowing the principal cannot directly assess their level of commitment. To address this issue, the principal can implement incentive structures and control mechanisms to align interests. They can also use informative signals to better assess the actions of the agent. As with hidden information, this scenario is also considered a form of moral hazard, where the agent may act in their own self-interest at the expense of the principal, because their actions are not directly observable (Küpper et al., 2024; Laffont & Martimort, 2002).

It is important to note that **agency conflicts** typically arise from the combination of conflict of interest and information asymmetry. The presence of only one of these conditions does not pose an incentive problem. If there were only a conflict of interest but information was symmetric (perfectly observable), then the principal could propose an incentive contract that enforces the optimal actions. Conversely, if there were information asymmetry but both parties had the same objectives, then the agent would naturally act in the principal's best interest without the need for incentive contracts (Laffont & Martimort, 2002).

2.2.2 The Agency Problem in Public companies

Typically, publicly traded companies exhibit dispersed ownership. This means that the owner and the manager of the company are not the same person, and shareholders are usually not involved in the day-to-day management of the firm. This situation is described as the **separation of ownership and control**. While executives are entrusted to manage the firm to

maximize shareholder wealth, this delegation grants them considerable discretion in decision-making, which they may utilize to pursue their own interests rather than those of the shareholders. Despite being the owners, shareholders have limited influence over managerial actions. Possible reasons for this are that direct monitoring is too costly and time-consuming, and that shareholders do not have the expertise to evaluate managerial decisions. Moreover, given the complexity and unpredictability of business environments, it is not feasible to write complete contracts that define desired managerial behaviour for every possible future contingency. Consequently, managerial actions are not directly observable. Shareholders must instead rely on observable outcomes, such as the share price, to infer executive effort. However, the share price is an imperfect signal as it depends not only on managers' actions but also on external factors such as market or industry shocks (Coenenberg et al., 2015; Ewert et al., 2023; Gaßner, 2012). This scenario is a classic example of the moral hazard problem resulting from hidden action, where the shareholders take the role of the principal, and the manager acts as the agent.

Principal-agent theory provides the fundamental theoretical framework for understanding corporate governance and the design of executive compensation contracts. It is used to study how to effectively incentivize managers in the presence of information asymmetry and moral hazard, ensuring that they undertake actions aligned with the principal's interests of maximizing shareholder value (Conyon & Peck, 2012; Gaßner, 2012). To mitigate agency conflicts, firms can implement incentive contracts that tie the executive's compensation to specific performance measures (Coenenberg et al., 2015; Gaßner, 2012). As the principal cannot directly observe or prescribe the agent's specific actions, they must instead influence the agent's decisions indirectly. The principal predicts how the agent will respond to various compensation structures and then selects the reward scheme that maximizes the principal's own wealth. However, this selection is subject to the constraint that the agent must be induced to choose the desired level of effort (Spremann, 1987). To effectively implement these reward schemes, companies utilize the variable compensation structures described in subsection **2.1.1**. Crucially, as outlined in subsection **2.1.2**, RPE aims to assist the principal in distinguishing the agent's actual performance from external factors.

2.2.3 The LEN Model

The Linear contracting, Exponential utility function, Normal distribution (LEN) model is a hidden action model that describes agency problems and offers valuable insights for the design of optimal incentive contracts. This subsection explains the peculiarities of the LEN model after first presenting the characteristics and structure of the general principal-agent model, which forms its foundation (Ewert et al., 2023; Küpper et al., 2024; Lambert, 2001; Spremann, 1987).

2.2.3.1 A General Agency Model

Within the principal-agent relationship, the following set-up is considered: The principal hires the agent to perform a task with outcome x . This outcome is driven by two factors: the action or effort exerted by the agent, a , and an exogenous risk factor outside the agent's control, ε . The output function is therefore $x = x(a, \varepsilon)$. The output could be interpreted as the company profit, revenue or shareholder value, which increases the principal's welfare. The principal, in turn, needs to compensate the agent for their contribution to this output (Ewert et al., 2023).

Without any information asymmetry, the principal could perfectly identify and distinguish the agent's actions from the external factor. Hence, they could design a compensation contract that depends only on a and thereby induces the agent to take the desired actions. In this case, the agent is exactly rewarded for their contribution. This scenario is described as the **first-best solution** and does, by definition, not include any incentive problems (Bolton & Dewatripont, 2005; Ewert et al., 2023).

However, under information asymmetry, the principal can only observe the outcome x , which includes the noise term ε . As only x is observable, compensation, w , must be based on this variable. A drawback of this measure is that it only provides partial information on a . Therefore, the agent could always claim that bad outcomes are due to the external factor, and good outcomes are entirely their contribution. Consequently, the principal cannot verify what the agent's real contribution to the outcome was. Because this creates a moral hazard problem, incentive contracts are needed to guide the agent's behaviour in a way that is desired by the principal (Ewert et al., 2023).

In designing the optimal incentive contract, it is crucial to consider the **agent's utility** function. Usually, it is assumed that it depends on the promised compensation, w , which can be seen as the agent's benefit, and the cost connected to their effort, a . In evaluating whether the contract should be accepted or not, the agent compares the expected utility from the contract with the expected utility from their next best alternative. This so-called **reservation utility** is exogenously given and provides a utility benchmark that the proposed incentive contract must

meet to be accepted. Another common assumption is the agent's risk aversion, which means they prefer certainty over uncertainty. Certain outcomes are therefore valued higher, and a premium is required for risky ones (Ewert et al., 2023).

As in most of the agency literature, the **agent's utility function** is assumed to be additively separable into a compensation and an action component, taking the following form:

$$H(w, a) = U(w) - k(a) \quad (1)$$

The action component, $k(a)$, captures the disutility of effort. Since effort is costly, this component reduces overall utility ($k'(a) > 0$ and $k''(a) > 0$). The compensation component, $U(w)$, reflects the monetary benefit derived from the contract. This function illustrates the agency conflict outlined in subsection 2.2.1. The agent faces a fundamental trade-off: while increasing effort raises the expected output for the principal, x , and in turn enhances the compensation for the agent, w , it also increases the personal cost of effort $k(a)$. Thus, the agent must balance the marginal gain in pay against the marginal disutility of effort (Küpper et al., 2024; Lambert, 2001).

The **principal's utility function**, G , is defined as the difference between the outcome, x , and the agent's compensation, w :

$$G(x - w) \quad (2a)$$

The function is increasing ($G' > 0$), reflecting the assumption that the principal prefers higher output. Moreover, the principal is assumed to be either risk neutral or risk averse ($G'' \leq 0$), with risk neutrality being the standard assumption in most models. The principal's utility is affected by the compensation function in two ways. First, an increase in compensation directly reduces the principal's benefit, as this represents a direct cost to them. Second, there is the incentive effect: since the agent's effort depends on their expected pay, w , the compensation contract indirectly influences the outcome, x . (Küpper et al., 2024; Lambert, 2001).

The **principal's decision problem** is to choose the optimal compensation function to maximize their expected utility:

$$\max_{w, a} E[G(x - w)] \quad (2b)$$

As the expected utility depends on the performance measure, x , which in turn depends on the agent's effort level, the agent's utility function needs to be considered as well. This introduces constraints into the optimization problem.

The first constraint ensures that the agent accepts the contract proposed by the principal. The **participation constraint** or **Individual Rationality (IR) constraint** requires the agent's

expected utility from the contract to be greater or at least equal to their reservation level of utility, $\bar{H}$.

$$E[H(w, a)] \geq \bar{H} \quad (3)$$

If this constraint is not satisfied, the agent will reject the contract, as they have a better alternative.

The second constraint captures the motivational or incentive function of the contract by ensuring that the agent will exert the effort level that is optimal from the principal's perspective. It anticipates the agent's choice of actions by assuming that they will choose them in a way that maximizes their expected utility H . This decision is defined by the **Incentive Compatibility (IC) constraint**:

$$a = \operatorname{argmax}_{a' \in A} E[H(w, a')] \quad (4)$$

The principal needs to incentivize the agent in a way that aligns the agent's self-interest with the principal's objectives. This reflects the central challenge in agency theory: both parties aim to maximize their own utility, but the agent's private action (effort) is not directly observable by the principal. By incorporating the IC constraint, the agency problem can be modelled as if the principal is selecting both the contract and the actions. However, the principal is constrained to choose them in a way that is compatible with the agent's incentives (Küpper et al., 2024; Lambert, 2001).

2.2.3.2 The Assumptions of the LEN Model

The previous subsection presented the general agency model in which the principal's problem is expressed as a constrained maximization problem. The compensation function is chosen to maximize the principal's expected utility while satisfying both the agent's participation and IC constraints. It is important to emphasize that the purpose of this model is not to derive a specific compensation contract, but rather to generate qualitative insights into the underlying incentive dynamics and trade-offs resulting from information asymmetry. To facilitate the analysis of the general model, further assumptions and simplifications are introduced through the **LEN framework** (Küpper et al., 2024; Lambert, 2001).

First, the utility function of the agent is assumed to be negative exponential, $U(w) = -e^{-r^A w}$, with r^A being the coefficient of risk aversion. This utility function exhibits constant absolute risk aversion (CARA), implying that the agent's attitude toward risk does not vary with wealth. This is an important assumption, as it ensures that the effect of incentives does not change with the level of wealth. Second, the performance measures (output) are assumed to follow a normal distribution, with error terms having 0 mean and a variance of σ^2 . A key property of the

normal distribution is that changes in the mean do not affect higher moments of the distribution (such as variance, skewness etc.). This simplifies the analysis, as a change in the mean does not produce any unintended effects on variance. Third, the compensation function is assumed to exhibit a linear relationship with the performance measures (output). This assumption reduces the complexity of the optimization problem and, in combination with the other assumptions, allows for closed-form first-order conditions. Even though this excludes non-linear compensation schemes, the linearity assumption is considered a useful approximation as it simplifies calculations and still captures the economically relevant comparative statics in many contexts. Together, these assumptions simplify the principal's problem significantly. They allow for a closed-form solution where the agent's certainty equivalent (CE) can be expressed as their expected compensation minus their cost of effort and a risk premium (half the variance times the coefficient of risk aversion). Consequently, the model can be solved analytically using first-order conditions and comparative statics can be derived more easily (Bolton & Dewatripont, 2005; Lambert, 2001; Spremann, 1987).

Under these assumptions, the output function is defined as $x(a) = \alpha a + \varepsilon$, representing marginal productivity times effort plus the normally distributed error term. The linear compensation contract is given by $w = vx + f$, with v representing the variable component and f the fixed compensation. The cost of effort is represented by the function $k(a) = \frac{1}{2}a^2$, and the risk-aversion levels of the principal and the agent are $r^P, r^A \geq 0$.

The principal's constrained maximization problem simplifies to maximizing their certainty equivalent, CE^P :

$$\max_{v,f,a} CE^P = (1-v)\alpha a - f - \frac{1}{2}r^P(1-v)^2\sigma^2 \quad (5a)$$

The CE of the principal can be decomposed into three components: their expected share of the total output, minus the fixed compensation paid to the agent, and minus a risk premium that captures their cost of bearing risk, which depends on their level of risk aversion and the variance of the outcome. This is subject to the agent's constraints:

$$a = \operatorname{argmax}_{a' \in A} \left\{ CE^A = v\alpha a' + f - \frac{1}{2}a'^2 - \frac{1}{2}r^A v^2 \sigma^2 \right\} \quad (5b)$$

$$CE^A = v\alpha a + f - \frac{1}{2}a^2 - \frac{1}{2}r^A v^2 \sigma^2 \geq \bar{H} \quad (5c)$$

Similarly, the CE of the agent reflects their expected total compensation, minus their personal cost of exerted effort, and minus their own risk premium, which depends on their risk aversion and the riskiness of their variable pay.

In the first-best solution (perfect information), the IC constraint (5b) is not considered, and effort can be directly contracted upon. Thus, no incentives are required to induce effort, and only risk-sharing considerations determine the optimal variable component of compensation. The solutions for v and a are:

$$v^{FB} = \frac{r^P}{r^A + r^P} \quad \text{and} \quad a^{FB} = \alpha \quad (6a)$$

The optimal level of effort, a^{FB} , equals marginal productivity. The variable compensation parameter, v^{FB} , depends on the risk aversion of both the principal and the agent. In this first-best scenario, optimal risk sharing is achieved: the more risk-averse the agent is, the lower their variable (risky) compensation share becomes. In the case of a risk-neutral principal ($r^P = 0$), this solution implies that the agent carries no risk ($v^{FB} = 0$) and the principal assumes all of it.

In the second-best solution, the IC constraint (5b) must be considered, which results in:

$$v^{SB} = \frac{\alpha^2 + r^P \sigma^2}{\alpha^2 + r^A \sigma^2 + r^P \sigma^2} > v^{FB} \quad \text{and} \quad a^{SB} = v^{SB} \alpha < a^{FB} \quad (6b)$$

As effort is not directly observable in the second-best case, the agent must be incentivized through a higher variable compensation, $v^{SB} > v^{FB}$. This forces the risk-averse agent to bear more risk, prompting them to demand a higher risk premium and ultimately inducing a lower effort level (Dierkes & Schäfer, 2008).

The second-best solution is characterized by a **risk-incentive trade-off**. If the principal is risk neutral, from a purely risk-sharing perspective it is optimal for the principal to bear all the risk, since bearing risk is costly for the risk-averse agent but costless for the risk-neutral principal. This results in $v^{FB} = 0$ and the agent receives a fixed wage. However, a fixed wage does not provide any incentive for the agent to increase effort. As effort is not observable, the agent could then shirk and still receive full pay. To mitigate this problem, the principal must include a variable compensation component that is tied to the observable outcome. While this strengthens incentives, it also exposes the agent to additional risk. Therefore, the principal needs to compensate the agent for bearing this extra risk by adding a risk premium to compensation. This creates a trade-off for the principal, who must balance the additional costs of exposing the agent to risk against the expected benefits of stronger incentive effects. Because of this trade-off, the variable compensation share needs to be larger in the second-best scenario $v^{SB} > v^{FB}$, but the resulting effort level will still be lower than in the first-best case, $a^{SB} < a^{FB}$ (Ewert et al., 2023; Küpper et al., 2024). Ultimately, information asymmetry reduces the CE of the principal. The reduction in the principal's utility when moving from the first-best to the second-best solution is described as **agency costs** (Dierkes & Schäfer, 2008; Ewert et al., 2023; Küpper et al., 2024).

2.3 Literature Review

The literature review focuses on two streams of research: the determinants of executive compensation through the lens of managerial power and the RPE puzzle with its possible explanations.

2.3.1 The Rise of CEO Compensation and Managerial Power

Executive compensation has become a subject of intense public debate, largely driven by criticism of perceived excessive pay levels and their steep rise in recent years. Critics often argue that these high pay levels are not justified by managerial or firm performance. However, data suggest that alongside executive compensation, the financial performance and market capitalization of listed companies have also increased (Göx, 2016).

As outlined in subsection 2.1.3, the main determinants of executive compensation and its growth are typically explained through the lenses of optimal contracting and managerial power theory. Another significant driver of rising CEO compensation is the practice of competitive benchmarking. Commonly, firms set executive pay in relation to a reference group of peers. This often leads firms to target above-average pay to remain competitive, which in turn triggers a continuous increase in compensation (DiPrete et al., 2010; Göx, 2016). Moreover, social-psychological factors within the interaction of the board and the CEO are also important and can help better understand the executive compensation setting process (O'Reilly & Main, 2010). Furthermore, Murphy and Sandino (2010) find that CEO pay is significantly higher when external compensation consultants provide additional services to the firm. This correlation suggests an agency conflict because consultants might recommend higher executive pay to ensure that they maintain their business relationships with management. However, the authors note that this correlation does not prove causation, as powerful managers might purposely hire compensation consultants that are known to recommend higher pay.

Empirical research typically measures managerial power using proxies related to corporate governance and ownership structure. Common indicators include the proportion of independent directors on the board, CEO duality, ownership stakes, CEO tenure and the presence of institutional investors. Weaker corporate governance is generally associated with greater managerial influence over compensation decisions (Bebchuk & Fried, 2004; Choe et al., 2014; van Essen et al., 2015). While these measures provide useful proxies, they capture only specific dimensions of managerial power and may not fully reflect the underlying bargaining position of the CEO. To better understand how executives influence their pay, Grabke-Rundell and Gomez-Mejia (2002) extend the traditional agency framework by explicitly incorporating different forms of managerial power. They criticize standard agency

theories for rarely measuring power directly and instead only acknowledging the existence of power imbalances. To address this gap, the authors utilize four dimensions of executive power originally identified by Finkelstein (1992) and propose that they directly increase compensation. The first is structural power, which stems from the manager's hierarchical position and authority within the firm. Next, ownership power is derived from the voting rights and influence granted by the shares owned by the executive. The third form is expert power that comes from the executive's organizational knowledge and expertise, allowing them to exploit information asymmetry to their advantage. Lastly, prestige power originates from the executive's social status and visibility. Ultimately, the authors argue that these specific behavioural dimensions of power must be incorporated as independent variables within agency theory to accurately predict executive pay.

Moreover, Bertrand and Mullainathan (2001) find that CEOs often have influence over their own pay, allowing them to exploit pay-for-performance structures to be rewarded for lucky outcomes. They define "pay for luck" as compensation linked to changes in firm performance that depend on external factors beyond the executive's control. As a result, when a firm performs well, CEOs receive higher pay even if the positive performance of the firm is not directly related to their actions. The authors demonstrate that this phenomenon is strongest among poorly governed firms. Furthermore, their evidence suggests an asymmetry in this reward system, where CEOs are paid for good luck but are not equally punished for bad luck.

2.3.2 RPE Puzzle and Managerial Power Hypothesis

One of the most debated areas where managerial power exerts its influence is the application of RPE. The term **RPE puzzle** describes the divergence between the theoretical optimality of RPE and its real-world application. As explained in subsection 2.1.2, RPE allows the design of more efficient and powerful incentive contracts by filtering out the "noise" of systematic risk from performance measures. Despite these theoretical benefits, empirical researchers have consistently failed to find widespread evidence of its implementation in practice. Numerous researchers have tried to solve this RPE puzzle by proposing possible explanations for the absence of RPE.

A prominent explanation of the puzzle is the **managerial power theory**, already explained in detail in subsection 2.1.3. Within this theory, powerful CEOs prevent or manipulate the use of RPE in their contracts to secure more favourable compensation outcomes. Because these CEOs can influence the design of their contracts, they secure a large share of their compensation as non-performance-based pay (Bebchuk & Fried, 2004). Expanding on this, Morse et al. (2011) examine the effect of CEO power on the design of incentive contracts. They find that when boards are weak or compromised, powerful CEOs can manipulate the

way their performance is evaluated. Specifically, executives use their influence to shift the weight of their incentive pay toward the performance metrics that they already know are performing well. This results in a distortion of the incentive mechanism induced by performance pay.

However, empirical evidence suggests that while managers can and do influence their compensation, managerial power theory alone cannot explain all executive pay outcomes. Even if managerial power can significantly impact the pay setting process, optimal contracting arrangements can simultaneously exist. Consequently, rather than being mutually exclusive, these two theoretical mechanisms likely coexist and describe different types of executive compensation contracts within agency theory (van Essen et al., 2015).

Furthermore, managerial power theory faces significant criticism. Weisbach (2007) agrees that optimal contracting theory does not perfectly reflect executive compensation practices because directors are rarely entirely independent and executives often hold some influence over their own pay. However, he argues that while Bebchuk and Fried (2004) successfully highlight the flaws of the current system, they fail to provide a practical alternative. Ultimately, Weisbach concludes that the extraction of rents by executives might not simply be a correctable governance flaw that can be fixed by regulation. Instead, it is likely an unavoidable cost of dispersed ownership and the modern form of organisation. Additionally, some research shows that weak corporate governance does not necessarily imply the absence of RPE. Under certain conditions, powerful managers may actually receive contracts with higher pay for performance sensitivity. Whenever a friendly board offers a contract with a higher pay-performance sensitivity, it also makes more intensive use of RPE. Therefore, the interpretation of the use or lack of RPE as an indicator for managerial power must be made carefully (Göx & Hemmer, 2020).

2.3.3 Alternative Explanations for the RPE Puzzle

Managerial power theory is not the only attempt to explain the RPE puzzle. Another stream of literature suggests that the RPE puzzle may stem from **imperfect testing methodologies** rather than an actual absence of RPE in practice. Theoretical work by Dikolli et al. (2013) demonstrates that because peer groups identified by researchers often differ from those used by boards, empirical models are biased toward finding no evidence of RPE.

Empirical studies support this view. Albuquerque (2009) argues that previous studies failed to detect RPE because they relied on misspecified peer groups, such as broad market indices or industry peers, which fail to accurately capture the specific common shocks a firm faces. She argues that firm size is an important characteristic to consider when constructing peer groups within an industry, as both the type of shocks and a firm's ability to respond to them

are functions of firm size. Consequently, by constructing peer groups matched on both industry and size, she finds robust evidence that CEO compensation is negatively associated with peer stock performance. This demonstrates that firms do utilize RPE when the benchmark is adapted to the firm's specific environment.

Similarly, Gong et al. (2011) argue that the RPE puzzle stems at least partly from the limitations of the traditional implicit tests used to detect RPE. They state that implicit tests, which regress CEO pay on a broad industry index, rely on simplified assumptions about the RPE contract. This simplification means that the tests fail to reflect the actual peer-group composition used by the firm, as they rely on a generic industry index instead. By leveraging the 2006 SEC disclosure requirements to identify explicit peer groups, their test shows a strong negative association between pay and peer performance. This negative association is an indicator of RPE, as it implies that common shocks affecting both the firm and its peers are being filtered out. Bettis et al. (2014) add that implicit RPE tests suffer from multiple forms of misspecification, beyond just the peer group. Such misspecifications may include incorrect performance metrics, incorrect benchmarks, or incorrect functional forms. Most recently, Jayaraman et al. (2021) claim that the inconsistent evidence on RPE stems from an imprecise categorization of peers. In contrast to previous researchers, they group peer firms according to product similarity, thereby ensuring that they face similar demand and supply shocks.

Beyond these methodological issues, RPE seems to be abandoned when an informative peer group cannot be formed. Albuquerque (2014) illustrates this by analysing the case of high growth-option firms. Because these companies operate in highly volatile environments and face fewer direct competitors, it is extremely difficult to find peer groups that share their exact risk exposure. Consequently, she finds that high growth-option firms rely significantly less on RPE. Thus, firms that are uncertain about the right peer group are less likely to use RPE (De Angelis & Grinstein, 2011).

Other authors explain the limited empirical evidence of RPE by arguing that managers can **privately hedge against systematic risk** using capital markets. Consequently, the effect of RPE is limited in an environment where managers privately trade on all assets except their own company's stock. By adjusting their private investment portfolios, managers can alter the risk exposure of their compensation packages and effectively neutralize undesired common risk. Because the manager creates this "home-made" benchmark, explicit RPE implemented by the firm becomes redundant. Therefore, even if RPE is not observed in the data, there is no loss in efficiency, as managers effectively filter out systematic risk themselves. This creates an implicit RPE that arises naturally when executives manage their private equity portfolios (Core et al., 2003; Jin, 2002; Maug, 2000). Consistent with this view, Garvey and Milbourn

(2003) find that while RPE is absent for the average wealthy executive who can self-hedge, it is strongly present for younger and less wealthy executives who face capital constraints and cannot easily trade to offset market risk in compensation. Furthermore, because managers cannot hedge against non-tradeable risks such as accounting profits, Maug's (2000) model predicts that explicit RPE should primarily be observed in contracts based on accounting benchmarks.

Additionally, De Angelis and Grinstein (2020) propose a **"Talent Retention Explanation"** for the use of RPE. Departing from the traditional view that RPE is designed primarily to provide incentives or filter risk, they argue that RPE serves as a commitment mechanism to retain talent by compensating CEOs based on their revealed ability relative to the market. If a CEO outperforms their peers, the labour market perceives them as having higher talent, which increases their outside employment opportunities. By conditioning compensation on peer performance, firms can automatically adjust pay to match the CEO's rising outside options and incentivize them to stay. However, the authors argue that this mechanism only holds when CEO talent is transferable. If talent is firm-specific, it holds less value to other firms and thus RPE is less necessary for retention. This is a possible explanation for the RPE puzzle as the authors find that firms are more likely to refrain from using RPE when the market for CEO talent is less efficient or when talent is not easily transferred.

Oyer (2004) agrees that the use of RPE depends on an executive's outside employment opportunities and the firm's need for CEO retention. However, he argues that an executive's outside value fluctuates with the broader market. Therefore, firms intentionally allow compensation to fluctuate with market trends rather than using RPE to reward executives strictly for their individual performance. This automatically aligns executive pay with outside job opportunities and ensures that the compensation package remains competitive. The drawback of this approach is that executives might receive excess pay in certain economic states. However, because adjusting or renegotiating contracts is assumed to be highly costly, firms accept this trade-off. This method allows them to automatically match outside offers and retain management without the need to renegotiate compensation agreements.

Ultimately, while both theories rely on the need to retain executives to explain compensation design, they differ in what drives an executive's market value. De Angelis and Grinstein (2020) link rising outside opportunities to individual relative performance, which is why they use RPE to match the CEO's market value. Conversely, Oyer (2004) links rising opportunities to general market booms, which leads firms to avoid RPE so that pay automatically increases with market movements.

More recently, DeMarzo and Kaniel (2023) consider a setting in which CEOs, in addition to their absolute wage, also care about how their compensation compares to that of their peers. They show that compensation benchmarking can offset performance benchmarking, potentially leading to a positive weight on peer performance and thereby explaining observed deviations from standard RPE.

Finally, industrial organization theory offers the “**Strategic Interaction Hypothesis**” proposed by Aggarwal and Samwick (1999). They argue that RPE is not only a risk-filtering tool, as standard agency theory suggests, but also an incentive mechanism that influences market competition. Specifically, when managers are rewarded for outperforming their peers, they are incentivized to take aggressive actions that undermine rival performance. While this improves the firm’s relative standing, it can erode overall industry profits. To mitigate this, Aggarwal and Samwick argue that firms that operate in markets with a need for softened competition such as strategic complements, will apply a positive weight on peer performance. This finding contradicts the predictions of standard agency theory, suggesting that firms in more competitive environments utilize less RPE and instead link pay positively to the performance of their rivals.

This explanation is empirically refined by Vrettos (2013), who differentiates between the types of strategic competition present in an industry. He finds that firms place a negative weight on peer performance when they compete as strategic substitutes but place a positive weight when they compete as strategic complements. Hence, Vrettos argues that the RPE puzzle arises because these directionally opposite weights cancel each other out when substitute and complement industries are aggregated in broad samples, thereby yielding inconsistent evidence for RPE.

Finally, more recent work by Bloomfield et al. (2023) also considers the effects of RPE on competition and finds that some firms avoid RPE to prevent incentivizing aggressive competitive behaviour, or “costly sabotage”. Exploiting data on cartel membership, they find that cartel members are significantly more likely to use RPE because the cartel agreement prevents such aggressive behaviour, allowing firms to safely utilize RPE for risk reduction.

In summary, the literature on the RPE puzzle highlights a gap between the theoretical advantages of RPE and its inconsistent empirical use in practice. Existing explanations range from managerial influence over contract design to methodological measurement issues, hedging opportunities, talent retention motives, and strategic competition effects. Overall, the literature suggests that executive compensation outcomes are driven by multiple interacting mechanisms rather than a single dominant theory. However, while managerial power is

frequently cited as a key explanation, less attention has been paid to how it affects the specific design of RPE contracts. In particular, there is limited theoretical analysis of how CEO influence alters key contract parameters such as peer-group selection or the weight placed on peer performance, and under which conditions RPE may become ineffective or abandoned altogether. This thesis contributes to this strand of the literature by analysing how managerial power affects the design and implementation of RPE within a formal agency model.

3. Model Analysis

This chapter presents the theoretical model adapted from Dikolli et al. (2018). It outlines the underlying assumptions, derives the main results, and discusses their implications for empirical analysis. The model's mechanisms and comparative statics are illustrated through a numerical example. The chapter concludes with a discussion of the model's key insights and limitations.

3.1 Model Assumptions

In this section, the main assumptions of the model are outlined. To implement RPE, the risk-neutral BoD uses a publicly available market index to filter out market risk from managerial compensation. In the standard setting, the risk-averse CEO can then accept or decline the contract offered by the BoD. In the case of acceptance, the CEO provides an unobservable effort, a . Lastly, the performance of the firm as well as that of the index-based peer group are observed, and the CEO is compensated accordingly. The decision timeline is represented in *Figure 1*.

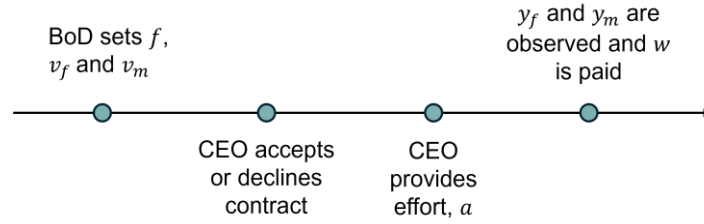


Figure 1 - Decision Timeline in the Case of a Powerful BoD

The **performance of the firm** is given by:

$$y_f = a + \varepsilon_f + \beta_f \varepsilon_m \quad (7)$$

implying that it depends on the unobservable effort of the CEO, a , the firm-specific risk, ε_f , and the market risk, ε_m , multiplied by the impact of market risk on firm performance, β_f . To be consistent with empirical data and for presentation clarity, it is assumed that $\beta_f > 0$. This implies a positive correlation between firm performance and market performance. Consistent with the assumptions of the LEN model, the noise terms are independently and normally distributed with $\varepsilon_f \sim N(0, \sigma_f^2)$, $\varepsilon_m \sim N(0, \sigma_m^2)$ with $\sigma_m^2 > 0$, and $Cov[\varepsilon_f, \varepsilon_m] = 0$.

The **performance of the index-based peer group** consists of two components, which are the expected performance of the peer group, μ_m , and the market risk, ε_m :

$$y_m = \mu_m + \varepsilon_m \quad (8)$$

with $\mu_m \geq 0$.

The BoD offers a linear compensation contract to the CEO, which in addition to a fixed wage component, f , depends on the performance of the peer group as well as the performance of the firm itself. v_f and v_m denote the incentive weights on firm and peer-group performance, respectively. The **CEO's compensation**, w , is expressed as follows:

$$w(y_f, y_m) = f + v_f y_f + v_m y_m \quad (9)$$

The CEO is assumed to be risk-averse and to exhibit a negative exponential utility function. Their utility depends on the received compensation, as well as the provided effort. Effort is costly for the CEO, and the cost of effort is denoted by the function $k(a)$, which is increasing ($k'(a) > 0$) and strictly convex ($k''(a) > 0$). An increasing function implies that the more effort is exerted, the more costly it gets, while strict convexity suggests that the marginal cost of effort increases with every additional unit of effort. Moreover, it is assumed that if the CEO does not exert any effort, then their cost is 0, so $k(0) = 0$, and $k'(0) = 0$ indicating that the marginal cost of effort at 0 is 0, such that the cost of an initial small increase in effort is negligible.

The **utility of the manager** is expressed by the following constant absolute risk aversion (CARA) function:

$$U(w, a) = -e^{-r(w-k(a))} \quad (10a)$$

where $r > 0$ is the coefficient of absolute risk aversion. CARA is an assumption of the LEN Model as explained in subsection 2.2.3. It means that the agent's risk aversion is constant for all levels of wealth, such that it does not depend on w .

The **CEO's certainty equivalent** (CE) represents the effective, risk-adjusted value of the compensation contract for the CEO. It therefore shows how much the contract is worth to the CEO after accounting for effort costs and risk. This is a standard result of the LEN Model (see appendix A.1a-e) and is given by:

$$CE(w, a) = E[w] - k(a) - \frac{r}{2} \text{Var}[w] \quad (10b)$$

That is, it equals the expected compensation, $E[w]$, minus the cost of effort, $k(a)$, minus a risk penalty which is proportional to the variance in compensation, $Var[w]$, and the risk aversion of the CEO, r . It therefore shows the total utility the CEO derives, considering their level of risk aversion, and not just the expected value.

Given this, the BoD **maximises the expected firm value**, Π , which can be expressed as the expected value of the firm's performance subtracting the costs of the CEO's compensation. This maximization includes choosing the fixed pay, the incentive weights on firm and peer performance and anticipating the effort the CEO will provide.

$$\max_{f, v_f, v_m, a} \Pi = E[y_f - w] \quad (11a)$$

The maximization needs to consider two constraints. The **IR constraint** ensures that the CEO accepts the contract offered by the BoD. This means that they must at least get their reservation utility, which is given by the CE (10b), normalized to 0.

$$s. t. (IR) : CE(w, a) \geq 0 \quad (11b)$$

The **IC constraint** ensures that the CEO chooses the intended effort level. The CEO chooses their effort level by equating the marginal benefit of effort with the marginal cost of effort. As shown in the appendix, this expression is derived from the effort level the CEO chooses when maximizing the CE from (10b).

$$(IC_a) : v_f = k'(a) \quad (11c)$$

3.2 Relative Performance Evaluation and Efficient Risk Filtering

The key theoretical justification for RPE lies in its ability to filter risk efficiently. One way to implement RPE in executive compensation is to compare the firm's performance to that of a predetermined peer group. Such groups typically consist of either a custom set of peers, a broad market index or an industry-specific index. According to Holmstrom's (1982) **informativeness principle**, managerial performance should be evaluated only using metrics that are informative about the manager's effort or talent. Therefore, efficient peer-group selection should remove common shocks to firm performance that are beyond managerial control. This improves the informativeness and accuracy of performance assessment. Moreover, it allows for better risk sharing between shareholders and managers, as the variability in compensation of the risk-averse managers is reduced. Lower pay variance also reduces the risk premium the BoD needs to offer, which in turn results in lower contracting costs and compensation. Finally, efficient risk filtering provides stronger incentives because managers' pay more directly reflects their actions and is less affected by uncontrollable factors.

In line with this logic, a well-designed RPE should rely on the peer group with the strongest filtering properties to remove systematic risk (Bakke et al., 2020; Bizjak et al., 2022; Ma et al., 2018).

Despite these theoretical advantages, the implementation of RPE in practice is often imperfect. A major challenge in peer-group benchmarking is the potential for bias in its implementation and peer-group selection. This form of performance measurement can be used opportunistically to inflate pay levels. Agency conflicts underlie this problem, as powerful CEOs may influence which firms are included in the peer group and choose those that are beneficial to them. Empirical evidence shows that peer groups are often chosen in a manner that biases compensation upward, particularly among firms outside the Standard & Poor's 500 (Bizjak et al., 2011). Hence, agency conflicts may explain why efficient risk filtering is not always observed in practice, despite its strong theoretical basis. Typically, the BoD determines the components of RPE compensation, including the peer group and payout structure. However, executives often participate in this process, which can lead to suboptimal peer groups with poorer filtering properties, higher costs and weaker incentives (Bizjak et al., 2022). Moreover, powerful CEOs may influence how performance measures are weighted, shifting incentives toward metrics that are more favourable to them (Morse et al., 2011).

3.3 CEO Power and Relative Performance Evaluation

This section first presents the benchmark case of a powerful BoD and then analyses how the results change when the CEO holds power over the contract design.

3.3.1 The Case of a Powerful Board of Directors

In the case of a powerful BoD, the BoD has absolute power over the design of RPE, and the CEO has no influence over it. Below, I characterize the solution to the board's decision problem (11a-c) under this assumption. The BoD chooses both incentive weights $v_f^\dagger$ and $v_m^\dagger$. The symbol “ $\dagger$ ” denotes the solution with a powerful BoD. Proofs are given in the appendix.

The CEO's choice of equilibrium effort, $a^\dagger$, is given by $v_f^\dagger = k'(a^\dagger)$ which represents the IC constraint in equation (11c). The IR constraint binds, such that the CEO gets exactly their reservation utility, which corresponds to the CE (10b).

OBSERVATION 1. *With a powerful BoD, the optimal incentive weights on firm and peer-group performance are:*

$$v_f^\dagger = \frac{1}{1 + rk''(a^\dagger)\sigma_f^2} \quad \text{and} \quad v_m^\dagger = -v_f^\dagger \beta_f \quad (12a)$$

And the resulting expected firm value is:

$$\Pi^\dagger = E[y_f | v_f^\dagger] - k(a^\dagger) - \frac{r}{2} (v_f^\dagger)^2 \sigma_f^2 \quad (12b)$$

$v_m^\dagger$ is the incentive weight on peer-group performance that minimizes the exposure of managerial compensation to market risk. With $v_m^\dagger < 0$, the CEO is offered an RPE-based contract that filters out market risk from compensation. As shown in the following equation, the CEO's compensation is only exposed to firm-specific risk ε_f :

$$w^\dagger = k(a^\dagger) + \frac{r}{2} (v_f^\dagger)^2 \sigma_f^2 + v_f^\dagger \varepsilon_f \quad (13)$$

The equation shows that executive compensation consists of three main components. First, the CEO's cost of effort. Second, a risk premium for bearing the firm-specific risk and the associated variance in wage, whose value depends on the risk aversion coefficient of the CEO. Third, the random component of pay that depends on firm-specific shocks. Note that the effect of market shocks has been fully eliminated.

An important insight is that the compensation contract filters market risk entirely from CEO pay, and therefore it is optimal to use RPE. The aggregate performance measure for the contract is given by:

$$y_f + y_m \cdot \frac{v_m^\dagger}{v_f^\dagger} \quad (14)$$

Note that this expression is equivalent to a weighted sum of firm and peer performance, given by $y_f v_f^\dagger + y_m v_m^\dagger$. This measure optimally uses the information about the CEO's effort by filtering out the noise from external factors, while preserving the incentive effect. In the context of the risk-incentive trade-off, this improves the compensation contract from a risk-sharing perspective, as it reduces the variance of pay, without weakening the incentive effect.

3.3.2 The Case of a Powerful CEO

In this subsection two forms of CEO power are considered, one where they can choose the incentive weight on a publicly observable aggregated measure of peer-group performance which is covered in 3.3.2.1 and one where they can influence the peer-selection process covered in 3.3.2.2.

3.3.2.1 CEO Choice of Peer-Group Incentive Weight

With a powerful CEO, the model setting changes slightly, because the BoD cannot choose the incentive weight on peer-group performance. Instead, the BoD proposes an incentive contract that solely specifies the fixed wage, f , and the incentive weight on firm performance, v_f . Contrary to the case of a powerful BoD, the CEO not only decides whether to accept the contract, but if it is accepted, also determines the incentive weight on peer-group performance, v_m . Then, the CEO exerts unobservable effort, a . Lastly, both performance measures are observed, and the CEO is compensated accordingly. *Figure 2* depicts the decision timeline for this case.

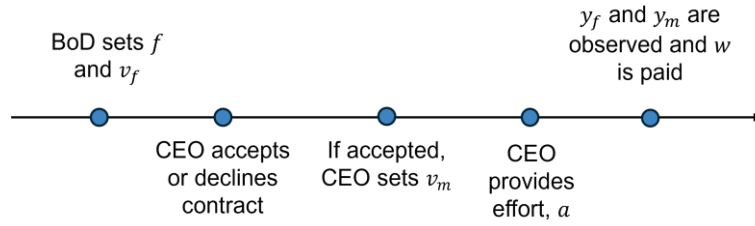


Figure 2 - Decision Timeline in the Case of a Powerful CEO (Choice of Peer-Group Incentive Weight)

The CEO chooses the weight on the peer group, v_m , such that it maximizes their CE given by the following equation. Note that this expression is derived by plugging the definitions of $E[w]$ and $Var[w]$ from (A.2) and (A.3) into (10b):

$$CE = f + v_f a + v_m \mu_m - k(a) - \frac{r}{2} \left(v_f^2 \sigma_f^2 + (v_f \beta_f + v_m)^2 \sigma_m^2 \right) \quad (15)$$

Unlike the case presented in subsection 3.3.1, the CEO does not necessarily choose the risk-minimizing peer weight, $v_m^\dagger = -v_f \beta_f$. This choice would eliminate all market risk from compensation, as σ_m^2 cancels from equation (15). However, it would also forgo potential benefits from positive expected peer-group performance, because when $v_m < 0$, the term $v_m \mu_m$ reduces the CEO's CE. Thus, the CEO faces a trade-off between minimizing the risk in

compensation by filtering out market shocks ($v_m < 0$) and increasing their expected pay by leveraging positive peer-group performance ($v_m > 0$).

The CEO's decision on v_m can be characterized through this first-order condition:

$$(IC_m): \frac{\partial CE}{\partial v_m} = \mu_m - r(v_f \beta_f + v_m) \sigma_m^2 = 0 \quad (16)$$

which yields $v_m = -v_f \beta_f + \frac{\mu_m}{r \sigma_m^2}$. It becomes clear that the CEO's choice depends on the weight on firm performance, v_f , determined by the BoD, as well as the expected performance of the peer group, μ_m . It is important to note that the CEO only makes the same choice as the powerful BoD (subsection 3.3.1) if expected peer-group performance is equal to $\mu_m = 0$. Then, their choice is the risk-minimizing incentive weight $v_m^\dagger = -v_f^\dagger \beta_f$ from (12a).

In the case of a powerful CEO, the BoD's decision problem from (11a) to (11c) is extended by an additional constraint, that reflects the CEO's choice of the peer incentive weight from equation (16). The symbol " $\ddagger$ " denotes the solution with a powerful CEO. In this case, the optimal incentive weights on firm performance and peer-group performance, $v_f^\ddagger$ and $v_m^\ddagger$, and the resulting expected firm value, $\Pi^\ddagger$, are given by the following equations. Note that equilibrium effort, $a^\ddagger$, is given by $v_f^\ddagger = k'(a^\ddagger)$.

PROPOSITION 1. *With a powerful CEO, the optimal incentive weights on firm performance and peer-group performance, $v_f^\ddagger$ and $v_m^\ddagger$, and the resulting expected firm value, $\Pi^\ddagger$, are given by:*

$$v_f^\ddagger = v_f^\dagger \quad \text{and} \quad v_m^\ddagger = v_m^\dagger + \frac{\mu_m}{r \sigma_m^2} \quad (17a)$$

$$\Pi^\ddagger = \Pi^\dagger - \frac{\mu_m^2}{2r \sigma_m^2} \quad (17b)$$

The powerful CEO introduces a distortion to the optimal peer incentive weight, $v_m^\dagger$. This distortion is represented by $v_m^\ddagger - v_m^\dagger = \frac{\mu_m}{r \sigma_m^2}$, and depends on expected peer performance, risk aversion and market risk. It is important to note that the deviation in the peer incentive weight is independent of the chosen firm incentive weight, v_f , and effort, a . Therefore, the BoD chooses the same optimal incentive weight, v_f , as in the benchmark case, (12a), so that $v_f^\ddagger = v_f^\dagger$. Adjusting v_f , does not mitigate the CEO's distortion in v_m . Instead, any increase in v_f introduces an additional inefficiency by increasing compensation risk and, thus, the required risk premium, while rent extraction through the distorted v_m remains possible.

Because the optimal v_f is unchanged, the induced effort level also remains unchanged. The level of effort is optimal, in that it balances the marginal cost of effort with the marginal benefit of inducing effort. Hence, the powerful CEO does not weaken incentives, and the distortion does not arise from inefficient effort provision. Rather, the inefficiency stems from distorted risk sharing. By choosing a higher (less negative) v_m , the CEO accepts additional market risk in order to benefit from positive expected peer performance. This increases the required compensation for bearing risk, resulting in a lump-sum cost to expected firm value equal to $\frac{\mu_m^2}{2r\sigma_m^2}$.

Moreover, it is important to point out that the CEO, despite being powerful, does not earn additional rents relative to their reservation utility, as the BoD adjusts the fixed wage downward by exactly the extra expected benefit the CEO could otherwise capture. This means that, anticipating the CEO's choice of v_m , the BoD ensures that the IR constraint is binding in equilibrium. In terms of the model, this implies that the CE from equation (15) is set to 0. In summary, the firm is worse off, because its expected value is reduced by the risk premium. Comparing equation (12b) to (17b), it follows that $\Pi^\ddagger < \Pi^\dagger$, but the CEO's utility does not improve as it remains at the reservation level ($CE = 0$).

As the powerful CEO's choice of $v_m^\ddagger$ deviates from the risk-minimizing incentive weight, $v_m^\dagger$, market risk, ε_m , is not completely removed from the CEO's compensation. Hence, the powerful CEO prevents market risk from being fully filtered out of compensation. As explained above, the CEO aims to benefit from potential positive peer performance.

COROLLARY 1. *With a powerful CEO, market risk, ε_m , is not eliminated from the CEO's compensation, that is:*

$$w^\ddagger = w^\dagger + \frac{\mu_m^2}{2r\sigma_m^2} + \frac{\mu_m}{r\sigma_m^2} \varepsilon_m \quad (18)$$

Analysing this expression shows that the CEO's compensation is higher compared to the compensation under a powerful BoD. This increase arises from two components. First, the CEO receives an additional risk premium of $\frac{\mu_m^2}{2r\sigma_m^2}$, which compensates them for the increased variance in compensation introduced by the distorted peer weight. Second, compensation now includes a random component linked to market shocks, $\frac{\mu_m}{r\sigma_m^2} \varepsilon_m$. This represents the unfiltered portion of systematic risk, whose magnitude increases with higher expected peer-group performance and decreases with market variance and the CEO's risk aversion. However, recall that the CEO's utility remains at the reservation level, as the BoD adjusts the fixed wage

such that the IR constraint is binding. This implies that while the CEO is fully compensated for the additional risk, they do not earn any additional rents.

These findings are consistent with empirical evidence showing that many firms implement imperfect forms of RPE that do not fully eliminate systematic risk. Empirical studies suggest that such imperfections are common in practice, although they often arise from factors such as benchmarking choices or peer-group design (Bizjak et al., 2022; Ma et al., 2018).

No use of RPE

If expected peer performance is sufficiently high, then there exists a threshold for μ_m above which the CEO will choose a positive weight on peer-group performance, even though firm and peer performance are positively correlated. Formally: $v_m^\dagger > 0$ if and only if $\mu_m > \frac{r\beta_f\sigma_m^2}{1+rk''(a^\dagger)\sigma_f^2}$.

A positive weight on peer performance implies that market risk is added to compensation rather than filtered out. This increases compensation risk and raises the required risk premium, thereby reducing expected firm value. Anticipating this outcome, the BoD may optimally refrain from using RPE altogether, preventing the CEO from choosing the peer weight. In that case, the CEO is evaluated solely on firm performance.

If RPE is not used, the BoD's decision problem corresponds to equations (11a-c), with the IC_m from (16) replaced by $v_m = 0$. The index "NR" denotes the solution without using RPE.

The optimal incentive weight on firm performance, v_f^{NR} , and the expected firm value, Π^{NR} , are represented by the subsequent equations. Note that from equilibrium effort a^{NR} , it follows that $v_f^{NR} = k'(a^{NR})$.

$$v_f^{NR} = \frac{1}{1 + rk''(a^{NR})(\sigma_f^2 + \beta_f^2\sigma_m^2)} \quad \text{and} \quad v_m^{NR} = 0 \quad (19a)$$

The incentive weight, v_f^{NR} , decreases with greater risk aversion, higher marginal cost of effort, and increased risk for both firm and peer performance. Without RPE, pay is exposed to both firm-specific risk, σ_f^2 , and market risk, σ_m^2 , increasing wage variance and the required risk premium. As a result, incentives are weaker than in the benchmark case where market risk is filtered out.

$$\Pi^{NR} = E[y_f|v_f^{NR}] - k(a^{NR}) - \frac{r}{2}(v_f^{NR})^2(\sigma_f^2 + \beta_f^2\sigma_m^2) \quad (19b)$$

Higher variance and/or greater risk aversion reduce expected firm value as they increase the compensation required to induce participation and effort.

From (9) it follows that the associated compensation for the CEO is given by $w^{NR} = f^{NR} + v_f^{NR}y_f$.

COROLLARY 2. *With a powerful CEO, the BoD optimally refrains from using RPE if the expected performance of the peer group is sufficiently large, that is, if:*

$$\mu_m > k = \frac{r\beta_f\sigma_m^2}{1+r k''(a^\pm)\sigma_f^2} \quad (20)$$

As shown in (A.12), this condition is equivalent to $v_m^\pm > 0$. When the CEO chooses a positive incentive weight on peer performance, the contract adds market risk to compensation rather than filtering it out, and therefore the BoD optimally refrains from using RPE. The threshold value, k , determines whether RPE is implemented. A higher threshold makes the use of RPE more likely, whereas a lower threshold makes its abandonment more likely. The threshold increases in CEO risk aversion, r , and in the firm's exposure to market risk, $\beta_f\sigma_m^2$, and decreases in firm-specific risk σ_f^2 .

These comparative statics have a clear economic intuition. When risk aversion, r , is low, the CEO requires only a small premium for bearing market risk. The benefit of filtering out market risk through RPE is therefore limited, making RPE less attractive. Similarly, if the firm's exposure to systematic risk, $\beta_f\sigma_m^2$, is small, there is little market risk to remove, reducing the value of RPE. Finally, when firm-specific risk, σ_f^2 , is high, overall compensation risk remains substantial even with RPE, weakening its relative benefit since RPE can only reduce market risk.

This finding aligns with empirical evidence on RPE adoption. For instance, studies by Bettis et al. (2014) and Gong et al. (2011) document that firms with greater exposure to systematic risk are more likely to use RPE-based compensation. This aligns with the model's prediction that RPE is most valuable when systematic risk is economically significant. At the same time, COROLLARY 2 provides a theoretical explanation for why RPE is not universally adopted. When expected peer performance is sufficiently high and CEO power allows manipulation of the peer weight, the BoD may optimally avoid RPE to prevent inefficient risk exposure introduced by the CEO.

3.3.2.2 CEO Choice in Peer Selection Process

In this subsection, the CEO's power to select the peer firm used for RPE is analysed. For this, the following setup is considered. First, it is assumed that there are two potential peer firms i , $i = 1, 2$ among which the CEO can choose. Second, as in the previous section, the CEO can also choose the incentive weight, v_i , on the performance of the selected peer firm. It is important to note that, different from the previous case, the firm's performance is now compared to a single peer firm and not to an index. The decision timeline for this case is represented in *Figure 3*.

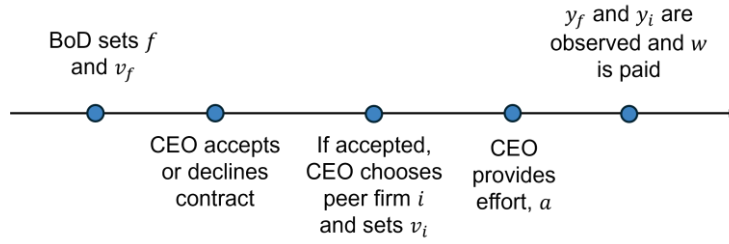


Figure 3 - Decision Timeline in the Case of a Powerful CEO (Choice in Peer Selection)

The performance of the peer firm i is given by:

$$y_i = \mu_i + \beta_i \varepsilon_m \quad (21)$$

With μ_i being the expected performance and β_i the impact of market risk on the peer firm. For the latter, a positive correlation between the market and the peer firm is assumed, hence $\beta_i > 0$. Note that it is presumed that the performance of the peer firm is only affected by market shocks, ε_m , but not by idiosyncratic shocks. In the following it is shown that a powerful CEO has strict preferences for which peer firm to choose, while the selection is irrelevant for a powerful BoD, as every peer firm can be used to eliminate market risk.

Once the peer firm is chosen, the CEO determines the incentive weight by solving the same decision problem as established in the previous case (16). While the detailed calculations are shown in the appendix, the resulting solution for the peer firm's incentive weight is given by:

$$v_i = -v_f \frac{\beta_f}{\beta_i} + \frac{\mu_i}{r\beta_i^2\sigma_m^2} \quad (22)$$

The first term represents the risk-minimizing incentive weight, as demonstrated in (A.13d). That is, if the CEO chooses $v_i = -v_f \frac{\beta_f}{\beta_i}$ then all market risk, ε_m , is eliminated from compensation. The second term captures the deviation from the risk-minimizing weight, that depends on the characteristics of the chosen peer firm. This term increases with expected

performance, μ_i , and decreases with risk aversion, r , the sensitivity of the peer firm to market shocks, β_i^2 , and market variance, σ_m^2 . The peer firm is chosen to maximize the CEO's CE, as explained below.

Given the expression for the peer incentive weight chosen by the CEO in (22), the CEO's CE for a selected peer firm i , denoted as CE_i , can be written as follows (see appendix A.14 for detailed derivation):

$$CE_i(w, a) = \left(f + v_f E[y_f | v_f] - k(a) - \frac{r}{2} v_f^2 \sigma_f^2 \right) + \frac{\mu_i}{\beta_i} \cdot \frac{1}{2r\sigma_m^2} \left(\frac{\mu_i}{\beta_i} - 2v_f r \beta_f \sigma_m^2 \right) \quad (23)$$

The CE consists of two components. The first term represents a “base” compensation component that expresses the CEO's welfare without the effect of peer selection. It includes the fixed wage, f , expected pay from firm performance, $v_f E[y_f | v_f]$, minus the cost of effort, $k(a)$, and the risk premium for firm-specific risk, $\frac{r}{2} v_f^2 \sigma_f^2$, that cannot be filtered out by RPE. The second term captures the effect of peer selection, that the CEO directly controls through their choice of peer firm i . As established in COROLLARY 2, the BoD only uses RPE if it fulfils its intended purpose of filtering out market risk, that is, if the incentive weight on the peer firm is negative, $v_i < 0$. When this condition holds, the second term in (23) is negative, and therefore reduces the CEO's utility.

For RPE to be used, the CEO must therefore choose a peer that results in a negative incentive weight, $v_i < 0$ which requires $-v_f \frac{\beta_f}{\beta_i} + \frac{\mu_i}{r\beta_i^2 \sigma_m^2} < 0$ and thus $\frac{\mu_i}{\beta_i} < v_f r \beta_f \sigma_m^2$. Hence, RPE is only implemented for peer firms whose ratio $\frac{\mu_i}{\beta_i}$ lies below this threshold. Given this constraint, the CEO selects the peer firm that maximizes their CE. Since the second term in (23) is negative under RPE, this is equivalent to selecting the peer that minimizes its absolute magnitude. This corresponds to choosing the firm with the lowest ratio $\frac{\mu_i}{\beta_i}$. This ratio reflects the trade-off between lower expected peer performance μ_i , which is beneficial under a negative weight, and the peer's exposure to market risk β_i . In summary, the second term in (23) captures the effect of peer selection on the CEO's utility under RPE. Because a negative incentive weight turns the peer into a performance benchmark, the CEO benefits from selecting a firm with relatively low expected performance, while taking into account its market risk exposure.

COROLLARY 3. *A powerful CEO selects the peer firm with the lowest ratio of the expected performance to the impact of market risk on the peer firm, $\frac{\mu_i}{\beta_i}$.*

This theoretical prediction is consistent with empirical evidence documenting strategic peer selection in RPE contracts. Specifically, Gong et al. (2011) provide evidence of “selection bias” in peer-group formation, documenting a significant negative relationship between a firm’s expected performance and its likelihood of being selected as a peer. Additionally, they find that firms are more likely to be selected as a peer if the impact of market risk on peer performance is higher. Similarly, Bakke et al. (2020) find that firms tend to select underperforming peers in RPE contracts and that peers added and retained over time are systematically weaker than those not selected. Such strategic selection of weaker peers lowers the performance benchmark against which the firm is evaluated, potentially inflating peer-adjusted performance and, in turn, increasing executive compensation. Overall, these findings suggest active management of peer groups and are consistent with managerial influence in the design of RPE contracts.

Using the definition for v_i from (22) the **expected firm value** with peer i can be expressed as:

$$\Pi_i = E[y_f | v_f^\dagger] - k(a) - \frac{r}{2} (v_f^2 \sigma_f^2 + (v_f \beta_f + v_i \beta_i)^2 \sigma_m^2) \quad (24a)$$

$$\Pi_i = \Pi^\dagger - \frac{\mu_i^2}{2r\beta_i^2\sigma_m^2} \quad (24b)$$

The second term in (24b) captures the reduction in firm value resulting from the distorted peer weight. Compared to expression (17b), this loss term now additionally depends on the peer firm’s sensitivity to market risk, β_i . Hence, the impact of CEO power on firm value now varies with the characteristics of the selected peer. A higher β_i reduces the magnitude of the loss, whereas a lower β_i amplifies it. This can be understood from the risk-filtering role of RPE. To remove market risk, the peer firm must exhibit sufficient exposure to market risk. A peer with a higher β_i allows a given peer weight to offset a larger portion of the firm’s market exposure, thereby reducing total compensation risk. Consequently, the risk premium and the associated reduction in firm value are smaller.

3.4 Implications for Empirical Analysis

Subsection 3.4.1 shows how to theoretically derive regression coefficients by modelling executive compensation regressions. Using these coefficients, the extent to which risk is filtered from CEO compensation can be analysed across the three cases. Then, subsection 3.4.2 presents the empirical implications that can be derived from this analysis.

3.4.1 Theoretical Derivation of Expected Regression Coefficients

In section 3.3, three cases have been analysed:

- I) a powerful CEO with an RPE-based contract in 3.3.2.1
- II) a powerful CEO without an RPE-based contract, as the BoD refrains from using it in 3.3.2.1
- III) a powerful BoD offering the CEO an RPE-based contract in 3.3.1

Each of them offers a different level of market risk filtering in CEO compensation. To empirically test for the presence of RPE in executive compensation, the equations derived in section 3.3 can be used as a regression model. In this section, I show how to model such executive compensation regressions and derive regression coefficients that can be used to develop empirical implications. For this, the expected relation between CEO compensation and measures of systematic and unsystematic firm performance is investigated.

First, by regressing firm performance on the performance of the index-based peer group, **firm performance** can be divided into systematic and unsystematic components:

$$y_f = \gamma_0 + \gamma_1 y_m + \theta_y \quad (25a)$$

The systematic component is correlated with the market, while the unsystematic component is unique to the firm. The terms γ_0 and γ_1 are the regression coefficients and θ_y is the uncorrelated error term. Recalling that $y_f = a + \varepsilon_f + \beta_f \varepsilon_m$ from (7) and $y_m = \mu_m + \varepsilon_m$ from (8), the estimated regression coefficient on the performance of the peer group is:

$$\hat{\gamma}_1 = \frac{\text{Cov}[y_f, y_m]}{\text{Var}[y_m]} = \beta_f \quad (25b)$$

This means that, under the model assumptions, the regression coefficient on the performance of the peer group, $\hat{\gamma}_1$, equals the market sensitivity of the firm, β_f , which measures how much firm performance moves with the market.

The systematic, y_s , and unsystematic, y_u , components of firm performance are therefore:

$$y_s = \hat{\gamma}_1 y_m = \beta_f y_m \quad \text{and} \quad y_u = y_f - y_s \quad (25c)$$

For the systematic part, variation is driven by the market. This means that y_s shows the portion of the firm's performance that can be linearly explained by the market's performance. The

unsystematic part is defined as the residual. This is, by definition, the part of firm performance, y_f , that cannot be explained by the market. By construction, $y_f = y_s + y_u$ and $Cov[y_s, y_u] = 0$. Second, the relation between the CEO's compensation, w , and the two components of performance is investigated. This is done by regressing executive pay on both the systematic and unsystematic component of firm performance from (25c):

$$w = \alpha_0 + \alpha_s y_s + \alpha_u y_u + \theta_w \quad (26a)$$

α_0 , α_s and α_u are the regression coefficients and θ_w is the uncorrelated error term. Given that the systematic and unsystematic components of firm performance are orthogonal, i.e. uncorrelated, the regression coefficients can be estimated by:

$$\hat{\alpha}_s = \frac{Cov[w, y_s]}{Var[y_s]} \quad \text{and} \quad \hat{\alpha}_u = \frac{Cov[w, y_u]}{Var[y_u]} \quad (26b)$$

PROPOSITION 2 summarizes the regression coefficients for each of the three cases. The table divides the cases of a powerful CEO in the use of RPE or its abandonment, by using the threshold derived in (20). Proofs are provided in the appendix (A.15a-f).

PROPOSITION 2. *For the cases of (I) a powerful CEO with an RPE-based contract, (II) a powerful CEO without an RPE-based contract, and (III) a powerful BoD offering the CEO an RPE-based contract, the expected regression coefficients on the systematic and unsystematic component of firm performance are:*

	$\mu_m \leq k = \frac{r\beta_f\sigma_m^2}{1 + rk''(a)\sigma_f^2}$	$\mu_m \geq k = \frac{r\beta_f\sigma_m^2}{1 + rk''(a)\sigma_f^2}$
Powerful CEO	Case (I) RPE-based contract $\hat{\alpha}_u^\dagger = v_f^\dagger$ and $\hat{\alpha}_s^\dagger = \frac{\mu_m}{r\beta_f\sigma_m^2}$	Case (II) no RPE-based contract $\hat{\alpha}_u^{NR} = \hat{\alpha}_s^{NR} = v_f^{NR} = \frac{1}{1 + rk''(a^{NR})(\sigma_f^2 + \beta_f^2\sigma_m^2)}$
Powerful BoD	Case (III) RPE-based contract $\hat{\alpha}_u^\dagger = v_f^\dagger$ and $\hat{\alpha}_s^\dagger = 0$.	

The results above can be used to analyse the relationship between CEO power and compensation design. The coefficient $\hat{\alpha}_u$ shows how much executive pay, w , changes with the unsystematic, firm-specific component of firm performance, y_u . In all three cases, this coefficient reflects the incentive rate on firm performance, v_f .

The coefficient $\hat{\alpha}_s$, on the other hand, shows how much the CEO's pay, w , changes with the systematic, market-driven component of firm performance, y_s . Thus, it reflects the extent to which risk is filtered from compensation. Consistent with the results in section 3.3, this

component differs from case to case. In the benchmark case of a powerful BoD (Case III) we observe “perfect” RPE as the contract efficiently filters out all market risk from pay, $0 = \hat{\alpha}_s^+ < \hat{\alpha}_u^+$. In Case (II), no RPE is implemented which means that no risk is filtered out from compensation, $0 < \hat{\alpha}_s^{NR} = \hat{\alpha}_u^{NR}$. The CEO is evaluated on total firm performance, which includes market movements. Hence, no distinction is made between systematic and unsystematic components of performance, and both coefficients are the same, $\hat{\alpha}_u^{NR} = \hat{\alpha}_s^{NR} = v_f^{NR}$.

An important implication arises from Case (I): The positive coefficient, $\hat{\alpha}_s^+ > 0$, indicates that compensation is still influenced by market performance, even when RPE is used. This suggests that a powerful CEO may be able to design an “imperfect” RPE system, that does not fully eliminate market risk, $0 < \hat{\alpha}_s^+ < \hat{\alpha}_u^+$, thereby allowing them to benefit from positive market movements. In this context, the coefficient can be interpreted as a measure of how much the CEO distorts their own performance evaluation. A higher coefficient reflects a greater dependence on market performance, signalling a more distorted application of RPE.

COROLLARY 4 shows comparative statics for the three cases, which are further illustrated in the following section.

COROLLARY 4. *The following comparative statics hold for the regression coefficients on the systematic component of firm performance:*

$$\text{Case (I): } \frac{\partial \hat{\alpha}_s^+}{\partial \mu_m} > 0; \quad \frac{\partial \hat{\alpha}_s^+}{\partial r} < 0; \quad \frac{\partial \hat{\alpha}_s^+}{\partial \beta_f} < 0; \quad \frac{\partial \hat{\alpha}_s^+}{\partial \sigma_m^2} < 0; \quad \frac{\partial \hat{\alpha}_s^+}{\partial \sigma_f^2} = 0;$$

$$\text{Case (II): } \frac{\partial \hat{\alpha}_s^{NR}}{\partial \mu_m} = 0; \quad \frac{\partial \hat{\alpha}_s^{NR}}{\partial r} < 0; \quad \frac{\partial \hat{\alpha}_s^{NR}}{\partial \beta_f} < 0; \quad \frac{\partial \hat{\alpha}_s^{NR}}{\partial \sigma_m^2} < 0; \quad \frac{\partial \hat{\alpha}_s^{NR}}{\partial \sigma_f^2} < 0;$$

$$\text{Case (III): } \frac{\partial \hat{\alpha}_s^+}{\partial \mu_m} = \frac{\partial \hat{\alpha}_s^+}{\partial r} = \frac{\partial \hat{\alpha}_s^+}{\partial \beta_f} = \frac{\partial \hat{\alpha}_s^+}{\partial \sigma_m^2} = \frac{\partial \hat{\alpha}_s^+}{\partial \sigma_f^2} = 0$$

In the case of imperfect RPE (Case I), systematic risk increases with expected peer performance, as higher peer performance strengthens the CEO’s incentive to distort the peer weight. In contrast, it decreases with CEO risk aversion, as this raises the cost of bearing risk. Similarly, higher market sensitivity and market risk make distortion more costly by increasing the resulting risk exposure, thereby limiting rent extraction. Firm-specific risk has no effect on risk filtering, as RPE can only eliminate systematic risk. In the absence of RPE (Case II), expected peer performance is irrelevant, as no peer group is used. Systematic risk decreases with CEO risk aversion, market exposure, market risk and firm-specific risk, as the board needs to reduce incentive intensity to limit the CEO’s risk exposure. Under perfect RPE (Case III), systematic risk is fully eliminated and therefore unaffected by these parameters.

3.4.2 Empirical Implications

Building on PROPOSITION 2, the model yields three main empirical implications for the analysis of RPE in the presence of CEO power. It generates testable hypotheses on compensation design, clarifies the interpretation of conventional RPE tests, and suggests a refined proxy for CEO power.

3.4.2.1 Hypotheses on the Role of Powerful CEOs in RPE Design

In this subsection, I show how the theoretical findings developed above can be translated into concrete and empirically testable hypotheses. Using observed compensation data, these hypotheses allow researchers to examine how managerial power affects the design and effectiveness of RPE.

The effect of managerial power on RPE design

An important question is whether managerial power influences the structure of RPE contracts. PROPOSITION 2 provides important insights into this matter, as it shows that pay sensitivity to systematic, market-driven performance differs across the three cases. Comparing firms with powerful CEOs (Cases I and II) to firms with powerful BoDs (Case III), the model predicts different compensation sensitivities to systematic performance. This leads to the following hypothesis:

HYPOTHESIS 1. *The sensitivity of compensation to systematic performance will differ for firms with powerful CEOs relative to firms with powerful BoDs, that is, $\hat{\alpha}_s^{CEO} \neq \hat{\alpha}_s^{BoD}$*

This entails that firms with powerful CEOs will use a different RPE design compared to firms with powerful BoDs.

The absence of RPE with a powerful CEO

The model further shows that when expected peer-group performance is sufficiently high, the BoD may optimally refrain from implementing RPE in the presence of a powerful CEO (Case II). Translating this into empirics, this means that a sample of firms with powerful CEOs may include firms that do not use RPE. In this case, PROPOSITION 2 shows that the regression coefficients on systematic and unsystematic performance should be identical.

HYPOTHESIS 2A. *For firms with powerful CEOs, the compensation sensitivity to systematic performance will be the same as the compensation sensitivity to unsystematic performance, that is, $\hat{\alpha}_s^{CEO} = \hat{\alpha}_u^{CEO}$.*

However, the sample might also include firms with powerful CEOs that employ RPE. This corresponds to Case (I). For those firms, expected peer performance is not sufficiently high to induce the BoD to refrain from using RPE.

HYPOTHESIS 2B. *For firms with powerful CEOs, the compensation sensitivity to systematic performance will differ from (be equal to) the compensation sensitivity to unsystematic performance if expected peer-group performance is sufficiently low (high), that is $\hat{\alpha}_s^{CEO} \neq \hat{\alpha}_u^{CEO}$ for sufficiently low μ_m and $\hat{\alpha}_s^{CEO} = \hat{\alpha}_u^{CEO}$ for sufficiently high μ_m .*

This insight enhances our understanding of the RPE puzzle, explaining that under a powerful CEO, it may be optimal for the BoD to refrain from implementing RPE if expected peer performance is high.

The effect of CEO power on the efficiency of risk filtering

The model also predicts that when RPE is employed under powerful CEOs (Case I), compensation does not fully filter out market risk (COROLLARY 1). Empirically, this can be tested by analysing the regression coefficient on systematic performance, $\hat{\alpha}_s^{CEO,RPE}$, across firms with powerful CEOs that use RPE. Identifying such firms requires reliable proxies for both CEO power and the application of RPE. If this can be achieved, PROPOSITION 2 suggests that compensation sensitivity to systematic performance should be different from 0, implying inefficient risk filtering.

HYPOTHESIS 3A. *For firms with powerful CEOs that use RPE, the compensation sensitivity to systematic performance will be nonzero, that is $\hat{\alpha}_s^{CEO,RPE} \neq 0$.*

Moreover, COROLLARY 4 highlights how economic determinants such as expected peer performance, the impact of market risk and CEO risk aversion affect pay sensitivity to systematic performance. These factors jointly determine the degree of risk-filtering efficiency. Building on this insight, researchers can consider the following hypothesis as a robustness test in a cross-sectional regression for a sample of firms with powerful CEOs that use RPE:

HYPOTHESIS 3B. *For firms with powerful CEOs that use RPE, the compensation sensitivity to systematic performance, that is, $\hat{\alpha}_s^{CEO,RPE}$,*

- (i) increases in the expected peer-group performance, μ_m*
- (ii) decreases in the expected impact of market risk, β_f*
- (iii) decreases in the CEO's risk aversion, r .*

3.4.2.2 Implications for Interpreting Conventional RPE Tests

Empirical research has primarily relied on two types of tests to examine the presence of RPE: strong-form and weak-form tests. The model provides important guidance for interpreting the results of these approaches.

Strong-form RPE

The **strong-form** test evaluates whether systematic risk is fully filtered from managerial compensation. It typically follows a two-step procedure (as in subsection 3.4.1). First, firm performance is decomposed into systematic and unsystematic components. Second, executive compensation is regressed on these two components to assess whether systematic performance affects pay (Kabitz, 2017). Strong-form RPE implies that compensation is unrelated to systematic performance, implying that the regression coefficient on the systematic component, $\hat{\alpha}_s$, equals 0.

The model shows that this prediction holds only in the benchmark case of a powerful BoD (Case III). When CEOs are powerful, compensation remains partially exposed to systematic performance even if RPE is formally implemented (Case I). If RPE is not used, compensation is based entirely on total firm performance (Case II). In both situations, the estimated sensitivity to systematic performance is positive, that is, $\hat{\alpha}_s > 0$.

This insight has important implications for empirical testing. In samples that mix firms with powerful CEOs and powerful BoDs, conventional strong-form tests may underestimate the existence of efficient RPE, as the presence of CEO power biases estimates of $\hat{\alpha}_s$ upward. To increase the power of strong-form tests, researchers should therefore restrict the sample to firms with powerful BoDs and test whether $\hat{\alpha}_s^{BoD} = 0$.

If it is not possible to construct a sample that only includes powerful BoD firms, the insights from COROLLARY 4 can be used to partition the sample according to moderating variables, thereby improving the power of the strong-form RPE test. In particular, compensation sensitivity to systematic performance increases with the expected peer-group performance and decreases with risk aversion and the impact of market risk on firm performance. Hence, a researcher is more likely to find evidence for strong-form RPE in a firm sample where:

- i) The expected peer-group performance, μ_m , is small
- ii) The CEO's risk aversion, r , is large
- iii) The impact of market risk on firm performance, $\beta_f \sigma_m^2$, is large

Weak-form RPE

Weak-form RPE tests examine whether systematic risk is at least partially filtered out from compensation. This is done by analysing the relationship between executive pay and the performance of the peer group. Assuming the correlation between firm and peer-group performance is positive, $\beta_f > 0$, filtering out systematic risk requires a negative incentive weight on peer-group performance. This weight creates an inverse relationship: for high peer performance, executive pay is adjusted downward, and vice versa (Kabitz, 2017). Empirically,

this can be tested by regressing CEO compensation on both firm performance and peer performance, which yields: $w = \tau_0 + \tau_f y_f + \tau_m y_m + \theta_w$ with regression coefficients τ_0 , τ_f and τ_m . The test for weak-form RPE is supported if the sign of the regression coefficient on peer-group performance, τ_m , is negative. Applying this approach to each of the three cases, $w \in \{w^\dagger, w^\ddagger, w^{NR}\}$, PROPOSITION 3 is obtained (see appendix A.17a-d).

PROPOSITION 3. *For the cases of (I) a powerful CEO with an RPE-based contract, (II) a powerful CEO without an RPE-based contract and (III) a powerful BoD offering the CEO an RPE-based contract, the expected regression coefficients on peer-group performance are as follows:*

	$\mu_m \leq k = \frac{r\beta_f\sigma_m^2}{1 + rk''(a)\sigma_f^2}$	$\mu_m \geq k = \frac{r\beta_f\sigma_m^2}{1 + rk''(a)\sigma_f^2}$
Powerful CEO	Case (I) RPE-based contract $\hat{t}_m^\ddagger = v_m^\ddagger = v_m^\dagger + \frac{\mu_m}{r\sigma_m^2} \leq 0$	Case (II) no RPE-based contract $\hat{t}_m^{NR} = v_m^{NR} = 0$
Powerful BoD	Case (III) RPE-based contract $\hat{t}_m^\dagger = v_m^\dagger = -v_f^\dagger\beta_f < 0$	

The proposition demonstrates that, in all three cases, the regression coefficient on peer performance coincides with the incentive weight on peer performance. In Cases (I) and (III), this coefficient is negative, which aligns with the test of weak-form RPE. By contrast, in Case (II) the coefficient is 0, such that the weak-form RPE test does not hold. This implies that empirical evidence for weak-form RPE is most likely to be found for firms with powerful BoDs (III) as well as for firms with powerful CEOs when expected peer performance is low (I). However, for firms with powerful CEOs and high expected peer performance (II), researchers may not find support for weak-form RPE. This limitation may help explain the mixed empirical findings in the literature on weak-form RPE. Moreover, weak-form tests cannot distinguish between efficient and distorted RPE. Both powerful boards and powerful CEOs generate negative peer weights when RPE is used, even though risk filtering is efficient only in the former case.

3.4.2.3 RPE-based Indicator Variable of CEO Power

The model results can be used to develop a more refined measure of CEO power. As conventional measures such as CEO duality typically do not show what type of power is held by CEOs, this new measure more specifically captures whether CEOs have power over RPE design. This makes it particularly useful for empirical analyses that seek to understand whether and how CEOs shape their own compensation structure.

The key idea is to estimate a firm-specific regression as in (25a) and examine the estimated regression coefficients on systematic and unsystematic performance. Using the logic of PROPOSITION 2, each firm can then be classified into one of the three theoretical cases. Firms for which the estimated coefficient on systematic performance is positive, $\hat{\alpha}_s > 0$, do not fully filter systematic risk from compensation. According to the model, this outcome is consistent with a CEO who has influence over RPE design, corresponding to either Case (I) or Case (II). By contrast, firms for which compensation is unrelated to systematic performance, that is, $\hat{\alpha}_s = 0$, correspond to Case (III), where the board retains control over contract design and implements efficient RPE.

Based on this classification, an RPE-based indicator variable of CEO power can be constructed. The variable takes the value one for firms in Cases (I) and (II), and 0 for firms in Case (III). This approach provides a simple and transparent proxy for CEO power that can be implemented even when detailed data on traditional governance characteristics are not publicly available.

The RPE-based indicator variable of CEO power is expected to be correlated with conventional measures of CEO power. However, the correlation may not be strong. The reason is that the proposed indicator captures power over a specific decision right, namely the ability to influence the design of RPE contracts, rather than general managerial dominance within the firm. As a result, it measures a more focused dimension of CEO power.

3.5 Illustration with a Numerical Example

This section provides a walkthrough of the model based on a numerical example and discusses the resulting comparative statics.

3.5.1 CEO Choice of Peer-Group Incentive Weight

3.5.1.1 Baseline Scenario

To illustrate the model's implications, a numerical example is constructed using the following parameterization. Firm performance is characterized by $y_f = a + \varepsilon_f + \beta_f \varepsilon_m$ and positively correlated with market performance, $\beta_f = 0.6$. The noise terms are normally distributed with $\varepsilon_f \sim N(0, 0.3)$, $\varepsilon_m \sim N(0, 0.2)$ and $Cov[\varepsilon_f, \varepsilon_m] = 0$. The manager's coefficient of absolute risk aversion is $r = 2$, and the cost of effort is given by $k(a) = \frac{a^2}{2}$, implying $k'(a) = a$ and $k''(a) = 1$.

1. Substituting these parameters into the utility function in (10a) yields $U(w, a) = -e^{-2\left(w - \frac{a^2}{2}\right)}$.

A critical determinant in the setting with a powerful CEO is the expected performance of the peer group, μ_m . According to COROLLARY 2, the BoD implements RPE only if expected peer performance is below a specific threshold level, k . When μ_m exceeds this threshold, the distortion introduced by the CEO's opportunistic weighting choice outweighs the risk-filtering benefits, and the BoD optimally refrains from RPE.

Substituting the parameter values into equation (20) yields:

$$k = \frac{r\beta_f\sigma_m^2}{1 + rk''(a^\dagger)\sigma_f^2} = \frac{2(0.6)(0.2)}{1 + 2(1)(0.3)} = 0.15 \quad (27)$$

Thus, RPE is implemented when $\mu_m \leq 0.15$ (Case I), whereas for $\mu_m > 0.15$, RPE is not used (Case II). *Figure 4* compares the variance of the CEO's performance measure when compensation is based solely on firm performance (as per equation 7) with the variance when RPE is applied. The variance of firm performance is constant, as it does not depend on the expected performance of the peer group. By contrast, the variance of the RPE-based performance measure increases with μ_m . This increase reflects the growing distortion in the contract when a powerful CEO assigns a more favourable weight to peer-firm performance. At $\mu_m = 0.15$, the two variance curves intersect. At this point, both evaluation measures are equally informative. For $\mu_m < 0.15$, the variance under RPE is lower than the variance of firm performance, implying that RPE provides a more precise signal of managerial effort. In this region, the board optimally implements RPE. For $\mu_m > 0.15$, however, the variance under RPE exceeds that of firm performance. The distortion becomes sufficiently strong that RPE no longer improves informativeness, and the board optimally abandons it. The figure therefore provides a numerical illustration of COROLLARY 2, based on the variance expressions

derived in appendix (A.11). The decision to implement RPE depends critically on expected peer performance because it determines whether risk filtering benefits dominate distortion costs.

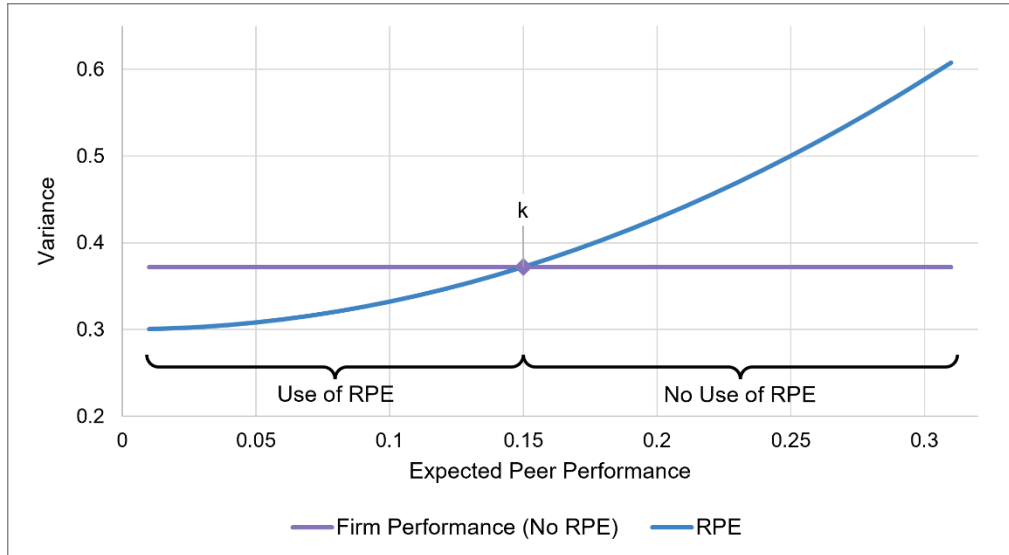


Figure 4 - Impact of Expected Peer Performance on the Decision to Implement RPE

Applying the formula derived in section 3.3 yields the results summarized in Table 1.

(I) Powerful CEO (RPE) for $\mu_m \leq k = 0.15$	(II) Powerful CEO (no RPE) for $\mu_m \geq k = 0.15$	(III) Powerful BoD
$v_f^\ddagger = 0.625$	$v_f^{NR} = 0.573$	$v_f^\dagger = 0.625$
$v_m^\ddagger = -0.375 + 2.5\mu_m$	$v_m^{NR} = 0$	$v_m^\dagger = -0.375$
$\Pi^\ddagger = 0.313 - 1.25\mu_m^2$	$\Pi^{NR} = 0.287$	$\Pi^\dagger = 0.313$
$w^\ddagger = 0.313 + 0.625\varepsilon_f + 1.25\mu_m^2$ $+ 2.5\mu_m\varepsilon_m$	$w^{NR} = f^{NR} + 0.573y_f$	$w^\dagger = 0.313 + 0.625\varepsilon_f$

Table 1 - Results for Baseline Scenario

Comparing expected firm values (Π) across the three cases shows that expected firm value is highest in Case (III), where the BoD retains control over contract design and selects the risk-minimizing incentive weight. In this case, no distortions arise. In Case (I), a powerful CEO introduces a deviation from the optimal peer-performance weight $v_m^\dagger = -0.375$ as it now

includes the term $2.5\mu_m$. As expected peer performance increases, the CEO increases the weight on peer performance, making it less negative. This allows the CEO to benefit opportunistically from positive peer performance and to load compensation on market risk rather than filtering it out. As a result, market risk is introduced into pay, and the firm must compensate the CEO for bearing this additional risk. The associated risk premium, $1.25\mu_m^2$, reduces expected firm value $\Pi^\dagger$. Notably, Case (III) and (I) coincide only when expected peer performance is $\mu_m = 0$.

The transition to Case (II) occurs at $\mu_m = 0.15$. At this point, the CEO's preferred weight on peer performance is 0 ($v_m^\dagger = -0.375 + 2.5(0.15) = 0$). For higher expected peer performance ($\mu_m > 0.15$), the CEO chooses a positive peer incentive weight. Since a positive peer incentive weight introduces additional market risk to compensation compared to when only firm performance is used for evaluation, the BoD refrains from using RPE altogether in that case.

A comparison of the firm incentive weights highlights the cost of abandoning RPE. In this numerical example, the incentive weight on firm performance equals the level of effort exerted by the CEO. From (IC_a) in (11c) it follows that $v_f = k'(a)$, and thus $v_f = a$. When RPE is not used (Case II), the incentive weight v_f^{NR} , and therefore the induced effort level, a , are strictly lower than in Cases (I) and (III) ($0.573 < 0.625$). This occurs because v_f^{NR} must account for total risk exposure, which includes both firm-specific risk and market risk. In contrast, under RPE the incentive weight $v_f^\dagger$ depends only on firm-specific risk, as market risk is filtered out.

Importantly, v_f is identical whether the BoD (Case III) or the CEO (Case I) holds power when RPE is used. One might argue that the BoD could offset the CEO's distortion of the peer-performance weight by increasing the incentive intensity on firm performance. However, as discussed in subsection 3.3.2, such an adjustment would further reduce expected firm value. To illustrate this numerically, consider an expected peer group performance of $\mu_m = 0.05$. Assume that the board anticipates the CEO's distortion of the peer-performance weight and therefore increases the incentive weight on firm performance to $v_f = 0.833$. This choice implies $v_m = -0.833(0.6) + 2.5(0.05) = -0.375 = v_m^\dagger$. Thus, the weight on peer performance coincides with the level that would arise under a powerful BoD. Despite this adjustment, expected firm performance becomes $\Pi = 0.278 - 1.25\mu_m^2$, which is lower than in Case (I) ($0.278 < 0.313$). The reason is that the distortion introduced by the powerful CEO is independent of v_f and increasing v_f therefore moves the contract away from the optimal

incentive level $v_f^\dagger = 0.625$. This deviation introduces additional inefficiencies in risk sharing, increases the CEO's risk premium and ultimately reduces expected firm value.

Finally, although expected compensation, w , is higher in Case (I) than in Case (III), the CEO is not better off in utility terms. The higher compensation merely compensates the CEO for the inefficient risk exposure created by the distorted peer weight, $v_m^\ddagger$. Because the BoD anticipates this behaviour, it adjusts the fixed wage downward such that the CEO's participation constraint remains binding. As a result, the CEO's certainty equivalent remains at reservation level ($CE = 0$) in all cases. Moreover, Case (III) is the only scenario in which market risk is fully filtered out from compensation. In Case (I), the term $2.5\mu_m\varepsilon_m$ remains in the compensation function, while in Case (II) the CEO is fully exposed to market risk as it is incorporated in y_f .

3.5.1.2 Comparative Statics: The Effect of CEO Risk Aversion

This subsection analyses how the model outcomes change when the CEO's degree of risk aversion varies. First, it should be noted that in the limiting case of a risk-neutral CEO ($r = 0$), the distinction between the three scenarios (Case I-III) disappears. The optimal incentive weight on firm performance becomes $v_f = 1$ across all three cases and compensation fully reflects firm performance. Since the CEO is indifferent to risk, the agency cost of risk-bearing vanishes and RPE no longer provides any benefit. When the CEO is risk averse ($r > 0$), compensation must balance incentives with risk sharing. Therefore, two alternative scenarios are examined: one with lower risk aversion ($r = 1$) and one with higher risk aversion ($r = 3$), while keeping all other parameters identical to the baseline case ($r = 2$). The results are summarized in *Table 2*. To illustrate the comparative statics graphically, an expected peer performance of $\mu_m = 0.08$ is assumed.

Parameter	Low Risk Aversion $r = 1$	Baseline Scenario $r = 2$	High Risk Aversion $r = 3$
k	0.09	0.15	0.19
$v_f^\dagger = v_f^\ddagger$	0.769	0.625	0.526
$v_m^\dagger$	-0.462	-0.375	-0.316
$\Pi^\dagger$	0.385	0.313	0.263
$v_m^\ddagger$	$-0.462 + 5\mu_m$	$-0.375 + 2.5\mu_m$	$-0.316 + 1.67\mu_m$
$\Pi^\ddagger$	$0.385 - 2.5\mu_m^2$	$0.313 - 1.25\mu_m^2$	$0.263 - 0.83\mu_m^2$
v_f^{NR}	0.729	0.573	0.473
Π^{NR}	0.364	0.287	0.236
$\Pi^\dagger - \Pi^{NR}$	0.021	0.026	0.027

Table 2 - Sensitivity Analysis for Low and High Risk Aversion

As risk aversion increases, expected firm value (Π) declines in all cases. This decline reflects the agency cost of risk aversion, since the CEO dislikes risk and the firm must therefore either reduce incentive intensity v_f or compensate the CEO with a higher risk premium. At the same time, the value of RPE increases with risk aversion. The loss from abandoning RPE, measured by the difference between Cases (II) and (III), increases slightly from 0.021 to 0.027 as risk

aversion rises. This reflects the fact that exposing the CEO to full market risk becomes increasingly costly as risk aversion increases. Consequently, filtering market risk through RPE becomes more valuable to the board. *Figure 5* illustrates these effects. With higher risk aversion, expected firm value decreases across all scenarios. At the same time, expected firm value under a powerful CEO ($\Pi^\ddagger$) approaches the benchmark case of a powerful BoD (Π^+). Moreover, the gap between Π^+ and Π^{NR} widens, highlighting that the benefits of risk filtering through RPE become more important when the CEO is more risk averse.

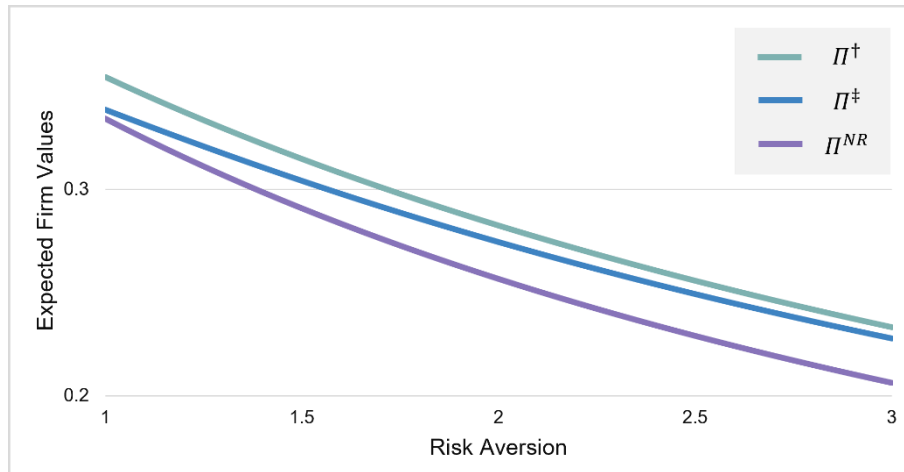


Figure 5 - Effect of Risk Aversion on Expected Firm Value

Higher risk aversion also reduces the optimal incentive intensity. The incentive weight on firm performance declines from 0.769 (when $r = 1$) to 0.526 (when $r = 3$). This implies that the CEO exerts less effort when risk aversion increases. The intuition is that stronger incentives expose the CEO to more income risk. When the CEO is highly risk averse, providing strong incentives requires a larger risk premium. To limit this cost, the board optimally reduces incentive intensity, accepting lower effort in exchange for less compensation risk. *Figure 6* illustrates this relationship by showing that effort (the firm incentive weights) declines as risk aversion increases.

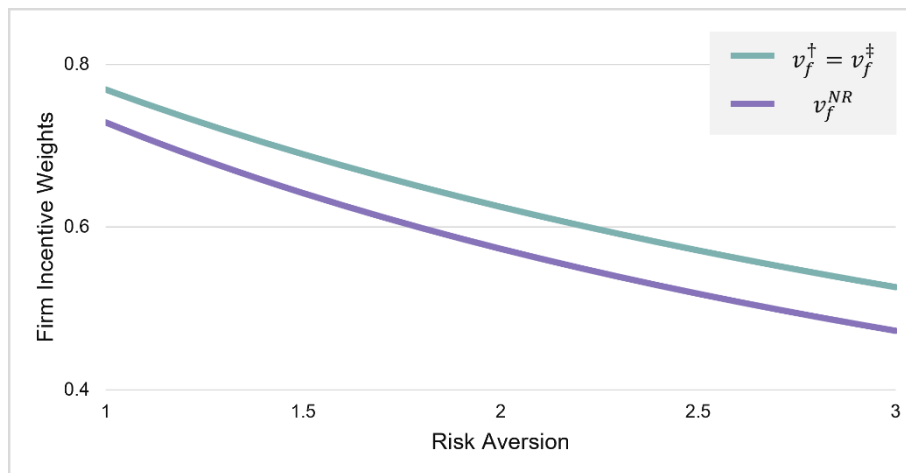


Figure 6 - Effect of Risk Aversion on Firm Incentive Weights

Higher risk aversion also affects the CEO's incentives to distort the peer performance weight. As r increases, the distortion term becomes smaller ($5\mu_m < 2.5\mu_m < 1.67\mu_m$). Thus, a more risk-averse CEO has less incentive to speculate on positive peer performance, since additional exposure to market risk becomes more costly. The chosen peer weight therefore moves closer to the risk-minimizing incentive weight $v_m^\dagger$. This also reduces the efficiency loss associated with CEO power. The reduction in firm value caused by the distorted peer weights equals $\frac{\mu_m^2}{2r\sigma_m^2}$ which declines from $2.5\mu_m^2$ to $0.83\mu_m^2$ as risk aversion increases. Economically, this implies that the agency cost of CEO influence over RPE design decreases when the CEO is more risk averse. As a consequence, the board finds it optimal to implement RPE for a larger range of expected peer performance. This is reflected in the threshold k , which increases from 0.15 in the baseline scenario to 0.19 when $r = 3$. *Figure 7* shows that as risk aversion increases, the distorted peer incentive weight under a powerful CEO moves closer to the benchmark case of a powerful BoD. Because $v_m^\dagger$ moves proportionally with $v_f^\dagger$, the risk-minimizing peer weight becomes less negative as risk aversion increases.

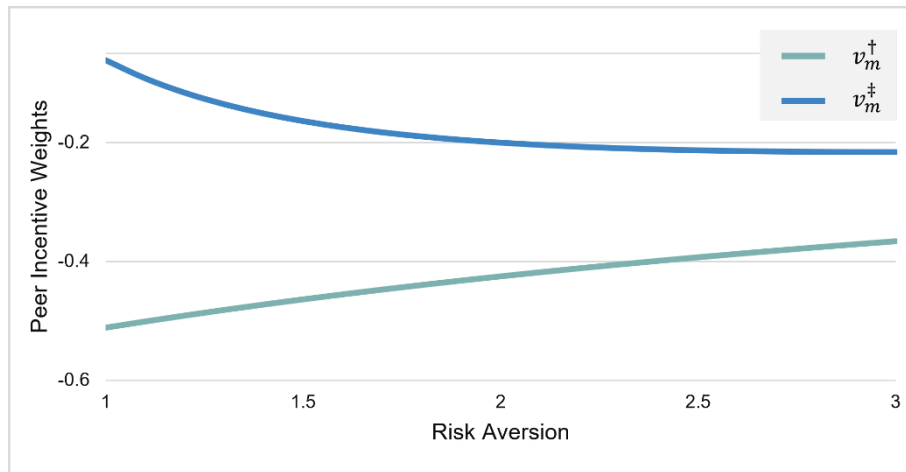


Figure 7 - Effect of Risk Aversion on Peer Incentive Weights

Finally, higher risk aversion also affects the structure of compensation. Total expected compensation decreases as risk aversion increases. This occurs for two reasons. First, the lower incentive intensity reduces the performance-based component of pay. Second, the CEO chooses a peer weight that is closer to the risk-minimizing level, which reduces the inefficiency costs associated with opportunistic distortion. Consequently, the compensation contract contains less variable pay and exposes the CEO to less market risk.

3.5.1.3 Comparative Statics: The Effect of Market Risk

This subsection analyses the effect of changes in market variance. To do so, a low market risk scenario ($\sigma_m^2 = 0.1$) and a high market risk scenario ($\sigma_m^2 = 0.4$) are compared with the baseline case ($\sigma_m^2 = 0.2$). The results are summarized in *Table 3*. Note that the parameters of the optimal RPE contract (Case III) remain unchanged ($v_f^\dagger = v_f^\ddagger = 0.625$; $v_m^\dagger = -0.375$; $\Pi^\dagger = 0.313$), since the RPE mechanism fully filters out market risk regardless of its magnitude. To illustrate the comparative statics graphically, an expected peer performance of $\mu_m = 0.08$ is assumed.

Parameter	Low Market Risk $\sigma_m^2 = 0.1$	Baseline Scenario $\sigma_m^2 = 0.2$	High Market Risk $\sigma_m^2 = 0.4$
k	0.075	0.15	0.3
$v_m^\ddagger$	$-0.375 + 5\mu_m$	$-0.375 + 2.5\mu_m$	$-0.375 + 1.25\mu_m$
$\Pi^\ddagger$	$0.313 - 2.5\mu_m^2$	$0.313 - 1.25\mu_m^2$	$0.313 - 0.625\mu_m^2$
v_f^{NR}	0.598	0.573	0.530
Π^{NR}	0.299	0.287	0.265
$\Pi^\dagger - \Pi^{NR}$	0.014	0.026	0.048

Table 3 - Sensitivity Analysis for Low and High Market Risk

An interesting observation is that lower market risk leads to higher inefficiency costs in Case (I). In the low-risk scenario, the deviation in the peer incentive weight, $v_m^\ddagger$, increases from $2.5\mu_m$ in the baseline to $5\mu_m$. A similar pattern is observed for expected firm value, where the inefficiency loss increases from $0.625\mu_m^2$ (high variance) to $2.5\mu_m^2$ (low variance). The economic intuition is that when market risk is low, the CEO faces little additional compensation risk from increasing exposure to peer performance. As a result, the CEO can distort the peer weight more aggressively in order to capture the positive expected peer performance, μ_m . In contrast, when market volatility is high, speculating on peer performance becomes more costly because it substantially increases the variance of the CEO's pay. Consequently, the CEO chooses a peer weight closer to the risk-minimizing level. Figure 8 illustrates this effect graphically. As market risk increases, the peer incentive weight under a powerful CEO moves closer to the benchmark weight chosen under a powerful BoD. The dashed line indicates the

RPE threshold. To the left of this threshold, the board optimally refrains from implementing RPE as the CEO would choose a positive peer incentive weight.

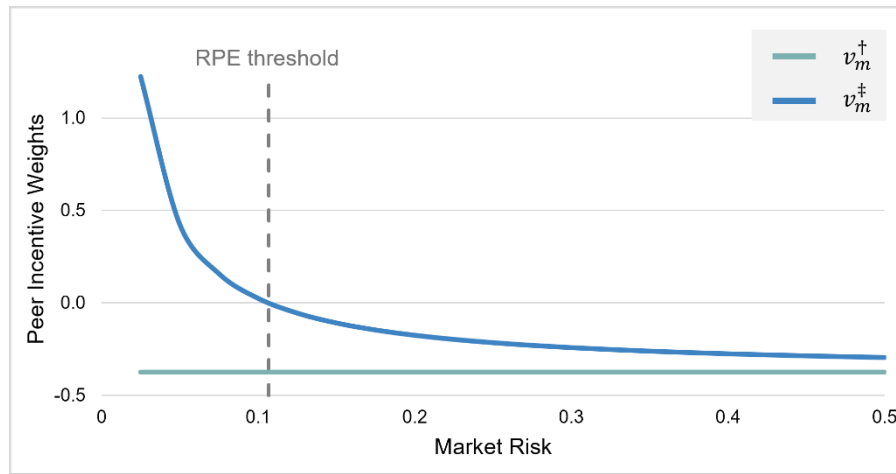


Figure 8 - Effect of Market Risk on Peer Incentive Weights

Moreover, the incentive weight v_f^{NR} decreases as market risk increases. When market volatility is low, total risk exposure without RPE ($\beta_f^2 \sigma_m^2$) is small, allowing the board to maintain stronger incentives without imposing excessive risk on the CEO. In the low-risk scenario, the incentive weight ($v_f^{NR} = 0.598$) is closer to the optimal level under RPE ($v_f^† = 0.625$). In contrast, when market risk is high, the unfiltered market risk adds substantial variance to the CEO's total pay. To limit the resulting risk premium, the board optimally reduces the incentive intensity. Figure 9 illustrates this relationship. When market risk is low, the two incentive weights are closest. As market risk increases, the incentive weight on firm performance without RPE declines, while the optimal RPE incentive weight remains constant.

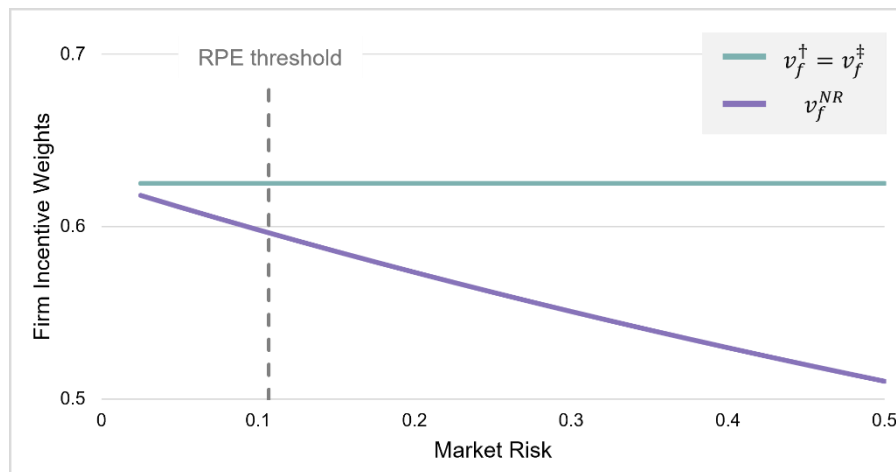


Figure 9 - Effect of Market Risk on Firm Incentive Weights

Finally, the value of RPE increases with market risk. This is reflected in the increase in the difference $\Pi^+ - \Pi^{NR}$ from 0.014 in the low-risk scenario to 0.048 in the high-risk scenario. When market variance is low, the benefits of filtering market risk are limited, so abandoning RPE leads to only a small efficiency loss. In contrast, when market risk is high, failing to filter out this risk significantly increases compensation volatility and the associated risk premium. Thus, RPE becomes more valuable as market risk increases. This is also reflected in the threshold k , which is lowest in the low-risk scenario. As a result, the board is more likely to refrain from implementing RPE because even small levels of expected peer performance may lead to inefficient weight distortions. Figure 10 illustrates this result. As market risk increases, expected firm value without RPE declines more strongly, widening the gap between the RPE and the no-RPE case. Moreover, $\Pi^\pm$ approaches Π^+ with higher market risk, as the powerful CEO has less incentives to distort RPE design.

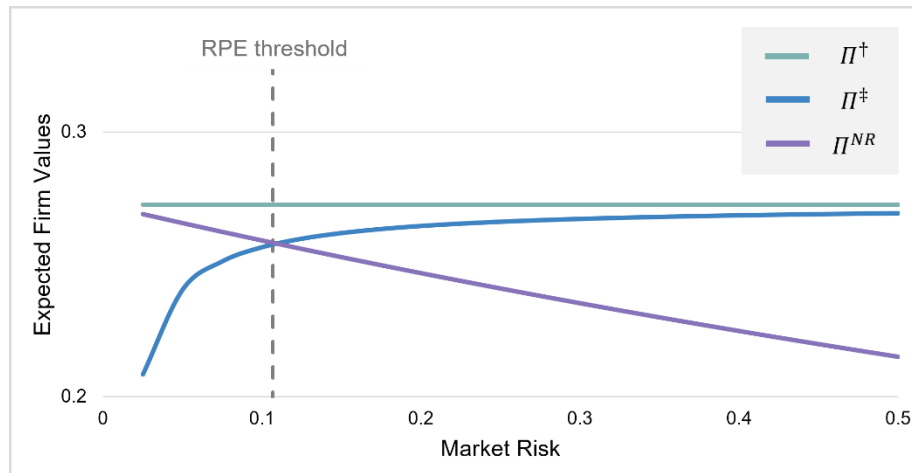


Figure 10 - Effect of Market Risk on Expected Firm Values

3.5.1.4 Comparative Statics: The Effect of Market Risk Exposure

To analyse the effect of market risk exposure, β_f , a low-beta scenario ($\beta_f = 0.4$) and a high-beta scenario ($\beta_f = 0.8$) are compared with the baseline case ($\beta_f = 0.6$). The results are summarized in *Table 4*. Note that the incentive weights $v_f^\dagger$ and $v_f^\ddagger$ as well as the expected firm values $\Pi^\dagger$ and $\Pi^\ddagger$ remain unchanged across all scenarios because RPE filters out the systematic risk component and the risk premium term $\frac{\mu_m^2}{2r\sigma_m^2}$ in $\Pi^\ddagger$ does not depend on β_f . Therefore, only the variables that are affected by changes in β_f are reported.

Parameter	Low Beta $\beta_f = 0.4$	Baseline Scenario $\beta_f = 0.6$	High Beta $\beta_f = 0.8$
k	0.10	0.15	0.20
$v_m^\dagger$	-0.25	-0.375	-0.50
$v_m^\ddagger$	$-0.25 + 2.5\mu_m$	$-0.375 + 2.5\mu_m$	$-0.5 + 2.5\mu_m$
v_f^{NR}	0.601	0.573	0.539
Π^{NR}	0.300	0.287	0.269
$\Pi^\dagger - \Pi^{NR}$	0.013	0.026	0.044

Table 4 - Sensitivity Analysis for Low and High Beta

An increase in β_f raises the potential benefit of RPE. When firm performance becomes more strongly correlated with the market, the systematic risk component $\beta_f^2 \sigma_m^2$ increases. Filtering this risk therefore becomes more valuable. This is reflected in the threshold k , which increases from 0.10 in the low-beta scenario to 0.20 in the high-beta scenario, indicating that RPE is more likely to be implemented when market exposure is high. In the following figures, this is illustrated by the RPE threshold: to the left of the line, RPE is not implemented (low β_f) while to the right it is used (high β_f). To calculate this threshold and develop the figures $\mu_m = 0.08$ is assumed.

The risk-minimizing weight $v_m^\dagger$ becomes more negative (from -0.25 to -0.50) as β_f increases. Intuitively, if firm performance is more sensitive to market movements, a stronger negative weight on the peer group is required to offset this risk. In the limiting case of perfect correlation with the market ($\beta_f = 1$) the optimal weight is $-v_f^\dagger$, effectively neutralizing market noise one-for-one. Figure 11 illustrates this relationship. As market exposure increases, both the optimal

and the distorted peer incentive weights decline. The lines remain parallel because the distortion term $\frac{\mu_m}{r\sigma_m^2} = 2.5\mu_m$ does not depend on β_f .

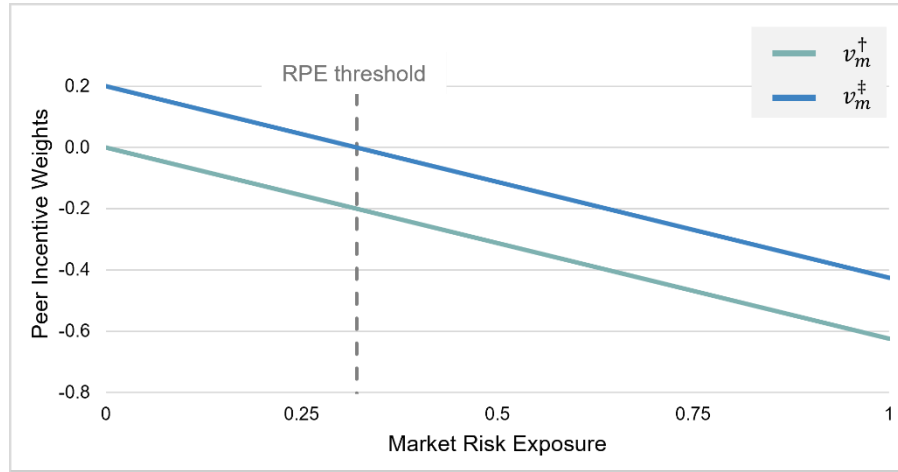


Figure 11 - Effect of Market Risk Exposure on Peer Incentive Weights

Moreover, the incentive weight v_f^{NR} decreases as β_f increases. Without RPE, the board cannot filter out market risk and therefore faces a trade-off between providing incentives and limiting compensation risk. As market exposure increases, firm performance becomes more volatile, which requires a higher risk premium for the CEO. To limit this cost, the board optimally reduces the incentive intensity and accepts a lower effort level. Figure 12 illustrates this result. When market exposure is low, all incentive weights are relatively close. As β_f increases, v_f^{NR} declines, while $v_f^+ = v_f^\dagger$ remain constant.

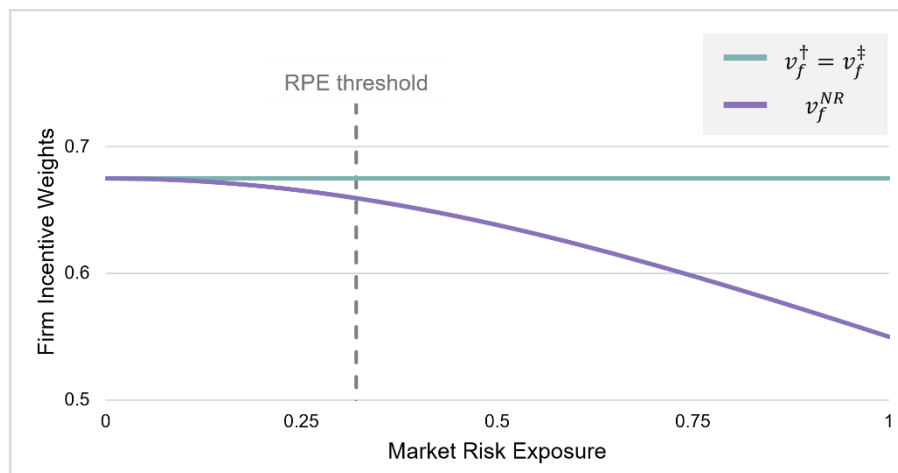


Figure 12 - Effect of Market Risk Exposure on Firm Incentive Weights

Finally, the loss in firm value from abandoning RPE increases with market risk exposure. The difference $\Pi^+ - \Pi^{NR}$ rises from 0.013 to 0.044, reflecting that higher market risk exposure increases both the compensation risk and the inefficiency from weaker incentives when RPE is not used. Figure 13 illustrates this effect. As β_f increases, expected firm value without RPE declines more strongly, widening the gap between the RPE and no-RPE cases.

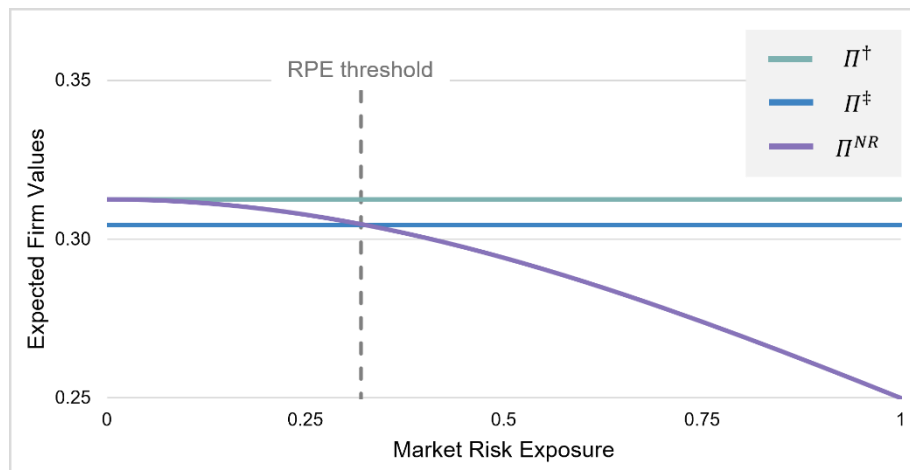


Figure 13 - Effect of Market Risk Exposure on Expected Firm Values

3.5.1.5 Comparative Statics: Effect of Firm-Specific Risk

Finally, the effect of inherent firm risk is analysed. Firm-specific risk is varied from a low-risk scenario ($\sigma_f^2 = 0.1$) to a high-risk scenario ($\sigma_f^2 = 0.5$), relative to the baseline scenario ($\sigma_f^2 = 0.3$). The results are summarized in *Table 5*.

Parameter	Low Firm Risk $\sigma_f^2 = 0.1$	Baseline Scenario $\sigma_f^2 = 0.3$	High Firm Risk $\sigma_f^2 = 0.5$
k	0.2	0.15	0.12
$v_f^\dagger = v_f^\ddagger$	0.833	0.625	0.5
$v_m^\dagger$	-0.5	-0.375	-0.3
$\Pi^\dagger$	0.417	0.313	0.25
$v_m^\ddagger$	$-0.5 + 2.5\mu_m$	$-0.375 + 2.5\mu_m$	$-0.3 + 2.5\mu_m$
$\Pi^\ddagger$	$0.417 - 1.25\mu_m^2$	$0.313 - 1.25\mu_m^2$	$0.25 - 1.25\mu_m^2$
v_f^{NR}	0.744	0.573	0.466
Π^{NR}	0.372	0.287	0.233
$\Pi^\dagger - \Pi^{NR}$	0.045	0.026	0.017

Table 5 - Sensitivity Analysis for Low and High Firm Risk

RPE is more likely to be implemented when firm risk is low. This is reflected by the threshold k , which decreases as firm-specific risk increases. The intuition is that the efficiency gains from filtering market risk are smaller for high-risk firms. When firm-risk dominates the total variance of performance, removing market risk through RPE provides only limited benefits. In this case, the potential efficiency gains from RPE are insufficient to offset the agency costs arising from the CEO's opportunistic distortion of the peer weight. Therefore, the board optimally refrains from implementing RPE if firm risk is high. This is illustrated in the following figures by the RPE threshold. To the left of the threshold line RPE is implemented, whereas to the right the board refrains from its use. To calculate this threshold and develop the figures $\mu_m = 0.08$ is assumed.

The optimal incentive weight on firm performance, $v_f^\dagger$, declines (from 0.833 to 0.5) as firm-specific risk increases. When firm-specific risk is low, firm performance provides a precise and informative signal of CEO effort. The board can therefore implement stronger incentives

without exposing the risk-averse CEO to excessive compensation risk. Consequently, the CEO exerts higher effort and expected firm value is highest in the low-risk setting ($\Pi^\dagger = 0.417$). Figure 14 illustrates this relationship. As firm-specific risk increases, the incentive weights on firm performance decline. At the same time, the incentive weights with and without RPE move closer together.

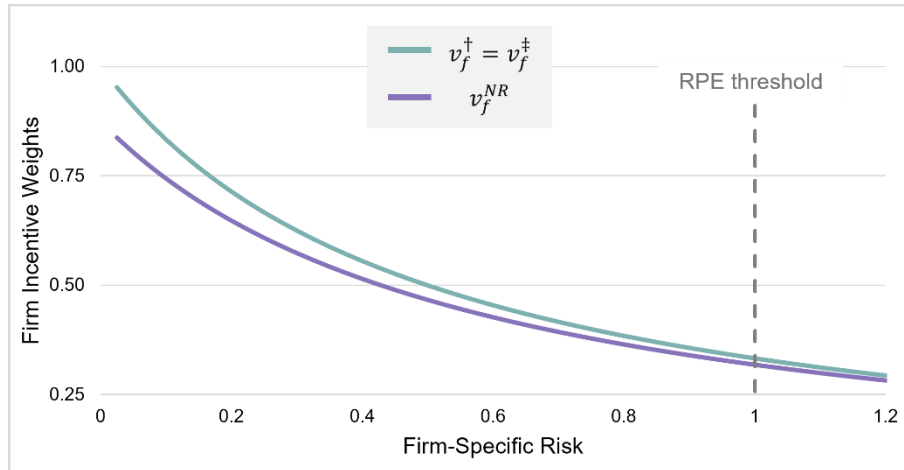


Figure 14 - Effect of Firm-Specific Risk on Firm Incentive Weights

Correspondingly, the optimal weight on peer performance, $v_m^\dagger$, becomes less negative as firm-specific risk increases. The risk-minimizing peer weight is proportional to the firm incentive weight ($-v_f^\dagger \beta_f$). When the CEO is strongly incentivized through firm performance, they are also indirectly exposed to market risk. A stronger negative peer weight is therefore required to offset this exposure. For the case with a powerful CEO (Case I), it is important to note that the CEO distortion term ($2.5\mu_m$) remains constant across these scenarios, because it depends only on market risk parameters and not firm-specific risk. Figure 15 illustrates this effect. As firm-specific risk increases, both peer weights become less negative and move in parallel because the distortion term remains constant.

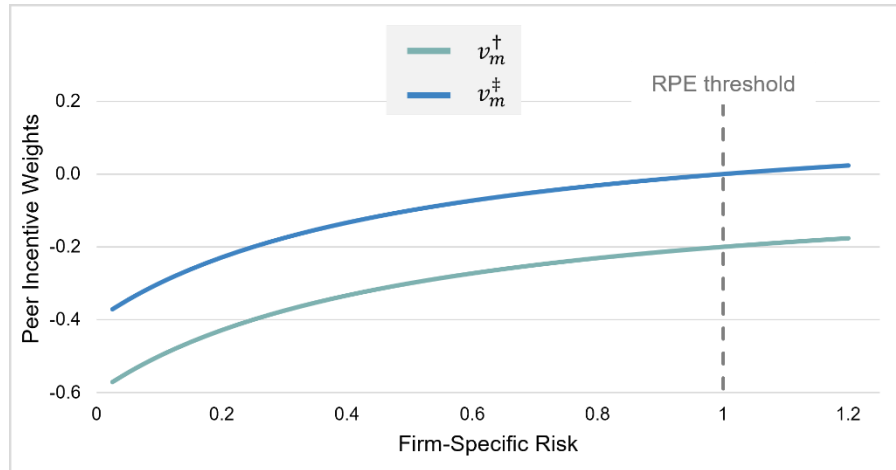


Figure 15 - Effect of Firm-Specific Risk on Peer Incentive Weights

Finally, the opportunity cost of abandoning RPE declines as firm-specific risk increases. In Table 5 the difference $\Pi^+ - \Pi^{NR}$ decreases from 0.045 to 0.017. Expected firm values decline because higher firm risk leads to lower incentive intensity on one hand and higher risk premia on the other. Figure 16 illustrates this result. As firm-specific risk increases, expected firm value declines across all cases and the gap between the RPE (Π^+) and no-RPE (Π^{NR}) scenario narrows, reflecting the reduced benefit of risk filtering through RPE.

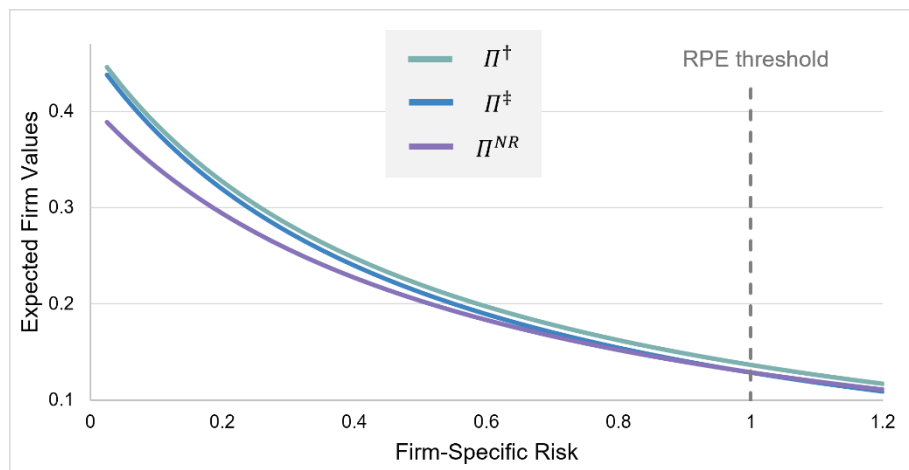


Figure 16 - Effect of Firm-Specific Risk on Expected Firm Values

3.5.2 CEO Choice in Peer Selection Process

To analyse the case of peer selection, it is assumed that the CEO can choose between two potential peer firms, a and b , with performance given by $y_a = \mu_a + \beta_a \varepsilon_m$ and $y_b = \mu_b + \beta_b \varepsilon_m$. For simplicity, both firms are assumed to have the same exposure to market risk, while their expected performance differs. The incentive weight on firm performance, v_f , remains the same as in the baseline scenario from subsection 3.5.1.1 ($v_f = 0.625$), as this parameter does not depend on the peer firm's characteristics. Recall from Equation (23) that the CEO's CE can be decomposed into a fixed component and a peer impact component. Since the fixed component is identical across both alternatives, the analysis focuses on the peer impact component. The relevant determinant of the CEO's choice is the ratio of expected peer performance to market exposure $\frac{\mu_i}{\beta_i}$. This ratio can increase either through higher expected performance or lower market risk. In the graphical illustration, the ratio is varied by increasing expected performance while keeping peer exposure ($\beta_i = 0.7$) constant. The comparative statics with respect to other parameters, such as CEO risk aversion, market risk, and firm risk, follow the same patterns as in the previous subsection. Therefore, they are not repeated here. Table 6 summarizes the results.

Parameters	Peer Firm a with $\mu_a = 0.05$ & $\beta_a = 0.7$	Peer Firm b with $\mu_b = 0.10$ & $\beta_b = 0.7$
$\frac{\mu_i}{\beta_i}$	0.071	0.143
v_i	-0.281	-0.026
CE (peer impact)	-0.020	-0.028
Π_i	0.306	0.287

Table 6 – Numerical Illustration of Strategic Peer Selection by a Powerful CEO

The table shows that peer firm a has a lower ratio $\frac{\mu_i}{\beta_i}$ than peer firm b . According to COROLLARY 3, the CEO therefore selects firm a . This choice minimizes the negative impact of peer selection on the CEO's CE, which equals -0.020 for firm a compared to -0.028 for firm b . The corresponding peer incentive weight for firm a , $v_a = -0.281$ is closer to the risk-minimizing incentive weight $-v_f \frac{\beta_f}{\beta_i} = -0.536$. Consequently, expected firm value is also higher when using peer firm a . The implementation of RPE depends on whether the CEO assigns a negative weight to peer performance. RPE is implemented only if $v_i < 0$ which requires the

ratio $\frac{\mu_i}{\beta_i}$ to remain below the threshold $v_f r \beta_f \sigma_m^2$. In the numerical example, this threshold equals 0.15. If the ratio exceeds this value, the CEO chooses a positive peer weight and the board optimally abandons RPE. *Figure 17* illustrates the relationship between expected peer performance and the chosen peer incentive weight. As performance increases, the incentive weight becomes less negative and eventually turns positive at the RPE threshold. To the right of this threshold, RPE is no longer implemented.

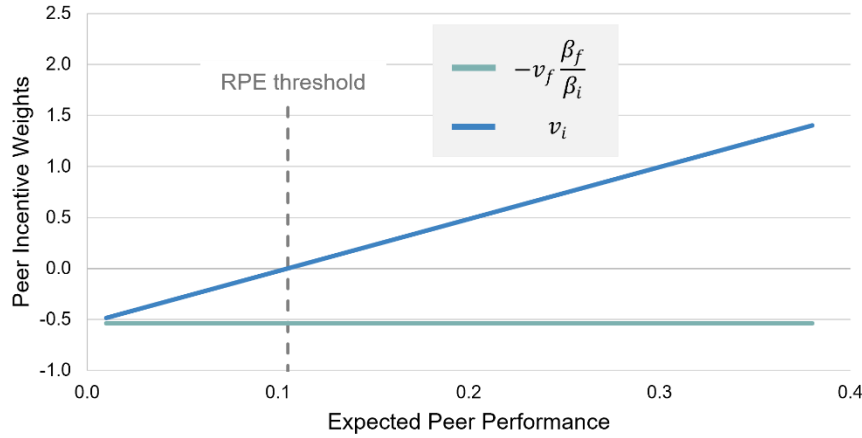


Figure 17 - Effect of Expected Peer Performance on Peer Incentive Weights

Figure 18 shows the corresponding effects on expected firm value and the peer impact on the CEO's CE. Higher peer performance increases the negative impact on the CEO's CE and reduces expected firm value. Consequently, selecting a peer with lower expected performance minimizes the inefficiency created by the distorted peer weight.

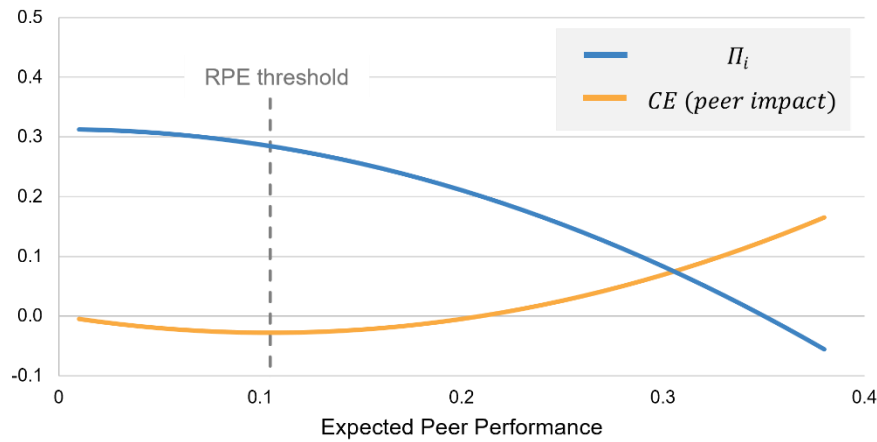


Figure 18 - Effect of Expected Peer Performance on Expected Firm Value and CE

Finally, an interesting implication is that the CEO's optimal choice of peer firm also benefits the board. By selecting the firm with the lowest $\frac{\mu_i}{\beta_i}$ ratio, the CEO chooses a peer weight that is closer to the risk-minimizing incentive weight. This reduces the compensation risk borne by the CEO, lowers the required risk premium, and ultimately results in higher expected firm value.

3.6 Discussion of Model Results

The model of Dikolli et al. (2018) extends optimal contract theory by incorporating aspects of managerial power. It offers a potential explanation for the RPE puzzle by showing that, in the presence of a powerful CEO, it may be optimal for the board to refrain from implementing RPE. At the same time, the model shows that the implementation of “imperfect RPE” may also arise as a consequence of CEO power. In some situations, such imperfect RPE can still be preferable to abandoning RPE altogether.

Within this framework, the CEO faces a trade-off between the benefits of rent extraction and the costs of increased risk exposure. On the one hand, the CEO can benefit from positive peer-group performance by distorting the peer incentive weight. On the other hand, this distortion increases compensation risk due to inefficient risk filtering. If expected peer performance is sufficiently high, the efficiency loss from distortion becomes so large that the BoD may optimally refrain from implementing RPE altogether. Importantly, the model shows that managerial power over RPE design does not reduce incentive intensity or effort provision. The optimal incentive weight on firm performance remains unchanged, implying that the efficient level of effort is still induced. Instead, the distortion primarily leads to a redistribution of wealth, as the powerful CEO captures part of the gains through the positive effect of peer-group performance on compensation.

An important insight of the model concerns the filtering of systematic risk from executive compensation. When RPE is implemented optimally under a powerful BoD, systematic risk is fully removed from compensation. In contrast, when RPE is not used, compensation remains fully exposed to systematic risk. When a powerful CEO distorts the RPE contract, market risk is no longer filtered efficiently, and systematic risk remains partially reflected in compensation. Furthermore, the model shows that when CEOs can influence peer selection, they tend to choose weaker peers, which is consistent with empirical evidence (Bakke et al., 2020; Gong et al., 2011). Although this behaviour appears opportunistic, this choice maximizes expected firm value conditional on RPE being implemented. This is because the CEO selects the peer with the lowest ratio of expected performance to market exposure, which minimizes the distortion of the peer incentive weight. Thus, while the resulting RPE contract remains imperfect, this choice represents the least costly outcome among the feasible options. Nevertheless, the first-best outcome remains the implementation of perfect RPE under a powerful board.

These insights have important implications for empirical research. Analysing the degree of systematic and unsystematic risk reflected in executive compensation can help researchers infer the role of managerial power in the design of RPE contracts. According to the model, the level of systematic risk present in compensation differs depending on whether the CEO has

influence over contract design. When a powerful CEO is present, either no systematic risk is filtered out, indicating the abandonment of RPE, or systematic risk is only partially filtered, indicating a distortion of RPE. These predictions can help improve the interpretation of conventional strong-form and weak-form tests of RPE. Moreover, the degree of risk filtering may itself serve as an indicator of CEO power over the design of RPE contracts. Nevertheless, these implications rely on the assumption that researchers can clearly distinguish between systematic and unsystematic risk in firm performance. In practice, this separation depends on the benchmark used and the construction of the peer group. As a result, measurement errors in the benchmark or peer selection may affect the estimated regression coefficients and complicate the empirical identification of risk filtering.

The comparative statics from the numerical analysis further illustrate how economic parameters affect the use and effectiveness of RPE. Higher risk aversion increases the value of risk filtering and therefore strengthens the benefits of RPE. At the same time, it reduces the optimal incentive intensity as the CEO is generally less tolerant of compensation risk. Moreover, higher market risk or greater exposure of firm performance to market shocks increases the potential gains from filtering systematic risk. Additionally, higher market risk, increases the CEO's cost of distorting RPE. By contrast, higher firm-specific risk reduces the relative benefits of RPE because this type of risk cannot be removed through benchmarking. These results highlight that the optimal use of RPE depends not only on CEO power but also on the underlying risk environment of the firm.

Limitations

While the model provides valuable insights into the relationship between managerial power and the design of RPE contracts, it also relies on several simplifying assumptions. The following discussion outlines some limitations of the framework.

In the model, managerial power is represented as a specific decision right. In one case, the CEO can influence the incentive weight on peer performance, and in another case, the CEO can select the peer firms used for comparison. This modelling choice is a strength of the framework, as it clearly identifies the source of managerial power and allows its direct effect on compensation design to be analysed. However, it may also represent a limitation, as CEO power can take more complex forms, and different dimensions of power may interact in ways that affect compensation outcomes (Choe et al., 2014). Consequently, other forms of managerial power that are not considered in the model may also influence the design of RPE, either directly or indirectly. For instance, instead of influencing the peer incentive weight, the CEO may manipulate reported firm performance through earnings management to increase their compensation (Infuehr, 2022).

Moreover, the model does not consider external constraints on CEO power. It assumes that the CEO can distort the incentive weights on peer performance without restriction, whereas the only constraint arises from the board's option to abandon RPE entirely. In practice, however, CEO power may be limited by external governance mechanisms. For instance, Kuhnen and Zwiebel (2008) highlight the disciplining role of shareholders. They argue that shareholders tolerate a certain level of rent extraction as long as the costs imposed by CEO power remain below the costs associated with replacing management. The threat of replacement therefore acts as a constraint on CEO influence over compensation design.

Similarly, institutional investors may exert pressure on management and influence compensation structures. Empirical evidence shows that stronger monitoring by institutional investors is associated with higher pay-performance sensitivity. This implies a more effective application of RPE, as CEO compensation becomes more sensitive to firm-specific performance, indicating that market risk is effectively filtered out from CEO pay (Liu & Yin, 2023).

Another limitation of the model relates to the structure of the compensation contract. The model assumes a linear compensation scheme consisting of a fixed wage and variable components tied to firm performance and peer performance. However, this simplified structure does not fully reflect more complex compensation arrangements that are commonly observed in practice. For instance, executive compensation packages often include instruments such as stock options, which introduce nonlinear payoff structures and asymmetric incentives (Conyon & Peck, 2012).

Moreover, the model does not consider other forms of executive compensation that are not necessarily monetary, such as perks or other non-transparent benefits. Powerful CEOs may use their influence to obtain such forms of hidden compensation in order to avoid attracting the attention of shareholders (Kuhnen & Zwiebel, 2008). These forms of rent extraction fall outside the scope of the model, which focuses exclusively on CEO influence over incentive weights or peer selection in RPE contracts. As a result, the model explains only a specific channel through which managerial power can affect executive compensation. It does not capture the full range of mechanisms through which CEOs may extract rents. However, constructing a model that incorporates all aspects of executive compensation would likely be infeasible. Thus, the strength of the model lies in isolating this particular mechanism and analysing its implications in a tractable framework.

A key limitation of the model is addressed by Diser and Hofmann (2018), who extend the framework to include the possibility of private hedging. If the CEO can trade in the stock of

peer firms, they gain an additional instrument to manage their risk exposure beyond the peer incentive weights. By privately hedging positions in peer firms, the CEO can further reduce the compensation risk arising from the RPE contract. Under a powerful BoD, private hedging and RPE act as complements. Their joint use can maximize firm value because the CEO can hedge any residual risk that cannot be fully eliminated through the RPE contract. Importantly, this risk reduction does not weaken the strength of incentives. However, this complementarity arises only when RPE is based on accounting measures. When RPE is based on stock performance, private hedging and RPE become perfect substitutes. In contrast, when the CEO holds power over contract design, the combination of RPE and private hedging may amplify distortions in the peer incentive weight. This is because private hedging reduces the cost of bearing additional risk, which makes it less costly for the CEO to manipulate the peer weight to their advantage. In response, the BoD may find it optimal to prohibit private hedging entirely. Consequently, the presence of private hedging opportunities can further reduce the likelihood that the board implements RPE when the CEO is powerful.

Another point to consider is that the model assumes that efficient RPE should completely remove systematic risk from managerial compensation. However, the full removal of common risk from CEO pay may not always be optimal (Wu, 2013, 2019). The reason is that RPE can introduce additional peer-specific risk into the CEO's compensation. Benchmarking against peer firms may expose the CEO to risks related to the peer's performance. To limit this additional risk, boards may choose not to filter out systematic risk completely. While the baseline model abstracts from this effect by assuming that peer firms are not exposed to idiosyncratic risk, Appendix B.1 of Dikolli et al. (2018) introduces peer-specific risk and shows that boards face a trade-off between removing common shocks and introducing additional noise. However, despite this extension, the model maintains that systematic risk should be fully filtered in the optimal contract ($\hat{\alpha}_s^\dagger = 0$). Thus, the model continues to attribute the presence of systematic risk in compensation primarily to distortions arising from CEO power, even though it may also reflect an optimal trade-off between removing market risk and avoiding excessive exposure to peer-specific risk (Wu, 2013, 2019). In line with this reasoning, Tice (2024) shows that firms are more likely to implement RPE when common risk is large and peer firms capture a substantial share of this common component. In such settings, the benefits of filtering out common shocks outweigh the potential costs of introducing additional noise through peer benchmarking.

A related argument is provided by Gopalan et al. (2010), who argue that it may not be optimal to fully filter out systematic shocks, such as sector movements, from executive compensation. Sector performance can influence firm performance and may therefore contain information

about managerial decisions. If sector conditions affect the firm's opportunities and strategic choices, completely removing sector performance from compensation could weaken incentives. In particular, when CEOs have greater flexibility to adjust the firm's strategy or its exposure to external factors, compensation should also reward the portion of firm performance driven by sector conditions. This is because the CEO may be able to influence how the firm responds to sector movements by adapting the firm's exposure to these external factors. In such cases, rewarding CEOs for the sector-driven component of firm performance can provide more appropriate incentives implying that the removal of systematic risk from compensation may not always be optimal. The presence of systematic risk in CEO pay therefore does not necessarily indicate distortions caused by CEO power, but may instead reflect an optimal compensation design.

Another limitation to consider is whether a powerful CEO can influence the compensation contract ex-post, after the performance has already been realized. The model focuses exclusively on the CEO's ex-ante influence and assumes that once the incentive weights or peer firms are chosen, the compensation contract cannot be adjusted. In practice, however, several studies suggest that powerful executives may engage in ex-post renegotiation or selectively apply benchmarking to their advantage. For example, Garvey and Milbourn (2006) show that powerful executives selectively apply benchmarking depending on whether performance benchmarks have turned out favourably or not. Specifically, they document asymmetric "pay for luck", whereby CEOs are rewarded for positive market shocks but protected from negative ones. Consequently, RPE is less likely to be used when the market is up, as the CEO then prefers to benefit from the positive market performance rather than filter it out. Conversely, RPE is more likely to be used when the market is down, allowing the CEO to use RPE to filter out this negative shock. These findings suggest that important aspects of compensation contracts may not always be determined ex ante but may instead be applied strategically ex post to transfer wealth from shareholders to executives. As a result, CEOs are rewarded for good luck and protected from bad luck. Similarly, Morse et al. (2011) develop a model in which CEOs manipulate incentive contracts after performance has been realized. In their framework, managers shift incentive weights toward the performance measures that turn out most favourably, thereby increasing their compensation and weakening the effectiveness of incentives. As a result, RPE may be applied selectively depending on whether filtering out external factors, such as market risk, increases compensation. Extending the model framework to include this dimension of ex-post manipulation would be valuable, as it highlights that distortions in incentive weights may arise not only from the ex-ante risk-return trade-off

decision faced by a powerful CEO, but also from opportunistic adjustments made after performance outcomes are observed.

Lastly, with regard to peer selection, the model assumes that any peer group can be used to eliminate market risk effectively. However, empirical evidence suggests that systematic risk filtering can be improved by using customized peer groups rather than broad industry indices (Bizjak et al., 2022; Ma et al., 2018). In practice, identifying suitable peers may be challenging depending on the firm's industry and competitive environment. Ideally, peer firms should be exposed to the same common shocks and have a similar ability to respond to them. This would require selecting firms with comparable characteristics such as industry, size, diversification, and financial constraints. Yet, choosing firms that are too similar may not be feasible, as it would substantially narrow the available pool and result in a peer group that is too small (Albuquerque, 2009). This issue is particularly relevant for firms operating across multiple business lines, for which only few peers reflect their diverse operations and revenue structures (Bakke et al., 2020). As a consequence, firms may rely on imperfect peer groups because sufficiently comparable peers are not always available. This limitation may help explain why RPE is often applied imperfectly in practice. In such cases, the partial filtering of common risk from compensation may not arise from distortions caused by CEO power, but rather from the difficulty of constructing an ideal peer group (Jayaraman et al., 2021).

Overall, the model provides a useful framework for understanding how managerial power can shape the design and effectiveness of RPE contracts. It highlights that distortions in compensation design may arise even when incentive intensity remains optimal, primarily through incomplete risk filtering and wealth redistribution. Moreover, it provides a possible explanation for the RPE puzzle. However, several of the identified limitations suggest that the observed imperfections in RPE may not be driven solely by CEO power, but also by optimal contracting considerations and practical constraints. Thus, its empirical implications should be interpreted with caution.

4. Summary and Conclusion

This thesis examines how and under which conditions managerial power affects the design of RPE in executive compensation contracts and how the BoD will respond to this. While RPE is theoretically attractive because it filters common shocks and improves incentive efficiency, empirical evidence shows that it is not consistently applied in practice. This discrepancy, often referred to as the RPE puzzle, motivates the analysis of how CEO power may influence compensation design.

Building on the framework of Dikolli et al. (2018), the analysis explores how managerial influence over compensation contracts can affect the implementation of RPE. The model allows the CEO to influence certain elements of the contract, including the weight placed on peer performance and the composition of the peer group used for benchmarking. The results show that when CEOs hold such influence, they may adjust these elements in ways that increase their expected compensation. As a result, the risk filtering properties of RPE may be weakened because market risk is no longer fully removed from compensation.

An important insight of the analysis is that managerial power does not necessarily weaken incentives. The incentive weight on firm performance remains unchanged, implying that the efficient level of effort is still induced. Instead, managerial influence primarily affects the distribution of gains between the firm and the CEO. The model also shows that under certain conditions, such as when expected peer performance is high, boards may choose not to implement RPE at all rather than accept a distorted version of the mechanism. This provides a possible explanation for the RPE puzzle. In addition, the analysis highlights how strategic peer selection may arise when CEOs have influence over the peer selection process, which is consistent with empirical findings on peer-group formation.

To illustrate the implications of the model, numerical examples and comparative statics analyses were conducted. These simulations show how the effectiveness of RPE depends on several economic parameters, including CEO risk aversion, exposure to market risk, and firm-specific risk. The results provide additional insight into the conditions under which RPE is more or less valuable as a component of executive compensation contracts.

Despite these insights, the model remains a simplified representation of executive compensation design and captures only one dimension of managerial power. In practice, compensation contracts are shaped by a wider range of factors. The model should therefore be interpreted as one contribution within a broader literature that seeks to explain the observed variation in compensation practices.

Future research could build on this framework by empirically testing the predictions derived from the model and by extending the theoretical setting to address some of the limitations discussed in this thesis. For example, future work could consider the possibility of ex-post influence on the RPE contract or incorporate constraints on managerial power. Such extensions may further improve the understanding of how managerial power affects the design and effectiveness of RPE in executive compensation.

5. List of References

- Aggarwal, R. K., & Samwick, A. A. (1999). Executive Compensation, Strategic Competition, and Relative Performance Evaluation: Theory and Evidence. *The Journal of Finance*, 54(6), 1999–2043.
- Albuquerque. (2014). Do Growth-Option Firms Use Less Relative Performance Evaluation? *The Accounting Review*, 89(1), 27–60. <https://doi.org/10.2308/accr-50574>
- Albuquerque, A. (2009). Peer Firms in Relative Performance Evaluation. *Journal of Accounting and Economics*, 48(1), 69–89. <https://doi.org/10.1016/j.jacceco.2009.04.001>
- Bakke, T., Mahmudi, H., & Newton, A. (2020). Performance Peer Groups in CEO Compensation Contracts. *Financial Management*, 49(4), 997–1027. <https://doi.org/10.1111/fima.12296>
- Bebchuk, L., & Fried, J. (2004). *Pay without Performance*. Harvard University Press. <https://doi.org/10.2307/j.ctv2jfvcp7>
- Bertrand, M., & Mullainathan, S. (2001). Are CEOs Rewarded for Luck? The Ones Without Principals Are. *The Quarterly Journal of Economics*, 116, 932. <https://doi.org/10.1162/00335530152466269>
- Bettis, J. C., Bizjak, J. M., Coles, J. L., & Young, B. (2014). The Presence, Value, and Incentive Properties of Relative Performance Evaluation in Executive Compensation Contracts. *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.2392861>
- Bizjak, J., Kalpathy, S., Li, Z. F., & Young, B. (2022). The Choice of Peers for Relative Performance Evaluation in Executive Compensation. *Review of Finance*, 26(5), 1217–1239.
- Bizjak, J., Lemmon, M., & Nguyen, T. (2011). Are All CEOs Above Average? An Empirical Analysis of Compensation Peer Groups and Pay Design. *Journal of Financial Economics*, 100(3), 538–555. <https://doi.org/10.1016/j.jfineco.2011.02.007>
- Bloomfield, M. J., Marvão, C., & Spagnolo, G. (2023). Relative Performance Evaluation, Sabotage and Collusion. *Journal of Accounting & Economics*, 76. <https://doi.org/10.1016/j.jacceco.2023.101608>
- Bolton, P., & Dewatripont, M. (2005). *Contract Theory*. The MIT Press.
- Choe, C., Tian, G. Y., & Yin, X. (2014). CEO Power and the Structure of CEO Pay. *International Review of Financial Analysis*, 35, 237–248. <https://doi.org/10.1016/j.irfa.2014.10.004>
- Clementi, G. L., & Cooley, T. (2023). CEO Compensation: Facts. *Review of Economic Dynamics*, 50, 6–27. <https://doi.org/10.1016/j.red.2023.07.006>

- Coenenberg, A. G., Salfeld, R., & Schultze, W. (2015). *Wertorientierte Unternehmensführung—Vom Strategieentwurf zur Implementierung* (3rd edn). Schäffer-Poeschel Verlag. www.wiso-net.de.
- Canyon, M. J., & Peck, S. I. (2012). Executive Compensation, Pay-for-Performance and the Institutions of Executive Pay Setting. In *The SAGE Handbook of Corporate Governance* (pp. 453–475). SAGE Publications. <http://ebookcentral.proquest.com/lib/univie/detail.action?docID=1000639>.
- Core, J. E., & Guay, W. R. (2010). Is CEO Pay Too High and Are Incentives Too Low? A Wealth-Based Contracting Framework. *Academy of Management Perspectives*, 24(1), 5–19. JSTOR.
- Core, J. E., Guay, W. R., & Larcker, D. F. (2003). Executive Equity Compensation and Incentives: A Survey. *FRBNY Economic Policy Review*, 27–50.
- De Angelis, D., & Grinstein, Y. (2011). Relative Performance Evaluation in CEO Compensation: Evidence from the 2006 Disclosure Rules. *SSRN Electronic Journal*.
- De Angelis, D., & Grinstein, Y. (2020). Relative Performance Evaluation in CEO Compensation: A Talent-Retention Explanation. *Journal of Financial and Quantitative Analysis*, 55(7), 2099–2123. <https://doi.org/10.1017/S0022109019000504>
- DeMarzo, P. M., & Kaniel, R. (2023). Contracting in Peer Networks. *The Journal of Finance*, 78(5), 2725–2778. <https://doi.org/10.1111/jofi.13260>
- Dierkes, S., & Schäfer, U. (2008). Prinzipal-Agenten-Theorie und Performance Measurement. *Controlling & Management*, 52(S1), 19–27. <https://doi.org/10.1365/s12176-012-0186-z>
- Dikolli, S. S., Diser, V., Hofmann, C., & Pfeiffer, T. (2018). CEO Power and Relative Performance Evaluation. *Contemporary Accounting Research*, 35(3), 1279–1296. <https://doi.org/10.1111/1911-3846.12316>
- Dikolli, S. S., Hofmann, C., & Pfeiffer, T. (2013). Relative Performance Evaluation and Peer-Performance Summarization Errors. *Review of Accounting Studies*, 18(1), 34–65. <https://doi.org/10.1007/s11142-012-9212-9>
- DiPrete, T. A., Eirich, G. M., & Pittinsky, M. (2010). Compensation Benchmarking, Leapfrogs, and the Surge in Executive Pay. *American Journal of Sociology*, 115(6), 1671–1712. <https://doi.org/10.1086/652297>
- Diser, V., & Hofmann, C. (2018). Hedging and Accounting-Based RPE Contracts for Powerful CEOs. *Journal of Business Economics*, 88(7–8), 941–970. <https://doi.org/10.1007/s11573-018-0907-7>

- Edmans, A., Gabaix, X., & Jenter, D. (2017). Executive Compensation: A Survey of Theory and Evidence. *NBER Working Paper 23596*. <https://doi.org/10.3386/w23596>
- Ewert, R., Wagenhofer, A., & Rohlfing-Bastian, A. (2023). *Interne Unternehmensrechnung*. Springer Berlin Heidelberg. <https://doi.org/10.1007/978-3-662-65283-1>
- Finkelstein, S. (1992). Power in Top Management Teams: Dimensions, Measurement, and Validation. *Academy of Management Journal*, 35(3), 505–538. <https://doi.org/10.2307/256485>
- Florin, B., Hallock, K. F., & Webber, D. (2010). *Executive Pay and Firm Performance: Methodological Considerations and Future Directions*. 49–86. [https://doi.org/10.1108/S0742-7301\(2010\)0000029004](https://doi.org/10.1108/S0742-7301(2010)0000029004)
- Frydman, C., & Jenter, D. (2010). *CEO Compensation*. National Bureau of Economic Research. <https://doi.org/10.3386/w16585>
- Garvey, G., & Milbourn, T. (2003). Incentive Compensation When Executives Can Hedge the Market: Evidence of Relative Performance Evaluation in the Cross Section. *The Journal of Finance*, 58(4), 1557–1582. <https://doi.org/10.1111/1540-6261.00577>
- Garvey, G. T., & Milbourn, T. T. (2006). Asymmetric Benchmarking in Compensation: Executives Are Rewarded for Good Luck But Not Penalized for Bad. *Journal of Financial Economics*, 82(1), 197–225. <https://doi.org/10.1016/j.jfineco.2004.01.006>
- Gaßner, F. (2012). *Wertorientierte Managementvergütung: Der Optimale Einsatz von Anreizsystemen für Manager*. Diplomica Verlag.
- Gibbons, R., & Murphy, K. J. (1990). Relative Performance Evaluation for Chief Executive Officers. *ILR Review*, 43(3), 30S-51S. <https://doi.org/10.2307/2523570>
- Gong, G., Li, L. Y., & Shin, J. Y. (2011). Relative Performance Evaluation and Related Peer Groups in Executive Compensation Contracts. *The Accounting Review*, 86(3), 1007–1043. <https://doi.org/10.2308/accr.00000042>
- Gopalan, R., Milbourn, T., & Song, F. (2010). Strategic Flexibility and the Optimality of Pay for Sector Performance. *Review of Financial Studies*, 23(5), 2060–2098. <https://doi.org/10.1093/rfs/hhp118>
- Göx, R. F. (2016). Wachstum und Höhe von Managementvergütungen. *Perspektiven Der Wirtschaftspolitik*, 17(4), 311–334. <https://doi.org/10.1515/pwp-2016-0025>
- Göx, R. F., & Hemmer, T. (2020). On the Relation between Managerial Power and CEO Pay. *Journal of Accounting and Economics*, 69(2–3), 101300–101300. <https://doi.org/10.1016/j.jacceco.2020.101300>

- Grabke-Rundell, A., & Gomez-Mejia, L. R. (2002). Power as a Determinant of Executive Compensation. *Human Resource Management Review*, 12(1), 3–23. [https://doi.org/10.1016/S1053-4822\(01\)00038-9](https://doi.org/10.1016/S1053-4822(01)00038-9)
- Holmstrom, B. (1982). Moral Hazard in Teams. *The Bell Journal of Economics*, 13(2), 324–340. <https://doi.org/10.2307/3003457>
- Infuehr, J. (2022). Relative Performance Evaluation and Earnings Management. *Contemporary Accounting Research*, 39(1), 607–627. <https://doi.org/10.1111/1911-3846.12731>
- Jayaraman, S., Milbourn, T., Peters, F., & Seo, H. (2021). Product Market Peers and Relative Performance Evaluation. *The Accounting Review*, 96(4), 341–366. <https://doi.org/10.2308/TAR-2018-0284>
- Jin, L. (2002). CEO Compensation, Diversification, and Incentives. *Journal of Financial Economics*, 66, 63. [https://doi.org/10.1016/S0304-405X\(02\)00150-2](https://doi.org/10.1016/S0304-405X(02)00150-2)
- Kabitz, L. (2017). Survey on Empirical RPE Tests. *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.2961190>
- Kuhnen, C. M., & Zwiebel, J. H. (2008). Executive Pay, Hidden Compensation and Managerial Entrenchment. *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.972622>
- Küpper, H.-U., Friedl, G., Hofmann, C., Hofmann, Y., & Pedell, B. (2024). *Controlling: Konzeption - Aufgaben - Instrumente* (7th edn). Schäffer-Poeschel. <https://doi.org/10.34156/9783791053486>
- Laffont, J.-J., & Martimort, D. (2002). *The Theory of Incentives: The Principal-Agent Model*. Princeton University Press. <https://doi.org/10.1515/9781400829453>
- Lambert, R. A. (2001). Contracting Theory and Accounting. *Journal of Accounting and Economics*, 32(1), 3–87. [https://doi.org/10.1016/S0165-4101\(01\)00037-4](https://doi.org/10.1016/S0165-4101(01)00037-4)
- Liu, C., & Yin, C. (2023). Institutional Investors' Monitoring Attention, CEO Compensation, and Relative Performance Evaluation. *Finance Research Letters*, 56, 104121. <https://doi.org/10.1016/j.frl.2023.104121>
- Ma, P., Shin, J.-E., & Wang, C. C. (2018). Relative Performance Benchmarks: Do Boards Follow the Informativeness Principle? *Harvard Business School Accounting & Management Unit Working Paper*, (17–039).
- Maug, E. (2000). The Relative Performance Puzzle. *Schmalenbach Business Review*, 52(1), 3–24. <https://doi.org/10.1007/BF03396608>
- Milgrom, P., & Roberts, J. (1992). *Economics, Management and Organization*. Englewood Cliffs.

- Morse, A., Nanda, V., & Seru, A. (2011). Are Incentive Contracts Rigged by Powerful CEOs? *The Journal of Finance*, 66(5), 1779–1821. <https://doi.org/10.1111/j.1540-6261.2011.01687.x>
- Murphy, K. J., & Sandino, T. (2010). Executive Pay and “Independent” Compensation Consultants. *Journal of Accounting and Economics*, 49(3), 247–262. <https://doi.org/10.1016/j.jacceco.2009.12.001>
- Murphy, K. J., & Zábojník, J. (2004). CEO Pay and Appointments: A Market-Based Explanation for Recent Trends. *American Economic Review*, 94(2), 192–196. <https://doi.org/10.1257/0002828041302262>
- O'Reilly, C. A., & Main, B. G. M. (2010). Economic and Psychological Perspectives on CEO Compensation: A Review and Synthesis. *Industrial and Corporate Change*, 19(3), 675–712. <https://doi.org/10.1093/icc/dtp050>
- Oyer, P. (2004). Why Do Firms Use Incentives That Have No Incentive Effects? *The Journal of Finance*, 59(4), 1619–1649. JSTOR.
- Spremann, K. (1987). Agent and Principal. In *Agency Theory, Information, and Incentives* (pp. 3–37). Springer Berlin Heidelberg. https://doi.org/10.1007/978-3-642-75060-1_2
- Tice, F. M. (2024). The Role of Common Risk in the Effectiveness of Explicit Relative Performance Evaluation. *Management Science*, 70(3), 1635–1655. <https://doi.org/10.1287/mnsc.2023.4764>
- van Essen, M., Otten, J., & Carberry, E. J. (2015). Assessing Managerial Power Theory. *Journal of Management*, 41(1), 164–202. <https://doi.org/10.1177/0149206311429378>
- Vrettos, D. (2013). Are Relative Performance Measures in CEO Incentive Contracts Used for Risk Reduction and/or for Strategic Interaction? *The Accounting Review*, 88(6), 2179–2212. <https://doi.org/10.2308/accr-50548>
- Weisbach, M. S. (2007). *Optimal Executive Compensation versus Managerial Power: A Review of Lucian Bebchuk and Jesse Fried's Pay without Performance: The Unfulfilled Promise of Executive Compensation*: XLV (Journal of Economic Literature, pp. 419–428).
- Wu, M. G. H. (2013). Common vs. Firm-Specific Risks in Relative Performance Evaluation. *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.2288042>
- Wu, M. G. H. (2019). Optimal Risk Trade-Off in Relative Performance Evaluation. *Journal of Management Accounting Research*, 31(1), 247–259. <https://doi.org/10.2308/jmar-52060>

6. Appendix – Proof of Results

Certainty Equivalent – LEN Model

The CE is defined as the certain wealth level that gives the same utility as the expected utility of the risky outcome:

$$U(CE) = E[U(w, a)] \quad (\text{A. 1a})$$

Assuming the exponential utility from (10a) and supposing that w is normally distributed yields:

$$-e^{-r(CE)} = E[-e^{-r(w-k(a))}] \quad (\text{A. 1b})$$

$$-e^{-r(CE)} = -e^{-r(E[w]-k(a))+\frac{r^2}{2}\text{Var}[w]} \quad (\text{A. 1c})$$

$$CE = E[w] - k(a) - \frac{r}{2}\text{Var}[w] \quad (\text{A. 1d})$$

Expected Value of Compensation

$$E[w] = f + v_f a + v_m \mu_m \quad (\text{A. 2})$$

Variance of Compensation

Variance of compensation is given by the variance of its components (9):

$$\begin{aligned} \text{Var}[w] &= v_f^2 \text{Var}[y_f] + v_m^2 \text{Var}[y_m] + 2v_f v_m \text{Cov}[y_f, y_m] \\ \text{Var}[w] &= v_f^2 (\sigma_f^2 + \beta_f^2 \sigma_m^2) + v_m^2 (\sigma_m^2) + 2v_f v_m \beta_f \sigma_m^2 \\ \text{Var}[w] &= v_f^2 \sigma_f^2 + \sigma_m^2 (v_f \beta_f + v_m)^2 \end{aligned} \quad (\text{A. 3})$$

Incentive Compatibility (IC) Constraint

The CE from (10b) can be rewritten using (A.2) and (A.3), which yields $CE(w, a) = f + v_f a + v_m \mu_m - k(a) - \frac{r}{2} (v_f^2 \sigma_f^2 + \sigma_m^2 (v_f \beta_f + v_m)^2)$. The CEO maximizes this expression through their choice of effort, a . Thus, by differentiating this expression with respect to a and setting it to 0 the IC constraint is obtained: $\frac{\partial CE}{\partial a} = v_f - k'(a) \xrightarrow{\text{yields}} v_f = k'(a)$.

OBSERVATION 1 - The Case of a Powerful BoD

The BoD's decision problem is the following optimization problem:

$$\begin{aligned} \max_{f, v_f, v_m, a} \quad & \Pi = E[y_f - w] \\ \text{s. t. (IR)} : \quad & E[w] - k(a) - \frac{r}{2} \text{Var}[w] \geq 0 \\ \quad \quad \quad & (IC_a) : v_f = k'(a) \end{aligned}$$

Using the expressions for $E[w]$ (A.2) and $Var[w]$ (A.3), the Lagrangian for the BoD's decision problem is given by:

$$L = a - (f + v_f a + v_m \mu_m) - \lambda_1 \left(f + v_f a + v_m \mu_m - k(a) - \frac{r}{2} (v_f^2 \sigma_f^2 + \sigma_m^2 (v_f \beta_f + v_m)^2) \right) - \lambda_2 (k'(a) - v_f) \quad (\text{A. 4})$$

The Kuhn-Tucker conditions are given by:

$$\frac{\partial L}{\partial a} = 1 - v_f - \lambda_1 (v_f - k'(a)) - \lambda_2 (k''(a)) = 0 \quad (\text{A. 5a})$$

$$\frac{\partial L}{\partial v_f} = -a - \lambda_1 (a - r v_f \sigma_f^2 - r \sigma_m^2 (v_f \beta_f + v_m) \beta_f) + \lambda_2 = 0 \quad (\text{A. 5b})$$

$$\frac{\partial L}{\partial v_m} = -\mu_m - \lambda_1 (\mu_m - r \sigma_m^2 (v_f \beta_f + v_m)) = 0 \quad (\text{A. 5c})$$

$$\frac{\partial L}{\partial f} = -1 - \lambda_1 = 0 \quad (\text{A. 5d})$$

$$\begin{aligned} \frac{\partial L}{\partial \lambda_1} &= f + v_f a + v_m \mu_m - k(a) - \frac{r}{2} (v_f^2 \sigma_f^2 + \sigma_m^2 (v_f \beta_f + v_m)^2) \geq 0 \text{ and } \lambda_1 \leq 0 \\ \text{and } \lambda_1 \frac{\partial L}{\partial \lambda_1} &= \lambda_1 \left(f + v_f a + v_m \mu_m - k(a) - \frac{r}{2} (v_f^2 \sigma_f^2 + \sigma_m^2 (v_f \beta_f + v_m)^2) \right) = 0 \end{aligned} \quad (\text{A. 5e})$$

$$\frac{\partial L}{\partial \lambda_2} = -(k'(a) - v_f) = 0 \quad (\text{A. 5f})$$

From (A.5d) it follows that $\lambda_1 = -1$ and from (A.5f) that $v_f = k'(a)$. (A.5e) reflects the complementary slackness condition which implies that $\lambda_1 \frac{\partial L}{\partial \lambda_1} = 0$. The negative Lagrange multiplier, λ_1 , means that the constraint is active and plugging it into (A.5e) the condition can be rewritten as $f + v_f a + v_m \mu_m - k(a) - \frac{r}{2} (v_f^2 \sigma_f^2 + \sigma_m^2 (v_f \beta_f + v_m)^2) = 0$.

Solving (A.5c) with λ_1 yields $-\mu_m - (-1) (\mu_m - r \sigma_m^2 (v_f \beta_f + v_m)) = 0 \xrightarrow{\text{yields}} r \sigma_m^2 (v_f \beta_f + v_m) = 0$. The incentive weight on peer-group performance is therefore given by $v_m^\dagger = -v_f \beta_f$.

Inserting λ_1 and v_m into (A.5b) gives $-a - (-1) (a - r v_f \sigma_f^2 - r \sigma_m^2 (v_f \beta_f - v_f \beta_f) \beta_f) + \lambda_2 = 0 \xrightarrow{\text{yields}} \lambda_2 = r v_f \sigma_f^2$. Using this and $v_f = k'(a)$ to solve (A.5a) yields the incentive weight on firm performance: $1 - v_f - \lambda_1 (v_f - v_f) - r v_f \sigma_f^2 (k''(a)) = 0 \xrightarrow{\text{yields}} v_f^\dagger = \frac{1}{1 + r \sigma_f^2 k''(a)}$.

Finally, **expected firm value** is obtained by substituting $E[y_f | v_f^\dagger]$ as well as $-E[w] = -k(a) - \frac{r}{2} \text{Var}[w]$ into the objective function: $\Pi = E[y_f] - E[w] = E[y_f | v_f^\dagger] - k(a) - \frac{r}{2} \text{Var}[w]$. Using (A.3) and the risk-minimizing incentive weight $v_m^\dagger$ this yields $\Pi^\dagger = E[y_f | v_f^\dagger] - k(a) - \frac{r}{2} (v_f^2 \sigma_f^2 + \sigma_m^2 (v_f \beta_f - v_f \beta_f)) \xrightarrow{\text{yields}} E[y_f | v_f^\dagger] - k(a) - \frac{r}{2} (v_f^2 \sigma_f^2)$.

Starting from equation (9) and inserting (7) and (8) the **CEO's compensation** can be rewritten as $w = f + v_f(a + \varepsilon_f + \beta_f \varepsilon_m) + v_m(\mu_m + \varepsilon_m) = f + v_f a + v_f \varepsilon_f + v_f \beta_f \varepsilon_m + v_m \mu_m + v_m \varepsilon_m$. Inserting the incentive weight on peer-group performance $v_m^\dagger$ from above and using the fact that $E[w] = f + v_f a + v_m \mu_m$ from (A.2) I obtain $w = E[w] + v_f \varepsilon_f$. Substituting $E[w] = k(a) + \frac{r}{2} \text{Var}[w]$ and the variance computed with $v_m = -v_f \beta_f$ the rewritten expression for CEO compensation is $w^\dagger = k(a) + \frac{r}{2} v_f^2 \sigma_f^2 + v_f \varepsilon_f$.

PROPOSITION 1 - The Case of a Powerful CEO

The BoD's decision problem is the same optimization problem as in the case above, but with one additional constraint:

$$\begin{aligned} \max_{f, v_f, v_m, a} \Pi &= E[y_f - w] \\ \text{s. t. (IR)} : E[w] - k(a) - \frac{r}{2} \text{Var}[w] &\geq 0 \\ (IC_a) : v_f &= k'(a) \\ (IC_m) : \mu_m - r(v_f \beta_f + v_m) \sigma_m^2 &= 0 \end{aligned}$$

Rewriting it in the same manner as in the case of a powerful BoD, the Lagrangian for the BoD's decision problem can be expressed as:

$$\begin{aligned} L = a - (f + v_f a + v_m \mu_m) - \lambda_1 \left(f + v_f a + v_m \mu_m - k(a) - \frac{r}{2} (v_f^2 \sigma_f^2 + \sigma_m^2 (v_f \beta_f + v_m)^2) \right) \\ - \lambda_2 (k'(a) - v_f) - \lambda_3 (\mu_m - r(v_f \beta_f + v_m) \sigma_m^2) \end{aligned} \quad (\text{A. 6})$$

The Kuhn-Tucker Conditions are given by:

$$\frac{\partial L}{\partial a} = 1 - v_f - \lambda_1 (v_f - k'(a)) - \lambda_2 (k''(a)) = 0 \quad (\text{A. 7a})$$

$$\frac{\partial L}{\partial v_f} = -a - \lambda_1 (a - r v_f \sigma_f^2 - r \sigma_m^2 (v_f \beta_f + v_m) \beta_f) + \lambda_2 + \lambda_3 r \beta_f \sigma_m^2 = 0 \quad (\text{A. 7b})$$

$$\frac{\partial L}{\partial v_m} = -\mu_m - \lambda_1 (\mu_m - r \sigma_m^2 (v_f \beta_f + v_m)) + \lambda_3 r \sigma_m^2 = 0 \quad (\text{A. 7c})$$

$$\frac{\partial L}{\partial f} = -1 - \lambda_1 = 0 \quad (\text{A. 7d})$$

$$\begin{aligned} \frac{\partial L}{\partial \lambda_1} &= f + v_f a + v_m \mu_m - k(a) - \frac{r}{2} (v_f^2 \sigma_f^2 + \sigma_m^2 (v_f \beta_f + v_m)^2) \geq 0 \text{ and } \lambda_1 \leq 0 \\ \text{and } \lambda_1 \frac{\partial L}{\partial \lambda_1} &= \lambda_1 \left(f + v_f a + v_m \mu_m - k(a) - \frac{r}{2} (v_f^2 \sigma_f^2 + \sigma_m^2 (v_f \beta_f + v_m)^2) \right) = 0 \end{aligned} \quad (\text{A. 7e})$$

$$\frac{\partial L}{\partial \lambda_2} = -(k'(a) - v_f) = 0 \quad (\text{A.7f})$$

$$\frac{\partial L}{\partial \lambda_3} = -(\mu_m - r(v_f \beta_f + v_m) \sigma_m^2) = 0 \quad (\text{A.7g})$$

As explained in the previous case, (A.7e) reflects the complementary slackness condition.

The Lagrangian is solved similarly to before. From (A.7d) it follows that $\lambda_1 = -1$ and from (A.7f) that $v_f = k'(a)$. As the Lagrange multiplier $\lambda_1 = -1$, the condition from (A.7e) can be rewritten as $f + v_f a + v_m \mu_m - k(a) - \frac{r}{2} (v_f^2 \sigma_f^2 + \sigma_m^2 (v_f \beta_f + v_m)^2) = 0$.

Rearranging (A.7g) the expression for the incentive weight on peer-group performance can be obtained: $v_m^\dagger = \frac{\mu_m}{r \sigma_m^2} - v_f \beta_f$. Inserting (A.7f) into (A.7a) yields $\lambda_2 = \frac{1-v_f}{k''(a)}$ and by using the expression for $v_m^\dagger$ in (A.7c) $\lambda_3 = \frac{\mu_m}{r \sigma_m^2}$ is obtained. Substituting the value of all three Lagrange multipliers, λ_1 , λ_2 , λ_3 , and $v_m^\dagger = \frac{\mu_m}{r \sigma_m^2} - v_f \beta_f$ into (A.7b) yields:

$$\frac{\partial L}{\partial v_f} = -a + \left(a - r v_f \sigma_f^2 - r \sigma_m^2 \left(v_f \beta_f + \frac{\mu_m}{r \sigma_m^2} - v_f \beta_f \right) \beta_f \right) + \frac{1-v_f}{k''(a)} + \frac{\mu_m}{r \sigma_m^2} r \beta_f \sigma_m^2 = 0 \quad (\text{A.8})$$

Simplifying this expression gives: $v_f^\dagger = \frac{1}{1+r k''(a) \sigma_f^2}$, which is equal to the expression for $v_f^\dagger$.

Expected firm value is obtained by substituting $E[y_f | v_f^\dagger]$ as well as $E[w^\dagger | a^\dagger]$ from COROLLARY 1 into the objective function: $\Pi = E[y_f] - E[w] = E[y_f | v_f^\dagger] - E[w^\dagger | a^\dagger] = E[y_f | v_f^\dagger] - k(a) - \frac{r}{2} v_f^2 \sigma_f^2 - \frac{\mu_m^2}{2 r \sigma_m^2}$. Recalling from (12b) that $\Pi^\dagger = E[y_f | v_f^\dagger] - k(a^\dagger) - \frac{r}{2} (v_f^\dagger)^2 \sigma_f^2$, this yields equation (17b) $\Pi^\dagger = \Pi^\dagger - \frac{\mu_m^2}{2 r \sigma_m^2}$.

COROLLARY 1

As demonstrated before, the **CEO's compensation** can be expressed as $w = f + v_f a + v_f \varepsilon_f + v_f \beta_f \varepsilon_m + v_m \mu_m + v_m \varepsilon_m$. Using the fact that $f + v_f a + v_m \mu_m = E[w] = k(a) + \frac{r}{2} \text{Var}[w]$ and inserting $v_m^\dagger = \frac{\mu_m}{r \sigma_m^2} - v_f \beta_f$ yields $w = k(a) + \frac{r}{2} (v_f^2 \sigma_f^2 + \sigma_m^2 (v_f \beta_f + \frac{\mu_m}{r \sigma_m^2} - v_f \beta_f)^2) + v_f \varepsilon_f + \frac{\mu_m \varepsilon_m}{r \sigma_m^2}$. Simplifying gives $w = k(a) + \frac{r}{2} v_f^2 \sigma_f^2 + v_f \varepsilon_f + \frac{\mu_m^2}{2 r \sigma_m^2} + \frac{\mu_m}{r \sigma_m^2} \varepsilon_m$. Finally, recalling the expression for $w^\dagger$ (13) from the previous case, I obtain $w^\dagger = w^\dagger + \frac{\mu_m^2}{2 r \sigma_m^2} + \frac{\mu_m}{r \sigma_m^2} \varepsilon_m$.

No use of RPE

The weight on peer-firm performance in the case where the BoD refrains from the use of RPE is obtained by solving the following optimization problem in the same manner as before

$$\begin{aligned} \max_{f, v_f, v_m, a} \quad & \Pi = E[y_f - w] \\ \text{s. t. (IR)} : \quad & E[w] - k(a) - \frac{r}{2} \text{Var}[w] \geq 0 \\ \text{(IC}_a\text{)} : \quad & v_f = k'(a) \\ & v_m = 0 \end{aligned}$$

Thus, the Lagrangian for the BoD's decision problem can be expressed as:

$$\begin{aligned} L = a - (f + v_f a + v_m \mu_m) - \lambda_1 \left(f + v_f a + v_m \mu_m - k(a) - \frac{r}{2} (v_f^2 o_f^2 + \sigma_m^2 (v_f \beta_f + v_m)^2) \right) \\ - \lambda_2 (k'(a) - v_f) - \lambda_3 (v_m) \end{aligned} \quad (\text{A. 9})$$

The Kuhn-Tucker Conditions are given by:

$$\frac{\partial L}{\partial a} = 1 - v_f - \lambda_1 (v_f - k'(a)) - \lambda_2 (k''(a)) = 0 \quad (\text{A. 10a})$$

$$\frac{\partial L}{\partial v_f} = -a - \lambda_1 (a - r v_f \sigma_f^2 - r o_m^2 (v_f \beta_f + v_m) \beta_f) + \lambda_2 = 0 \quad (\text{A. 10b})$$

$$\frac{\partial L}{\partial v_m} = -\mu_m - \lambda_1 (\mu_m - r \sigma_m^2 (v_f \beta_f + v_m)) - \lambda_3 = 0 \quad (\text{A. 10c})$$

$$\frac{\partial L}{\partial f} = -1 - \lambda_1 = 0 \quad (\text{A. 10d})$$

$$\begin{aligned} \frac{\partial L}{\partial \lambda_1} = f + v_f a + v_m \mu_m - k(a) - \frac{r}{2} (v_f^2 o_f^2 + \sigma_m^2 (v_f \beta_f + v_m)^2) \geq 0 \text{ and } \lambda_1 \leq 0 \\ \text{and } \lambda_1 \frac{\partial L}{\partial \lambda_1} = \lambda_1 \left(f + v_f a + v_m \mu_m - k(a) - \frac{r}{2} (v_f^2 o_f^2 + \sigma_m^2 (v_f \beta_f + v_m)^2) \right) = 0 \end{aligned} \quad (\text{A. 10e})$$

$$\frac{\partial L}{\partial \lambda_2} = -(k'(a) - v_f) = 0 \quad (\text{A. 10f})$$

$$\frac{\partial L}{\partial \lambda_3} = -v_m = 0 \quad (\text{A. 10g})$$

Substituting $v_m = 0$ and $\lambda_1 = -1$ in (A.10b) I derive $\lambda_2 = r v_f \sigma_f^2 - r \sigma_m^2 v_f \beta_f^2$. Then, inserting that and $k'(a) = v_f$ in (A.10a) and simplifying yields $v_f^{NR} = \frac{1}{1 + r k''(a^{NR}) (\sigma_f^2 + \beta_f^2 \sigma_m^2)}$ from (19a).

Expected firm performance is obtained by $\Pi = E[y_f] - E[w] = E[y_f | v_f^{NR}] - k(a) - \frac{r}{2} \text{Var}[w]$ using $v_m = 0$ and (A.3) this yields $\Pi^{NR} = E[y_f | v_f^{NR}] - k(a^{NR}) - \frac{r}{2} (v_f^{NR})^2 (\sigma_f^2 + \beta_f^2 \sigma_m^2)$ from (19b).

COROLLARY 2 – RPE Threshold

It is assumed that the BoD chooses the performance measure that is more informative. This corresponds to the measure with the smallest variance. To establish the threshold from COROLLARY 2, the aggregated performance measures with and without the use of RPE are compared. Without RPE, the CEO is only evaluated based on firm performance, $y_f = a + \varepsilon_f + \beta_f \varepsilon_m$ from equation (7), with $Var[y_f] = \sigma_f^2 + \beta_f^2 \sigma_m^2$. If RPE is used under a powerful CEO, the aggregated performance measure corresponds to $y^\ddagger = y_f + y_m \cdot \frac{v_m^\ddagger}{v_f^\ddagger}$ from (14). Substituting

the expressions for y_f from (7), y_m from (8) and $v_m^\ddagger$ from (17a) this yields $y^\ddagger = a + \varepsilon_f + \beta_f \varepsilon_m + (\mu_m + \varepsilon_m) \cdot \left(-v_f^\ddagger \beta_f + \frac{\mu_m}{r\sigma_m^2}\right) \frac{1}{v_f^\ddagger} = a + \varepsilon_f + \beta_f \varepsilon_m - \mu_m \beta_f - \varepsilon_m \beta_f + \frac{\mu_m^2}{r\sigma_m^2 v_f^\ddagger} + \frac{\varepsilon_m \mu_m}{r\sigma_m^2 v_f^\ddagger}$. Simplifying, substituting $v_f^\ddagger$ from (17a) and considering only the noise terms that are relevant for the

variance yields: $y^\ddagger = \varepsilon_f + \varepsilon_m \frac{\mu_m(1+rk''(a^\ddagger)\sigma_f^2)}{r\sigma_m^2}$ with $Var[y^\ddagger] = \sigma_f^2 + \frac{\mu_m^2(1+rk''(a^\ddagger)\sigma_f^2)^2}{r^2\sigma_m^2}$.

To evaluate whether to implement RPE or not, the BoD chooses the option with the most informative performance measure, hence the one with the smallest variance. Therefore, the BoD will not implement RPE if $Var[y_f] \leq Var[y^\ddagger]$:

$$\sigma_f^2 + \beta_f^2 \sigma_m^2 \leq \sigma_f^2 + \frac{\mu_m^2(1+rk''(a^\ddagger)\sigma_f^2)^2}{r^2\sigma_m^2} \quad (\text{A.11})$$

Simplifying and solving for μ_m yields the threshold from COROLLARY 2 in (20).

By substituting for $v_f^\ddagger$ and $v_m^\ddagger$ in the just derived threshold, this condition implies $v_m^\ddagger > 0$:

$$\begin{aligned} \mu_m > \frac{r\beta_f\sigma_m^2}{1+rk''(a^\ddagger)\sigma_f^2} &\xrightarrow{\text{yields}} \mu_m - r\beta_f\sigma_m^2 \frac{1}{1+rk''(a^\ddagger)\sigma_f^2} > 0 &\xrightarrow{\text{yields}} \mu_m - r\beta_f\sigma_m^2 v_f^\ddagger > 0 \\ &\xrightarrow{\text{yields}} \frac{\mu_m}{r\sigma_m^2} - \beta_f v_f^\ddagger = v_m^\ddagger > 0 & (\text{A.12}) \end{aligned}$$

COROLLARY 3 - CEO choice in peer selection process

If a single peer firm is used as compensation benchmark instead of an index, the CE of the CEO changes to:

$$CE = f + v_f a + v_i \mu_i - k(a) - \frac{r}{2} \left(v_f^2 \sigma_f^2 + (v_f \beta_f + v_i \beta_i)^2 \sigma_m^2 \right) \quad (\text{A. 13a})$$

And the decision on v_i can be modelled through this first order condition:

$$(IC_i): \frac{\partial CE}{\partial v_i} = \mu_i - r(v_f \beta_f + v_i \beta_i) \beta_i \sigma_m^2 = 0 \quad (\text{A. 13b})$$

Rearranging for v_i yields the expression from (22):

$$v_i = -v_f \frac{\beta_f}{\beta_i} + \frac{\mu_i}{r \beta_i^2 \sigma_m^2} \quad (\text{A. 13c})$$

The **risk-minimizing incentive weight** is given by:

$$\frac{\partial Var[w]}{\partial v_i} = 2(v_f \beta_f + v_i \beta_i) \beta_i \sigma_m^2 = 0 \xrightarrow{\text{yields}} v_i = -v_f \frac{\beta_f}{\beta_i} \quad (\text{A. 13d})$$

Inserting the expression for v_i from (A.13c) into (A.13a) yields:

$$\begin{aligned} CE_i &= f + v_f a + \left(-v_f \frac{\beta_f}{\beta_i} + \frac{\mu_i}{r \beta_i^2 \sigma_m^2} \right) \mu_i - k(a) \\ &\quad - \frac{r}{2} \left(v_f^2 \sigma_f^2 + \left(v_f \beta_f + \left(-v_f \frac{\beta_f}{\beta_i} + \frac{\mu_i}{r \beta_i^2 \sigma_m^2} \right) \beta_i \right)^2 \sigma_m^2 \right) \\ &= f + v_f a - k(a) - \frac{r}{2} v_f^2 \sigma_f^2 - v_f \mu_i \frac{\beta_f}{\beta_i} + \frac{\mu_i^2}{r \beta_i^2 \sigma_m^2} - \frac{\mu_i^2}{2r \beta_i^2 \sigma_m^2} \\ &= \left(f + v_f E[y_f | v_f] - k(a) - \frac{r}{2} v_f^2 \sigma_f^2 \right) - v_f \mu_i \frac{\beta_f}{\beta_i} + \frac{\mu_i^2}{2r \beta_i^2 \sigma_m^2} \\ &= \left(f + v_f E[y_f | v_f] - k(a) - \frac{r}{2} v_f^2 \sigma_f^2 \right) + \frac{\mu_i}{\beta_i} \cdot \frac{1}{2r \sigma_m^2} \left(\frac{\mu_i}{\beta_i} - 2v_f r \beta_f \sigma_m^2 \right) \end{aligned} \quad (\text{A. 14})$$

PROPOSITION 2 - Expected Regression Coefficients

Case (I) - Powerful CEO

Starting from the compensation equation (9) and inserting $v_m^\dagger$, it follows that $w^\dagger = f^\dagger + v_f^\dagger y_f + v_m^\dagger y_m = f^\dagger + \frac{\mu_m}{r\sigma_m^2} y_m + v_f^\dagger (y_f - \beta_f y_m)$. Using $y_u = y_f - y_s = y_f - \beta_f y_m$ and $y_s = \beta_f y_m$ from (25c) so that $y_m = \frac{y_s}{\beta_f}$, I obtain $w^\dagger = f^\dagger + \frac{\mu_m}{r\sigma_m^2 \beta_f} y_s + v_f^\dagger y_u$. Substituting into (26b), using $Cov[y_s, y_u] = 0$ and simplifying yields:

$$\hat{\alpha}_u^\dagger = \frac{Cov[w^\dagger, y_u]}{Var[y_u]} = \frac{Cov\left[f^\dagger + \frac{\mu_m}{r\sigma_m^2 \beta_f} y_s + v_f^\dagger y_u, y_u\right]}{Var[y_u]} = \frac{v_f^\dagger Var[y_u]}{Var[y_u]} = v_f^\dagger \quad (A.15a)$$

$$\hat{\alpha}_s^\dagger = \frac{Cov[w^\dagger, y_s]}{Var[y_s]} = \frac{Cov\left[f^\dagger + \frac{\mu_m}{r\sigma_m^2 \beta_f} y_s + v_f^\dagger y_u, y_s\right]}{Var[y_s]} = \frac{\frac{\mu_m}{r\sigma_m^2 \beta_f} Var[y_s]}{Var[y_s]} = \frac{\mu_m}{r\sigma_m^2 \beta_f} \quad (A.15b)$$

Case (II) - No RPE

If no RPE is used, compensation is equal to $w^{NR} = f^{NR} + v_f^{NR} y_f$. Using $y_f = y_u + y_s$ from (25c) yields $w^{NR} = f^{NR} + v_f^{NR} (y_u + y_s)$. Substituting into (26b) yields:

$$\hat{\alpha}_u^{NR} = \frac{Cov[w^{NR}, y_u]}{Var[y_u]} = \frac{Cov[f^{NR} + v_f^{NR} (y_u + y_s), y_u]}{Var[y_u]} = \frac{v_f^{NR} Var[y_u]}{Var[y_u]} = v_f^{NR} \quad (A.15c)$$

$$\hat{\alpha}_s^{NR} = \frac{Cov[w^{NR}, y_s]}{Var[y_s]} = \frac{Cov[f^{NR} + v_f^{NR} (y_u + y_s), y_s]}{Var[y_s]} = \frac{v_f^{NR} Var[y_s]}{Var[y_s]} = v_f^{NR} \quad (A.15d)$$

Case (III) - Powerful BoD

Starting from (9) and inserting $v_m^\dagger$, it follows that $w^\dagger = f^\dagger + v_f^\dagger y_f + v_m^\dagger y_m = f^\dagger + v_f^\dagger (y_f - \beta_f y_m)$. Using $y_u = y_f - \beta_f y_m$ from (25c) this simplifies to $w^\dagger = f^\dagger + v_f^\dagger y_u$. Substituting into (26b) yields:

$$\hat{\alpha}_u^\dagger = \frac{Cov[w^\dagger, y_u]}{Var[y_u]} = \frac{Cov[f^\dagger + v_f^\dagger y_u, y_u]}{Var[y_u]} = \frac{v_f^\dagger Var[y_u]}{Var[y_u]} = v_f^\dagger \quad (A.15e)$$

$$\hat{\alpha}_s^\dagger = \frac{Cov[w^\dagger, y_s]}{Var[y_s]} = \frac{Cov[f^\dagger + v_f^\dagger y_u, y_s]}{Var[y_s]} = \frac{v_f^\dagger Cov[y_u, y_s]}{Var[y_s]} = 0 \quad (A.15f)$$

COROLLARY 4 – Comparative Statics

Case (I) - Powerful CEO

$$\begin{aligned}\frac{\partial \hat{\alpha}_s^{\dagger}}{\partial \mu_m} &= \frac{1}{r\beta_f\sigma_m^2} > 0; \quad \frac{\partial \hat{\alpha}_s^{\dagger}}{\partial r} = -\frac{\mu_m}{r^2\beta_f\sigma_m^2} < 0; \quad \frac{\partial \hat{\alpha}_s^{\dagger}}{\partial \beta_f} = -\frac{\mu_m}{r\beta_f^2\sigma_m^2} < 0; \\ \frac{\partial \hat{\alpha}_s^{\dagger}}{\partial \sigma_m^2} &= -\frac{\mu_m}{r\beta_f(\sigma_m^2)^2} < 0; \quad \frac{\partial \hat{\alpha}_s^{\dagger}}{\partial \sigma_f^2} = 0;\end{aligned}\tag{A.16a}$$

Case (II) - No RPE

$$\begin{aligned}\frac{\partial \hat{\alpha}_s^{NR}}{\partial \mu_m} &= 0; \quad \frac{\partial \hat{\alpha}_s^{NR}}{\partial r} = -\frac{k''(a^{NR})(\sigma_f^2 + \beta_f^2\sigma_m^2)}{[1 + rk''(a^{NR})(\sigma_f^2 + \beta_f^2\sigma_m^2)]^2} < 0; \\ \frac{\partial \hat{\alpha}_s^{NR}}{\partial \beta_f} &= -\frac{rk''(a^{NR})(2\beta_f\sigma_m^2)}{[1 + rk''(a^{NR})(\sigma_f^2 + \beta_f^2\sigma_m^2)]^2} < 0; \\ \frac{\partial \hat{\alpha}_s^{NR}}{\partial \sigma_m^2} &= -\frac{rk''(a^{NR})\beta_f^2}{[1 + rk''(a^{NR})(\sigma_f^2 + \beta_f^2\sigma_m^2)]^2} < 0; \\ \frac{\partial \hat{\alpha}_s^{NR}}{\partial \sigma_f^2} &= -\frac{rk''(a^{NR})}{[1 + rk''(a^{NR})(\sigma_f^2 + \beta_f^2\sigma_m^2)]^2} < 0;\end{aligned}\tag{A.16b}$$

PROPOSITION 3 – Expected Regression Coefficients on Peer-Group Performance

Applying standard regression techniques for the linear function $w(y_f, y_m) = f + v_f y_f + v_m y_m$ from (9), yields:

$$\hat{\tau}_m = \frac{Cov[w, y_m] Var[y_f] - Cov[w, y_f] Cov[y_f, y_m]}{Var[y_m] Var[y_f] - Cov[y_f, y_m]^2}\tag{A.17a}$$

For all cases $\{\dagger, \ddagger, NR\}$ it holds that:

$$Cov[w, y_m] = v_f Cov[y_f, y_m] + v_m Var[y_m]\tag{A.17b}$$

$$Cov[w, y_f] = v_f Var[y_f] + v_m Cov[y_f, y_m]\tag{A.17c}$$

Inserting (A.17b) and (A.17c) into (A.17a) yields:

$$\hat{\tau}_m = \frac{\{v_f Cov[y_f, y_m] + v_m Var[y_m]\} Var[y_f] - \{v_f Var[y_f] + v_m Cov[y_f, y_m]\} Cov[y_f, y_m]}{Var[y_m] Var[y_f] - Cov[y_f, y_m]^2}\tag{A.17d}$$

Simplifying, I obtain $\hat{\tau}_m = v_m$.