

収穫逓減を考慮した最適インセンティブ契約

— 線形契約ベンチマークと会計・ガバナンスへの含意 —

鈴木 孝則

早稲田大学商学学術院教授

2026 年 6 月 10 日

論文要旨

標準的なプリンシパル・エイジェント・モデルでは、成果が努力に対して線形に反応する、すなわち収穫一定であることがしばしば仮定される。しかし、研究開発、市場開拓、業務改善、管理会計上の KPI 達成などの場面では、努力の限界効果が逓減することが多い。本稿は、成果生成過程に収穫逓減性を導入したうえで、線形報酬契約クラスにおける最適インセンティブ係数を分析する。具体的には、成果を $y = 1 - e^{-a} + \varepsilon$, $\varepsilon \sim N(0, \sigma^2)$ とし、エイジェントは CARA 効用を持つと仮定する。エイジェントの確実性等価は努力に関して厳密に凹であるため、一階条件アプローチが本モデル内では正当化され、努力反応はランベルト W 関数により $a(\beta) = W_0(\beta/k)$ と表される。主要な結果は次の通りである。第一に、最適インセンティブ係数 $\beta^*(\sigma)$ は一意に存在し、リスク σ に対して単調減少する。第二に、標準モデルの係数 $\beta_{\text{std}} = 1/(1 + rk\sigma^2)$ と比較すると、収穫逓減モデルの β^* は任意の $\sigma > 0$ で小さい。第三に、両者の絶対乖離 $\Delta(\sigma) \equiv \beta_{\text{std}} - \beta^*$ は中程度のリスク領域で最大化する一峰性を示す。基準ケース $k = r = 1$ では、 $\Delta(\sigma)$ が最大となる点は $\sigma^\dagger \simeq 0.570$ であり、そこでの下方修正率は約 30.8% である。なお、相対的な下方修正率 $(\beta_{\text{std}} - \beta^*)/\beta_{\text{std}}$ 自体の最大値は $\sigma \simeq 0.957$ 付近でおよそ 35.0% に達する。 (k, r) パラメータに関する数値検討では、一峰性は頑健に観察される一方、最大乖離の大きさは努力コスト係数 k に依存する： k が小さいほど、すなわち努力が相対的に安価で高努力領域に入りやすいほど、標準モデルとの乖離は大きくなる。リスク回避度 r は主として σ 軸をリスケールする。本稿では正規化された凹型成果関数族 $f_\theta(a) = (1 - e^{-\theta a})/\theta$ および対数族 $\log(1 + \theta a)/\theta$ による頑健性確認を行う。さらに、本稿の下方修正命題は収穫逓減型または凹型の成果関数に関するベンチマークであり、研究開発・基礎研究のような収穫逓増・閾値型業務には機械的に適用できないことを明示したうえで、日本のコーポレートガバナンス・コードへの実務的含意を示す。

キーワード：プリンシパル・エイジェント理論、インセンティブ契約、収穫逓減、CARA 効用、一階条件アプローチ、ランベルト W 関数、会計情報品質、コーポレートガバナンス・コード

Optimal Incentive Contracts under Diminishing Returns:

A Linear-Contract Benchmark and Implications for Accounting and Governance

Takanori Suzuki

Professor, Faculty of Commerce, Waseda University

June 10, 2026

Abstract

This paper studies a linear incentive contract under a diminishing-return technology, $y = 1 - e^{-a} + \varepsilon$, $\varepsilon \sim N(0, \sigma^2)$, with a CARA agent. The agent's effort response is given by $a(\beta) = W_0(\beta/k)$, and the optimal incentive coefficient $\beta^*(\sigma)$ is unique, decreases monotonically in risk, and lies strictly below the constant-return benchmark $\beta_{\text{std}} = 1/(1 + rk\sigma^2)$ for every $\sigma > 0$. The absolute gap $\Delta(\sigma) \equiv \beta_{\text{std}} - \beta^*$ is unimodal. In the benchmark $k = r = 1$, the absolute gap is maximized at $\sigma^\dagger \simeq 0.570$, where the relative downward adjustment is about 30.8%; the relative adjustment rate itself is maximized at $\sigma \simeq 0.957$ at about 35.0%. Numerical analysis shows that while the unimodality is robust over a (k, r) grid, the magnitude of the gap depends on the effort-cost coefficient k : smaller k (cheaper effort) yields a larger gap, while the risk-aversion coefficient r primarily rescales the σ -axis. Robustness is verified using the normalized concave families $f_\theta(a) = (1 - e^{-\theta a})/\theta$ and $\log(1 + \theta a)/\theta$. The paper also clarifies that the downward-adjustment result is a benchmark for diminishing-return or concave technologies and should not be mechanically applied to increasing-return or threshold-type tasks such as basic research and early-stage R&D. Implications for Japan's Corporate Governance Code are discussed.

Keywords: principal-agent theory, incentive contracts, diminishing returns, CARA utility, Lambert W function, accounting quality, Corporate Governance Code

1 はじめに

企業が持続的に成長し、株主をはじめとするステークホルダーの期待に応えるためには、経営者や従業員の行動を企業価値の向上へ方向付けるインセンティブ設計が不可欠である。成果連動型報酬は努力を促す一方で、成果指標に含まれるノイズをエイジェントに負担させる。インセンティブの強さとリスク負担とのトレードオフは、コーポレート・ガバナンス、管理会計、業績評価制度における基本問題であり続けている。

この問題に理論的基盤を与えてきたのがプリンシパル・エイジェント（principal-agent; P-A）理論である。Holmström (1979) は、観測可能な成果指標がエイジェントの行動に関する情報をどの程度含むかに応じて報酬契約を設計すべきであることを示した。Holmström and Milgrom (1987) は、CARA 効用、正規ノイズ、線形契約の枠組みにおいて、線形報酬契約が自然に導かれる条件を明らかにした。この標準的な線形・正規・CARA モデルでは、努力コスト係数を k 、絶対的リスク回避度を r 、成果ノイズの標準偏差を σ とすると、標準的なインセンティブ係数は

$$\beta_{\text{std}}(\sigma) = \frac{1}{1 + rk\sigma^2} \quad (1)$$

で与えられる。式 (1) は、成果指標のリスクが高いほど報酬感応度を下げるべきであるという、直観的かつ実務的に有用なベンチマークである。

もっとも、この標準ベンチマークは、成果が努力に対して線形に反応するという強い仮定に依存している。現実の経済活動では、努力の限界効果が逓減する場面が多い。研究開発では初期努力が大きな知識蓄積をもたらしても、追加的なブレークスルーの難度は高まる。市場開拓では、初期顧客の獲得は相対的に容易でも、シェア拡大が進むにつれて限界顧客獲得コストが上昇する。管理会計上の KPI 改善においても、初期の改善余地が大きい段階を過ぎると、追加努力による改善幅は小さくなりやすい。したがって、収穫一定を前提にした式 (1) をそのまま用いると、成果の限界効果が小さくなっているにもかかわらず、過度に強いインセンティブを与える可能性がある。

本稿は、この問題を扱うため、成果生成関数を

$$y = f(a) + \varepsilon = 1 - e^{-a} + \varepsilon, \quad \varepsilon \sim N(0, \sigma^2) \quad (2)$$

と特定化した P-A モデルを分析する。 $f'(a) = e^{-a} > 0$ 、 $f''(a) = -e^{-a} < 0$ であり、成果は努力に対して厳密に凹である。エイジェントは CARA 効用を持ち、契約は線形契約 $w(y) = \alpha + \beta y$ に限定する。

本稿の位置づけ。 努力に関する成果分布の凹性や一階条件アプローチ（first-order approach; FOA）の妥当性は、すでに Holmström (1979) や Jewitt (1988) などでも一般的に検討されている。本稿はこれらの一般理論を置き換えるものではない。また、ランベルト W 関数が現れること自体を独自の貢献として主張するものでもない。本稿の貢献は次の三点に

整理できる。第一に、線形・正規・CARA という実務的に頻用される枠組みの中に収穫逓減を導入し、標準ベンチマーク式 (1) がどのように修正されるかを明示する。第二に、最適インセンティブ係数 $\beta^*(\sigma)$ の存在・一意性・単調性、および標準モデルとの大小関係を命題として示し、両者の絶対乖離 $\Delta(\sigma) \equiv \beta_{\text{std}} - \beta^*$ が中程度のリスク領域で最大化する一峰性を数値的に確認する。第三に、正規化された代替成果関数族による頑健性確認、および日本のコーポレートガバナンス・コードへの応用的含意を示す。

同時に、本稿の射程は収穫逓減型またはより一般に凹型の成果関数に限定される。研究開発の初期段階や基礎研究のように、一定の努力水準を超えて初めて成果が立ち上がる収穫逓増・閾値型業務では、努力反応や最適契約の構造が本稿のモデルとは異なり得る。そのような非凹型の状況は、本稿の命題を機械的に適用すべき対象ではなく、第 6 節および第 8 節で限界と今後の課題として明示する。

構成は以下の通り。第 2 節で先行研究を整理し、第 3 節でモデルを定式化する。第 4 節は均衡の主要命題を示し、第 5 節は基準ケースと (k, r) パラメータ頑健性の数値分析を行う。第 6 節は正規化された代替凹型関数による頑健性を確認し、あわせて収穫逓増・閾値型成果関数への適用範囲を整理する。第 7 節は日本のコーポレートガバナンス制度への含意を、第 8 節は結論をまとめる。会計品質投資および複数 KPI への拡張は付録 A および付録 B で論じる。

2 先行研究と本稿の貢献

P-A 理論の基本問題は、エイジェントの努力がプリンシパルから直接観測できない状況で、観測可能な成果指標に基づく契約をどのように設計すべきかにある。Mirrlees (1976) はモラルハザード契約問題の基礎を形成し、Holmström (1979) は、成果指標が隠れた努力に関する情報を提供する程度に応じてインセンティブを与えるべきであるという情報原理を提示した。Grossman and Hart (1983) は、より一般的な契約設計問題を定式化した。

本稿に特に関係するのは、一階条件アプローチと成果関数の凹性である。一般のモラルハザード問題では、エイジェントの誘因制約を一階条件で置き換えることが常に正当化されるわけではない。Jewitt (1988) は FOA が妥当となる十分条件を整理した。本稿のモデルでは、成果が可分ノイズ型 $y = f(a) + \varepsilon$ であり、契約が線形で、 f が凹、努力コストが凸であるため、エイジェントの確実性等価は努力 a に関して厳密に凹となる。したがって、本稿内では FOA を補題として直接確認できる。これは本モデルの特定化に依存する性質であり、一般理論上の新規命題ではない。

線形契約の実務的有用性については、Holmström and Milgrom (1987) が代表的である。CARA 効用と正規ノイズを組み合わせると、確実性等価が期待報酬から分散に比例したリスクプレミアムを差し引く形になり、式 (1) のような明快な報酬感応度が得られる。Bolton and Dewatripont (2004) および Laffont and Martimort (2002) もこの枠組みを契約理論

の標準的ベンチマークとして整理している。

会計研究との接点も重要である。成果指標のノイズ σ^2 は、会計情報や KPI の測定誤差、外部環境リスク、成果指標が真の貢献をどの程度反映しているかを表す。Baker (1992) は測定可能な業績指標が真の目的と完全に一致しない場合のインセンティブ歪みを示し、Holmström and Milgrom (1991) は複数タスク環境では強いインセンティブが努力配分を歪め得ることを明らかにした。近年の Ewert and Wagenhofer (2021) は、会計品質投資が経営者の誘因に与える影響を分析している。

本稿は、これらの研究と同じく会計情報の品質や制度設計をインセンティブ問題として捉えるが、焦点は「成果指標がノイズを含む」ことだけでなく、「努力による成果改善が逓減する」ことに置かれる。ここで強調すべき点は、本稿の結果が既存の P-A 理論の主命題を置き換えるものではないという点である。むしろ本稿は、既存理論が与える線形契約ベンチマークを出発点とし、成果関数が線形から凹型へ変わると、報酬感応度の水準だけでなく、努力の限界生産性を通じたリスク負担項も変化することを明示する。したがって、本稿の発見事項の意義は、「関数形を変えれば数値が変わる」という当然の事実にあるのではなく、標準モデルが実務上参照されやすい状況において、収穫逓減型 KPI では標準係数が系統的に過大となり得ることを、命題と数値例で確認した点にある。実務上は、 k, r, σ や成果関数 f を厳密に推定することは難しいため、本稿の数値表は特定企業の報酬係数を直接算定する公式というより、報酬委員会、内部監査、管理会計部門が対象 KPI のリスク、収穫逓減性、閾値性を検討するための参照フレームとして位置づけられる。

3 モデル

3.1 基本設定

プリンシパルはリスク中立、エイジェントはリスク回避的である。時点 0 でプリンシパルが線形報酬契約 $w(y) = \alpha + \beta y$, $\beta \geq 0$ を提示する。時点 1 でエイジェントが努力 $a \geq 0$ を選択する。努力はプリンシパルから観測不能である。時点 2 で成果 y が実現し、契約に基づいて報酬が支払われる。

成果生成過程は式 (2) で与えられる。期待成果は $f(a) = 1 - e^{-a}$ であり、限界成果 $f'(a) = e^{-a}$ は努力とともに逓減する。努力コストは $c(a) = (k/2)a^2$, $k > 0$ であり、エイジェントの効用関数は CARA 型 $u(\omega) = -\exp(-r\omega)$, $r > 0$ である。ここで $\omega = w(y) - c(a)$ である。

3.2 エージェント問題

CARA 効用と正規ノイズの下で、エージェントは確実性等価を最大化する。 $\text{Var}(\beta\varepsilon) = \beta^2\sigma^2$ に基づき、

$$\text{CE}(a; \alpha, \beta) = \alpha + \beta(1 - e^{-a}) - \frac{k}{2}a^2 - \frac{r}{2}\beta^2\sigma^2. \quad (3)$$

最後の項は CARA・正規環境におけるリスクプレミアムである。

補題 1 (エージェント問題とランベルト W 関数). 任意の $\beta \geq 0$ に対して、確実性等価 (3) は努力 a に関して厳密に凹である。 $\beta > 0$ のとき最適努力は一意的に存在し、

$$a(\beta) = W_0\left(\frac{\beta}{k}\right). \quad (4)$$

$\beta = 0$ のとき $a(0) = 0$ である。

Proof. $\partial^2 \text{CE} / \partial a^2 = -\beta e^{-a} - k < 0$ より厳密に凹。一階条件 $\beta e^{-a} - ka = 0$ を $ae^a = \beta/k$ と書き換えると、ランベルト W 関数の定義 $xe^x = z \Leftrightarrow x = W(z)$ より $a = W_0(\beta/k)$ を得る。 $\beta/k \geq 0$ より主枝 W_0 で一意に定まる。□

ランベルト W 関数の数学的性質や数値計算法については Corless et al. (1996) を参照。補題 1 の意義は特殊関数の導入自体ではなく、努力反応 $a(\beta)$ を明示的な単調関数として分離できる点にある。これによりプリンシパルの問題を 1 次元の凹問題として扱える。

3.3 プリンシパルの問題

エージェントの留保確実性等価を 0 に正規化する。プリンシパルは固定給 α で参加制約を等号で満たし、期待利得は

$$\Pi(\beta; \sigma) = 1 - e^{-a(\beta)} - \frac{k}{2}a(\beta)^2 - \frac{r}{2}\beta^2\sigma^2. \quad (5)$$

補題 1 より $\beta = kae^a$ であるため、努力 a を直接の選択変数として

$$\max_{a \geq 0} \tilde{\Pi}(a; \sigma) = 1 - e^{-a} - \frac{k}{2}a^2 - \frac{r}{2}\sigma^2 k^2 a^2 e^{2a}. \quad (6)$$

4 均衡分析

4.1 最適解の存在と一意性

式 (6) を a で微分すると、

$$H(a; \sigma) \equiv \frac{\partial \tilde{\Pi}}{\partial a} = e^{-a} - ka - r\sigma^2 k^2 a(1+a)e^{2a}, \quad (7)$$

であり、内点解は

$$e^{-a} - ka = r\sigma^2 k^2 a(1+a)e^{2a} \quad (8)$$

を満たす。

命題 1 (最適解の存在・一意性). 任意の $\sigma \geq 0$ について、式 (6) の最大化問題は一意な解 $a^*(\sigma)$ を持つ。対応する最適インセンティブ係数は $\beta^*(\sigma) = ka^*(\sigma)e^{a^*(\sigma)}$ である。 $\sigma = 0$ のとき $\beta^*(0) = 1$ 、 $\sigma > 0$ のとき $0 < \beta^*(\sigma) < 1$ である。

Proof. $\partial H / \partial a = -e^{-a} - k - r\sigma^2 k^2 e^{2a}(1+4a+2a^2) < 0$ より H は a について厳密に減少する。 $H(0; \sigma) = 1$ かつ $\lim_{a \rightarrow \infty} H(a; \sigma) = -\infty$ より $H(a; \sigma) = 0$ を満たす $a^*(\sigma)$ が一意に存在し、これが (6) の大域的最大点である。 $\sigma = 0$ では $e^{-a} = ka$ であり $\beta = kae^a$ を代入して $\beta^*(0) = 1$ 。 $\sigma > 0$ では (8) の右辺が正であるから $e^{-a} > ka$ 、したがって $\beta^* = kae^a < 1$ 。 $\square$

4.2 リスクに対する単調性

命題 2 (リスクに対する単調減少). $\sigma > 0$ において、 $a^*(\sigma)$ および $\beta^*(\sigma)$ は σ の厳密な減少関数である。さらに $\sigma \rightarrow \infty$ のとき $\beta^*(\sigma) \rightarrow 0$ である。

Proof. 陰関数定理より $da^*/d\sigma = -H_\sigma/H_a$ である。 $H_\sigma = -2r\sigma k^2 a^*(1+a^*)e^{2a^*} < 0$ 、 $H_a < 0$ より $da^*/d\sigma < 0$ 。 $d\beta/da = ke^a(1+a) > 0$ より $d\beta^*/d\sigma < 0$ 。 $\sigma \rightarrow \infty$ で a^* が正の下限を持つと (8) の右辺が発散して矛盾するため $a^* \rightarrow 0$ 、したがって $\beta^* \rightarrow 0$ 。 $\square$

4.3 標準モデルとの大小関係

命題 3 (標準モデルとの比較). 任意の $\sigma > 0$ について、

$$\beta^*(\sigma) < \beta_{\text{std}}(\sigma) = \frac{1}{1 + rk\sigma^2}. \quad (9)$$

Proof. $\sigma > 0$ では命題 1 より $0 < \beta^* < 1$ 。(1) より $\beta^* < \beta_{\text{std}}$ は $(1 - \beta^*)/\beta^* > rk\sigma^2$ と同値。(8) を $\beta^* = kae^a$ を用いて整理すると

$$rk\sigma^2 = \frac{k^2 a^2 (1 - \beta^*)}{(\beta^*)^3 (1 + a)}, \quad (10)$$

したがって主張は

$$(\beta^*)^2 (1 + a) > k^2 a^2 \iff e^{2a} (1 + a) > 1 \quad (11)$$

に帰着する (最後は $\beta^* = kae^a$ より)。 $a = 0$ で等号、 $a > 0$ で $e^{2a} > 1$ かつ $1 + a > 1$ により狭義不等号 $e^{2a} (1 + a) > 1$ が成立。 $\sigma > 0$ では $a^* > 0$ であるため結論が従う。 $\square$

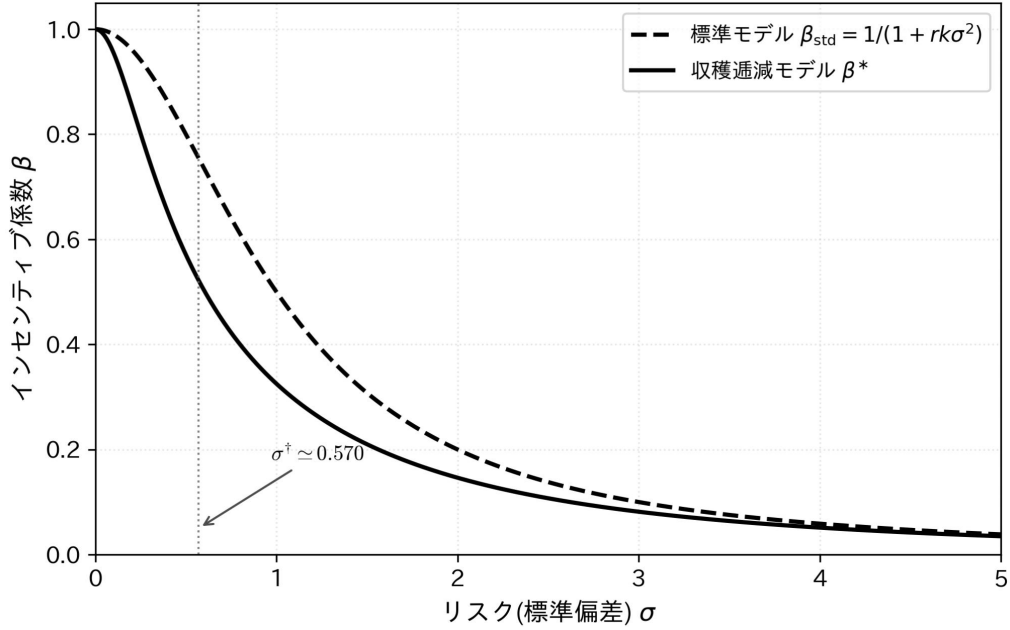


図1 最適インセンティブ係数の比較 ($k = r = 1$)。実線は収穫逓減モデルの β^* ，破線は標準モデルの β_{std} 。 $\sigma^\dagger \simeq 0.570$ は両者の絶対乖離 $\Delta(\sigma) = \beta_{\text{std}} - \beta^*$ が最大となる点。

5 数値分析と (k, r) パラメータ頑健性

本節では式 (8) を数値的に解き，最適インセンティブ係数の挙動を検討する。すべての数値および図表は同梱の Python スクリプトから生成しており，本文・表・図の整合性を保証している。

5.1 基準ケース ($k = r = 1$) の挙動

図1は，収穫逓減モデルの $\beta^*(\sigma)$ と標準モデルの $\beta_{\text{std}}(\sigma) = 1/(1 + \sigma^2)$ を比較したものである。命題2-3と整合的に， β^* は σ に関して単調減少し，任意の $\sigma > 0$ で標準モデルを下回る。

5.2 乖離の一峰性

標準モデルとの差を $\Delta(\sigma) = \beta_{\text{std}}(\sigma) - \beta^*(\sigma)$ と定義する。命題3より $\Delta(\sigma) > 0$ である。 $\Delta(0) = 0$ かつ高リスク極限では両モデルとも0に近づくため $\Delta(\sigma) \rightarrow 0$ である。したがって乖離幅は中間のリスク水準で最大となる可能性がある。

観察1 (絶対乖離の一峰性 (基準ケース))。 $k = r = 1$ における数値計算では $\Delta(\sigma)$ は一峰

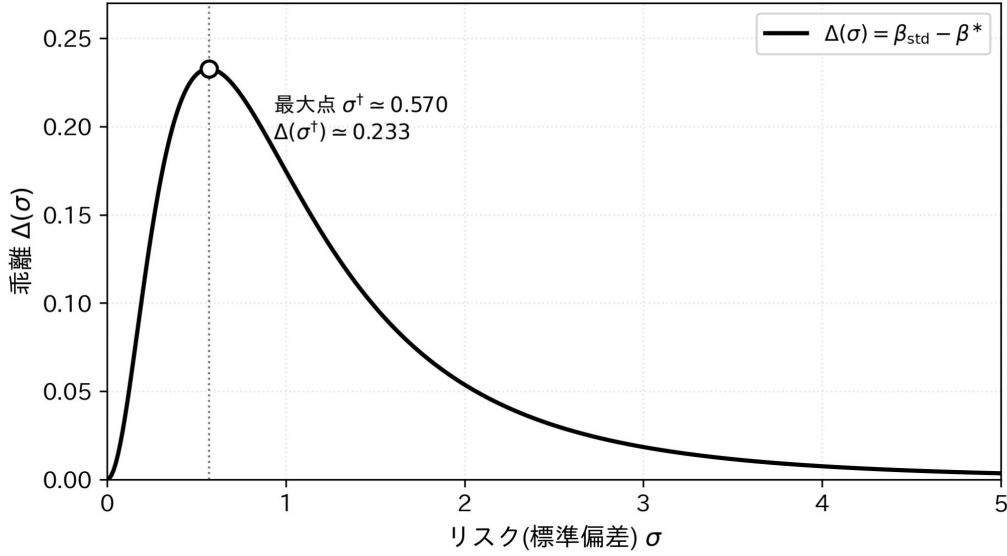


図2 標準モデルとの絶対乖離 $\Delta(\sigma) = \beta_{\text{std}} - \beta^*$ ($k = r = 1$)。乖離は中程度のリスク領域で最大化する。

性を示し、最大点と最大値は

$$\sigma^\dagger \simeq 0.5700, \quad \Delta(\sigma^\dagger) \simeq 0.2327. \quad (12)$$

このとき $\beta_{\text{std}}(\sigma^\dagger) \simeq 0.7548$, $\beta^*(\sigma^\dagger) \simeq 0.5220$ であり、下方修正率は約 30.83% である。

絶対乖離と相対修正率の区別。本稿で「下方修正率」と呼んでいる量は $(\beta_{\text{std}} - \beta^*)/\beta_{\text{std}}$ である。観察 1 は Δ (絶対乖離) の最大点を扱っているが、相対比率自体の最大点とは異なる点に注意が必要である。 $k = r = 1$ における数値計算では、相対修正率が最大化する σ は約 0.957 であり、その値は約 35.0% に達する。報酬調整の絶対額が問題となる場面では $\sigma^\dagger$ を、設計時の相対的な修正幅が問題となる場面では相対修正率の最大点を参照することが望ましい。

図 2 は $\Delta(\sigma)$ を示し、表 1 は代表的な数値を与える。

5.3 (k, r) パラメータ頑健性

最大乖離が (k, r) にどう依存するかは、本稿の中心的主張の頑健性を判断するうえで重要である。 $(k, r) \in \{0.5, 1, 2, 5\}^2$ の格子上で式 (8) を数値的に解いた結果を表 2 に示す。

表 2 から以下の三点が読み取れる。

1. **一峰性は (k, r) にわたって頑健に成立する。**試験した全 16 パラメータ点で $\Delta(\sigma)$ は一峰性を示し、乖離が中リスク領域で最大化するという定性的特徴は頑健である。
2. **リスク回避度 r は σ 軸をリスケールする。**行を固定して r を変えると $\sigma^\dagger$ は変化するが、 $\Delta(\sigma^\dagger)$ および下方修正率は不変である。これは数学的に、標準モデルの分母が

表 1 代表的な数値結果 ($k = r = 1$)。下方修正率は $(\beta_{\text{std}} - \beta^*)/\beta_{\text{std}}$ と定義する。

σ	$\beta^*(\sigma)$	$\beta_{\text{std}}(\sigma)$	$\Delta(\sigma)$	下方修正率
0.000	1.0000	1.0000	0.0000	0.00%
0.100	0.9554	0.9901	0.0347	3.50%
0.200	0.8557	0.9615	0.1058	11.01%
0.300	0.7479	0.9174	0.1695	18.48%
0.400	0.6517	0.8621	0.2104	24.41%
0.500	0.5706	0.8000	0.2294	28.68%
0.570	0.5220	0.7548	0.2327	30.83%
0.700	0.4468	0.6711	0.2243	33.42%
0.957	0.3392	0.5220	0.1827	35.01%
1.500	0.2100	0.3077	0.0977	31.77%
2.000	0.1462	0.2000	0.0538	26.88%
5.000	0.0349	0.0385	0.0036	9.24%

太字は Δ の最大点 ($\sigma^\dagger \simeq 0.570$) と相対修正率の最大点 ($\sigma \simeq 0.957$) を示す。

表 2 (k, r) 格子上での最大乖離分析。各セルは $\sigma^\dagger / \Delta(\sigma^\dagger) /$ 下方修正率 を示す (Δ の最大化基準)。

$k \backslash r$	0.5	1	2	5
0.5	1.028 / 0.333 / 42.1%	0.727 / 0.333 / 42.1%	0.514 / 0.333 / 42.1%	0.325 / 0.333 / 42.1%
1.0	0.806 / 0.233 / 30.8%	0.570 / 0.233 / 30.8%	0.403 / 0.233 / 30.8%	0.255 / 0.233 / 30.8%
2.0	0.618 / 0.149 / 20.6%	0.437 / 0.149 / 20.6%	0.309 / 0.149 / 20.6%	0.195 / 0.149 / 20.6%
5.0	0.419 / 0.074 / 10.6%	0.296 / 0.074 / 10.6%	0.210 / 0.074 / 10.6%	0.133 / 0.074 / 10.6%

$rk\sigma^2$ という積で現れるため、 r の変化を $\sigma \mapsto \sigma\sqrt{r}$ という変数変換で吸収できることに対応する。実際、 (k, r, σ) における (8) の解は、 $(k, 1, \sigma\sqrt{r})$ における解と一致する。

3. 努力コスト係数 k は最大乖離の大きさを変える。列を固定して k を変えると、 $\Delta(\sigma^\dagger)$ および下方修正率は明確に変化し、 k が小さいほど乖離は大きくなる： $k = 0.5$ では下方修正率約 42.1%， $k = 5$ では約 10.6% である。

経済的解釈は次の通りである。 k が小さい（努力が相対的に安価な）環境では、最適努力 a^* は大きくなり、エイジェントは収穫逡減関数 $f(a) = 1 - e^{-a}$ の高努力領域へ移動する。この領域では $f'(a) = e^{-a}$ が小さく、すなわち追加努力の限界生産性が大幅に低下している。このとき標準モデル（線形成果）は限界生産性 1 を仮定しているため、過大なインセンティブを推奨する度合いも大きい。逆に k が大きい（努力が高価な）環境では a^* は小さく、 $f'(a) \simeq 1$ に近いため、標準モデルとの乖離も小さい。したがって本稿の含意は、「収穫逡減

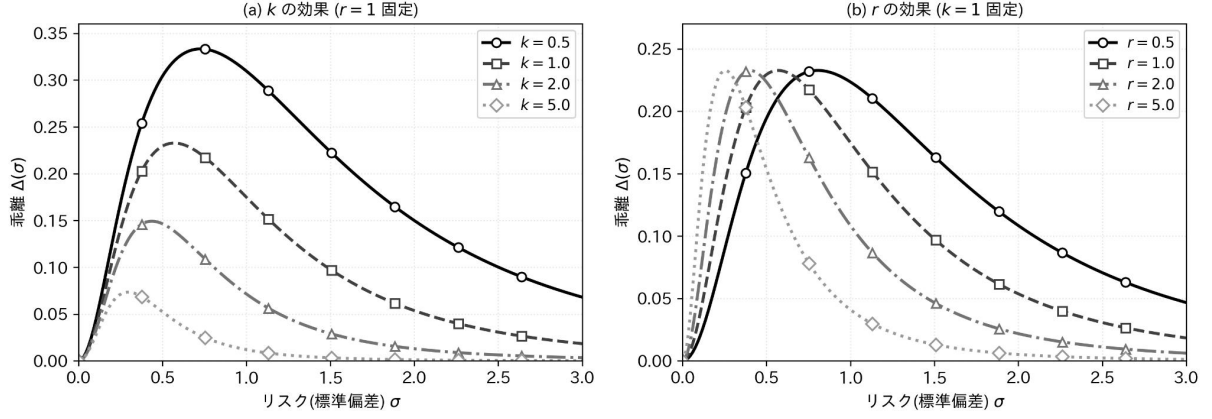


図3 (k, r) 別の $\Delta(\sigma)$ 形状。(a) $r = 1$ 固定で k を変化 (最大値が変化)。(b) $k = 1$ 固定で r を変化 (σ 軸のリスケールのみ)。

性が顕在化する高努力領域への参入が容易な業務」(例：相対的に努力コストが低く成果に近づける業務)でこそ強く現れる。

注. 表 2 の r 軸のリスケール構造は、無次元化 $\xi \equiv rk\sigma^2$ で書き換えれば、標準モデルが $\beta_{\text{std}} = 1/(1+\xi)$ となり ξ のみの関数になることに対応する。一方、収穫逓減モデルの FOC は $e^{-a}/k - a = \xi \cdot a(1+a)e^{2a}$ と整理でき、 k と ξ が独立に作用する。このため r 単独の変化は ξ を通じてのみ作用するが、 k は ξ に加えて係数 $-a/k$ を介して独立に作用する。

5.4 (k, r) 別の Δ 形状

図 3 は左パネルで $r = 1$ 固定、 k を変化させた場合の $\Delta(\sigma)$ 、右パネルで $k = 1$ 固定、 r を変化させた場合の $\Delta(\sigma)$ を示す。左パネルでは曲線の最大値が k により変化し、右パネルでは形状は同じで σ 軸方向に縮尺が変わるのみであることが視覚的に確認できる。

6 代替的成果関数による頑健性と適用範囲

本稿の特定化 $f(a) = 1 - e^{-a}$ は、 $\lim_{a \rightarrow \infty} f(a) = 1$ で飽和する関数であり、 $f'(0) = 1$ である。後者は標準モデルの線形成果 $f(a) = a$ における $f'(0) = 1$ と一致しており、初期限界生産性において両モデルは公平に比較されている。本節では、初期限界生産性 $f'(0) = 1$ を保ちながら凹性の強さを変える正規化された関数族で頑健性を検証する。あわせて、査読者の指摘を踏まえ、収穫逓増・閾値型成果関数に対する本稿の適用範囲を整理する。

一般に成果関数を $y = f(a) + \varepsilon$ とし、 $f'(a) > 0$ を仮定すると、線形契約 $w = \alpha + \beta y$ のもとでエイジェントの一階条件は

$$\beta f'(a) = ka \quad (13)$$

表 3 正規化指数族 $f_\theta(a) = (1 - e^{-\theta a})/\theta$ における乖離最大点 ($k = r = 1$)。

θ	$\sigma^\dagger$	$\beta_{\text{std}}(\sigma^\dagger)$	$\beta_\theta^*(\sigma^\dagger)$	下方修正率
0.5	0.6177	0.7238	0.5745	20.62%
1.0 (中心モデル)	0.5700	0.7548	0.5220	30.83%
2.0	0.5140	0.7910	0.4577	42.13%
3.0	0.4795	0.8131	0.4169	48.72%

である。したがって、努力水準 a を実装するために必要なインセンティブ係数は

$$\beta_f(a) = \frac{ka}{f'(a)} \quad (14)$$

と表される。参加制約を等号で満たすとき、プリンシパルの誘導目的関数は

$$\Pi_f(a; \sigma) = f(a) - \frac{k}{2}a^2 - \frac{r\sigma^2}{2} \left\{ \frac{ka}{f'(a)} \right\}^2 \quad (15)$$

となる。式 (15) が示すように、成果関数の変更は、期待成果 $f(a)$ そのものだけでなく、 $a/f'(a)$ を通じてリスク負担項にも影響する。本節の数値分析は、この一般式のうち、凹型かつ初期限界生産性を標準化した関数族で、標準モデルに対する下方修正がどの程度頑健に観察されるかを確認するものである。非凹型の関数では、式 (13) が大域的な誘因両立性を保証しない場合があるため、別途注意が必要である。

6.1 正規化された指数族

成果関数族

$$f_\theta(a) = \frac{1 - e^{-\theta a}}{\theta}, \quad \theta > 0 \quad (16)$$

を考える。 $f'_\theta(0) = 1$ で初期限界生産性は線形モデルと一致し、 $f'_\theta(a) = e^{-\theta a}$ より凹性の強さは θ で制御される。 $\theta \rightarrow 0$ の極限で $f_\theta(a) \rightarrow a$ となり、標準線形モデルに連続接続する。 $\theta = 1$ が本稿の中心モデル (2) に対応する。

エイジェントの一階条件は $\beta e^{-\theta a} = ka$ より $\beta = ka e^{\theta a}$ であり、プリンシパルの一階条件は

$$e^{-\theta a} - ka = rk^2\sigma^2 \cdot a(1 + \theta a)e^{2\theta a}. \quad (17)$$

表 3 は $k = r = 1$, $\theta \in \{0.5, 1, 2, 3\}$ における結果を示す。

図 4 は $\beta_\theta^*(\sigma)$ を θ 別に示し、 $\theta \rightarrow 0$ で標準モデル (破線) に近づくことを視覚的に確認できる。

数値検証として $\theta = 0.1, 0.05, 0.01$ における最大乖離を計算したところ、それぞれ $\Delta_{\max} = 0.040, 0.021, 0.004$ であり、 $\theta \rightarrow 0$ で標準モデル ($\Delta = 0$) に連続的に収束することが確認できた。

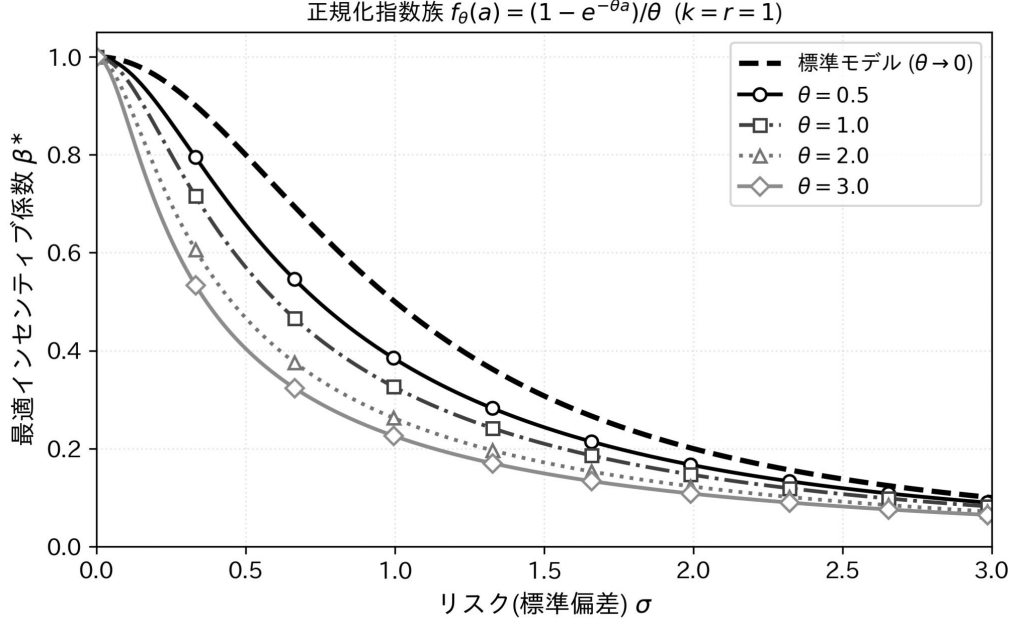


図4 正規化指数族における $\beta_\theta^*(\sigma)$ ($k = r = 1$)。 θ が大きいほど凹性が強く，標準モデル（破線）からの乖離が大きい。 $\theta \rightarrow 0$ で標準モデルに収束する。

表4 正規化対数族 $f_\theta(a) = \log(1 + \theta a)/\theta$ における乖離最大点 ($k = r = 1$)。

θ	$\sigma^\dagger$	$\beta_{\text{std}}(\sigma^\dagger)$	$\beta_\theta^*(\sigma^\dagger)$	下方修正率
0.5	0.6435	0.7071	0.5723	19.07%
1.0	0.6091	0.7294	0.5257	27.93%
2.0	0.5676	0.7563	0.4715	37.65%

6.2 対数族による補足確認

正規化対数族 $f_\theta(a) = \log(1 + \theta a)/\theta$ ($f'_\theta(0) = 1$, $\theta \rightarrow 0$ で線形モデル) でも数値計算を行った。表4に結果を示す。指数族の中心モデル（中程度の凹性で約30.8%）と同オーダーの下方修正率が観察され，定性的結論は頑健である。

小括。 初期限界生産性 $f'(0) = 1$ を揃えた正規化族で検証した範囲では，「中リスク領域における下方修正余地」という本稿の中心メッセージは関数の特定化に依存しない。下方修正率の具体的な大きさは凹性の強さ θ により変化するが（指数族で20–50%弱，対数族で19–38%），凹型成果関数の範囲では定性的特徴は頑健である。

6.3 収穫逦増・閾値型成果関数への適用範囲

もっとも、本稿の分析はすべての非線形成果関数に対して同じ結論を与えるものではない。とりわけ、研究開発や基礎研究のように、一定の努力水準を超えて初めて成果が立ち上がる閾値型業務では、成果関数に収穫逦増的な領域が含まれる可能性がある。この場合、エイジェントの目的関数 $\beta f(a) - ka^2/2$ は大域的に凹であるとは限らず、一階条件は十分条件ではなくなる。また、努力反応は連続的・単調的とは限らず、複数の局所解やジャンプを伴う可能性がある。したがって、命題 3 で示した標準モデルに対する下方修正結果は、収穫逦減型または凹型成果関数のベンチマークに限定して解釈すべきである。

収穫逦増・閾値型業務を表す単純な例として、次のような正規化関数を考えることができる。

$$f_{IR}(a) = a + \frac{\lambda}{\eta} \log\{1 + \exp[\eta(a - \bar{a})]\} - \frac{\lambda}{\eta} \log\{1 + \exp[-\eta\bar{a}]\}, \quad \lambda > 0, \eta > 0. \quad (18)$$

この関数は $f_{IR}(0) = 0$ となるよう正規化され、 $a = \bar{a}$ 付近で限界生産性が立ち上がる。こうした非凹型の成果関数では、低努力領域と高努力領域で誘因の役割が異なり、標準係数からの単純な下方修正ではなく、段階的評価、マイルストーン型報酬、あるいは非線形契約を含む別の契約設計が必要となり得る。

この区別は実務的にも重要である。成熟事業の改善活動や定型的 KPI 改善のように成果が早期に飽和する業務では、本稿の収穫逦減型ベンチマークが過大インセンティブの検討に役立つ。一方、基礎研究、新規事業開発、プラットフォーム立ち上げのように成果が閾値を超えて初めて顕在化する業務では、本稿の下方修正命題を機械的に適用することは避けるべきである。むしろ本稿の実務的含意は、対象業務が収穫逦減型か、それとも閾値型かを識別する必要性を明示する点にもある。

本節のまとめ。凹型成果関数の範囲では、標準モデルに対する下方修正という中心メッセージは頑健である。しかし、収穫逦増・閾値型成果関数では、努力反応や最適契約の構造が根本的に異なり得るため、本稿の結果は直接適用できない。この点を明示することで、本稿の結論は、収穫逦減型 KPI に対する線形契約ベンチマークとして位置づけられる。

7 ガバナンス上の含意：日本のコーポレートガバナンス・コードとの接続

本節では本稿の理論的結果を日本の実務環境に接続する。特に、金融庁および東京証券取引所が 2015 年に策定（2018 年・2021 年改訂）した株式会社東京証券取引所（2021）におけるコーポレートガバナンス・コード（以下 CG コード）の関連原則との関係を論じる^{*1}。

^{*1} なお、本稿執筆時点（2026 年 4 月末）では、2026 年 4 月 10 日に金融庁および東京証券取引所がコーポレートガバナンス・コード改訂案を公表し、同日より同年 5 月 15 日までパブリックコメント手続中である。

7.1 関連する CG コード原則の整理

経営者報酬および業績連動報酬に関しては、まず原則 4-2 が、経営陣の報酬について中長期的な会社の業績や潜在的リスクを反映させ、健全な企業家精神の発揮に資するインセンティブ付けを行うべきことを規定する。これを敷衍した補充原則 4-2 ① は、経営陣の報酬が持続的成長に向けた健全なインセンティブとして機能するよう、客観性・透明性ある手続に従い、報酬制度を設計し、具体的な報酬額を決定すべきこと、さらに中長期業績連動報酬と現金報酬・自社株報酬の割合を適切に設定すべきことを定めている。さらに補充原則 4-2 ② は人的資本・知的財産への投資、事業ポートフォリオに関する取締役会の監督に触れ、補充原則 4-10 ① は独立した諮問委員会としての報酬委員会の設置・活用を求めている。

業績連動報酬の「強さ」と「設計根拠」については、2021 年改訂で説明責任が一段と重視された。本稿の結果は、この説明責任に対して具体的な数理的フレームワークを提供する。

7.2 中リスク・収穫逡減型業務における過大インセンティブの可能性

第 5 節および第 6 節の結果は、標準的な線形ベンチマーク式 (1) に従うと、中リスク・収穫逡減型業務で下方修正余地が生じることを示している。本稿の数値例では、下方修正率は設定に応じておおむね 10% 台から 50% 弱まで変化する。中心的な設定 ($k = r = \theta = 1$) および中程度の凹性では、20–40% 程度の下方修正が観察される。日本企業の実務に照らすと、このような業務領域には以下が含まれる。

- デジタル・トランスフォーメーション (DX) の全社展開フェーズ
- R&D の応用研究フェーズ
- 中規模 M&A 後のシナジー実現
- 成熟市場での追加シェア獲得

これらの業務に対して、標準モデルに基づいて高い業績連動係数を設定すると、理論的には経営者の合理的な努力配分から見て過大であり得る。

7.3 業務類型の識別と応用可能性

本稿の実務的応用可能性は、特定企業の最適報酬係数を機械的に算出することよりも、対象業務の成果関数の形状を報酬設計上の論点として可視化する点にある。報酬委員会や内部監査部門は、KPI のノイズ水準だけでなく、努力に対する成果の反応が収穫逡減的か、閾値

る (株式会社東京証券取引所, 2026)。今回の改訂は「プリンシプル化」「スリム化」を方針としており、補充原則の多くが原則または解釈指針へ再配置される予定である。本稿では現行の 2021 年改訂コードを基準として議論する。改訂後のコードでは原則の付番が変更される可能性があるが、本稿の論旨自体には影響しない。

型か、あるいは複数段階の性質を持つかを検討する必要がある。

収穫逓減型業務では、努力の限界効果が低下するため、標準的な線形成果モデルに基づく高い報酬感応度は過大になりやすい。これに対して、収穫逓増・閾値型業務では、努力が一定水準に達するまで成果が現れにくい可能性があり、単純な感応度の引き下げではなく、研究開発マイルストーン、段階的な評価指標、中長期的な非財務 KPI との組合せが重要になる。したがって、本稿の枠組みは、収穫逓減型業務に対する係数調整の参照フレームであると同時に、収穫逓増型業務を別の契約設計問題として切り分けるための診断枠組みでもある。

7.4 報酬委員会への実務的提案

本稿の結果は、補充原則 4-2 ① が要請する適切な報酬制度設計に対して、以下の三点の具体化を提供する。

1. **業務特性の収穫逓減度評価。**対象業務が「初期段階」「成長段階」「成熟段階」のいずれにあるかを定期的に評価し、報酬感応度の調整を検討する。これは原則 4-2 の「中長期的な業績」を反映するインセンティブ付けという要請に対し、業務サイクルに応じた動学的な調整の根拠となる。
2. **中リスク領域での慎重な係数設計。**成果指標のノイズが極端に大きい場合や極端に小さい場合は標準モデルとの違いは小さい。問題は中間領域である。本稿は、この中間領域こそ報酬委員会が独立した検討を要する領域であることを示している。
3. **設計根拠の説明枠組み。**本稿の枠組みは、「成果指標のノイズ」と「成果関数の凹性」の両方を考慮した係数設計の説明枠組みを提供する。表 2 のように業務の特性パラメータ（努力コスト水準、リスク水準）から下方修正率を導出する手順は、報酬委員会の議論に対する一つの参照フレームとなり得る。

もっとも、これらのパラメータ (k, r, σ) を実務上直接推定することは容易ではない。したがって本稿の数値表は、厳密な算定式というより、業務特性に応じた報酬感応度調整の方向性を示す参照フレームとして位置づけられる。具体的な水準を決定する際には、報酬委員会が同業他社のベンチマーク水準、過去の達成度分布、業務の段階性（初期・成長・成熟）といった追加情報と組み合わせて判断することが現実的である。

7.5 内部監査への含意

内部監査の観点からは、標準モデルに近い高い業績連動係数が設定されている中リスク・収穫逓減型プロジェクトは、過大インセンティブによる行動歪み（短期成果偏重、利益調整、過度なリスクテイク、長期投資の抑制）の監査対象となり得る。監査資源配分においては、「中リスクでありながら報酬感応度が高い案件」を重点領域として把握することが望ましい。

役員報酬の個別開示については、1 億円以上の役員報酬に関する個別開示は 2010 年（平成 22 年）の開示府令改正で導入されており、2022 年プライム市場移行と直接結び付くものではない。むしろプライム市場移行や 2021 年 CG コード改訂で強化されたのは、独立社外取締役比率や報酬決定方針の説明責任である。本稿の枠組みは、この説明責任の論拠の一つとなり得る。

8 結論

本稿は、成果が努力に対して収穫逓減する状況を CARA 効用、正規ノイズ、線形報酬契約という標準的な P-A モデルの枠組みに組み込み、最適インセンティブ係数の性質を分析した。成果関数を $y = 1 - e^{-a} + \varepsilon$ とした場合、エイジェントの努力反応はランベルト W 関数により $a(\beta) = W_0(\beta/k)$ と表される。

主要結果は以下の通りである。第一に、最適インセンティブ係数 $\beta^*(\sigma)$ は一意に存在し、成果ノイズの標準偏差 σ に対して単調に減少する（命題 1, 2）。第二に、標準モデルの係数 $\beta_{\text{std}} = 1/(1 + rk\sigma^2)$ と比較すると、収穫逓減モデルの係数は任意の $\sigma > 0$ で小さい（命題 3）。第三に、両者の絶対乖離 $\Delta(\sigma)$ は中程度のリスク領域で最大化する一峰性を示す。基準ケース $k = r = 1$ では Δ の最大点は $\sigma^\dagger \simeq 0.570$ で下方修正率約 30.8% であり（観察 1）、この一峰性自体は (k, r) パラメータの広い範囲で頑健に成立する。なお相対修正率の最大点は $\sigma \simeq 0.957$ 付近で約 35.0% となり、参照すべき最適化基準により「中リスク領域」の具体的位置は若干異なる。第四に、最大乖離の大きさは努力コスト係数 k に依存し、 k が小さい（努力が安価な）環境ほど標準モデルとの乖離は大きくなる。第五に、初期限界生産性を揃えた正規化された凹型関数族 $f_\theta(a) = (1 - e^{-\theta a})/\theta$ および $\log(1 + \theta a)/\theta$ で頑健性を確認した。

本稿の貢献は、FOA や凹性条件に関する一般理論の新規性ではなく、収穫逓減を含む線形契約ベンチマークを、命題・観察・図表によって実務的に利用しやすい形で示した点にある。また、代替的な凹型成果関数および一般成果関数の式を用いることで、本稿の結論がどの範囲で頑健であり、どの範囲では慎重な解釈を要するかを明示した。日本のコーポレートガバナンス・コード（原則 4-2 および補充原則 4-2 ①）が求める業績連動報酬の客観性・透明性ある設計に対し、本稿は「成果指標のノイズ」と「成果関数の凹性」を統合した数理的枠組みを提供する。本稿の数値例では下方修正率は設定に応じて 10% 台から 50% 弱まで変化し、中心的な設定（中程度の凹性、努力コスト・リスク回避度が中位）では 20–40% 程度 of 下方修正が観察された。実務的には努力コストが相対的に低い高努力業務でこの調整が顕著になる。

本稿には限界もある。第一に、契約を線形契約に限定しており、非線形契約全体の中での最適性は扱っていない。第二に、単一エイジェント・単一期間モデルであり、動学的誘因、マルチタスク環境における努力配分は十分に扱っていない。第三に、付録で扱う会計品質投資・複数 KPI への拡張は研究ノートの議論に留めており、本格的な実証検証は今後の課

題である。第四に、本稿の中心命題は収穫逓減型または凹型の成果関数を前提としており、研究開発、基礎研究、新規事業立ち上げのような収穫逓増・閾値型業務には一般には直接適用できない。このような非凹型環境では、複数均衡、不連続な努力反応、非線形契約、動学的マイルストーン契約が重要となる可能性がある。今後の研究では、日本企業の役員報酬データを用いて、業績連動報酬比率と業務の収穫逓減度との関係を定量的に評価することに加え、収穫逓減型業務と閾値型業務を識別したうえで、両者に異なる報酬設計が採用されているかを検証することが期待される。

付録 A 会計品質投資との相互作用

成果指標のノイズ分散 σ^2 は、外部環境ショックや本質的不確実性に由来する成分と、会計情報品質に依存する測定誤差成分に分解できる。本付録では簡略な分析を通じて、収穫逓減環境下での会計品質投資の役割を論じる。

A.1 ノイズの分解と統合問題

ノイズを

$$\sigma^2(q) = \sigma_0^2 + s(q) \quad (19)$$

と分解する。ここで $\sigma_0^2 \geq 0$ は不可避なノイズ、 $s(q) \geq 0$ は会計品質投資 $q \geq 0$ により削減可能な部分で、 $s'(q) < 0$ 、 $s''(q) > 0$ 、 $s(0) = \bar{s}$ 、 $\lim_{q \rightarrow \infty} s(q) = 0$ とする。投資コスト $C(q)$ は $C'(q) > 0$ 、 $C''(q) \geq 0$ とする。プリンシパルは (β, q) を同時に選択し、

$$\max_{\beta \geq 0, q \geq 0} 1 - e^{-a(\beta)} - \frac{k}{2} a(\beta)^2 - \frac{r}{2} \beta^2 [\sigma_0^2 + s(q)] - C(q). \quad (20)$$

q に関する一階条件は

$$-\frac{r}{2} \beta^2 s'(q) = C'(q). \quad (21)$$

左辺は $-s'(q) > 0$ であるためリスクプレミアム削減の限界便益、右辺は限界費用である。

A.2 比較静学

(21) を β について陰関数微分すると、

$$\frac{\partial q^*}{\partial \beta} = \frac{-r\beta s'(q)}{(r/2)\beta^2 s''(q) + C''(q)} > 0. \quad (22)$$

分子は $s'(q) < 0$ より正、分母は二階条件で正であるため $\partial q^* / \partial \beta > 0$ である。すなわち、より強いインセンティブ β が設定されるほど、会計品質投資 q^* も大きくなる。直観的には、 β が大きいほどリスクプレミアム $(r/2)\beta^2 \sigma^2$ が β^2 に比例して増大し、ノイズ削減の限界価値が高まるためである。

A.3 傾向としての含意

収穫逦減モデルでは標準モデルより β^* が小さい (命題 3)。同一のノイズ構造 $s(q)$ および同一の投資コスト $C(q)$ の下では, (22) の関係を通じて, 会計品質投資の最適水準 q_{DR}^* も標準モデルのそれよりも小さくなる傾向がある。ただし (β, q) は同時決定であり, $\sigma^2(q)$ を通じて両者は内生的に結び付くため, これは「他条件一定」のもとでの局所的傾向を示す比較静学であり, 独立した命題として強く主張するには追加の正則条件が必要である。

実務的含意としては, β^* の調整と q^* の調整は互いに補完的に動くため, 収穫逦減環境では両者を同時に下方修正する整合的設計が望ましい。標準モデルを盲目的に適用し β を高く設定しつつ会計品質投資を行うと, 過大インセンティブと過大投資の二重歪みが生じうる。これは Ewert and Wagenhofer (2021) が分析した会計品質投資の役割を, 収穫逦減環境において補完的視点から再評価する材料となる。

付録 B 複数 KPI ポートフォリオへの拡張

実務では複数 KPI を組み合わせた業績評価が一般的である。本付録では本稿のモデルを複数 KPI 環境に拡張した際の基本構造を示す。

B.1 独立 KPI モデル

m 個の KPI を考え,

$$y_j = f_j(a_j) + \varepsilon_j, \quad \varepsilon_j \sim N(0, \sigma_j^2), \quad j = 1, \dots, m \quad (23)$$

で生成されたとする。 f_j は厳密に凹で増加し, $\{\varepsilon_j\}$ は互いに独立とする。エイジェントは KPI 別の努力 $(a_1, \dots, a_m)$ を選択し, 努力コストは可分 $c = \sum_j (k_j/2) a_j^2$ とする。線形契約は $w = \alpha + \sum_j \beta_j y_j$ である。確実性等価は

$$CE = \alpha + \sum_j \beta_j f_j(a_j) - \sum_j \frac{k_j}{2} a_j^2 - \frac{r}{2} \sum_j \beta_j^2 \sigma_j^2. \quad (24)$$

エイジェントの一階条件 $\beta_j f'_j(a_j) - k_j a_j = 0$ は KPI 別に分離され, $f_j(a) = 1 - e^{-a}$ かつ $k_j = k$ の場合 $a_j(\beta_j) = W_0(\beta_j/k)$ となる。プリンシパルの問題も KPI 別に分解され, 各 β_j^* は単一 KPI モデルにおける (k_j, σ_j) に対応する最適係数である。したがって, 各 KPI に対して命題 1–3 がそのまま適用される。

B.2 相関ノイズの効果

$\varepsilon \sim N(0, \Sigma)$, $\Sigma_{ij} = \rho_{ij} \sigma_i \sigma_j$ の場合, リスクプレミアムは $(r/2) \beta^\top \Sigma \beta$ となる。プリンシパル最適解は分解できず, 多変量最適化となる。包絡定理を用いてエイジェントの誘因制

約 $\beta_j f'_j(a_j) = k_j a_j$ を考慮した一階条件は

$$(1 - \beta_j) f'_j(a_j) \frac{da_j}{d\beta_j} - r(\Sigma\beta)_j = 0, \quad j = 1, \dots, m \quad (25)$$

である。(25) の左辺第一項は、 β_j を増やしたときの直接的な期待成果増加と固定報酬調整の差、第二項はリスクプレミアム増加（相関を含む）である。負の相関 $\rho_{ij} < 0$ は $\Sigma\beta$ の成分を縮小し、より強いインセンティブを正当化する。正の相関は逆である。

B.3 設計上の含意

KPI ポートフォリオ設計では、(i) 各 KPI の収穫逨減性に応じた個別下方修正、(ii) ノイズ相関構造の活用（負相関 KPI の組合せでリスクを部分相殺）、(iii) 努力コストが非可分な場合の歪み（測定容易な KPI への努力偏重；Holmström and Milgrom, 1991）が同時に考慮される。本稿の枠組みではこれらの最初の点を中心に論じたが、後二者の本格的分析は今後の課題である。

参考文献

- Baker, George P. (1992), “Incentive Contracts and Performance Measurement,” *Journal of Political Economy*, 100(3), 598–614. DOI: 10.1086/261831.
- Bolton, Patrick and Mathias Dewatripont (2004), *Contract Theory*, MIT Press. ISBN: 978-0-262-02576-8.
- Corless, Robert M., Gaston H. Gonnet, D. E. G. Hare, David J. Jeffrey, and Donald E. Knuth (1996), “On the Lambert W Function,” *Advances in Computational Mathematics*, 5, 329–359. DOI: 10.1007/BF02124750.
- Ewert, Ralf and Alfred Wagenhofer (2021), “Motivating Managers to Invest in Accounting Quality: The Role of Conservative Accounting,” *Contemporary Accounting Research*, 38(3), 2000–2033. DOI: 10.1111/1911-3846.12664.
- Grossman, Sanford J. and Oliver D. Hart (1983), “An Analysis of the Principal-Agent Problem,” *Econometrica*, 51(1), 7–45. DOI: 10.2307/1912246.
- Holmström, Bengt (1979), “Moral Hazard and Observability,” *The Bell Journal of Economics*, 10(1), 74–91. DOI: 10.2307/3003320.
- Holmström, Bengt and Paul Milgrom (1987), “Aggregation and Linearity in the Provision of Intertemporal Incentives,” *Econometrica*, 55(2), 303–328. DOI: 10.2307/1913238.
- Holmström, Bengt and Paul Milgrom (1991), “Multitask Principal-Agent Analyses: Incentive Contracts, Asset Ownership, and Job Design,” *The Journal of Law, Eco-*

- nomics, and Organization*, 7 (Special Issue), 24–52. DOI: 10.1093/jleo/7.special_issue.24.
- Jewitt, Ian (1988), “Justifying the First-Order Approach to Principal-Agent Problems,” *Econometrica*, 56(5), 1177–1190. DOI: 10.2307/1911363.
- Laffont, Jean-Jacques and David Martimort (2002), *The Theory of Incentives: The Principal-Agent Model*, Princeton University Press. ISBN: 978-0-691-09184-6.
- Mirrlees, James A. (1976), “The Optimal Structure of Incentives and Authority within an Organization,” *The Bell Journal of Economics*, 7(1), 105–131. DOI: 10.2307/3003192.
- 株式会社東京証券取引所 (2021), 『コーポレートガバナンス・コード～会社の持続的な成長と中長期的な企業価値の向上のために～』, 2021 年 6 月 11 日改訂版.
- 株式会社東京証券取引所 (2026), 「コーポレートガバナンス・コードの改訂について（パブリックコメント）」, 2026 年 4 月 10 日公表, 意見募集期間 2026 年 4 月 10 日–5 月 15 日（改訂案・パブリックコメント段階）.