

Article

Not peer-reviewed version

Study on Coordination Mechanism of n -Level Principal-Agent with Multi-Branches at the Chain Terminal Under Fairness Preferences

[Simo Sun](#), Yuandong Wang, [Jianjun Jiao](#), [Yadong Shu](#) *

Posted Date: 12 March 2026

doi: 10.20944/preprints202603.0948.v1

Keywords: fairness preference; n -level principal-agent; incentive mechanism; farm-supermarket direct purchase; terminal synergy



Preprints.org is a free multidisciplinary platform providing preprint service that is dedicated to making early versions of research outputs permanently available and citable. Preprints posted at Preprints.org appear in Web of Science, Crossref, Google Scholar, Scilit, Europe PMC.

Copyright: This open access article is published under a [Creative Commons CC BY 4.0 license](#), which permit the free download, distribution, and reuse, provided that the author and preprint are cited in any reuse.

Disclaimer/Publisher's Note: The statements, opinions, and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions, or products referred to in the content.

Article

Study on Coordination Mechanism of n -Level Principal-Agent with Multi-Branches at the Chain Terminal Under Fairness Preferences

Simo Sun, Yuandong Wang, Jianjun Jiao and Yadong Shu *

School of Mathematics and Statistics, Guizhou University of Finance and Economics, Guiyang, Guizhou 550025, P.R. China

* Correspondence: shuyadong@mail.gufe.edu.cn

Abstract

Based on the principal-agent theory, this paper constructs a n -level principal-agent model with multi-branches at the chain terminal, introduces the improved Fehr-Schmidt fairness preference theory, and establishes a n -level incentive framework covering the horizontal and vertical fairness preferences of terminal agents. By treating intermediate agents as a unified whole through the "black-box" approach, this paper focuses on analyzing the influence mechanism of terminal agents' fairness preferences on effort level, incentive contract design and the operational efficiency of the entire principal-agent chain. The research results confirm that incorporating fairness preferences into the incentive mechanism design of green supply chain principal-agent relationships can effectively stimulate the production and operation enthusiasm of terminal agents, improve their effort levels; at the same time, it can significantly reduce agency costs in the principal-agent chain, optimize resource allocation efficiency, and thereby enhance green production efficiency. The conclusions and model design of this study provide a new theoretical path and practical reference for the collaborative operation and sustainable development of actual green supply chains with terminal multi-branch principal-agent structures.

Keywords: fairness preference; n -level principal-agent; incentive mechanism; farm-supermarket direct purchase; terminal synergy

1. Introduction

Principal-agent relationship is a core analytical framework for interpreting incentive mechanism design, resource allocation efficiency and cooperation maintenance. And the evolution of principal-agent theory has always been closely aligned with the complexity of real economic scenarios. Early principal-agent research, rooted in the Holmstrom & Milgrom(H-M) model, was based on the classic assumptions of complete rationality and risk aversion, focusing on the single principal-agent bilateral relationship and providing theoretical support for basic incentive contract design [1]. However, with the development of experimental economics, the limitations of the traditional pure self-interest assumption have become increasingly evident. The ultimatum game experiment conducted by Güth et al. pioneered the discovery that responders would reject non-zero but unfair allocation proposals, which violates the logic of purely self-interested decision-making [2]. Fehr et al.'s Gift Exchange Game experiments show that when employers offer wages above the market-clearing level, approximately 40% to 50% of employees will reciprocate with an effort level that exceeds contractual requirements, which breaks through the classic assumptions of pure self-interest and complete contract enforcement [3]. The dictator game experiment conducted by Forsythe et al. further confirmed that even when endowed with absolute allocation power, decision-makers would still cede part of their gains to others, which highlights the core role of fairness concerns in human behavior and lays a solid empirical foundation for the development of fairness preference theory [4].

The integration of fairness preference theory with traditional incentive theory has driven breakthrough progress in principal-agent research, and foreign scholars have accumulated rich achievements around this topic. Grund & Sliwka found in their research on tournament mechanisms that hiring agents with envy and compassion preferences can improve corporate payoffs [5]. Demougin & Fluit pointed out that when performance evaluation incurs costs, agents with envy preferences are more favored by enterprises [6]. The Reciprocity Model developed by Rabin and the F-S Model proposed by Fehr & Schmidt incorporate social preferences into the utility function, break through the constraints of the assumption of complete rationality from two distinct dimensions of intentional fairness and outcome fairness, and jointly form the core framework for the study of social preferences, thus providing a systematic explanation for relevant experimental phenomena [7,8]. Englmaier & Wambach embedded the F-S model into the H-M framework, and proposed that an optimal contracts need to achieve a three-dimensional equilibrium of incentives, insurance, and fairness [9]. Bartling confirmed that fairness preferences can effectively reduce agency costs in multi-task agency scenarios, which offers new for complex contract design [10].

Domestic research has kept pace with international frontiers while deepening and expanding in combination with local contexts. Wei Guangxing was the first to systematically introduce the game experiments and model construction of Western fairness preference theory, laying a foundational groundwork for domestic research in this field [11]. Xue Qiuzhi et al. were among the early representatives in China who systematically introduced fairness preference theory, expounding its guiding significance for management practices such as incentive mechanism optimization and payoffs distribution coordination [12]. Addressing the applicability limitations of traditional single principal-agent theory, Feng Genfu proposed dual agency theory, providing a new framework for the governance analysis of equity-concentrated enterprises [13]. Pu Yongjian implanted the concept of motivational fairness into the model, revealing the substitution relationship between material utility and fairness motivation, and enriching principal-agent theory from the perspective of behavioral economics [14,15]. Li Xun et al. incorporated the F-S Model into multi-level principal-agent analysis and found that fairness preferences alter the optimal contract structure [16]. Ding Chuan studied the fairness preference-based cooperation mechanism of marketing channels and found that fairness preferences can reduce free-riding behavior and promote cooperation among game players [17]. Fu Qiang and Zhu Hao considered the fairness preferences between agents and between principals and agents, and expanded the research framework of the principal-agent model [18]. Zhao Chenyuan pioneeringly introduced the fairness preference theory into the chain-type multi-level principal-agent model, and compared the incentive efficiency under the scenarios of complete rationality and procedural fairness preference [19]. Bai Chunkai verified through dynamic models that fairness preferences can optimize payoffs-sharing contracts in venture capital, which further expands the theoretical application scenarios [20].

Despite the existing research having broken through the traditional single rationality assumption [21], two core gaps still remain. First, there is insufficient research on the interaction mechanisms both within each level and across-level of multi-level principal-agent relationships, making it difficult to depict the mechanisms of incentive transmission and conflict coordination in complex chains. Second, there is a scarcity of quantitative modeling and empirical testing on fairness preferences in such structures with multi-branches at the chain terminal as Principal $\rightarrow$ Intermediary $\rightarrow$ Multiple Agents configuration. In real economic activities, such complex structures are widely present in fields such as supply chain collaboration, financial intermediation services, and public service provision. Taking farm-supermarket direct purchase as an example, because supermarkets(principals) connect with scattered farmers(multiple agents) via suppliers(intermediary) [22] and farmers' planting effort level directly determines agricultural product quality, farmers' fairness perception of the price gap between the purchase price and the supermarket's selling price often acts as a key factor affecting their effort level. Another example in the internet platform ecosystem, where platform parties(principals) manage a large number of settled merchants(multiple agents) through regional agents(intermediary), merchants' judgments on the fairness of platform

rules and the rationality of payoffs distribution directly affect their operational investment and compliance willingness. Individual decisions are generally driven by fairness preferences as people not only care about their absolute gains but also attach great importance to the fairness of payoffs distribution and the reciprocity of behavioral motives, which leads to significant deviations in traditional pure self-interest assumption models while explaining such scenarios [23]. More crucially, in the structures of multi-branches at the chain terminal, intermediary have a dual identity as both principals and agents, which results in intertwined issues such as incentive conflicts, risk transmission, and cross-level information asymmetry [24]. So simply superimposing single-level models cannot accurately depict their dynamic characteristics, which poses severe challenges for incentive mechanism design.

Based on this, targeting the limitations of existing theories, this paper adopts the three-dimensional research method of Theoretical Modeling-Numerical Simulation-Case Verification, and extends fairness preference theory from bilateral principal-agent relationships to more realistic chain-type n -level principal-agent scenarios. In the structural dimension, it horizontally expands the terminal agent layer of the chain structure to construct a multi-branch agency model. In the theoretical dimension, it breaks through the homo economicus assumption, incorporating horizontal(among agents) and vertical(between agents and principals) fairness preferences into the utility function, and focuses on the collaborative decision-making mechanism of terminal agents with intermediate levels treated as a black-box. In the empirical dimension, it quantifies the impact of fairness preferences on effort levels and agency costs through numerical simulation, and verifies the practical applicability of the theoretical model with the case of farm-supermarket direct purchase. By revealing the influence mechanism of fairness preferences on the collaborative decision-making of terminal agents and designing collaborative mechanisms and incentive contracts adapted to complex structures, this paper not only enriches the analytical depth and application boundaries of principal-agent theory but also provides practical theoretical guidance and practical solutions for incentive optimization in fields such as supply chain collaboration, organizational management, and platform ecosystem governance.

2. Model Construction

2.1. Model Definitions and Assumptions

Definition 1: Let the set of all participants be P . Among them, the set of pure principals is $\{p_0\}$, and the set of terminal agents is $A = \{A_1, ..., A_m\}$ ($|A| = m \geq 1$). The remaining participants who act as both principals and agents are termed as intermediary, and their set is $B = \{B_2, ..., B_n\}$ ($|B| = n \geq 1$). The model satisfies $|P| = n + m$ and has n -level consecutive principal-agent relationships ($n \geq 2$). This structure is defined as a n -level principal-agent with multi-branches at the chain terminal, as shown in Figure.1.

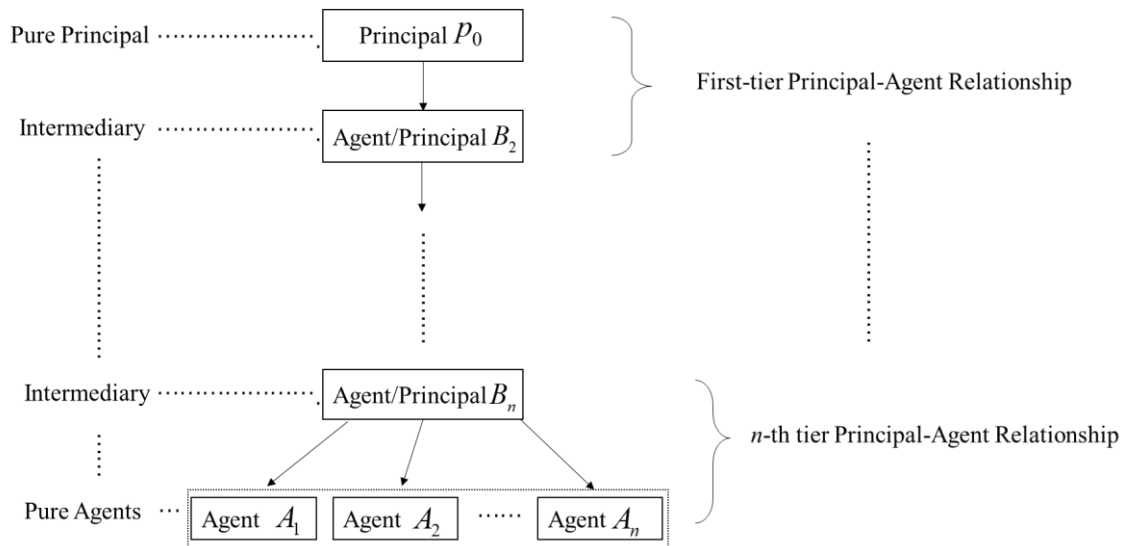


Figure. 1: n -level Principal-Agent Structure with Multi-branches at the Chain Terminal

To investigate the impact of pure agents' fairness preferences, the following assumptions are made.

- 1) (1) The principal p_0 is risk-neutral, while both the intermediary and pure agents are strictly risk-averse.
- (2) Pure agents are homogeneous, and their output function is $\pi_i = a_i + \theta$, ($i = 1, 2, \dots, m$), where a_i represents the effort level, and θ is an exogenous shock following a normal distribution $N(0, \sigma^2)$.
- 2) (3) The wage of a pure agent takes a linear form as $s_i = \alpha + \beta\pi_i$, where α is fixed payoffs and β is the output sharing coefficient. And the effort cost function is $c_i = \frac{1}{2}ba_i^2$ ($b > 0$).
- 3) (4) The output function of an intermediary is as following. When $1 < j < n$, $\tau_j = \mu_j + \tau_{j+1} + \theta$. While when $j = n$, $\tau_n = \sum_{i=1}^m \pi_i + \mu_n + \theta$. Where, μ_j denotes the effort level of the intermediary, and θ follows a normal distribution $N(0, \sigma^2)$.
- 4) (5) The expected wage of an intermediary is $q_j = \lambda + \beta_j\tau_j$, $j \in \{2, 3, \dots, n\}$, where λ is fixed payoffs and β_j is the output sharing coefficient. And the effort cost function is $e_j = \frac{1}{2}C\mu_j^2$ ($C > 0$).

2.2. Principal-Agent Model with Multi-branches at the Chain Terminal

Considering the preceding $n-1$ levels principal-agent relationships, the intermediary participates as a dual-identity participant whose payoffs is affected by multiple parties. Referring to the research method of Zhao Chenyuan [24], the analysis of the preceding $n-1$ levels relationship is as follows.

The intermediary B_j , as the direct principal of the intermediary B_{j+1} , gains actual payoffs equal to the difference between the wages of intermediary B_j and intermediary B_{j+1} . Thus the actual payoffs is as follows:

$$\begin{aligned}
\xi_j &= q_j - q_{j-1} = \lambda + \beta_j \tau_j - (\lambda + \beta_{j+1} \tau_{j+1}) \\
&= \beta_j \mu_j + \beta_j \theta + (\beta_j - \beta_{j+1}) \tau_{j+1} - \frac{1}{2} C \mu_j^2 \\
&\dots\dots \\
&= \beta_j \mu_j + (\beta_j - \beta_{j+1}) \left(\sum_{k=j+1}^n \mu_k + \sum_{i=1}^m a_i \right) + [\beta_j + (\beta_j - \beta_{j+1})(n+m-j)] \theta
\end{aligned}$$

When $j \in \{2, 3, \dots, n-1\}$, the actual net payoffs of intermediary B_j is the wage difference from B_{j+1} :

$$\begin{aligned}
\psi_j &= \xi_j - e_j = \beta_j \mu_j + (\beta_j - \beta_{j+1}) \left(\sum_{k=j+1}^n \mu_k + \sum_{i=1}^m a_i \right) \\
&\quad + [\beta_j + (\beta_j - \beta_{j+1})(n+m-j)] \theta - \frac{1}{2} C \mu_j^2
\end{aligned}$$

When $j = n$, the actual payoffs of intermediary B_n is

$$\begin{aligned}
\xi_n &= \lambda + \beta_n \tau_n - \sum_i s_i \\
&= \lambda + \beta_n \mu_n + (\beta_n - \beta) \sum_{i=1}^m a_i - m\alpha + [(\beta_n - \beta)m + \beta_n] \theta
\end{aligned}$$

$$\psi_n = \xi_n - e_n$$

$$\text{Its net payoffs is } \psi_n = \lambda + \beta_n \mu_n + (\beta_n - \beta) \sum_{i=1}^m a_i - m\alpha + [(\beta_n - \beta)m + \beta_n] \theta - \frac{1}{2} C \mu_n^2$$

Its expected value and variance are respectively as follows:

$$E\psi_n = \lambda + \mu_n \beta_n + (\beta_n - \beta) \sum_{i=1}^m a_i - m\alpha - \frac{1}{2} C \mu_n^2, \quad D\psi_n = [(\beta_n - \beta)m + \beta_n]^2 D\theta$$

The actual net payoffs of the pure principal p_0 is

$$\xi_1 = \tau_2 - q_2 = -\lambda + (1 - \beta_2) \left[\sum_{j=2}^n \mu_k + \sum_{i=1}^m a_i + (n-1+m)\theta \right]$$

In the n -th level of the principal-agent relationship, the net payoffs of the pure agent is

$$\omega_i = \alpha + \beta \pi_i - \frac{1}{2} b a_i^2$$

Its expected value and variance are respectively as follows:

$$E\omega_i = \alpha + \beta a_i - \frac{1}{2} b a_i^2, \quad D\omega_i = \beta^2 \sigma^2$$

Because both pure agents and intermediary are strictly risk-averse, so the certainty equivalent payoffs is used to quantify risk preferences, with its core expression shown below [25]:

$$CE = E(R) - \frac{1}{2} \rho D^2(R)$$

Where CE denotes the certainty equivalent payoffs, $E(R)$ denotes the expected payoffs, $\frac{1}{2}\rho D^2(R)$ denotes the risk premium, ρ denotes the risk aversion coefficient and $D^2(R)$ denotes the variance.

Based on this, the certainty equivalent payoffs for intermediary B_j ($2 \leq j < n$) can be derived as follows.

$$U_j = \beta_j \mu_j + (\beta_j - \beta_{j+1}) \left(\sum_{k=j+1}^n \mu_k + \sum_{i=1}^m a_i \right) - \frac{1}{2} C \mu_j^2 - \frac{1}{2} [\beta_j + (\beta_j - \beta_{j+1})(n + m - j)] \rho \sigma^2.$$

When $j = n$, the certainty equivalent payoffs of the intermediary B_n is

$$U_n = \lambda + \beta_n \mu_n + (\beta_n - \beta) \sum_{i=1}^m a_i - m\alpha - \frac{1}{2} C \mu_n^2 - \frac{1}{2} [(\beta_n - \beta)m + \beta_n] \rho \sigma^2.$$

The certainty equivalent payoffs of the pure agent A_i is

$$x_i = \alpha + \beta a_i - \frac{1}{2} b a_i^2 - \frac{1}{2} \rho \beta^2 \sigma^2.$$

The pure principal p_0 is risk-neutral, so its certainty equivalent payoffs is

$$U_1 = E\xi_1 = -\lambda + (1 - \beta_2) \left(\sum_{k=2}^n \mu_k + \sum_{i=1}^m a_i \right).$$

Thus, the basic model under the assumption of complete rationality is derived as follows:

$$\max U_1 = -\lambda + (1 - \beta_2) \left(\sum_{k=2}^n \mu_k + \sum_{i=1}^m a_i \right),$$

$$IR : U_{2 \leq j < n} = \beta_j \mu_j + (\beta_j - \beta_{j+1}) \left(\sum_{k=j+1}^n \mu_k + \sum_{i=1}^m a_i \right) - \frac{1}{2} C \mu_j^2 - \frac{1}{2} [\beta_j + (\beta_j - \beta_{j+1})(n + m - j)] \rho \sigma^2 \geq y_j,$$

$$U_n = \lambda + \beta_n \mu_n + (\beta_n - \beta) \sum_{i=1}^m a_i - m\alpha - \frac{1}{2} C \mu_n^2 - \frac{1}{2} [(\beta_n - \beta)m + \beta_n] \rho \sigma^2 \geq y_n$$

$$x_{1 \leq i \leq m} = \alpha + \beta a_i - \frac{1}{2} b a_i^2 - \frac{1}{2} \rho \beta^2 \sigma^2 \geq x_0,$$

$$IC : \max_{a_i} x_{1 \leq i \leq m}, \max_{\mu_j} U_{2 \leq j \leq n}.$$

3. Research on the Terminal Incentive Mechanism

3.1. Analysis of the Black-Box Model

To focus on the terminal coordination mechanism, the complex internal interactions of the intermediary is ignored and the intermediary can be treated as an integrated intermediary B_M . The set of participants is then $W = \{p_1, B_M, X_i\}$, ($i = 1, \dots, m$), where p_1 represents the pure principal and X_i ($i = 1, \dots, m$) represents the pure agent. The additional assumptions for the intermediary are as follows.

(1) The intermediary is risk-averse.

(2) The wage form is $M = \lambda + \beta_M \tau_M$, where τ_M denotes output and β_M denotes the output sharing coefficient.

(3) The effort cost function is $e_M = \frac{1}{2} C \mu_M^2$, ($C > 0$). And the output function is $\tau_M = \mu_M + S_A + \theta$ where S_A is the sum of outputs from pure agents and θ follows a $N(0, \sigma^2)$ distribution.

Based on the above assumptions, the core variables are derived as follows.

The output function of the intermediary is

$$\tau_M = \sum_{i=1}^m \pi_i + \mu_M + \theta = \sum_{i=1}^m a_i + \mu_M + (m+1)\theta.$$

The actual payoffs of the intermediary is

$$\begin{aligned} s_M &= \lambda + \beta_M \tau_M - \sum_{i=1}^m s_i \\ &= \lambda + \beta_M \mu_M + (\beta_M - \beta) \sum_{i=1}^m a_i - m\alpha + [(\beta_M - \beta)m + \beta_M] \theta. \end{aligned}$$

The actual net payoffs of the intermediary is

$$\begin{aligned} s'_M &= \lambda + \beta_M \tau_M - \sum_{i=1}^m s_i - e_M \\ &= \lambda + \beta_M \mu_M + (\beta_M - \beta) \sum_{i=1}^m a_i - m\alpha + [(\beta_M - \beta)m + \beta_M] \theta - \frac{1}{2} C \mu_M^2. \end{aligned}$$

The expectation value and variance of the intermediary's net payoffs are respectively as follows:

$$Es'_M = \lambda + \beta_M \mu_M + (\beta_M - \beta) \sum_{i=1}^m a_i - \frac{1}{2} C \mu_M^2 - m\alpha \quad \text{and}$$

$$Ds'_E = [(\beta_M - \beta)m + \beta_M]^2 D\theta.$$

The certain equivalent payoffs of the intermediary is

$$U_M = \lambda + (\beta_M - \beta) \sum_{i=1}^m a_i + \beta_M \mu_M - \frac{1}{2} C \mu_M^2 - m\alpha - \frac{1}{2} [(\beta_M - \beta)m + \beta_M]^2 \rho \sigma^2.$$

The payoffs of the pure principal p_1 equals the difference between output and wage from the intermediary B_M . The payoffs of the principal is

$$\begin{aligned} U &= \tau_M - M = \tau_M - (\lambda + \beta_M \tau_M) \\ &= -\lambda + (1 - \beta_M) \left[\sum_{i=1}^m a_i + \mu_M + (m+1)\theta \right]. \end{aligned}$$

Given the assumption of the principal's risk neutrality, the certain equivalent payoffs of the principal is

$$EU = -\lambda + (1 - \beta_M) \left(\mu_M + \sum_{i=1}^m a_i \right).$$

3.2. Fairness Preference Theory

Drawing on Fehr-Schmidt theory, this paper introduces the mechanisms of "relative deprivation" and "relative satisfaction" to reflect the sense of contentment [14,24,26]. From the perspective of bounded rationality, the horizontal and vertical fairness preferences of pure agents are incorporated into the n-level principal-agent structure with multiple terminal branches.

About horizontal fairness preference, pure agent X_i pays attention to the payoffs difference with other agents, and the perceived utility of pure agent X_i is

$$\omega_i = n_i + k'_i \sum_{(i \neq j)} \max\{s_i - s_j, 0\} - k'_j \sum_{i \neq j} \max\{s_j - s_i, 0\}.$$

In this equation $i, j \in \{1, 2, \dots, m\}$, s_i and s_j denote the actual payoffs of pure agents X_i and X_j respectively, and n_i denotes the actual net payoffs of pure agent X_i . k'_i and k'_j respectively measure the positive utility (a sense of pride) derived from having a higher payoffs than others and the negative utility (a sense of envy) arising from having a lower payoffs than others with $k'_i, k'_j \geq 0$. Under the assumption of homogeneous pure agents, let $k'_i = k'_j = k'$ ($k' > 0$). The larger the value of k , the stronger the intensity of preference for horizontal fairness of pure agent X_i . The actual net payoffs of pure agent X_i is

$$\omega_i = n_i + k' \sum_{i \neq j} (s_i - s_j) = n_i + k' \beta \sum_{i \neq j} (\pi_i - \pi_j).$$

About vertical fairness preference, considering the payoffs comparison with intermediary and the pure principal, the revised net payoffs is

$$\begin{aligned} \omega_i &= n_i + \gamma_i \{ \max[s_i - s_M, 0] + \max[s_i - U, 0] \} \\ &\quad - \eta_i \{ \max[\xi_j - s_M, 0] + \max[s_i - U, 0] \}. \end{aligned}$$

where s_i is the actual payoffs of pure agent X_i , ξ_j is the actual payoffs of intermediary B_M (when $j=1$, ξ_j represents the actual payoffs of the pure principal p_0), and n_i is the actual net payoffs of pure agent X_i . $\gamma_i \{\max[s_i - s_M, 0] + \max[s_i - U, 0]\}$ and $\eta_i \{\max[\xi_j - s_M, 0] + \max[s_i - U, 0]\}$ respectively represent the vertical pride-preference utility and envy-preference utility of the pure agent, with $\gamma_i, \eta_i \geq 0$. Under the assumption of homogeneous pure agents, let $\gamma_i = \eta_i = k''$ to measure the intensity of vertical fairness preference. The larger the value of k'' , the more sensitive the pure agent is to such fairness issues, leading to the more significant adjustment effect of payoffs distribution fairness on actual net payoffs.

At this point, the actual net wage of pure agent X_i is

$$\begin{aligned}\omega_i &= n_i + k''[(s_i - s_M) + (s_i - U)] \\ &= n_i + k''[(m+2)\alpha + 2\beta\pi_i + (\beta-1)\sum_{i=1}^m \pi_i - \theta - \mu_M] .\end{aligned}$$

Referring to the research by Fu Qiang et al. [18], by integrating horizontal and vertical fairness preferences, the actual net wage of pure agent X_i is

$$\omega_i = n_i + k'\beta \sum_{i \neq j} (\pi_i - \pi_j) + k''[(m+2)\alpha + 2\beta\pi_i + (\beta-1)\sum_{i=1}^m \pi_i - \theta - \mu_M] \quad (3.1)$$

About comprehensive fairness preference, Let $k' = k'' = k$ denote the combined effect of horizontal and vertical fairness preferences. Equation (3.1) can be simplified as

$$\omega_i = n_i + k[\beta \sum_{i \neq j} (\pi_i - \pi_j) + (m+2)\alpha + 2\beta\pi_i + (\beta-1)\sum_{i=1}^m \pi_i - \theta - \mu_M] .$$

Here s_i is the actual payoffs of pure agent X_i , s_M is the actual payoffs of the intermediary, n_i is the actual net payoffs of pure agent X_i , and k is the intensity of the pure agent's fairness preference. The expectation and variance of the actual net payoffs of pure agent X_i are respectively

$$\begin{aligned}E\omega_i &= \alpha + \beta a_i - \frac{1}{2} b a_i^2 + k\beta \sum_{i \neq j} (a_i - a_j) \\ &\quad + k[(m+2)\alpha + 2\beta a_i + (\beta-1)\sum_{i=1}^m a_i - \mu_M] , \\ D\omega_i &= [(km + 2k + 1)\beta - k(m+1)]^2 D\theta .\end{aligned}$$

When $k=0$, the pure agent is fully rational and no longer values fairness preferences. The certainty equivalent payoffs of pure agent X_i is

$$\begin{aligned}x_i &= \alpha + \beta a_i - \frac{1}{2} b a_i^2 + k[\beta \sum_{i \neq j} (a_i - a_j) + (m+2)\alpha + 2\beta a_i] \\ &\quad + k[(\beta-1)\sum_{i=1}^m a_i - \mu_M] - \frac{1}{2} [(km + 2k + 1)\beta - k(m+1)]^2 \rho \sigma^2 .\end{aligned}$$

The model degenerates into the traditional form under the assumption of full rationality, achieving consistency with the classical model.

3.3. Model Analysis Under Information Symmetry

When information is symmetric, the effort levels of all parties being directly observed and without information barriers or moral hazard, the model only needs to satisfy the participation constraint (IR), that is, the payoffs of the intermediary B_M and the pure agent X_i is not lower than their opportunity cost. The pure principal p_1 aims to maximize the expected net payoffs EU . Assuming the reservation payoffs of both parties is x_0 , the optimization problem is constructed as follows:

$$\begin{aligned} \max EU &= -\lambda + (1 - \beta_M)(\mu_M + \sum_{i=1}^m a_i) \\ IR: U_M &= \lambda + (\beta_M - \beta) \sum_{i=1}^m a_i + \beta_M \mu_M - \frac{1}{2} C \mu_M^2 - m\alpha \\ &\quad - \frac{1}{2} [\beta_M(m+1) - \beta m]^2 \rho \sigma^2 \geq x_0 \\ x_i &= \alpha + \beta a_i - \frac{1}{2} b a_i^2 + k [\beta \sum_{i \neq j} (a_i - a_j) + (m+2)\alpha + 2\beta a_i] \\ &\quad + k [(\beta - 1) \sum_{i=1}^m a_i - \mu_M] - \frac{1}{2} [(km + 2k + 1)\beta - k(m+1)]^2 \rho \sigma^2 \geq x_0 \end{aligned}$$

When setting the participation constraint to equality, solve the system of equations to obtain the fixed payoffs

$$\begin{aligned} \alpha &= \frac{2km(1 - \beta) \sum_{i=1}^m a_i + 2m(k\mu_M + T_2 + x_0) - 2\beta(2k + 1) \sum_{i=1}^m a_i + b \sum_{i=1}^m a_i^2}{2m(km + 2k + 1)} \quad \text{and} \\ \lambda &= \frac{2km(\mu_M + \sum_{i=1}^m a_i) + 2m(T_2 + x_0) + b \sum_{i=1}^m a_i^2}{2(km + 2k + 1)} \\ &\quad + \frac{C\mu_M^2 + 2(T_1 + x_0) - 2\beta_M(\sum_{i=1}^m a_i + \mu_M)}{2} \end{aligned}$$

$$\text{where } T_1 = \frac{1}{2} [\beta_M(m+1) - \beta m]^2 \rho \sigma^2 \quad \text{and} \quad T_2 = \frac{1}{2} [\beta + k(\beta - \beta_M)(m+1)]^2 \rho \sigma^2.$$

Substitute λ into the objective function, take the partial derivatives with respect to β , β_M , μ_M , a_1 , a_2 , $\dots$, a_m respectively, set the first-order partial derivatives as zero, and rearrange to obtain as follows:

$$\beta = \frac{k(m+1)}{km+2k+1}, \quad \beta_M = \frac{km}{km+2k+1},$$

$$a_1 = a_2 = \dots = a_m = \frac{2k+1}{b} \text{ and } \mu_M = \frac{2k+1}{C(km+2k+1)}.$$

Substitute the above expressions into the objective function, then

$$EU = \frac{(g+m)bx_0}{2Cbg^3} - \frac{(2k+1)^2 m}{4Cbg^3},$$

$$\alpha = \frac{(2m+4b)k^2 + [2bCx_0(m+2) + (m+4)C]k + C(2bx_0+1)}{2Cbg^2} \text{ and}$$

$$\lambda = \frac{Cg[2bx_0(m+g) + m(2k+1)^2] + b(2k+1)}{2bCg^2}, \text{ where } g = km+2k+1.$$

Core Conclusions are as follows.

Theorem 1: The optimal output sharing coefficient for the pure agent is $\beta = \frac{m+1}{m+2}$, which is positively correlated with the fairness preference intensity k .

When $k=0$, $\beta=0$, optimal risk sharing is achieved. When $k \rightarrow \infty$, β tends to $\frac{m+1}{m+2}$, which is consistent with the research conclusion of Zhao Chengyuan [24].

Proof: Since $\beta = \frac{k(m+1)}{km+2k+1}$, $\frac{\partial \beta}{\partial k} = \frac{m+1}{(km+2k+1)^2} > 0$, which means β is positively correlated with k and $\lim_{k \rightarrow \infty} \beta = \frac{m+1}{m+2}$. When $k=0$, $\beta=0$.

Theorem 2: The effort level of the pure agent $X_i (i=1, \dots, m)$ is $a_1 = a_2 = \dots = a_m = \frac{2k+1}{b}$, which is positively correlated with the fairness preference intensity k and negatively correlated with the effort cost coefficient b . Fairness preference can enhance the incentive effect.

Proof: When $k=0$, $a_i = \frac{1}{b}$, thus $a_i - a_i(k=0) = \frac{2k}{b} > 0$, indicating that the incentive effect on the agent is enhanced. Furthermore, $\frac{\partial a_i}{\partial k} = \frac{2}{b} > 0$, so the agent's effort is positively correlated with the fairness preference intensity k and negatively correlated with the effort cost coefficient b .

Theorem 3: The overall payoffs of the model is jointly determined by the number of pure agents m , the effort cost coefficient, the reservation payoffs x_0 and the fairness preference intensity k . It is independent of the risk aversion degree and negatively correlated with the effort cost coefficient b .

Proof: Since $EU = \frac{(g+m)bx_0}{2Cbg^3} - \frac{(2k+1)^2 m}{4Cbg^3}$, where $g = km+2k+1$, therefore

$$\frac{\partial EU}{\partial b} = -\frac{m(2k+1)^2}{2b^2g} < 0 \quad \text{and} \quad \frac{\partial EU}{\partial C} = -\frac{(2k+1)^2}{2C^2g} < 0.$$

3.4. Incentive Model under Asymmetric Information

Under asymmetric information, because the principal cannot directly observe the effort level, and the behaviors adopted in the principal-agent relationship are incompletely cooperative. A contract should be designed through the incentive compatibility constraint(IC) [27]. The intermediary and the pure agents choose the optimal effort level under the participation constraint, and the optimization problem is constructed as follows:

$$\begin{aligned} \max EU &= -\lambda + (1-\beta_M)(\mu_M + \sum_{i=1}^m a_i) \\ IR: U_M &= \lambda + (\beta_M - \beta) \sum_{i=1}^m a_i + \beta_M \mu_M - \frac{1}{2} C \mu_M^2 - m\alpha \\ &\quad - \frac{1}{2} [\beta_M(m+1) - \beta m]^2 \rho \sigma^2 \geq x_0 \\ x_i &= \alpha + \beta a_i - \frac{1}{2} b a_i^2 + k[\beta \sum_{i \neq j} (a_i - a_j) + (m+2)\alpha + 2\beta a_i] \\ &\quad + k[(\beta - 1) \sum_{i=1}^m a_i - \mu_M] - \frac{1}{2} [(km + 2k + 1)\beta - k(m+1)]^2 \rho \sigma^2 \geq x_0 \end{aligned}$$

The "intermediary" selects their own effort level μ_M to maximize their own utility, while the principal designs the contract parameter β_M to maximize their own utility. Therefore, when solving for the maximum value of the objective function, the maximum values of μ_M and a_i should be derived first under the incentive compatibility constraint by backward induction. Set $\frac{\partial U_M}{\partial \mu_M} = \frac{\partial x_i}{\partial a_i} = 0$, from which it can be derived that $a_i = \frac{(km + 2k + 1)\beta - k}{b}$, $\mu_M = \frac{\beta_M}{C}$.

When setting the participation constraint to equality, α and λ can be calculated as follows:

$$\begin{aligned} \alpha &= \frac{2km(1-\beta)S_1 + 2m(k\mu_M + T_2 + x_0) - 2\beta S(2k+1) + bS_2}{2m(km + 2k + 1)}, \\ \lambda &= \frac{2km(\frac{\beta_M}{C} + S_1) + 2m(T_2 + x_0) + bS_2}{2(km + 2k + 1)} + \frac{\beta_M^2}{2C} + T_1 + x_0 - \beta_M(S_1 + \frac{\beta_M}{C}), \\ \text{where, } S_1 &= \frac{(km + 2k + 1)\beta - k}{b}, S_2 = \frac{[(km + 2k + 1)\beta - k]^2}{b^2}, \\ T_1 &= \frac{1}{2} [\beta_M(m+1) - \beta m]^2 \rho \sigma^2 \quad \text{and} \quad T_2 = \frac{1}{2} [\beta + k(\beta - \beta_M)(m+1)]^2 \rho \sigma^2. \end{aligned}$$

Substitute λ , μ_M and a_i into the objective function, take partial derivatives with respect to β and β_M , and after rearrangement the following equations can be obtained.

$$\frac{\partial EU}{\partial \beta} = \frac{[(m+g)\beta - (m+1)(k+\beta_M)]bm\rho\sigma^2 + m(g\beta - 3k - 1)}{b},$$
$$\frac{\partial EU}{\partial \beta_M} = \frac{Cg\rho\sigma^2(m+1)[m\beta - (m+1)\beta_M] - \beta_M g + 2k + 1}{Cg}.$$

Meanwhile, $\frac{\partial^2 EU}{\partial \beta^2} = -\frac{m(m+g)b\rho\sigma^2 + g}{b} < 0$, $\frac{\partial^2 EU}{\partial \beta_M^2} = -\frac{1 + (m+1)^2 C\rho\sigma^2}{C} < 0$, where $g = km + 2k + 1$. So the objective function has a local maximum. Set the first-order partial derivatives to zero and solve the system of equations simultaneously.

$$\beta = \frac{(ft^2 + t)bg^2 + [(1 - k - ft(k+1))bt + (3k+1)(ft+1)g]}{(ft^2 + t)bg^2 + bgmt + (ft+1)g^2},$$
$$\beta_M = \frac{bfg(g-2k-1)t^2 + [3(m+1)Cg^2 + (2k+3)b + (6k+1)f]gt}{(ft+1)(bt+1)g^2 + bgmt} + \frac{2(m+1)Cgt + (2k+1)g - (2k-m+2)bt}{(ft+1)(bt+1)g^2 + bgmt},$$

where $g = km + 2k + 1$, $f = C(m+1)^2$ and $t = \rho\sigma^2$. Substitute β and β_M into a_i , u_M .

$$a_i = \frac{bt(ft+1)g^2 - [(2k+1)(bft^2 - ft - 1) + 2btk]g + bmt}{(ft+1)(bt+1)bg + bmt},$$
$$\mu_M = \frac{bft^2g^2 + Ctg(m+1)(2-3g) - (2k+1)(bft^3 - 1)g}{(ft+1)(bt+1)Cg^2 + bgmt} + \frac{[(2k+3)b + (6k+1)f]tg - (4k-m+2)tb}{(ft+1)(bt+1)Cg^2 + bgmt}.$$

The core conclusions are as follows.

Theorem 4: Under full rationality ($k=0$), the output sharing coefficient β is determined by the effort cost coefficient and $\rho\sigma^2$. Moreover, the sharing coefficient of pure agents is higher than that of the intermediary, achieving a balance between incentives and stability.

Proof: When $k=0$, there is $\beta = \frac{1 + (m+1)(Cm + C + b)\rho\sigma^2}{1 + bC\rho^2\sigma^4(m+1)^2 + (m+1)(Cm + C + b)\rho\sigma^2},$

$$\beta_M = \frac{1 + (m+1)(Cm + b)\rho\sigma^2}{1 + bC\rho^2\sigma^4(m+1)^2 + (m+1)(Cm + C + b)\rho\sigma^2}.$$

Thus, $\beta - \beta_M = \frac{(m+1)C\rho\sigma^2}{1 + bC\rho^2\sigma^4(m+1)^2 + (m+1)(Cm + C + b)\rho\sigma^2} > 0.$

Theorem 5: After introducing fairness preference, the optimal effort level of pure agents is higher than that under the assumption of complete self-interest. Fairness perception forms implicit incentives and reduces moral hazard.

Proof: When $k = 0$, $a_i' = \frac{(m+1)(mC + C + b)\rho\sigma^2 + 1}{b^2C\rho^2(m+1)^2\rho^4 + (m+1)(mC + C + b)b\rho\sigma^2 + b}$, Construct the function $Q = a_i - a_i'$, where a_i is the effort level of pure agents when $k \neq 0$ and a_i' is the effort level of pure agents when $k = 0$. In this case,

$$Q = \frac{k\{mgb^2f^2t^4 + ft^2[kmtb^2(m+2) + gf(bt+2)]\}}{[fb^2t^2 + bt(bm+b+f)][bgtf(t+1) + tf(bm+g) + g]} + \frac{kg[bt^2(m+1)(bm+4f) + bt(3m+2) + 4ft+2]}{[fb^2t^2 + bt(bm+b+f)][bgtf(t+1) + tf(bm+g) + g]}$$

where $g = km + 2k + 1$, $f = C(m+1)^2$ and $t = \rho\sigma^2$. Obviously, $Q > 0$, thus $a_i > a_i'$.

Theorem 6: The effort level of pure agents is positively correlated with the fairness preference intensity k , and maintaining a fair payoffs distribution structure is the key to improving incentive efficiency.

Proof: Since

$$\frac{\partial a_i}{\partial k} = \frac{bt[g^2(m+1) + 2(2mg - m)] + 4tfg^2 + 2g^2}{[bfgt^2 + (m+g)bt + fgt]^2 b} + \frac{[b^2m(g^2 + 2mg - m) + 2(g^2m + 2g^2 + 2mg - m)]fbt^2 + 2f^2g^2t^2}{[bfgt^2 + (m+g)bt + fgt]^2 b} + \frac{[b^2m(g^2 + 2mg - m) + 2(g^2m + 2g^2 + 2mg - m)]fbt^2 + 2f^2g^2t^2}{[bfgt^2 + (m+g)bt + fgt]^2 b}$$

where $g = km + 2k + 1$, $f = C(m+1)^2$, $t = \rho\sigma^2$ and $2mg - m = m(2km + 4k + 1) > 0$. Therefore $\frac{\partial a_i}{\partial k} > 0$, which implies that the effort level of pure agents is positively correlated with the degree of fairness preference.

Theorem 7: The effectiveness of the incentive contract is jointly determined by k , C , b and $\rho\sigma^2$, and the optimal sharing coefficient of pure agents is greater than that of the intermediary.

When k approaches infinity, the sharing coefficient is determined by m and $b\rho\sigma^2$.

Proof: When $k \neq 0$, there is

$$\beta - \beta_M = \frac{(k+1)(gC + k)}{Cbg^2 + [Ck^2m^3 + 2Ckm^2 + Cm + (C+b)(km+k+1)]g + (km+k+1)^2}$$

where $g = (m+1)(k+1)\rho\sigma^2$. Obviously $\beta - \beta_M > 0$, i.e., $\beta > \beta_M$. As k approaches infinity,

$$\text{it follows that } \lim_{k \rightarrow \infty} \beta = \frac{mb\rho\sigma^2 + b\rho\sigma^2 + 3}{(m+2)(b\rho\sigma^2 + 1)}.$$

Theorem 8: The agency cost $AC = EU^* - EU$ (where EU^* denotes the principal's expected payoffs under symmetric information, and EU denotes the expected payoffs under asymmetric information) is positively correlated with ρ and σ^2 . When $k \neq 0$, a reasonable combination of m and k can significantly reduce the agency cost.

$$AC = EU^* - EU$$

Proof: Since
$$= \frac{\{(m+1)^2[(Cmg+b)t+mg+1]\}(km-2k-1)^2t}{2g\{bfgt^2+[m^2(g+2k)+(5Ck+bk+2C+b)m+(2k+1)(b+C)]t+g\}},$$

Where $g = km + 2k + 1, f = C(m+1)^2, t = \rho\sigma^2$, Therefore

$$\frac{\partial AC}{\partial \rho} = \frac{(km-2k-1)^2\sigma^2[(m+1)^2(f_1+f_2)\rho^2\sigma^4+f_3]}{2g[(m+1)^2bgC\rho^2\sigma^4+(f_4\rho\sigma^2+g)]^2} > 0,$$

$$\frac{\partial AC}{\partial \sigma^2} = \frac{(km-2k-1)^2\rho[(m+1)^2(f_1+f_2)\rho^2\sigma^4+f_3]}{2g[(m+1)^2bgC\rho^2\sigma^4+(f_4\rho\sigma^2+g)]^2} > 0.$$

Where $f_1 = C^2k^2m^5 + 2kC^2m^4(3k+1) + [13C^2k^2 + 2kC(4C+b) + C^2]m^3,$

$$f_2 = 2Cm^2(3k+1)(2Ck+C+b) + 4C^2k^2m + [(2C+b)^2k + (C+b)^2]m + (2k+1)b^2,$$

$$f_3 = 2g(m+1)^2(Cmg+b)\rho\sigma^2 + (mg+1)g \text{ and}$$

$$f_4 = Ckm^3 + (4k+1)Cm^2 + [(5C+b)k + 2C+b]m + (2k+1)(b+C).$$

When $k = 0$, we have

$$\frac{\partial AC}{\partial m} = \frac{\{1+(m+1)^4Cb\rho^3\sigma^6 + [(Cm+b)^2 + 2m+1](m+1)^2\rho^2\sigma^4 + 2(m+1)(Cm+b)\rho\sigma^2\}}{2[1+(m+1)^2Cb\rho^2\sigma^4 + \rho\sigma^2(m+1)(Cm+C+b)]^2} > 0.$$

4. Numerical Simulation and Case Verification

4.1. Numerical Simulation

Set the parameters as $\rho\sigma^2 = 1, C = 1, b = 1, m = 3$ and $x_0 = 0$. Simulate the impact of different fairness preference intensities k on the key variables of the model, and the results are shown in Table 1.

Table 1. payoffs of each level under different fairness preferences.

Symmetric information	Asymmetric information
-----------------------	------------------------

Fairness preferenc e intensity	Effort level of pure agents	Expected payoffs of pure agents	Net payoffs of principal	Effort level of pure agents	Expected payoffs of pure agents	Net payoffs of principal
0.00	1.00	0.5000	2.0000	0.5676	2.1622	0.3221
0.50	2.00	1.5102	1.8776	1.7602	2.5835	1.2766
1.00	3.00	2.3333	2.3750	3.0000	2.9679	2.3333
1.50	4.00	3.1419	2.9343	4.2457	3.3293	3.4010
2.00	5.00	3.9463	3.5124	5.4934	3.6756	4.4720
2.50	6.00	4.7490	4.0988	6.7419	4.0116	5.5446
3.00	7.00	5.5508	4.6895	7.9909	4.3403	6.6179
3.50	8.00	6.3521	5.2827	9.2401	4.6637	7.6916
4.00	9.00	7.1531	5.8776	10.4895	4.9830	8.7657
4.50	10.00	7.9538	6.4735	11.7391	5.2992	9.8400
5.00	11.00	8.7544	7.0703	12.9887	5.6130	10.9144

Comparison of Table 1 and Figure.2 reveals the following findings.

- (1) Regardless of whether information is symmetric or not, the effort level and expected payoffs of pure agents, as well as the net payoffs of pure principals, increase significantly with the rise in the fairness preference intensity k .
- (2) The incentive effect under symmetric information is overall superior to that under asymmetric information, while fairness preference exerts a positive regulatory effect on both states.
- (3) When $k = 5.00$, the effort level of pure agents under symmetric information is 10 times higher than that when $k = 0$ (complete rationality), and the net payoffs of pure principals increases by 4.47 times, which verifies the significant incentive effect of fairness preference.

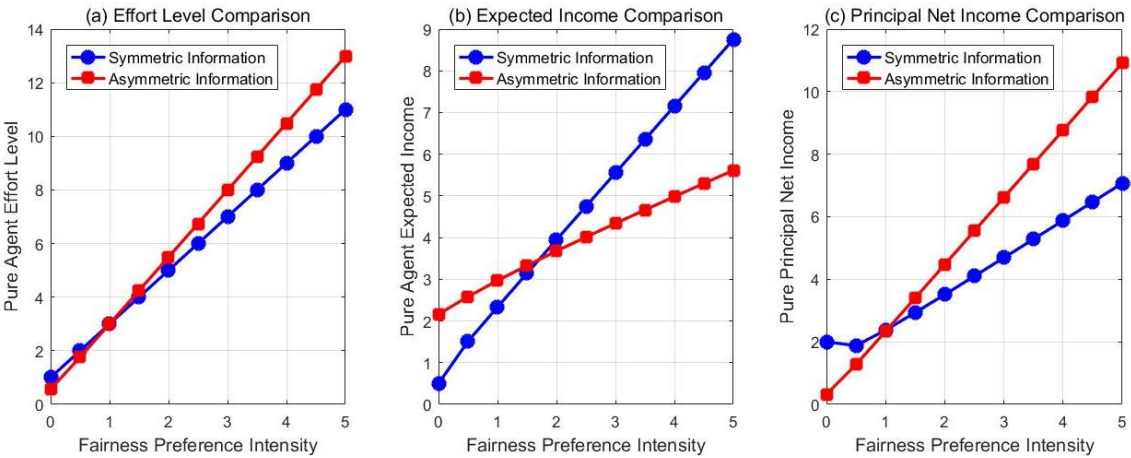


Figure. 2 The Impact of Fairness Preference on Key Variables

A further analysis is conducted on the relationship between agency cost, the number of agents m , and fairness preference intensity k in Table 2.

Table 2. Relationship between Agency Cost, Number of Agents and Fairness Preference

m k AC	$k=0$	$k=1$	$k=2$	$k=3$	$k=4$	$k=5$
$m=1$	0.4545	0.3049	0.3622	0.4356	0.5138	0.5939
$m=2$	0.6818	0.1078	0.0580	0.0396	0.0301	0.0243
$m=3$	0.9189	0.0000	0.0697	0.1904	0.3252	0.4659
$m=4$	1.1607	0.1464	0.7015	1.3277	1.9735	2.6275
$m=5$	1.4051	0.6353	2.1016	3.6582	5.2393	6.8305
$m=6$	1.6509	1.5177	4.3506	7.2909	10.2603	13.2416
$m=7$	1.8978	2.8248	7.4954	12.2889	17.1153	21.9554
$m=8$	2.1453	4.5768	11.5655	18.6909	25.8528	33.0296
$m=9$	2.3934	6.7873	16.5804	26.5223	36.5039	46.5016
$m=10$	2.6417	9.4657	22.5530	35.8002	49.0896	62.3963

It can be concluded from Table 2 that

- (1) When $k=0$, the agency cost increases monotonically with the rise in the number of agents m .
- (2) When $k \neq 0$, the agency cost is jointly determined by m and k , and there exist multiple value combinations that significantly reduce the agency cost (e.g., the agency cost is only 0.0580 when $m=2$ and $k=2$).
- (3) When $m=3$ and $k=1$, the agency cost drops to 0, achieving the optimal synergy effect, which verifies the feasibility of reducing agency costs by optimizing the combination of m and k .

4.2. Case Verification: Quality Incentives under the Farm-Supermarket Direct Purchase Model

The farm-supermarket direct linkage green supply chain is an important channel for farmers to increase their payoffs, with a principal-agent structure as supermarkets(pure principals) → suppliers(intermediary) → farmers(pure agents). Supermarkets need to incentivize farmers to provide high-quality agricultural products through suppliers, to achieve a win-win outcome for all three parties.

Studies by Fei have confirmed that the sum of the price premium for high-quality agricultural products and the planting investment subsidy provided by supermarkets should equal the cost difference for farmers between planting high-quality and low-quality agricultural products [28]. This is highly consistent with the conclusions of the model in this paper. Farmers have a strong fairness preference, and their effort level(investment in agricultural product quality) is significantly affected by their perception of payoffs fairness.

In practice, farmers not only focus on their own absolute payoffs but also attach importance to the relative fairness of payoffs distribution. If they perceive that supermarkets obtain excess payoffs through channel advantages while their own labor is not reasonably rewarded, they will experience a strong sense of unfairness and thus reduce quality investment. Conversely, when the gap between the acquisition price and the supermarket retail price is reasonable and subsidies are sufficient, farmers' fairness preferences are satisfied and they are willing to exert greater effort to ensure the quality of agricultural products.

Based on the model in this paper, supermarkets can optimize the incentive mechanism through the following measures.

(1) Increase the purchase price of high-quality agricultural products and raise the output sharing coefficient to directly enhance incentives for farmers.

(2) Increase training in planting technology and subsidies for production materials to reduce farmers' effort costs, indirectly improve their effort levels, and thereby enhance green production efficiency.

(3) Establish a transparent price announcement mechanism to narrow the perceived gap between the purchase price and retail price, thereby strengthening the positive incentive of vertical fairness preference.

(4) Optimize the scale of farmers managed by suppliers according to the fairness preference characteristics of farmers in different regions, so as to reduce agency costs.

This case verifies the practical applicability of the model in this paper and provides an operable practical path for the design of incentive mechanisms in complex principal-agent structures.

5. Conclusions

Based on principal-agent theory, this paper introduces fairness preference to construct a n -level principal-agent incentive model with multi-branches at the chain terminal, and systematically explores the impact of fairness preference on the effort level of terminal agents, the design of incentive contracts, and the overall efficiency of the principal-agent chain. The core finding is that the fairness preference of terminal agents is a key variable affecting the efficiency of the principal-agent chain. This finding remedies the limitation of traditional incentive theory that neglects individual fairness concerns and reveals the core cause of the deviation of contracts from reality.

Model analysis shows that in a complex chain-type multi-branch agency structure, fairness preference profoundly affects the design of optimal contracts by altering risk sharing and incentive costs. With the contract matching fairness preference characteristics, it motivates agents' intrinsic motivation and achieves incentive compatibility, while reducing agency costs and improving green production efficiency. Numerical simulation and The farm-supermarket direct linkage green supply chain case further verify the validity and practicality of the theoretical model, which confirms that fairness preference has a significant effect on controlling agency costs and improving the effort level of agents.

The theoretical contribution of this paper lies in extending fairness preference theory to the complex principal-agent structure with multi-branches at the chain terminal, and constructing a n -level incentive model integrating horizontal and vertical fairness preferences. The practical value lies in providing a quantitative reference for the design of incentive mechanisms in fields such as enterprise management and supply chain collaboration. Incentive efficiency can be improved by regulating factors related to fairness preference (e.g., optimizing the fairness of payoffs distribution, matching the number of agents with the intensity of fairness preference).

Future research can be further expanded in the following directions. First, introduce heterogeneous fairness preferences of agents to enhance the realistic explanatory power of the model. Second, consider a dynamic evolutionary perspective to analyze the long-term interaction between fairness preference and incentive mechanisms. Third, combine blockchain and quantum technologies to explore the regulatory effect of information transparency on the incentive effect of fairness preference.

Author Contributions: Conceptualization, S.S. and Y.W.; methodology, S.S.; software, Y.S.; validation, S.S., Y.S. and J.J.; formal analysis, J.J.; investigation, S.S.; resources, S.S.; data curation, Y.S.; writing-original draft preparation, S.S.; writing-review and editing, J.P.; visualization, Y.W.; supervision, Y.W.; project administration, S.S.; funding acquisition, S.S. All authors have read and agreed to the published version of the manuscript.

Funding: This research was funded by Universities Key Laboratory of System Modeling and Data Mining in Guizhou Province of funder grant number No.2023013.

Data Availability Statement: All data included in this study are available upon request by contact with the corresponding author.

Conflicts of Interest: The authors declare no conflict of interest.

References

1. Holmstrom B, Milgrom P. Aggregation and linearity in the provision of intertemporal incentives. *Econometrica: Journal of the Econometric Society*, 1987, 303-328.
2. Güth W, Schmittberger R, Schwarze B. An experimental analysis of ultimatum bargaining. *Journal of economic behavior & organization*, 1982,3(4), 367-388.
3. Fehr E, Kirchsteiger G, Riedl A. Does fairness prevent market clearing? An experimental investigation. *The quarterly journal of economics*, 1993, 108(2), 437-459.
4. Forsythe R, Horowitz J L, Savin N E, et al. Fairness in simple bargaining experiments. *Games and Economic behavior*, 1994, 6(3), 347-369.
5. Grund C, Sliwka D. Envy and compassion in tournaments. *Journal of Economics & Management Strategy*, 2005, 14(1), 187-207.
6. Demougin D, Fluet C, Helm C. Output and wages with inequality averse agents. *Canadian Journal of Economics*, 2006, 399-413.
7. Rabin M. Incorporating fairness into game theory and economics. *The American economic review*, 1993, 1281-1302.
8. Fehr E, Schmidt KM. A theory of fairness, competition and cooperation. *Quarterly Journal of Economics*, 1999, 114(3), 817- 868.
9. Englmaier F, Wambach A. Optimal incentive contracts under inequity aversion. *Games and economic Behavior*, 2010, 69(2), 312-328.
10. Bartling B. Relative performance or team evaluation? Optimal contracts for other-regarding agents. *Journal of Economic Behavior & Organization*, 2011, 79(3), 183-193.
11. Wei, G. X. A literature review on game experiments and theoretical models of fairness preferences. *The Journal of Quantitative Economics and Technical Economics*, 2006, (8), 152-161.
12. Xue, Q. Z., Huang, P. Y., Lu, Z., et al. *Behavioral economics. Theories and applications*. Wise Books. 2005.
13. Feng, G. F. Dual principal-agent theory: Another analytical framework for the corporate governance of listed companies-with a discussion on new ideas for further improving the corporate governance of listed companies. *Economic Research Journal*, 2004, (12), 16-25.
14. Pu, Y. J. Incorporating the "fairness game" into the principal-agent model-A contribution from behavioral economics. *Contemporary Finance & Economics*, 2007a, (3), 5-11.
15. Pu, Y. J. A principal-agent model based on behavioral economics theory: The trade-off between material utility and motivational fairness. *China Economic Quarterly*, 2007b, 7(1), 297-318.
16. Li, X., & Cao, G. H. A study on incentive mechanisms based on fairness preference theory. *Journal of Industrial Engineering and Engineering Management*, 2008, (2), 107-111, 116.
17. Ding, C., Wang, K. H., & Ran, R. A study on cooperation mechanisms of marketing channels based on fairness preferences. *Journal of Management Sciences in China*, 2013, 16(8), 80-94.
18. Fu, Q., & Zhu, H. A study on incentive mechanisms based on public preference theory —Considering both horizontal and vertical fairness preferences. *Journal of Industrial Engineering and Engineering Management*, 2014, (3), 190-195.

19. Zhao, C. Y., Pu, Y. J., & Pan, L. W. Incentives in chain-type multi-level principal-agent relationships-Based on a comparison of the complete rationality and procedural fairness preference models. *Journal of Chinese Management Science*, 2017, (6), 121-131.
20. Bai, C. K. A study on the multi-level principal-agent model of venture capital based on fairness preferences. Southwestern University of Finance and Economics, 2023.
21. Liu, Y. G., & Jiang, N. Y. A literature review of principal-agent theory. *Academics in China*, 2006, (1), 69-78.
22. Li, B. Y., Ma, D. Q., Dai, G. X., et al. Dynamic quality improvement strategies for the Farm-Supermarket Docking supply chain considering the reference quality effect. *Operations Research and Management Science*, 2021, 30(5), 52-59.
23. Lan, C. F. A study on supply chain decision-making and coordination considering behavioral factors. Beijing University of Posts and Telecommunications, 2017.
24. Zhao, C. Y. A study on multi-level principal-agent incentive mechanisms based on fairness preferences. Chongqing University, 2016.
25. Pu, Y. J. Incorporating the fairness game into the principal-agent model-A contribution from behavioral economics. *Contemporary Finance & Economics*, 2007, (3), 5-11.
26. Fehr E, Goette L, Zehnder C. A behavioral account of the labor market: The role of fairness concerns. Discussion Paper No.3901, 2008.
27. Cooper R, Ross T W.. Product warranties and double moral hazard. *The RAND Journal of Economics*, 1985, 103-113.
28. Fei, W. A study on the effort behaviors of subjects in typical agricultural product circulation channels-From the perspective of agricultural product quality and safety. *Journal of Beijing Technology and Business University*, 2013, 28(4), 1-9.

Disclaimer/Publisher's Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.