

Competing for the Future of AI: The Economic Incentives of Open-Source Foundation Models*

Mark A. Jamison[†] Byoungmin Yu[‡]

April 20, 2026

Abstract

This paper studies incentives for openness in foundation models (FMs). Using data from Hugging Face with model documentations, we develop and estimate a two-stage dynamic structural model in which the FM owners choose the degree of openness for their models, and downstream AI developers decide which FM to adopt for application development. Our estimates indicate that the short-term benefits of openness account for less than half of its contemporaneous costs, implying that immediate returns alone cannot rationalize observed openness choices of FM owners. Instead, longer-term considerations play a central role. By increasing openness, owners gain knowledge spillovers from downstream developers, which expands later adoption by improving the quality of subsequent model generations. This intertemporal feedback mechanism raises long-term continuation value and helps explain why FM owners adopt open-source strategies despite limited short-run monetization benefits.

Keywords: Open-source foundation model, knowledge spillovers, long-term payoff, innovation

JEL classification: L20, L17, O31, O33, O36

*We thank the Public Utility Research Center for financial support. All errors are our own.

[†]Public Utility Research Center and Digital Markets Initiative, Warrington College of Business, University of Florida, Gainesville, Florida, United States. (mark.jamison@warrington.ufl.edu).

[‡]Corresponding author. Public Utility Research Center and Digital Markets Initiative, Warrington College of Business, University of Florida, Gainesville, Florida, United States. (byoungminyu9@gmail.com).

1 Introduction

At least since the time of Adam Smith, economists and other scholars have been interested in innovation. Smith’s observations on division of labor and improvements in specialization emphasized innovation in production. Schumpeter’s creative destruction considered both product innovation and process innovation, and Aghion and Howitt (1990) modeled how this can lead to economic growth. Romer (1990) examined how intentional innovation drives economic growth, and Mokyr (1990, 1992) explained the role of scientific advancements in sustained technological innovation. Utterback (1994) studied how new products go through evolutionary phases, including novelty creation, standardization, and process improvement. Hughes et al. (1987) examined interdependencies across innovations in large technological systems.

Recent innovations in artificial intelligence (AI) present new opportunities to study how firms choose their innovation paths. Innovations in AI have been occurring at least since the 1950s, but the advent of foundation models (FMs) represents a structural break in this evolution. FMs are almost always large neural networks using the self-attention mechanism developed at Google (Vaswani et al., 2017) that enable a wide variety of new AI applications to build downstream of the FM. The first major, widely recognized FMs emerged around 2018, with BERT (Bidirectional Encoder Representations from Transformers) by Google and GPT (Generative Pre-trained Transformer) by OpenAI being the most prominent early examples. Evidence from GitHub highlights the importance and rapidity of this shift in AI: In 2023 alone, there were more than 65,000 generative AI projects, many of which are built on FMs, representing a 248% year-over-year increase.¹ The economic impact of generative AI is estimated to range between \$2.6 and \$4.4 trillion annually, positioning it as a major new driver of economic growth (Chui et al., 2023).

A diversity of firms—from large technology incumbents such as Meta to new entrants like DeepSeek—have entered the race to develop FMs. One feature of this competition is a firm choosing its openness features, i.e., the degree to which the model’s components—such as weights, training data, code, and documentation—are accessible, modifiable, and transparent for others allowing downstream developers to build generative AI applications at minimal cost.² The proliferation of open-source models has substantially lowered barriers to entry and democratized access to advanced AI capabilities, accelerating innovation well

¹<https://github.blog/news-insights/research/the-state-of-open-source-and-ai/>

²Weights are numerical parameters learned during training that determine how the model uses its various neural network nodes to generate output. Also, while there is ongoing debate over the precise definition of *open source* in the AI industry, we use the term *open-source* to broadly encompass foundation models whose parameters are made publicly accessible. For a detailed discussion of terminology, see the Open Source Initiative at <https://opensource.org/ai/open-source-ai-definition>.

beyond the confines of closed models such as proprietary GPT systems. This expansion is particularly evident on platforms like Hugging Face, where the number of hosted models has grown rapidly, enabling developers to adapt frontier technologies to niche use cases and generate a vast array of new applications.³

That open-source FMs differ in their choices of degrees of openness implies differences in their beliefs about costs and the potential revenue returns. It is generally accepted that FMs incur substantial sunk costs in training their large-scale models. Closed FMs in general receive revenue from Application Programming Interfaces (API) usage (Per-Token Fees); subscriptions to consumer-facing services, such as OpenAI’s ChatGPT Plus; and fees charged for allowing companies to customize and fine-tune the FM model using proprietary data. All other things being equal, choosing to close a FM allows the owner to receive more short-run revenue from selling API fees to downstream developers.

But openness involves tradeoffs for FMs. Some downstream developers who value openness will adopt open-source FM for customization and product-specific fine-tuning. The FM owner can then learn from such developers’ use and leverage that knowledge to subsequently improve the next-generation FM, consequently increasing the future demand.⁴ Xu et al. (2025) describe this as a “data flywheel effect.” And openness comes at a cost for an FM: Openness could benefit rivals, i.e., the process by which the generalized knowledge embedded within the FM spreads, leaks, or is repurposed by downstream developers, researchers, and other firms that may compete with the FM owner’s offerings. Xu et al. (2025) examine these tradeoffs. They consider situations where an FM’s openness has above effects and impacts rivalry. More specifically, rivals may benefit from the knowledge spillovers, leading to less investment and lower welfare.

Despite the central role of openness in shaping competition among foundation models, the trade-off between its short-run costs and long-run payoffs remains largely underexplored in the empirical literature. This paper helps fill this gap by developing and estimating a two-stage dynamic structural model of FM competition. In the first stage, FM owners choose the degree of openness of their models. In the second stage, downstream AI developers choose which FM to adopt for application development. Openness enters both stages. On the demand side, it affects developers’ adoption choices by changing accessibility, customization costs, and deployment flexibility, shifting their profit which we estimate in a reduced-form way. On the supply side, openness shapes the future state of competition by influencing

³<https://hai.stanford.edu/ai-index-report>

⁴In line with this view, Meta’s CEO Mark Zuckerberg has emphasized that the primary motivation for releasing the Llama series as open-source models is to expand market share and cultivate a new ecosystem around Meta. See <https://www.cnbc.com/2023/10/16/metaspopen-source-approach-to-ai-puzzles-wall-street-techies-love-it.html>.

the magnitude of knowledge spillovers and the evolution of subsequent model quality. The model therefore allows the private value of openness to depend not only on contemporaneous revenue, but also on the discounted value of future ecosystem rents.

This mechanism parallels the evolution of Linux as a viable competitor to Microsoft’s Windows ecosystem through an open-source strategy (Lerner and Tirole, 2002). However, in contrast to traditional open-source software—where contributors directly improve the core codebase—downstream developers in the AI ecosystem primarily contribute indirectly by generating fine-tuning data, applications, and usage patterns. The FM owner subsequently internalizes this ecosystem-generated knowledge to improve subsequent generations of the model rather than the current release. Accordingly, our framework captures a key dynamic feature of open-source strategies: the openness of the current model shapes the future trajectory of innovation in FMs and market share.

The empirical analysis combines data from Hugging Face with information collected from model documentation. These sources are used to construct measures of openness, model quality, and downstream adoption. We address identification in two ways. First, because FM markets are nascent, the time dimension of the data is relatively short. However, the rapid pace of innovation and frequent turnover across model generations provide substantial variation in relevant state variables over a compressed time horizon. Second, openness choices may be endogenous if firms respond to unobserved demand shocks. The paper addresses this concern by exploiting the production cycle of FMs: because pre-training, safety alignment, and release decisions require meaningful lead times, openness is plausibly predetermined relative to contemporaneous short-run shocks. The empirical specification also includes owner and time fixed effects to absorb persistent heterogeneity across firms and aggregate market trends. In addition, the paper adopts a conditional two-step calibration-and-recovery strategy where we first calibrate parsimonious openness cost parameters to match observed openness choices under a maintained normalization, and then recover the relative importance of short-run monetization and continuation value from the optimality condition.

We find that short-run returns alone are insufficient to rationalize observed openness choices. Preliminary estimates suggest that the immediate benefits of openness account for less than half of its contemporaneous costs, implying that a static profit framework cannot explain why firms open their models to the extent observed in the data. Instead, firms appear to place substantial value on the long-run benefits of openness. The predicted long-run value of openness is driven by the ecosystem formed around downstream developers who fine-tune and deploy the model. These activities generate knowledge spillovers that enhance the quality of subsequent foundation models. Higher-quality future models, in turn, attract additional downstream adopters, reinforcing ecosystem expansion and increasing short-run

revenues in later periods. This feedback mechanism explains why firms may rationally sacrifice short-run profits in exchange for future market expansion and suggests that firms choose open-source strategies as a long-run pathway to innovation.

1.1 Related Literature

This paper relates to three strands of research: the emerging literature on the economics of open-source foundation models, the broader literature on firms’ open-source strategies, and classic work on innovation. Recent work on open-source FMs identifies several factors that incentivize FM owners to release open-source models, including the profitability of complementary services, downstream competition, and innovation spillovers. Leisten (2025) develops a theoretical framework in which open-sourcing can be privately optimal when complementary services (such as hosting or enterprise partnerships) are sufficiently profitable and when downstream application development expands the ecosystem, thereby strengthening demand for such services. Relatedly, Xu et al. (2025) emphasize a trade-off whereby greater openness lowers fine-tuning and deployment frictions for downstream developers but may intensify competition by strengthening rival developers in upstream model markets. Habibi (2025) further empirically shows that incentives to release open-source FMs may decline as model quality approaches the technological frontier, while also highlighting the connection between open-sourcing and long-run innovation activity through a theoretical model. From a complementary perspective, Azoulay et al. (2025) argue that weak appropriability in AI increases the importance of complementary assets for commercialization. Our paper builds on these insights but shifts the focus from short-term profitability conditions to the long-term valuation of openness through downstream adoption and future quality evolution.

Our analysis also contributes to the literature on open-source strategies by for-profit firms in software and platform markets. This literature shows that proprietary firms may adopt open-source strategies not only to monetize the released product directly, but also to attract external contributions (Haese and Peukert, 2025), stimulate complementary innovation (Petrailia, 2025), and strengthen demand for complementary offerings (Iansiti and Richards, 2006; Fosfuri et al., 2008). In a similar vein, Nagle (2019) empirically shows that open-source software adoption significantly enhances firm productivity, particularly for firms that possess the necessary complementary IT capabilities. In these views, openness is a strategic instrument for ecosystem formation and innovation rather than a pure licensing choice. This perspective aligns closely with our framework, in which firms may rationally incur short-run costs of openness to increase downstream participation and generate future benefits through ecosystem feedback. At the same time, our setting differs from standard software

contexts because foundation model development features rapid generation cycles and a tight link between downstream use and subsequent model improvement. These features make the intertemporal consequences of openness particularly central in AI markets.

More broadly, the paper relates to the classic tension in innovation economics between appropriability and diffusion. Traditional theories of innovation emphasize that firms require sufficient control and appropriability to recover large R&D investments and sustain innovation incentives (Schumpeter, 1942; Romer, 1990). By contrast, the economics of open-source software stresses that relaxing control can stimulate distributed experimentation, user customization, and collaborative development, thereby generating external knowledge that may accelerate innovation (Lerner and Tirole, 2002). The source of innovation is thus distributed across all the players in the field, including the competitors (Von Hippel, 1988). The foundation-model context brings these forces into especially sharp relief: open-sourcing may reduce direct control over current monetization, yet it can expand the downstream ecosystem that generates feedback and experimentation relevant for future model quality. Our paper formalizes this trade-off in a setting where the strategic value of openness depends on firms’ expectations about future ecosystem rents and innovation gains, not only on contemporaneous revenues.

The existing literature provides important insights into the benefits and costs of openness, but leaves limited evidence on how firms dynamically value openness as a pathway to future ecosystem growth and subsequent model improvement. We address this gap by developing and estimating a dynamic structural model in which openness affects downstream adoption, generates knowledge spillovers, and feeds into later-period model quality and market share. In doing so, the paper brings a structural, intertemporal measurement approach to a literature that has largely emphasized theoretical or reduced-form analyses of AI openness. In addition, this study provides a unified explanation for heterogeneity in open-source strategies across FM owners and clarifies when short-run losses from openness can be privately rational. More broadly, our results inform debates in information economics and AI policy by suggesting that observed openness choices may reflect dynamic competition for future ecosystem control, rather than static pricing or licensing incentives alone.

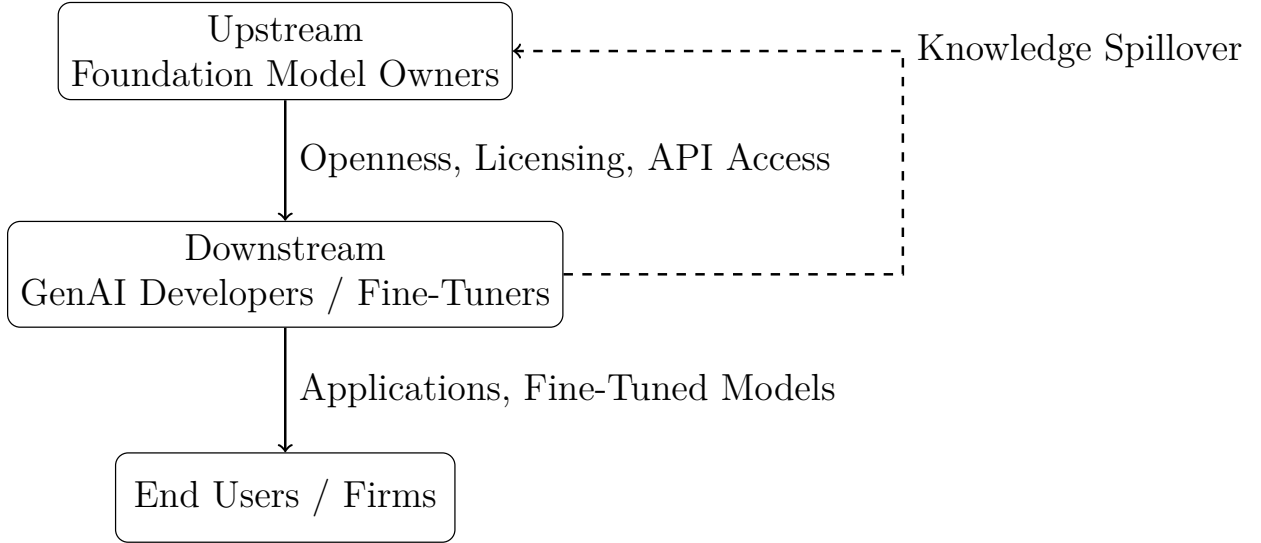
The remainder of this paper is organized as follows. In Section 2, we present a two-stage structural model for the developers’ choices on the openness level of their foundation models. Section 3 presents a dataset we use. Section 4 reports the estimated parameters and payoff parameters of FM owners. Concluding remarks follow in Section 5.

2 Model

This section develops a simple vertical-structure model that captures the strategic interaction between upstream foundation model owners and downstream Generative AI (GenAI) application developers. The goal is to understand what are the factors that drive upstream owners to increase the openness level of their foundation model. We exclude time subscripts t for the ease of exposition, and recover it in Section 2.3.

2.1 Industry Structure

Figure 1: Value-chain of foundation model ecosystem



We consider a vertical chain illustrated in Figure 1. In specific:

- **Upstream FM Owners** $u \in \mathcal{U} = \{1, \dots, U\}$: Owners such as Meta, Mistral, Alibaba (Qwen), or DeepSeek that produce open-source foundation models $j \in \mathcal{J} = \{1, 2, \dots, J\}$.
- **Downstream Developers** $i \in \mathcal{I} = \{1, \dots, I\}$: startups, enterprise developers, and independent researchers that fine-tune or integrate FMs into applications.

We construct a two-stage game. In Stage 1, multi-product upstream owners $u \in \mathcal{U}$ decide openness level $O_j \in (0, 1]$ of a foundation model j , given a set of model quality $\vec{q} = \{q_1, q_2, \dots, q_J\}$. $O_j = 1$ indicates a fully opened model, and O_j affects the quality of its next generation model, q_{j_u} , through knowledge spillovers from expanded downstream ecosystem. In Stage 2, downstream developers choose which foundation model to adopt for application development.

2.2 Stage 2: Downstream Adoption

We specify a reduced-form latent profit index for downstream developers adopting FM j as:⁵

$$\pi_{ij} = Aq_j + \beta O_j + \gamma q_j \times O_j + \alpha p_j + \mathbf{X}_j \lambda + \xi_j + \varepsilon_{ij}. \quad (1)$$

Here, q_j denotes the base quality of model j that affects the downstream developer’s profit. O_j represents its degree of openness. In open-source environments, users can directly customize software and fix bugs to suit their own needs, which can facilitate decentralized adaptation and user-driven improvements (Lerner and Tirole, 2002). In the context of foundation models, this is analogous to fine-tuning and other forms of downstream model adaptation.

For downstream developers, however, the marginal effect of openness on profit is theoretically ambiguous. On the one hand, greater openness can reduce adaptation costs by improving access and transparency, thereby facilitating fine-tuning (Xu et al., 2025). Openness may also generate strategic benefits such as flexibility, transparency, and reliability in deployment (Varian and Shapiro, 2003; Hermansen, Cailean and Osborne, Michael, 2025). On the other hand, greater openness may increase implementation and maintenance burdens—especially under local deployment—because integration, debugging, and system upkeep must be handled by the developer rather than the platform provider. Accordingly, β captures the net marginal effect of openness on downstream profit, reflecting the balance between adaptation benefits and implementation costs. Furthermore, we include an interaction term between FM quality and openness, $q_j \times O_j$. The coefficient γ captures the hypothesis that the marginal impact of openness is conditional on the underlying model’s capabilities.

Finally, p_j denotes the effective per-use cost of accessing model j . For closed models (e.g., GPT-4), this cost can be interpreted as the price paid to the foundation model owner for API access. For open-source models (e.g., Llama 3), by contrast, access to model need not require payment to the model developer; instead, the relevant marginal cost is better interpreted as inference infrastructure cost—such as payments to a third-party cloud provider or the opportunity cost of self-hosting.⁶ We also include observed determinants of profit when

⁵We would like to carefully point out that equation (1) represents a latent profit index that we estimate with reduced-form parameters, rather than profit itself. However, for the sake of brevity and to maintain consistency with the empirical IO literature (e.g., Seim (2006)), we will refer to equation (1) as a *profit* for downstream developers for the rest of the paper.

⁶This interpretation is consistent with recent theoretical work distinguishing between closed and open-source model monetization. In Leisten (2025), open models are offered for free, and the founder monetizes through a complementary service, whereas in the closed paradigm, the founder sells access to the model and its applications. Habibi (2025) similarly models open-sourcing as a trade-off between faster technology growth and forgone immediate returns from licensing or API sales. Our use of p_j extends this distinction by treating it as an effective access cost: for open-source models, it captures the cost of compute and deployment rather than a direct licensing payment to the model owner. In the similar vein, the CEO of Meta stated

adopting model j in a vector $\mathbf{X}_j$. The unobserved quality (to the econometricians) of FM is captured by ξ_j , which also affects downstream profit. Downstream developers may also choose an outside option by not using any foundation model. Let the indirect profit from the outside option be $\pi_{i0} = \varepsilon_{i0}$.

If ε_{ij} is assumed to follow a type-I extreme value distribution, then the implied choice probability of adopting model j takes the multinomial logit form:

$$s_j = \frac{\exp(D_j)}{1 + \sum_{k \in J} \exp(D_k)}, \quad (2)$$

where $D_j = Aq_j + \beta O_j + \gamma q_j \times O_j + \alpha p_j + \mathbf{X}_j \lambda + \xi_j$ represents the deterministic component of reduced-form profit, not as accounting profit itself. As noted in the footnote, since the scale of ε_{ij} is normalized for identification, D_j should be interpreted as a dimensionless latent index that moves monotonically with accounting profit, rather than a direct profit measure.

2.3 Stage 1: Openness Strategy

2.3.1 Dynamic quality evolution

We incorporate a structural feature of AI development cycles: the time lag between model iterations. A foundation model initiated at time t requires a training and safety alignment period τ before release. Consequently, the quality of the next-generation model evolves according to the following law of motion:

$$q_{j'_u, t+\tau} = \rho \cdot q_{j_u, t} + \theta_s \cdot \text{Ecosystem}_{j_u, t} + \theta_p \cdot \text{Size}_{j'_u, t+\tau} + \omega_{j'_u, t+\tau}, \quad (3)$$

where:

- $q_{j_u, t}$ denotes the quality of the deployed model j by owner u at time t , and $q_{j'_u, t+\tau}$ represents the quality of the subsequent version j'_u (e.g., the transition from Llama 2 to Llama 3).
- $\rho \in (0, 1)$ is the *persistence parameter*, capturing the extent to which current architectural capabilities carry over to the next generation.
- $\theta_s > 0$ captures the *knowledge spillover (net) efficiency*. It quantifies how effectively the developer harvests downstream innovations—such as fine-tuning recipes, bug reports, and usage data—to enhance the pre-training and lower the effective marginal cost of

that: “a key difference between Meta and closed model providers is that selling access to AI models isn’t our business model” (Meta, 2024). For more details, see Eastwood (2026).

achieving a higher quality level of the next model. Here, $\text{Ecosystem}_{j_u,t}$ is the cumulative number of downstream developers for each model.

- $\text{Size}_{j'_u,t+\tau}$ is the size of parameters used for pre-training the model.
- $\omega_{j'_u,t+\tau} \sim \mathcal{N}(0, \sigma^2)$ is an independent and idiosyncratic innovation shock realized at the time of launch.

This specification implies a trajectory for innovation via open-sourcing FMs and the subsequent strategic gains: Openness (O_j) increases market share ($s_{j,t}$) by lowering adoption barriers, which in turn accelerates quality improvement ($\theta_s \cdot \text{Ecosystem}_{j_u,t}$) for the future period $t + \tau$, eventually expanding future market share via A as well.

We conceptually distinguish this knowledge spillover mechanism from traditional open-source software such as Linux, where value accrues through direct code contributions (e.g., GitHub). For the open-source model, a larger installed base of downstream developers generates various forms of feedback that enhance the quality (or reduce the cost) of training the next-generation model ($q_{j',t+\tau}$). For instance, FM owners could incorporate high-quality synthetic data generated from successful downstream fine-tuning recipes and preference datasets, which often become public, when pre-training their next model. Also, the downstream community contributes to identifying edge-case failures and safety vulnerabilities, sometimes more efficiently than internal teams. In our structural model, θ_s captures the reduced-form efficiency in taking advantage of this feedback loop.

2.3.2 Frontier proximity penalty

Another key implication of our framework is that the incentive to increase openness (O_j) is non-monotonic in the quality gap from the technological frontier, defined as $\Delta_t = \bar{q}_t - q_{j,t}$, consistent with the mechanism in Habibi (2025). When a model lags behind the frontier, openness promotes downstream adoption and ecosystem expansion, thereby enhancing quality improvement and market penetration. However, as model quality approaches the frontier, the marginal benefit of openness may be limited because the opportunity cost of openness increases, as high-quality models can be monetized more effectively through proprietary deployment and API-based revenue.

We propose a *frontier proximity penalty* in the upstream developer profit function:

$$\max_{O_j} \Pi_{j,t}(\mathbf{q}_t, O_j, \mathbf{O}_{-j}^*(\mathbf{q}_t)) = \underbrace{R(\mathbf{O}_t, \mathbf{q}_t)}_{\text{Revenue}} - \underbrace{\left[\frac{\kappa_O}{2} O_j + \mu \cdot e^{-z(\bar{q}_t - q_{j,t})} \right] O_j}_{\text{Strategic Cost of Openness}} - \underbrace{C(q_{j,t})}_{\text{Production Cost}}, \quad (4)$$

where:

- $R(\mathbf{O}_t, \mathbf{q}_t) = \phi \Lambda_t s_{j,t}$ is an ecosystem value, with ϕ capturing the strategic value per adopter, capturing indirect monetization through ecosystem spillovers (e.g., selling complementary services) (Leisten, 2025), and Λ_t is the market.
- $\frac{\kappa_O}{2} O_j^2$ captures a baseline convex exposure cost of openness: as a firm releases more of the model's capabilities, the scope for imitation, appropriation, and unintended competitive spillovers rises at an increasing rate.
- $\mu e^{-z\Delta_t} O_j$ represents the frontier proximity penalty; as the gap narrows ($\Delta_t \rightarrow 0$), the term rises exponentially. If the gap falls below a critical implicit threshold (determined by parameters μ and z), the cost dominates the ecosystem benefits (ϕ), driving the optimal openness $O_{j,t}^*$ towards zero.
- $C(q_{j,t}) = \kappa q_{j,t}^2$ is the convex cost of producing a model with quality level $q_{j,t}$.

2.3.3 Foundation model owner's problem

The FM owner maximizes the discounted sum of profits. Since market shares $s_{j,t}$ depend on the openness levels of all firms through the denominator of the logit demand system, the profit function depends on the vector of industry openness. The openness decision at time t incurs an immediate cost (and opportunity cost) but enhances future model quality at $t + \tau$ through the quality evolution process described in (3). Consequently, upstream developers solve for a Markov Perfect Equilibrium conditional on competing owners' openness strategies. The value function for model j for owner u , given the state vector of qualities $\mathbf{q}_t$, is defined as:

$$V_j(\mathbf{q}_t) = \max_{O_j \in (0,1]} \left\{ \Pi_j(\mathbf{q}_t, O_j, \mathbf{O}_{-j}^*(\mathbf{q}_t)) + \delta^\tau \mathbb{E}_t [V_j(\mathbf{q}_{t+\tau}) \mid \mathbf{q}_t, O_j, \mathbf{O}_{-j}^*(\mathbf{q}_t)] \right\} \quad (5)$$

where $\mathbf{O}_{-j}^*(\mathbf{q}_t)$ denotes the vector of competitors' equilibrium openness strategies. The future state vector $\mathbf{q}_{t+\tau} = (q_{j'_u, t+\tau}, \mathbf{q}_{-j'_u, t+\tau})$ evolves according to the law of motion in (3), conditional on the current state $\mathbf{q}_t$, the developer's choice O_j , and competitors' choices $\mathbf{O}_{-j}^*$. The discount factor δ reflects time preferences over the development lag τ .

Taking the first-order condition of the Bellman equation (5) with respect to O_j , we obtain the necessary condition for optimality:

$$\frac{\partial \Pi_j}{\partial O_j} + \delta^\tau \mathbb{E}_t \left[\frac{\partial V_j(\mathbf{q}_{t+\tau})}{\partial q_{j'_u, t+\tau}} \frac{\partial q_{j'_u, t+\tau}}{\partial O_j} \right] = 0. \quad (6)$$

The effect of openness on future quality operates through downstream adoption. From the quality evolution equation (3), $\partial q'_{j_u,t+\tau}/\partial O_j = \theta_s \Lambda_t \cdot \partial s_{j,t}/\partial O_j$, so openness improves next-generation quality only insofar as it expands the current ecosystem from which the owner draws knowledge spillovers. Substituting this expression together with the logit share derivative and the linear approximation $E_t[\partial V_j/\partial q'_{j_u,t+\tau}] \approx \nu_q q'_{j_u,t+\tau}$, we rewrite the first-order condition (6) as:

$$\underbrace{\phi \cdot \Lambda_t s_{j,t} (1 - s_{j,t}) (\gamma q_{j,t} + \beta)}_{\text{Marginal Short-run Revenue } (X_{1j,t})} - \underbrace{(\kappa_O O_j + \mu \cdot e^{-z(\bar{q}_t - q_{j,t})})}_{\text{Marginal Cost } (Y_{j,t})} + \underbrace{\nu_q \cdot \delta^\tau \theta_s \Lambda_t \frac{\partial s_{j,t}}{\partial O_j} q'_{j_u,t+\tau}}_{\text{Marginal Long-run Revenue } (X_{2j,t})} = 0. \quad (7)$$

The three components have distinct economic interpretations. The first term captures marginal short-run revenue: greater openness shifts current downstream adoption, generating contemporaneous ecosystem value at rate ϕ per adopter.⁷ The second term captures the marginal cost of openness, combining a convex exposure cost with a frontier proximity penalty that rises as the quality gap narrows. The third term captures marginal long-run revenue: the discounted value of future quality improvement enabled by today's expanded ecosystem. This continuation-value term summarizes the dynamic benefits of openness, including future gains in complementary-service revenue, market share, and technological leadership.

In the baseline specification, we assume that complementary-service monetization does not directly depend on the contemporaneous openness level:

$$\frac{\partial \phi}{\partial O_j} = 0.$$

This restriction should be interpreted as a baseline abstraction rather than as a claim that openness can never affect complementary-service demand. The economic rationale is that the monetized complementary service layer—such as Qwen's managed deployment and API access through Alibaba Cloud Model Studio, DeepSeek's API platform, and Meta's Llama API developer platform—depends more naturally on model capability and enterprise readiness than on contemporaneous openness per se.⁸ We therefore abstract away from a direct current-period effect of openness on ϕ . Any indirect effect of openness on complementary-service revenue is instead captured dynamically through the continuation value: higher

⁷Examples include enterprise deployments (Mistral) and agent-building tools such as Llama API (i.e., a developer platform for Llama application development).

⁸Alibaba Cloud Model Studio provides official API access to Qwen models and enterprise-facing deployment support; DeepSeek operates its own API platform with model access and pricing; and Meta introduced Llama API as a developer platform in limited preview.

openness increases downstream adoption, strengthens knowledge spillovers, improves future model quality, and thereby expands future monetization opportunities.⁹

Since the observed history of competition in foundation models is short, we adopt a parsimonious approximation in which the expected marginal continuation value of quality is proportional to future quality:

$$\mathbb{E}_t \left[\frac{\partial V_j(\mathbf{q}_{t+\tau})}{\partial q_{j_u, t+\tau}} \right] \approx \nu_q q_{j_u, t+\tau},$$

where ν_q is a structural parameter to be estimated.¹⁰ If $\nu_q > 0$, the marginal continuation value of quality is increasing in model quality, consistent with a convex continuation value and with a superstar-type mechanism in which technological leadership near the frontier yields disproportionately large dynamic rewards (Rosen, 1981). We can then collapse equation (7) into the linear decomposition

$$Y_{j,t} = \phi X_{1j,t} + \nu_q X_{2j,t}, \tag{8}$$

where $Y_{j,t}$ denotes marginal static cost, and ϕ and ν_q are estimated using the data.

3 Data

This section describes the data used to estimate the model parameters. Our primary data source is the Hugging Face platform, from which we collect monthly data at the foundation-model level on the number of downstream developers (also known as children) and new developers. The *N of Children* measures the number of developers who built upon a given foundation model in a given month, and therefore serves as a direct proxy for monthly adoption demand. We also collect model-specific characteristics—including parameter size, an intelligence benchmark index, API price, and model type (e.g., instruct or reasoning)—from *Artificial Analysis*.¹¹

Because no standardized metric exists to quantify the degree of openness of foundation models, we construct a composite openness index using a structured scoring procedure based on official model cards provided by the developers on Hugging Face or GitHub. For each

⁹In the same spirit, Meta’s strategy on allowing a free use of Llama for developers relies on yielding return to Meta in the form of *innovative products*. See Reuters (2025) for details.

¹⁰A higher openness level today facilitates quality improvement in period $t + \tau$, which may also imply higher future production or training costs. These future costs are internalized in the continuation value function $V(\cdot)$. Accordingly, ν_q captures the net marginal continuation value of quality improvement, rather than gross future revenue alone.

¹¹<https://artificialanalysis.ai/>

Table 1: Summary Statistics for Estimation Samples

Statistic	N	Mean	St. Dev.	Min	Median	Max
<i>Panel A: Demand-side Variables</i>						
N of Children	475	22.644	47.299	1	8	534
Intelligence Index	475	0.211	0.143	0.010	0.169	0.575
OpenIndex	475	0.758	0.174	0.166	0.867	0.916
API Price	475	0.783	1.140	0.050	0.250	4.190
Instruct	475	0.305	0.461	0	0	1
<i>Panel B: Supply-side Variables</i>						
Model Parameters	70	62.826	100.303	0.600	–	480.000
Ecosystem	82	24.439	34.945	1	–	178
Intelligence Index	62	0.246	0.154	0.010	–	0.575

Notes: Panel A reports summary statistics for the estimation sample used in the downstream adoption regression. Panel B reports summary statistics for the estimation sample used in the quality evolution regression.

foundation model j , we define a normalized openness index, $\text{OpenIndex}_j \in (0, 1]$, which captures the extent to which the model can be independently accessed, audited, adapted, and deployed by third parties. We evaluate six dimensions of openness, each scored as $s_{d,j} \in \{0, 0.5, 1\}$ for $d \in \{\text{weights, license, data, architecture, finetune, inference}\}$, where higher values indicate fewer restrictions and greater transparency. The final index is the arithmetic mean of these scores: $\text{OpenIndex}_j = \frac{1}{6} \sum_d s_{d,j}$.

The six dimensions are defined as follows. Weight openness measures the public availability of pretrained model weights. License openness reflects the permissiveness of the governing license, distinguishing between fully open licenses (e.g., Apache 2.0, MIT) and restrictive commercial or non-commercial alternatives. Data disclosure captures transparency regarding pre-training and post-training datasets. Architecture openness evaluates the disclosure of model design and training recipes. Fine-tuning and inference openness reflect the technical and legal feasibility of adapting and deploying the model on independent infrastructure. This framework aligns with the Model Openness Framework (MOF) classifications proposed by White et al. (2024). Detailed scoring criteria are provided in Appendix A.

We construct the estimation sample by merging three sources: monthly adoption counts (number of children and new children) collected from Hugging Face at the foundation-model level; model-level attributes, including parameter size, intelligence benchmark index, API price, and model type, drawn from *Artificial Analysis*; and the openness index scored manu-

ally from official model cards following the rubric described above. The merged panel covers 95 unique open-source foundation models from five owners (Meta, Mistral, DeepSeek, Qwen, and OpenAI) over July 2023 to October 2025. Because model attributes are time-invariant, within-model variation comes exclusively from adoption counts. The demand-side estimation uses 475 model-month observations with complete covariates; the quality evolution regression uses 62 successive-generation pairs for which both current and next-generation quality can be matched; and the supply-side FOC decomposition uses 324 interior observations where calibrated openness lies strictly between the bounds. Table 1 reports summary statistics.

4 Estimation

4.1 Empirical Model

In this section, we present our estimation strategy for identifying demand- and supply-side parameters and the results. We first estimate the downstream adoption parameters in equation (1). In specific, our estimating equation is as follows:

$$\begin{aligned} \ln(s_{jt}) - \ln(s_{0t}) = & Aq_j + \beta O_j + \gamma q_j \times O_j + \alpha p_j + \lambda_1 \text{Instruct}_j + \lambda_2 \text{Children}_{j,t-1} \quad (9) \\ & + \xi_{jt} + \delta_u + \delta_t + \varepsilon_{ij}. \end{aligned}$$

On the left-hand side, we employ the monthly number of children (i.e., downstream developers) using the foundation model j in month t to compute the market share s_{jt} . s_{0t} captures the outside options available for the downstream deployers who may not choose to fine-tune, or rather use closed models.¹² On the right-hand side, we include the quality variable (q_j) measured by the time-invariant intelligence index, the openness level (O_j), their interaction term, and model-specific covariates such as API price (p_j), an indicator for whether the model is instruction-oriented (Instruct_j), and the number of children last month ($\text{Children}_{j,t-1}$) to capture the ecosystem effect.¹³ We also capture model owner and year-month fixed effects by δ_u and δ_t , respectively. We then estimate the dynamic quality evolution equation (3).

¹²We estimate a conditional logit model, as we define the market size as the total number of children $\times$ 1.5.

¹³We get similar estimates for other variables without $\text{Children}_{j,t-1}$ in the unreported regression result.

4.2 Identification

We address the identification challenge in this section before we move on to the results table. Identifying the causal impact of openness on market performance requires addressing the potential endogeneity of the firm’s decision. In our demand specification equation (9), we explicitly control for the fundamental quality of the model (q_j) using observable benchmark scores. Consequently, the structural error term, ξ_{jt} , represents unobserved, transient shocks to model demand, such as temporary fluctuations in user sentiment, viral trends, or idiosyncratic preference shifts that are not captured by objective quality metrics.

The primary concern is that these unobservable demand shocks (ξ_{jt}) might be correlated with the owner’s decision regarding the degree of openness (O_j), the parameter of our interest. For instance, a firm facing a negative demand shock might strategically increase openness level to regain attention. However, we argue that the decision on openness is predetermined relative to these temporary shocks for two reasons.

First, the development and safety alignment of foundation models require a substantial lead time. The strategic choice of O_j is therefore made well in advance of the realization of ξ_{jt} , during a period when the market for any given model has yet to materialize. Given the rapid and largely unpredictable evolution of the foundation model market since 2023—marked by sudden capability jumps, the entry of new competitors, and volatile developer sentiment—owners face considerable uncertainty about future demand at the time of the openness decision. Under such conditions, O_j is more plausibly driven by the firm’s long-term architectural and monetization roadmap than by anticipatory responses to unforeseen demand. Second, openness is effectively irreversible once implemented. Releasing model weights publicly cannot be undone, as the weights immediately become freely reproducible and redistributable by third parties. This irreversibility implies that owners cannot treat openness as a flexible instrument to be adjusted in response to realized or anticipated demand shocks, which further supports its treatment as a predetermined variable in our framework. Short-term idiosyncratic shocks, by construction, are therefore orthogonal to strategic openness decisions.

We further mitigate endogeneity concerns by incorporating a robust set of fixed effects. The inclusion of the quality index and owner fixed effects absorbs an important share of persistent differences in model capability and provider-specific unobserved quality, while year-month fixed effects capture common cost shocks and market-wide pricing movements. In a similar spirit, we treat p_j as conditionally exogenous to the idiosyncratic demand shock ξ_{jt} after controlling for observed model quality and fixed effects. The identifying intuition is that inference access costs are largely determined by supply-side considerations—such as inference compute requirements, hosting costs, and broader pricing benchmarks in model-

serving markets—rather than being adjusted in response to temporary adoption shocks in the downstream market. This argument is particularly relevant in the foundation-model setting, where access prices often reflect not only the model developer’s pricing choice but also external hosting and deployment channels. Accordingly, our estimate of α should be interpreted under a conditional exogeneity assumption.

4.3 Results

The first column in Table 2 reports the estimates from the downstream adoption equation (9). As expected, higher model quality and openness induce greater downstream adoption ($A > 0$ and $\beta > 0$). $\beta > 0$ also implies that a net marginal benefit of a higher degree of openness on downstream profit is positive. We also see that the API price coefficient (α) is negative, consistent with our assumption that API price represents the cost-side of running model j .

Interestingly, the interaction coefficient, γ , is estimated to be negative (-14.923), implying that the marginal effect of openness on downstream adoption declines with model quality. Formally,

$$\frac{\partial \ln(s_{jt}/s_{0t})}{\partial O_j} = \beta + \gamma q_j = 6.453 - 14.923 \cdot q_j,$$

so the adoption benefit of greater openness becomes smaller as q_j increases.¹⁴ In our dataset, the maximum value of q_j is 0.575 while its median is 0.169 (Table 1). Thus, though most of the open-source FM owners expect an increase in market share in O_j , some owners of relatively high-quality FMs would expect lower market share with higher O_j .

One economic interpretation is that deployment constraints reduce the effective benefit of openness for high-quality models. While lower-quality open-source models can often be deployed and fine-tuned locally at relatively low cost, higher-quality models require substantially greater computational resources, making local deployment infeasible for many downstream users. This shifts deployment toward cloud-based inference and increases effective usage costs, reducing the marginal demand response to openness. This mechanism, however, should be interpreted as a reduced-form explanation of the negative interaction term, rather than as a directly identified causal channel.

This result alone may not imply that highly open models are chosen primarily for long-run strategic reasons. In fact, because $\beta > 0$ and $\gamma < 0$, the static adoption incentive for openness

¹⁴Symmetrically,

$$\frac{\partial \ln(s_{jt}/s_{0t})}{\partial q_j} = A + \gamma O_j,$$

implying that the marginal effect of quality is weaker for more opened models.

Table 2: Demand and Quality Evolution Estimates

Model:		Downstream Adoption	Quality Evolution
		(1)	(2)
<u>Variables</u>	<u>Parameters</u>		
q_j (Intelligence)	A	11.111*** (2.599)	
O_j (Openness)	β	6.453*** (1.064)	
$q_j \times O_j$	γ	-14.923*** (2.883)	
p_j (API Price)	α	-0.1288** (0.0425)	
$Instruct_j$	λ_1	0.3501* (0.1303)	
$\ln(\text{Children}_{j,t-1})$	λ_2	0.0676 (0.0460)	
$q_{j_u,t}$ (Persistence)	ρ		0.5137*** (0.0246)
$\text{Ecosystem}_{j_u,t}$	θ_s		0.0005* (0.0002)
$\ln(\text{Size}_{j'_u,t+\tau})$	θ_p		0.0335** (0.0077)
Owner fixed effects		Yes	Yes
Year-Month fixed effects		Yes	—
Observations		475	62

Notes: Standard errors clustered by month in (1) and by owner in (2) are reported in parentheses. In column (1), the dependent variable is $\log(s_{jt}) - \log(s_{0t})$. * $p < 0.1$; ** $p < 0.05$; *** $p < 0.01$.

can be strongest for lower-quality entrants. Accordingly, evidence for dynamic motives—such as future ecosystem rents, quality improvement, or standard-setting advantages—should come from the continuation-value component of the structural model, rather than from the demand estimates alone.

The second column shows supply-side parameters estimated with the dynamic quality evolution equation (3). Notably, the estimated knowledge spillover parameter, $\hat{\theta}_s = 0.0005$, implies that having a hundred more downstream developers is associated with an approximate 0.05 unit increase in the next generation model's quality ($q_j \in [0, 1]$), on average. Given the rapid exponential growth of the developer ecosystem, this spillover effect represents a substantial contribution to quality improvement. We use these estimated parameters to back

out the supply-side monetization parameters, ϕ and ν_q .

4.4 Simulation

In this section, we evaluate whether the model can recover the true monetization parameters from simulated data. We generate panel data from the structural environment described above, incorporating the quadratic openness cost and the dynamic continuation value. In each simulated market-period, firms choose openness by maximizing the model-implied profit function over the feasible interval $O_{j,t} \in (0, 1]$. Using the simulated equilibrium choices, we then construct the marginal openness cost and estimate the monetization coefficients from

$$Y_{j,t} = \phi X_{1j,t} + \nu_q X_{2j,t}.$$

Because this estimating equation is derived from the first-order condition, parameter recovery is based on interior observations. Our primary goal is to verify that the procedure accurately recovers the true monetization parameter ϕ used in the data-generating process, while also assessing recovery of ν_q in the joint specification. We repeat this exercise across Monte Carlo replications ($ns = 100$) and summarize the mean estimate, bias, and root mean squared error of the recovered parameters in Table 3. We note that the estimated $\hat{\phi}$ and $\hat{\nu}_q$ are closer to the true parameters, validating the functionality of our model.

4.5 Monetization Parameters

We now turn to recovering the monetization parameters that rationalize openness choices in the data. Identification of long-run revenue components typically requires long time-series variation, but the foundation model industry exhibits unusually rapid strategic cycles: major model generations often arrive within roughly 6 to 9 months. As a result, our 28-month panel spans multiple technological generations, generating meaningful variation in openness choices, quality gaps relative to the frontier, and subsequent model quality. This variation allows us to decompose the marginal return to openness into short-run and long-run monetization components.

4.5.1 Two-step recovery process

Our empirical strategy proceeds in two steps. In the first step, we calibrate the marginal cost schedule of openness using the observed openness data while maintaining the normalization $\phi = 2.5$. As described in Appendix B, this step should be interpreted as a conditional cost-side calibration rather than as a full-information estimator of the dynamic openness

Table 3: Data-Generating Process and Monte Carlo Recovery of Structural Parameters

Category	Parameter	Description	DGP Value	Recovered Value
<i>Panel A: Fixed Estimated Parameters</i>				
Demand Side	A	Quality Effect on Adoption	11.111	—
	β	Marginal Effect of Openness on Adoption	6.453	—
	α	Price Sensitivity	-0.1288	—
Quality Dynamics	γ	Quality–Openness Interaction	-14.923	—
	ρ	Quality Persistence	0.5137	—
	θ_s	Knowledge Spillover Parameter	0.0005	—
<i>Panel B: Simulation Parameters</i>				
Cost Structure	κ_O	Convex Openness Cost Parameter	100	—
	μ	Frontier-Proximity Penalty	5,000	—
	z	Decay Rate of Frontier Penalty	20.0	—
	Λ	Market Size	300	—
Dynamics	δ	Discount Factor	0.95	—
	σ	Innovation Shock Standard Deviation	0.0673	—
	ν_q	Continuation-Value Coefficient	50,000	50,011
	ϕ	Monetization per Adopter	2.500	2.498
Initial Conditions	$q_{j,0}$	Initial Model Quality	$\sim U(0.2, 0.4)$	—
	$p_{j,t}$	API Price	$\sim U(0.5, 3.0)$	—
<i>Panel C: Monte Carlo Recovery Performance</i>				
Recovery of ϕ	$\mathbb{E}[\hat{\phi}]$	Mean Estimate Across Replications	2.500	2.498
	$\text{RMSE}(\hat{\phi})$	Root Mean Squared Error	—	0.0099
Recovery of ν_q	$\mathbb{E}[\hat{\nu}_q]$	Mean Estimate Across Replications	50,000	50,011
	$\text{RMSE}(\hat{\nu}_q)$	Root Mean Squared Error	—	71.92

Notes: This table reports the parameterization of the data-generating process (DGP) and the Monte Carlo recovery performance of the structural parameters under the quadratic openness-cost specification. Panel A shows the estimated demand and quality-dynamics parameters that are reported in Table 2, and Panel B illustrates other baseline parameters that are chosen for the simulation. With these parameters fixed, we generate firms’ openness choices over $T = 50$ periods with $N = 100$ firms in each period. In each firm-period observation, openness is chosen on a discrete grid over $(0, 1]$ to maximize the model-implied profit function. Parameter recovery is based on interior observations only, defined as $O_{j,t} \in (0.02, 0.98)$, since the first-order condition need not hold at corner solutions. Reported recovery statistics are computed from 100 Monte Carlo replications.

problem. Specifically, we solve a static approximation of the openness choice problem at the observation level and choose the cost parameters to minimize the distance between observed openness and model-implied openness.

This first step yields $\hat{\kappa}_O = 105.41$ and $\hat{\mu} \approx 0$, implying that the empirical cost schedule is driven primarily by the convex openness-cost term rather than by an additional frontier-proximity penalty. We may not interpret the near-zero estimate of μ literally as evidence that frontier proximity imposes no strategic cost. In our maintained specification, the demand

system already allows the marginal revenue from openness to decline with model quality through the negative interaction term $q_{j,t} \times O_{j,t}$ in equation (9). Since models closer to the frontier also tend to be higher-quality models, this interaction absorbs much of the same empirical variation that would otherwise help identify an additional frontier-specific penalty. Thus, the estimate $\hat{\mu} \approx 0$ is more appropriately viewed as reflecting limited separate identifying power for μ once the quality–openness interaction is included, rather than a literal absence of frontier-related strategic cost.

In the second step, we treat the calibrated cost parameters as given and estimate the decomposition implied by the marginal openness condition:

$$Y_{j,t} = \phi X_{1j,t} + \nu_q X_{2j,t},$$

where $Y_{j,t}$ is the calibrated marginal cost of openness, $X_{1j,t}$ is the short-run marginal revenue component, and $X_{2j,t}$ is the long-run marginal revenue component. Separate recovery of ϕ and ν_q comes from the fact that the two components in the marginal openness condition are not mechanically proportional in the data. The short-run component ($X_{1j,t}$) varies with the contemporaneous sensitivity of downstream adoption to openness, whereas the continuation-value component ($X_{2j,t}$) additionally depends on future model quality. Because future quality and current adoption sensitivity do not move one-for-one across model-period observations, the data generate independent variation in the two components.

Table 4: Estimated Openness Marginal Revenue Parameters

	Parameter	Estimate
$X_{1j,t}$ (Marginal Short-run Revenue)	ϕ	0.0303 (0.0163)
$X_{2j,t}$ (Marginal Long-run Revenue)	ν_q	1,724.1** (209.3)
Observations		324

Notes: This table reports the second-step estimates from the FOC decomposition

$$Y_{j,t} = \phi X_{1j,t} + \nu_q X_{2j,t}.$$

Standard errors clustered by owner are reported in parentheses. The first-step calibration is conducted under the normalization $\phi = 2.5$, and yields calibrated cost parameters $\hat{\kappa}_O = 105.41$ and $\hat{\mu} \approx 0$. Then we estimate payoff parameters, ϕ and ν_q , with the cost parameters given. More details are in Appendix B. * $p < 0.10$, ** $p < 0.05$, and *** $p < 0.01$.

4.5.2 Results

Table 4 reports the resulting monetization estimates. The estimates reveal a pronounced asymmetry between the short-run and long-run components. The short-run monetization coefficient is small and statistically imprecise, whereas the long-run monetization coefficient is positive and statistically significant. This pattern suggests that openness in the data is not primarily rationalized by contemporaneous monetization gains. Instead, the dominant force behind openness is the continuation value associated with future quality improvement and ecosystem expansion.¹⁵ This result is consistent with the model’s interpretation of $\nu_q > 0$ as a reduced-form winner-takes-all mechanism. Openness can expand downstream adoption and ecosystem activity, which then improves future model quality through knowledge spillovers and raises continuation value. Accordingly, even when the contemporaneous marginal return to openness is limited, firms may still optimally choose high openness if the expected longer-run return to ecosystem formation is sufficiently large.

4.5.3 Short-run monetization parameter across owners

To examine heterogeneity in the short-run monetization parameter, we plug the estimated long-run revenue coefficient $\hat{\nu}_q$ back into the second-step decomposition and recover owner-specific short-run revenue parameters from

$$\tilde{Y}_{j,t} = \phi_j X_{1j,t} + \epsilon_{j,t}, \quad (10)$$

where

$$\tilde{Y}_{j,t} = Y_{j,t} - \hat{\nu}_q X_{2j,t}.$$

We estimate equation (10) separately by owner, with the standard errors computed using a bootstrap procedure with 1,000 replications.

Table 5 reports the resulting owner-specific short-run monetization parameters. The estimates reveal substantial heterogeneity. Among the four major model owners, only Qwen exhibits a statistically significant positive coefficient. By contrast, the coefficients for DeepSeek, Meta, and Mistral are statistically indistinguishable from zero. This pattern suggests that Qwen’s openness strategy is more closely tied to immediate market-penetration incentives, whereas the openness choices of Meta, DeepSeek, and Mistral are not well explained by contemporaneous short-run revenue gains alone.

¹⁵Although one could in principle iterate on ϕ by feeding the second-step estimate back into the first-step calibration until convergence (e.g., Singh et al. (2024)), in our data this fixed-point procedure is numerically unstable. We therefore adopt the more transparent conditional two-step procedure as our baseline. We delineate more details in the Appendix B.

Table 5: Owner-Specific Short-run Monetization Parameters

Owner	Estimate ($\hat{\phi}_j$)	Boot. SE	95% CI	Z-stat	P-value
Qwen	0.622***	0.114	[0.422, 0.889]	5.45	0.000
DeepSeek	0.379	0.493	[-0.191, 1.700]	0.77	0.441
Meta	0.069	0.066	[-0.032, 0.224]	1.03	0.301
Mistral	0.031	0.078	[-0.042, 0.131]	0.40	0.693

Notes: This table reports owner-specific estimates of the short-run revenue parameter ϕ_j , obtained from

$$\tilde{Y}_{j,t} = \phi_j X_{1j,t} + \epsilon_{j,t}, \quad \tilde{Y}_{j,t} = Y_{j,t} - \hat{\nu}_q X_{2j,t}.$$

Standard errors are computed using a bootstrap procedure with 1,000 replications to account for the generated-regressor problem induced by the first-stage estimate of $\hat{\nu}_q$. The confidence intervals are percentile bootstrap intervals. * $p < 0.1$; ** $p < 0.05$; *** $p < 0.01$.

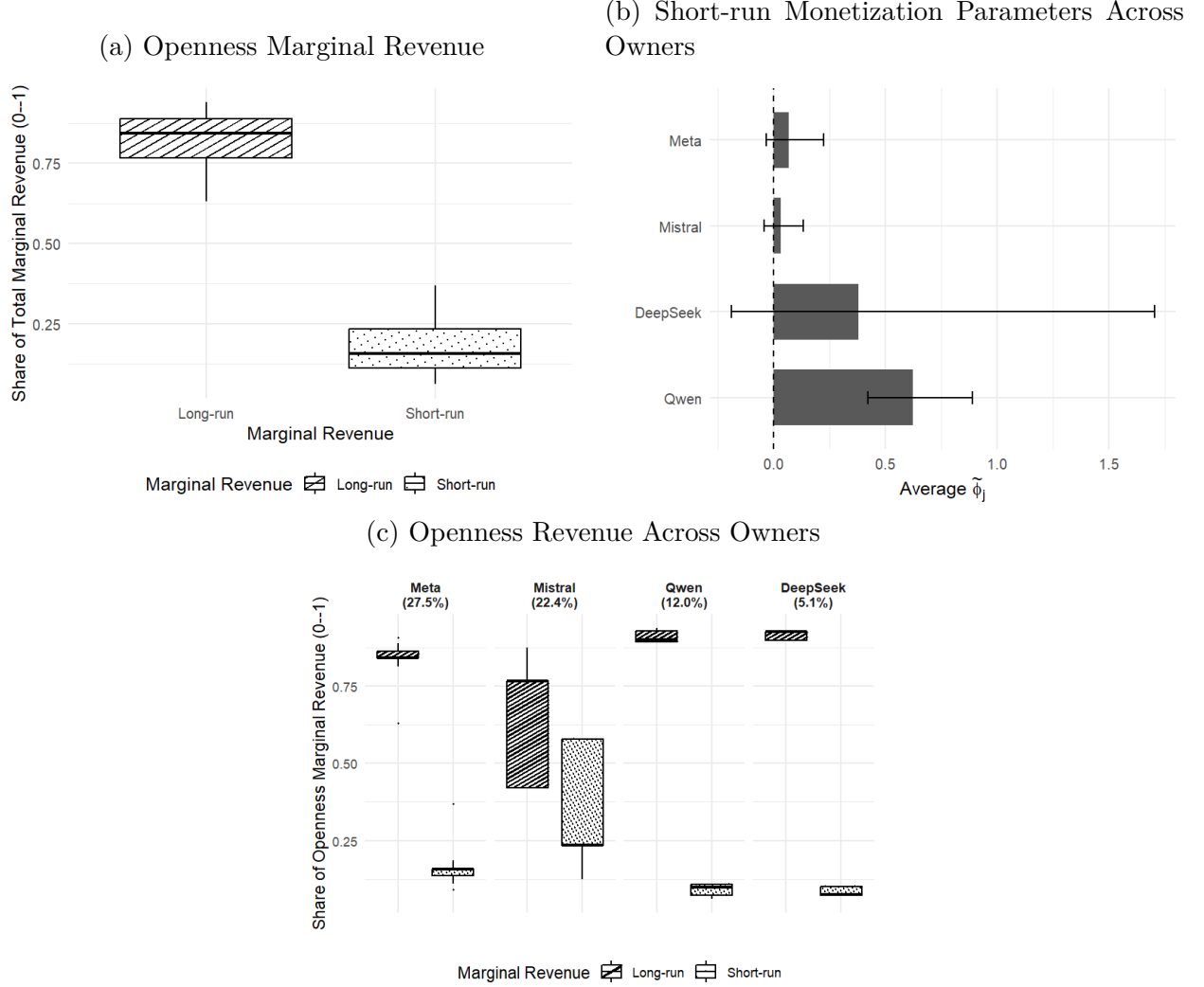
This interpretation is also consistent with the broader structure of the model. At the aggregate level, the evidence points to future ecosystem value as the dominant driver of openness. At the owner level, however, firms differ in the extent to which current monetization matters at the margin. Qwen stands out as the only owner for which the static component is estimated precisely and positively, whereas the remaining major owners appear to rely more heavily on future ecosystem value.

Figure 2 summarizes the implied monetization decomposition. Panel (2a) shows that the long-run component accounts for the majority of total openness revenue in the aggregate distribution. Panel (2b) reports the owner-specific short-run monetization coefficients, showing that Qwen is the only major owner with a precisely estimated positive short-run component. Panel (2c) illustrates the distribution of the revenue shares across owners. Meta and DeepSeek appear strongly future-oriented, while Mistral exhibits a more balanced revenue composition even though its short-run coefficient is estimated imprecisely. Qwen combines a statistically significant short-run monetization coefficient with a revenue composition that remains tilted toward future value, indicating that its openness strategy reflects both short-run monetization incentives and long-run ecosystem considerations.

Figure 3 further examines how openness varies with the underlying revenue structure as models move closer to the frontier. Panels (3a) and (3b) show that openness is generally higher for models with larger revenue components, with the relationship appearing slightly stronger for long-run revenue than for short-run revenue. This pattern is consistent with the regression results, which indicate that long-run revenue is the more important source of variation in openness.

The relationship is also state-dependent. Differences in openness across revenue groups

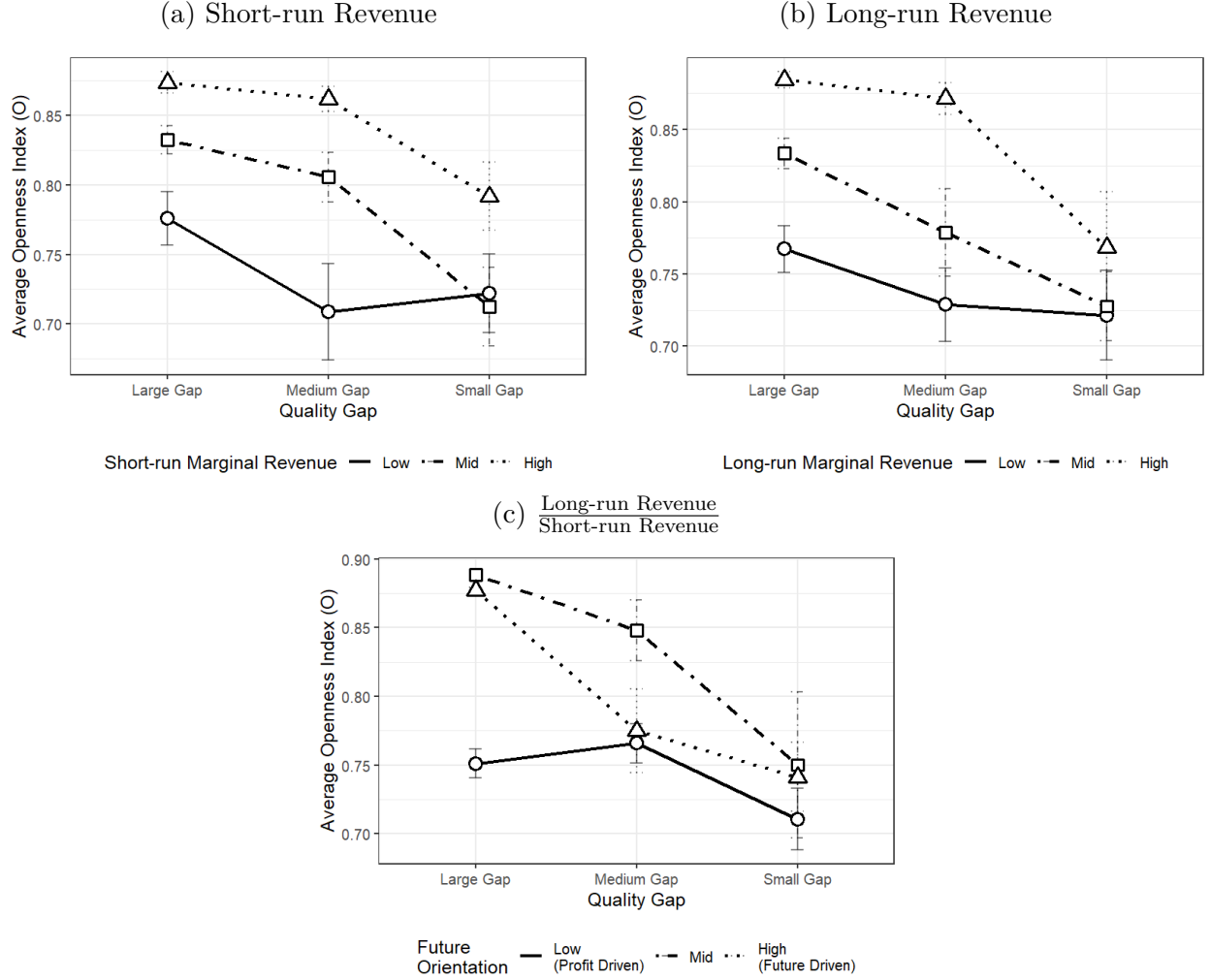
Figure 2: Estimated Monetization Parameters



Notes: We calibrate the openness cost parameters in a first step under the maintained normalization $\phi = 2.5$, which yields $\hat{\kappa}_O = 105.41$ and $\hat{\mu} \approx 0$. Panel (a) reports the distribution of the long-run and short-run shares of total openness revenue. Panel (b) reports the owner-specific short-run monetization coefficients $\hat{\phi}_j$ with 95% bootstrap confidence intervals. Panel (c) reports the distribution of payoff shares by owner. Details of the first-step cost calibration and second-step payoff recovery are provided in Appendix B.

are largest when the quality gap from the frontier FM is wide, suggesting that revenue heterogeneity matters most for the Large Gap group. As models approach the frontier, openness declines overall and the differences across groups become smaller. Panel (3c) further shows that a high degree of future orientation alone does not necessarily imply higher openness near the frontier. The mid-orientation group overall shows higher degree of openness compared to the high-orientation group, somewhat different patterns from what we observe in other panels.

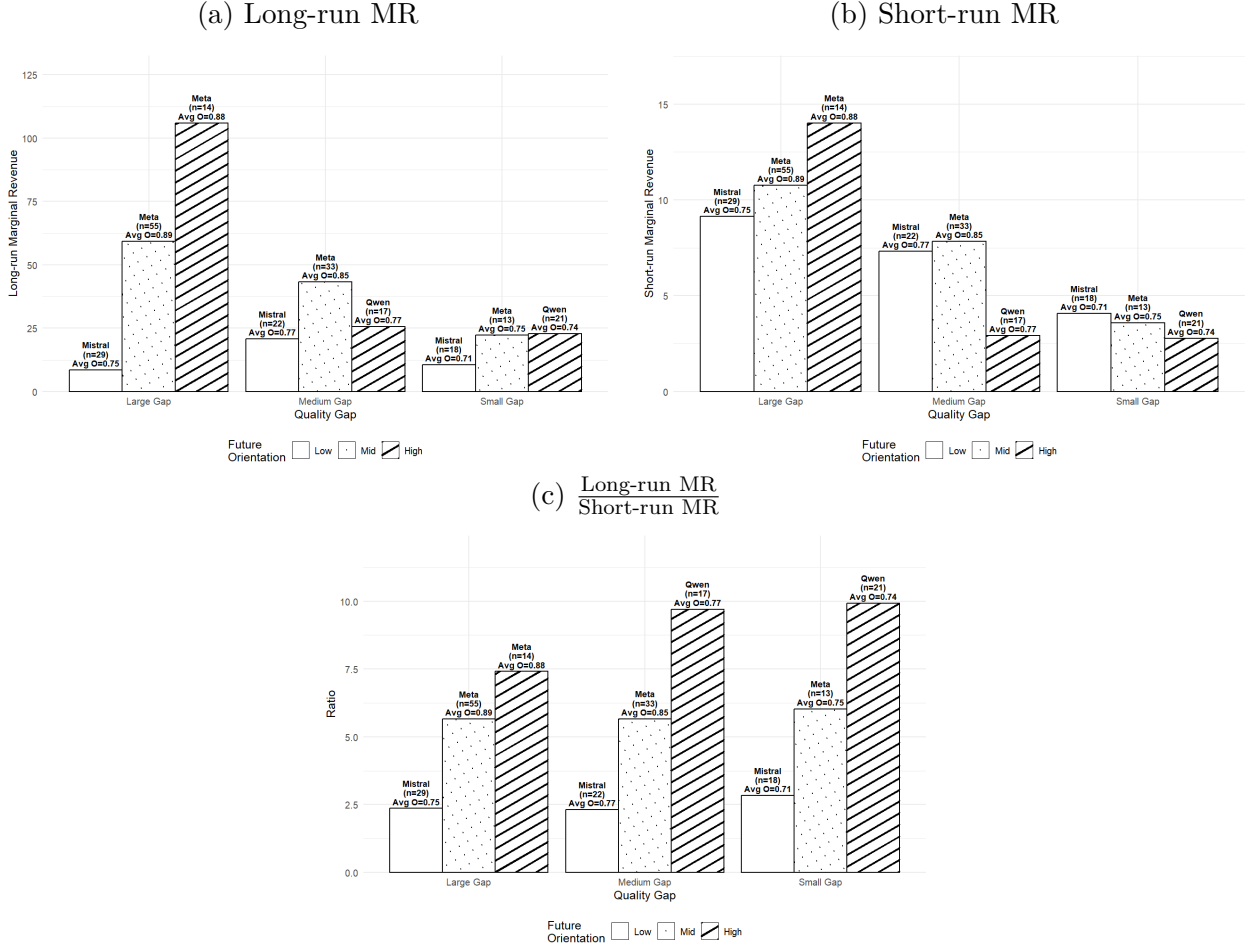
Figure 3: Openness Across Revenue Structures and the Quality Gap



Notes: The figure documents heterogeneity in openness across alternative revenue structures. Observations are partitioned into tertiles based on (a) the estimated short-run revenue, $|\hat{\phi}X_{1,j,t}|$, (b) the estimated long-run revenue, $|\hat{\nu}_qX_{2,j,t}|$, and (c) the future-orientation ratio, $|\hat{\nu}_qX_{2,j,t}|/|\hat{\phi}X_{1,j,t}|$. The horizontal axis denotes the quality gap relative to the frontier model, ordered from a large gap to a small gap. Each marker represents the group mean of the Openness Index, and vertical bars denote standard errors.

Figure 4 clarifies the mechanism behind these patterns by decomposing revenues into long-run and short-run levels. A central implication is that openness is not determined by future orientation alone, but by the interaction between future orientation and the scale of absolute revenues. When the quality gap is large, the highly future-oriented group combines strong relative orientation with substantial long-run revenue, and this is associated with high openness. As the quality gap narrows, however, both long-run and short-run revenues decline in absolute terms. In the medium- and small-gap regions, the highly future-oriented group is increasingly dominated by challengers such as Qwen, whose high orientation ratios

Figure 4: Marginal Revenue (MR) Across the Quality Gap



Notes: The figure illustrates heterogeneity in marginal revenue structures across different levels of the quality gap relative to the frontier model. Panels (a), (b), and (c) report group means of the absolute long-run MR, the absolute short-run MR, and the future-orientation ratio (long-run MR/short-run MR), respectively. Observations are first partitioned into three groups according to the quality gap—Large, Medium, and Small. Within each gap category, observations are further divided into tertiles based on the future-orientation ratio. The height of each bar represents the average value of the corresponding metric within each subgroup. Labels above the bars indicate the dominant owner within the subgroup, together with the number of models (n) and the average openness level (Avg O).

arise mainly because their short-run revenues are small. Yet their absolute long-run revenues remain below the levels observed for Meta-dominated groups in the large-gap region. This explains why a high future-orientation ratio does not mechanically translate into greater openness near the frontier.

More broadly, Figure 4 suggests that openness is strongest when future-oriented incentives are supported by sufficiently large long-run revenue bases. Relative orientation matters, but only insofar as it operates through scale. This interpretation helps reconcile

the aggregate and owner-level results: long-run revenue is the main driver of openness in the aggregate, while heterogeneity in short-run monetization still shapes how different firms implement open-source strategies.

5 Conclusion

This paper sheds light on the economic forces underlying developers’ strategic decisions on the open-source foundation models by disentangling their short-run monetization and continuation value parameters. Our findings indicate that, on average, firms place at least as much weight on the long-run value generated by openness as on contemporaneous revenue. In particular, newer entrants such as Qwen and DeepSeek appear more future-oriented, willingly incurring the short-run costs of openness in exchange for accelerated ecosystem growth and long-run competitive advantages. By contrast, Mistral, an established developer, balances current monetization with ecosystem considerations, reflecting heterogeneity in strategic positioning across firms.

We show that the long-run revenue of openness is rooted in the ecosystem formed around downstream deployers. By releasing open-source FMs, developers enable a large pool of downstream players to fine-tune models and develop applications at negligible marginal cost. Although these activities do not directly augment the quality of the current model, they generate valuable knowledge spillovers—such as fine-tuning data, usage patterns, and application-level feedback—that can be internalized by the model owner when developing subsequent generations, providing a trajectory for innovation. Higher-quality future models, in turn, attract additional downstream adopters, reinforcing ecosystem expansion and raising future revenue. This self-reinforcing mechanism creates a virtuous cycle in which openness today shapes innovation and market share tomorrow.

Taken together, the results underscore that the economic rationale for open-source strategies cannot be understood through a purely static lens. While openness erodes short-run competitive advantages and intensifies imitation risks, especially for frontier models, dynamic considerations—driven by ecosystem formation and winner-takes-all forces in the race toward advanced AI capabilities—substantially offset these costs. Our analysis provides evidence that open-source foundation models could be viewed as intertemporal strategic investments rather than a short-run monetization paradigm.

There are caveats to discuss. First, due to limited data, we assume a parsimonious functional form for the value function. We leave it for the future agenda with a more longitudinal dataset, and the continuation value parameter would be more precisely identified. Second, we are examining competition among open-source foundation models only, while

competition between closed models (e.g., Claude and GPTs) is also important. We partially capture this by including the effect of the gap between the top-performing model’s quality on the strategic cost of openness. However, due to limited data, this paper instead focuses on the economic incentives of open-source model owners, which are absent for closed model owners. In a similar vein, integrating the entry model of foundation model owners would be an interesting extension.

References

- Aghion, P. and Howitt, P. (1990). A model of growth through creative destruction. NBER Working Paper.
- Azoulay, P., Krieger, J., and Nagaraj, A. (2025). Old moats for new models: Openness, control, and competition in generative artificial intelligence. Entrepreneurship and Innovation Policy and the Economy, 4(1):7–46.
- Bajari, P., Benkard, C. L., and Levin, J. (2007). Estimating dynamic models of imperfect competition. Econometrica, 75(5):1331–1370.
- Chui, M., Hazan, E., Roberts, R., Singla, A., and Smaje, K. (2023). The economic potential of generative ai. McKinsey Report.
- Eastwood, B. (2026). Ai open models have benefits. so why aren’t they more widely used? <https://mitsloan.mit.edu>. Accessed March 12, 2026.
- Fosfuri, A., Giarratana, M. S., and Luzzi, A. (2008). The penguin has entered the building: The commercialization of open source software products. Organization Science, 19(2):292–305.
- Habibi, M. (2025). Open sourcing gpts: Economics of open sourcing advanced ai models. arXiv preprint arXiv:2501.11581.
- Haese, J. and Peukert, C. (2025). Open at the core: Moving from proprietary technology to building a product on open source software. Management Science.
- Hermansen, Cailean and Osborne, Michael (2025). The economic and workforce impacts of open source ai. Linux Foundation. Accessed: 2026-01-08.
- Hughes, T. P. et al. (1987). The evolution of large technological systems. The Social Construction of Technological Systems, 82:51–82.
- Iansiti, M. and Richards, G. L. (2006). The Business of Free Software: Enterprise Incentives, Investment, and Motivation in the Open Source Community. Harvard Business School.
- Leisten, M. (2025). Open (?) ai. SSRN Working Paper.
- Lerner, J. and Tirole, J. (2002). Some simple economics of open source. Journal of Industrial Economics, 50(2):197–234.

- Mokyr, J. (1990). Punctuated equilibria and technological progress. American Economic Review, 80(2):350–354.
- Mokyr, J. (1992). The Lever of Riches: Technological Creativity and Economic Progress. Oxford University Press.
- Nagle, F. (2019). Open source software and firm productivity. Management Science, 65(3):1191–1215.
- Pakes, A., Ostrovsky, M., and Berry, S. (2007). Simple estimators for the parameters of discrete dynamic games. RAND Journal of Economics, 38(2):373–399.
- Petralia, S. (2025). Open source software as digital platforms to innovate. Journal of Economic Behavior & Organization, 237:107109.
- Romer, P. M. (1990). Endogenous technological change. Journal of Political Economy, 98(5):S71–S102.
- Rosen, S. (1981). The economics of superstars. American Economic Review, 71(5):845–858.
- Schumpeter, J. A. (1942). Creative destruction.
- Seim, K. (2006). An empirical model of firm entry with endogenous product-type choices. RAND Journal of Economics, 37(3):619–640.
- Singh, A., Hosanagar, K., and Nevo, A. (2024). Network externalities and app store fees in mobile platforms. SSRN Working Paper.
- Utterback, J. (1994). Mastering the dynamics of innovation: How companies can seize opportunities in the face of technological change. Historical Research Reference in Entrepreneurship.
- Varian, H. R. and Shapiro, C. (2003). Linux adoption in the public sector: An economic analysis. Manuscript.
- Vaswani, A. et al. (2017). Attention is all you need. Advances in Neural Information Processing Systems, 30.
- Von Hippel, E. (1988). The sources of innovation. In The Sources of Innovation, pages 111–120.

White, M., Osborne, C., Abdelmonsef, A., et al. (2024). The model openness framework: Promoting completeness and openness for reproducibility, transparency, and usability in ai. [arXiv:2403.13784](#).

Xu, F., Wang, X., Chen, W., and Xie, K. (2025). The economics of ai foundation models: Openness, competition, and governance. [arXiv preprint arXiv:2510.15200](#).

A Details on Computing the Openness Degree

To ensure an objective and reproducible quantification of model openness, we evaluated each foundation model based on the official Model Cards and documentation provided by the developers on public repositories, primarily Hugging Face and GitHub.¹⁶ The scoring follows a structured rubric across six dimensions, where each dimension is assigned a score $s_{d,j} \in \{0, 0.5, 1\}$. The criteria for each dimension were manually verified against the latest available Model Cards as follows:

1. **Weights (W):** A score of 1.0 is assigned if the model weights are publicly available for download (e.g., in .safetensors format). All the models in our dataset are assigned 1.0 for this dimension.
2. **License (L):** 1.0 for standard open-source licenses (Apache 2.0, MIT). 0.5 for custom commercial licenses with usage-based restrictions (e.g., Meta Llama or Qwen licenses). 0 for restrictive, non-commercial research licenses (e.g., Mistral Research License).
3. **Data Disclosure (D):** 1.0 for models providing detailed dataset mixtures and sampling strategies. 0.5 for models that disclose general data categories and token counts but withhold specific sources. 0 for no meaningful data documentation.
4. **Architecture (A):** 1.0 if a comprehensive technical report or whitepaper is provided, detailing model configurations (e.g., GQA, MoE routing, RoPE). 0.5 for basic architectural specifications without an in-depth technical report.
5. **Fine-tuning (F):** 1.0 if the model supports standard fine-tuning procedures on consumer/research-grade hardware with provided instructions. 0.5 if only parameter-efficient methods (PEFT/LoRA) are feasible or if instructions are limited.
6. **Inference (I):** 1.0 for models natively supported by mainstream inference engines (e.g., vLLM, Ollama, Hugging Face Transformers). 0.5 for models requiring proprietary kernels or extreme hardware configurations for basic deployment.

Table A.1 summarizes the scoring for representative models from major firms. The final *OpenIndex_j* is the arithmetic mean of these six dimensions, providing a granular view of how different open-source models vary in their actual degree of openness.

¹⁶All model documentation was verified as of February, 2026.

Table A.1: Representative Model Openness Scores by Dimension

Model	Openness Dimensions						OpenIndex
	W	L	D	A	F	I	
<i>Meta</i>							
Llama-3.2-3B	1.0	1.0	0.5	1.0	1.0	1.0	0.916
Llama-3.1-70B	1.0	0.5	0.5	1.0	1.0	1.0	0.833
Llama-2-7B	1.0	0.5	0.0	1.0	1.0	1.0	0.750
<i>DeepSeek</i>							
DeepSeek-V3	1.0	1.0	0.0	1.0	0.5	1.0	0.750
DeepSeek-R1-Distill	1.0	0.5	0.0	0.5	1.0	1.0	0.667
DeepSeek-R1 (Base)	1.0	0.5	0.0	0.5	0.0	0.5	0.416
<i>Mistral</i>							
Mistral-7B-v0.3	1.0	1.0	0.0	0.5	1.0	1.0	0.750
Mistral-Large-2411	1.0	0.0	0.0	0.5	0.5	0.0	0.416
<i>Qwen</i>							
Qwen3-8B	1.0	1.0	0.5	1.0	1.0	1.0	0.916
Qwen2.5-72B	1.0	0.5	0.5	1.0	1.0	1.0	0.833

Notes: This table reports openness scores across six dimensions: Weights (W), License (L), Data Disclosure (D), Architecture (A), Fine-tuning (F), and Inference (I). Each dimension is scored as 0, 0.5, or 1 based on publicly available model documentation. The overall OpenIndex is computed as the unweighted mean across the six dimensions. All model documentation was verified as of February 2026 using official Model Cards and developer repositories (Hugging Face and GitHub).

B Calibration of the Openness Cost Parameters

In this appendix, we describe how we calibrate the openness cost parameters, κ_O and μ . In principle, one could attempt to solve for the openness cost parameters and the payoff parameters jointly from the full dynamic first-order condition in equation (7). In our data, however, such a joint procedure is numerically unstable. The main difficulty is that the short-run monetization parameter ϕ and the openness cost parameters (κ_O, μ) enter the optimality condition with similar scaling roles, while the long-run payoff $X_{2j,t}$ is mechanically close to a scaled version of the short-run payoff term $X_{1j,t}$. As a result, a full joint fixed-point procedure is weakly identified and prone to unstable updates.

We therefore adopt a transparent two-step conditional calibration-estimation strategy. In the first step, we discipline the cost-side parameters conditional on a maintained normalization for the short-run revenue coefficient, setting $\phi = \phi_0 = 2.5$. Rather than imposing the full dynamic first-order condition, we solve a static approximation of the openness problem at the observation level. For each model-period observation (j, t) , define

$$\Pi_{j,t}^{\text{static}}(O) = \phi_0 \Lambda_{j,t} s_{j,t}(O) - \frac{\kappa_O}{2} O^2 - \mu \exp(-z \Delta_{j,t}) O, \quad (11)$$

where

$$\Delta_{j,t} = \max\{0, \bar{q}_t - q_{j,t}\}, \quad (12)$$

$\Lambda_{j,t}$ denotes market size, and $s_{j,t}(O)$ is the model-implied adoption share from equation (2).

For each observation, the model-implied openness choice is defined by the bounded optimization problem

$$\hat{O}_{j,t}(\kappa_O, \mu; \phi_0) = \arg \max_{O \in [0,1]} \Pi_{j,t}^{\text{static}}(O). \quad (13)$$

We solve this problem by direct constrained maximization rather than by imposing only the first-order condition. This is numerically preferable because it remains valid under corner solutions and does not require the optimum to lie in the interior of the action space.

Given the implied choice rule in (13), we choose (κ_O, μ) to minimize the observation-level discrepancy between observed openness and model-implied openness:

$$(\hat{\kappa}_O, \hat{\mu}) = \arg \min_{\kappa_O, \mu} \frac{1}{N} \sum_{j,t} \left(O_{j,t} - \hat{O}_{j,t}(\kappa_O, \mu; \phi_0) \right)^2. \quad (14)$$

Accordingly, the first step should be interpreted as a nested minimum-distance calibration. The inner problem delivers the model-implied openness choice at the observation level, and the outer problem selects the cost parameters to minimize the average squared discrepancy

between observed and predicted openness choices.

This first-step procedure is best viewed as a conditional cost-side discipline rather than as a full-information structural estimator of the dynamic openness problem. Conditional on the maintained normalization $\phi = \phi_0$, it yields a cost schedule that rationalizes the cross-sectional pattern of observed openness choices. In principle, variation in the quality gap $\Delta_{j,t}$ helps discipline the frontier-proximity term μ , while variation in observed openness levels helps pin down the convexity parameter κ_O . In our empirical implementation, however, the estimated contribution of the frontier-proximity component is limited, so the data are most informative about the overall curvature of the openness cost schedule.

Given the calibrated cost parameters, we then construct the implied marginal openness cost:

$$Y_{j,t} = \hat{\kappa}_O O_{j,t} + \hat{\mu} \exp(-z \Delta_{j,t}). \quad (15)$$

In the second step, we estimate the decomposition implied by the marginal openness condition:

$$Y_{j,t} = \phi X_{1j,t} + \nu_q X_{2j,t}, \quad (16)$$

where

$$X_{1j,t} = \Lambda_{j,t} \frac{\partial s_{j,t}(O_t)}{\partial O_j}, \quad X_{2j,t} = \delta^{\tau_{j,t}} \theta_s \Lambda_{j,t} \frac{\partial s_{j,t}(O_t)}{\partial O_j} q'_{j_u,t+\tau}. \quad (17)$$

Under the logit adoption specification,

$$\frac{\partial s_{j,t}(O_t)}{\partial O_j} = s_{j,t}(O_t) (1 - s_{j,t}(O_t)) (\beta + \gamma q_{j,t}). \quad (18)$$

This second step separates the short-run monetization component from the continuation-value component while taking the calibrated cost schedule from the first step as given.

Our empirical strategy is related in spirit to the two-step literature in empirical IO and dynamic games, which separates the recovery of policy-relevant objects from the subsequent recovery of structural parameters using equilibrium optimality conditions (Bajari et al., 2007; Pakes et al., 2007). Our implementation, however, differs from these estimators in an important respect. Rather than estimating policy functions and transition rules in the first step, we impose an initial normalization for the short-run monetization coefficient and calibrate the openness cost parameters to match the observed openness choices under a static approximation. In the second step, we use the marginal openness condition to recover the monetization parameters. The resulting procedure is therefore best interpreted as a conditional two-step calibration-and-recovery strategy, rather than as a direct implementation of the two-step estimation strategy used in empirical IO literature (e.g., Bajari et al. (2007)).

B.1 Robustness to Alternative Initial Values of ϕ_0

Because the first-step calibration is conducted conditional on a maintained normalization for the short-run monetization parameter, it is useful to examine whether the results are sensitive to alternative values of ϕ_0 . In the baseline analysis, we set $\phi_0 = 2.5$. As a robustness check, we repeat the two-step procedure under two alternative maintained values, $\phi_0 = 1$ and $\phi_0 = 5$.

Table B.1 reports the resulting first-step calibrated openness-cost parameters and the second-step revenue decomposition estimates. The results show that the initial values of ϕ mainly affect the scale of the calibrated cost parameters and the levels of the recovered coefficients, rather than the qualitative interpretation of the model. When $\phi_0 = 1$, both calibrated cost parameters are estimated to be very small, and the recovered second-step coefficients are also small in absolute value. When $\phi_0 = 5$, by contrast, the calibrated curvature parameter $\hat{\kappa}_O$ rises sharply, while the frontier-proximity term $\hat{\mu}$ is pushed to its lower bound, indicating that this component contributes little once the short-run revenue normalization is increased.

At the same time, a common pattern emerges across both exercises. In particular, the recovered short-run coefficient $\hat{\phi}$ is far smaller than the maintained value used in the first step under both alternative normalizations. This is consistent with a strong dampening pattern in the static component of openness marginal revenue. Most importantly, the continuation-value term remains necessary in the second-step decomposition under both normalizations. Although the level of $\hat{\nu}_q$ changes substantially across specifications, this should be interpreted as a rescaling property of the decomposition rather than as evidence against the main mechanism. Overall, these exercises suggest that changing the maintained value of ϕ_0 affects the scale of the decomposition, but does not overturn the paper’s main qualitative conclusion that observed openness choices cannot be rationalized solely by contemporaneous short-run revenue and instead require a forward-looking component.

In unreported calculations, the continuation-value component continues to account for the majority of implied total marginal revenue, with the long-run share remaining substantially larger than the short-run share on average. Taken together, these exercises show that alternative maintained values of ϕ_0 affect the scale of the decomposition, but leave the main qualitative result unchanged: observed openness choices are not well explained by contemporaneous monetization incentives alone and instead require a forward-looking component.

Table B.1: Results under Different ϕ_0

	Maintained value in the first step	
	$\phi_0 = 1$	$\phi_0 = 5$
<i>Panel A. First-step calibration</i>		
$\hat{\kappa}_O$	0.00155	210.87
$\hat{\mu}$	0.00144	0.000
<i>Panel B. Second-step decomposition: $Y_{j,t} = \phi X_{1j,t} + \nu_q X_{2j,t}$</i>		
$\hat{\phi}$	9.23×10^{-7} (4.90×10^{-7})	0.0606 (0.0325)
$\hat{\nu}_q$	0.0464 (0.0061)	3449.27 (418.78)
Observations	324	324