

Article

An Auditable Pricing Reference Framework for Medical Data Products: An Early Proof-of-Concept Study

Junwei Wang ¹, Wei Dai ¹, Konglin Zhu ^{2,*}  and Bo Qu ³

¹ Beijing International Big Data Exchange Co., Ltd., Beijing 100012, China; wangyiming0811@gmail.com (J.W.); daiwei@mail.bjdex.com (W.D.)

² School of Artificial Intelligence, Beijing University of Posts and Telecommunications, Beijing 100876, China

³ School of Computer Science and Technology, University of Science and Technology of China, Hefei 230026, China; qubo@hxb.com.cn

* Correspondence: klzhu@bupt.edu.cn

Abstract

Exchange-listed medical data products create information and pricing asymmetry because sellers, buyers, and governance reviewers do not observe the same product boundaries, scenario permissions, processing depth, or compliance costs. This study proposes SM-DPF, an auditable pricing reference framework that makes these asymmetric conditions explicit through a reproducible three-layer chain: a public investment anchor, a locked structural scoring model, and a proposed future market learning governance interface that is not implemented or evaluated in the present study. Using 23 de-identified transaction-descriptive records from a single exchange context across insurance claims, model pretraining, and pharmaceutical R&D, the pricing reference outputs $y_{0,i}$ show in-sample diagnostic consistency with de-identified transaction prices y_i , with an overall mean absolute percentage error of 4.82% and median absolute percentage error of 4.55%. The same 23 records informed the initial scenario-response calibration and the subsequent diagnostics; the reported errors and correlations are, therefore, in-sample diagnostics only. SM-DPF is a governance-oriented reference tool for transparent listing decisions and audit replay with no transaction price prediction claim.

Keywords: information asymmetry; pricing symmetry; medical data product; data exchange; pricing reference; scenario-based pricing; auditability; proof-of-concept

1. Introduction

Medical data circulation has become a high-value use case for research, insurance claims, model pretraining, pharmaceutical R&D, and health system governance. Health data reuse requires lawful access and purpose limitation [1,2]. It also depends on privacy protection, institutional accountability, and documented data quality [3,4]. These requirements become more demanding when hospital-derived data are converted into exchange-listed products [5,6]. A listed medical data product is defined by data scope, permitted use, de-identification, processing state, delivery conditions, and compliance boundaries. Medical data holders, therefore, need pricing evidence that can withstand stewardship, auditability, and underpricing risk scrutiny, while buyers need to understand whether a reference price is proportionate to product scope, scenario, processing depth, and data quality.

This setting creates an exchange-facing pricing reference problem. Data element reforms and data asset valuation guidance support compliant circulation and traceable valuation ev-



Academic Editor: Vasilis K. Oikonomou

Received: 17 June 2026

Revised: 20 July 2026

Accepted: 22 July 2026

Published: 24 July 2026

Copyright: © 2026 by the authors.

Licensee MDPI, Basel, Switzerland.

This article is an open access article distributed under the terms and

conditions of the [Creative Commons Attribution \(CC BY\)](https://creativecommons.org/licenses/by/4.0/) license.

idence [7,8]. They do not, however, specify how a medical data exchange should construct a replayable listing reference before dense comparable transactions exist [9]. Cost-based and income-based approaches provide useful but assumption-sensitive valuation evidence [10,11]. Market-comparable approaches become more persuasive only when transactions are sufficiently dense and comparable [12,13]. Data contribution and machine learning data pricing methods mainly address task-specific value or marginal contribution, leaving exchange listing-reference formation less specified [14,15]. Shapley value and coalition value approaches further clarify contribution allocation, but they do not themselves provide an exchange-ready cold-start rule chain [16,17]. An exchange-facing reference, therefore, needs to explain the calculation path to suppliers, buyers, compliance reviewers, and governance committees without substituting itself for negotiated transaction prices.

This study proposes SM-DPF, an early proof-of-concept framework for scenario-based pricing references for exchange-listed medical data products. The framework follows a three-layer lifecycle logic: a public investment anchor provides a cold-start monetary scale, a locked structural scoring model explains product-level differentiation across hospital source, application value, intrinsic data attributes, processing intensity, and scenario response, and a market learning governance interface defines how comparable transaction evidence can inform future versions. The empirical demonstration uses 23 de-identified descriptive transaction records from a single exchange context across insurance claims, model pretraining, and pharmaceutical R&D. The study evaluates executability, disclosed-output consistency, and mechanism-level scenario response, with no claim of external validation, final transaction-pricing authority, or market-wide predictive performance.

2. Theoretical Background and Related Work

2.1. Health Data Secondary Use as an Institutional Governance Problem

Health data secondary use has moved from a narrow technical sharing issue to a broader institutional governance problem. The HealthData@EU pilot shows that secondary use depends on metadata quality, common access procedures, cross-border infrastructure, data-holder coordination, and operational governance capacity, not merely on the existence of datasets [18]. The European Health Data Space (EHDS) formalizes this shift by creating a legal framework for electronic health data access, primary use, and secondary use across healthcare delivery, research, innovation, public policy, and regulatory activities [1]. Recent commentaries on the EHDS and OECD health data governance work similarly argue that cross-border research value requires lawful processing, institutional trust, practical standards, and governance mechanisms that can operate beyond one data holder [2,3]. Precision medicine data sharing reaches the same conclusion from a clinical and genetic data perspective: the scientific value of reuse depends on clear access rights, privacy safeguards, consent boundaries, reuse procedures, and accountability across institutions [4].

These governance requirements are directly relevant to exchange-listed medical data products. A listed product is a regulated offering defined by data scope, permitted use, processing state, de-identification, delivery conditions, contractual restrictions, and audit records. Work on the EHDS warns that secondary use frameworks must be fit for purpose in real operational settings, because abstract rights or infrastructures do not automatically solve questions of access review, proportionality, data quality, or accountability [5]. Similarly, decentralized analysis and federated reuse can reduce some transfer risks while still requiring governance of metadata, data quality, computational procedures, and institutional responsibility [6]. Pricing reference formation should therefore be understood as part of the governance record attached to a listed data product.

China's data element reforms add a second institutional context. National policy documents have framed data as a factor of production and encouraged market-oriented allocation,

data exchanges, and service institutions that support compliant circulation [7,8]. Data asset valuation guidance also recognizes cost, income, and market evidence as valuation inputs under different evidentiary conditions [9]. These policy instruments help explain why an exchange-facing pricing reference is needed, but they do not themselves provide a medical data-specific rule chain that converts product scope, scenario, processing depth, and quality into a replayable listing reference.

2.2. Data Markets, Pricing, and Valuation Methods

The economics and computational literature on data pricing provides the broader methodological background. Recent reviews describe data pricing as a fragmented field spanning economic theory, market design, privacy economics, machine learning contribution measurement, and platform implementation [10]. A foundational difficulty is that data are reusable and often non-rival: use by one party does not necessarily exhaust future uses, and the same data basis can support different decisions, models, or research tasks [11]. Data markets therefore differ from ordinary goods markets: A buyer's value can be task-specific, the seller may not observe downstream benefit, and excessive or poorly governed data access may create inefficiencies relative to social welfare [12].

Existing pricing and valuation approaches each address part of the problem. Cost-based valuation is transparent and auditable when collection, storage, processing, compliance, and infrastructure expenditures can be traced. Its limitation is that cost alone cannot explain why the same clinical data basis should receive different references when prepared for insurance claims, model pretraining, or pharmaceutical R&D. Income-based valuation connects data to expected revenue, cost savings, or downstream benefit, although it requires cash-flow and discount-rate assumptions that are difficult to defend for early-stage exchange listings. Market-comparable methods become powerful once transaction density and comparability are sufficient; however, exchange-listed medical data products are often sparse, confidential, and heterogeneous in disease scope, field composition, processing version, de-identification method, and contractual restriction. Computational data pricing work offers additional perspectives. Model-based pricing for machine learning data marketplaces links price to model performance and buyer tasks [13], while mobile health and federated learning pricing studies show that data valuation can be embedded in dynamic or distributed learning environments [14,15]. Shapley value approaches further estimate the marginal contribution of data to model performance or coalition value [16,17]. These methods are valuable for training data assessment and benefit allocation. Exchange listing raises a more operational question: how to form a governance-ready, version-replayable pricing reference for a specific medical data product before or during transaction preparation.

Privacy and trust research also matters. Individuals and institutions may attach value to control, confidentiality, and expected misuse risk, and small monetary incentives may not fully capture privacy preferences or trust effects [19,20]. For medical data products, this means that a pricing reference cannot be treated as a purely technical output. It is embedded in compliance boundaries, data-holder willingness, buyer requirements, and institutional legitimacy.

2.3. Scenario Dependence, Product Differentiation, and Health Data Quality

Scenario dependence is central to medical data product pricing. The same underlying clinical records may become different products when used for insurance claims, model pretraining, pharmaceutical research, public health analysis, or hospital operations. Each scenario changes the required fields, de-identification level, linkage depth, documentation burden, processing intensity, and buyer value logic. Product differentiation theory provides a conceptual bridge: buyers may pay different amounts for different product attributes

when those attributes affect utility [21]. Medical data products extend this logic because quality is multidimensional. It includes provenance integrity, completeness, standardization, modality, field match, temporal coverage, linkage depth, de-identification, annotation, institutional credibility, and fitness for intended use.

Health data quality literature reinforces this point. Studies of electronic health record reuse have long shown that data quality assessment must consider dimensions such as completeness, correctness, concordance, plausibility, conformance, and fitness for use, with volume forming only one part of quality assessment [22,23]. More recent work on trustworthy artificial intelligence (AI) in medicine similarly argues that medical AI requires data quality assessment across documented dimensions, including the suitability of data for the target task and the risk created by missing, biased, or inconsistent inputs [24]. In genomic medicine, value assessment also depends on the way evidence is generated, interpreted, and translated into clinical or research use, with raw data existence providing insufficient evidence on its own [25]. These studies support a rule chain approach in which observable product and delivery attributes enter a pricing reference through explicit mappings that replace informal narrative adjustments.

2.4. Research Positioning

The reviewed literature indicates that the problem addressed here is narrower than general data valuation and broader than statistical price fitting. Governance literature explains why secondary use of health data requires lawful access and safeguards [1,3]. It also stresses accountability, operational trust, and cross-institutional coordination [4,5]. The data market and pricing literature explains why data value is task-dependent and non-rival [10,11]. It also shows why privacy sensitivity and sparse transactions make medical data prices difficult to infer directly from observed deals [12,19]. Product differentiation and health data quality studies explain why the same clinical data basis may support different value references when delivery specifications and intended uses differ [21,22]. Later work on AI and genomic data quality reinforces the need to document task fitness, provenance, and evidence generation conditions [23,24]. However, these strands do not themselves provide an exchange-ready, auditable, and scenario-conditioned pricing reference rule chain for cold-start medical data product listing [25].

SM-DPF is positioned in this gap. It constructs a transparent pricing reference chain in which a public investment anchor supplies a cold-start monetary scale, locked product indicators represent hospital source characteristics, application value, intrinsic data attributes, and processing intensity, a scenario-response operator captures use-specific payment logic, and a market learning governance interface defines future transaction-informed updates. Its expected contribution is governance-oriented: each pricing reference should be traceable to its inputs, parameter versions, scenario definition, and calculation path. This positioning also sets the claim boundary for the proof-of-concept. The study can demonstrate executability, in-sample consistency, and mechanism-level scenario differentiation. A single-exchange sample of 23 records cannot establish market-wide predictive validity or external generalizability.

This positioning is also linked to the theme of symmetry. In the present setting, pricing reference formation is distorted by information asymmetry between data suppliers, data buyers, and governance reviewers. SM-DPF seeks to make product attributes, scenario permissions, and calculation versions symmetrically inspectable, so that different parties can replay the same reference logic before negotiated transactions occur.

3. Materials and Methods

3.1. Study Design and Claim Boundary

This is an early proof-of-concept methodology study. The objective is to test whether a version-bound rule chain can generate auditable pricing references and support in-sample consistency checks on de-identified transaction metadata. The study does not claim market-wide predictive accuracy, external validation across exchanges, or statistical estimation of scenario coefficients.

The proof-of-concept dataset comprises 23 de-identified transaction price records from a single exchange context and its partners: nine insurance claim records, eight model pre-training records, and six pharmaceutical R&D records. Medical data have value only in a concrete permitted application, so the framework calculates a unit reference for one person's application-specific data group rather than for an isolated visit record. Buyer demand may initiate a pricing inquiry, whereby the seller requests a reference, the exchange provides a target pricing reference, and the buyer and seller negotiate before the final transaction price is formed. For a contract covering multiple persons, total contract value equals the number of persons multiplied by the negotiated unit price. The same 23 records were used to inform the initial scenario-response values and to compute the MAPE and rank correlation diagnostics reported below. Consequently, the transaction comparison is an in-sample diagnostic replay, not an independent validation or an out-of-sample prediction test. Supplementary Table S2 consolidates H , A , S , D , K , γ_s , $y_{0,i}$, y_i , and every row-level error measure used in the study.

The operational specification contains two evidently distinct parameter groups. The indicator mappings, H/A aggregation weights, intrinsic data mappings, and processing mappings were fixed through structured expert consultation and were not estimated by regressing the 23 transaction prices. The scenario-response coefficients were developed differently: experts identified discount, neutral, and premium ranges for the three permitted-use scenarios, and the available 23 proof-of-concept transactions were then used iteratively to select the initial versioned values 0.60, 1.00, and 1.60. These scenario values are, therefore, small-sample calibration parameters, not official tariffs or externally validated market coefficients. The public investment anchor was constructed separately from public-informatization inputs. Supplementary Table S3 reports the complete numerical construction, while Supplementary Tables S4–S6 report in-sample comparator and sensitivity analyses.

3.2. SM-DPF Architecture: Anchor, Scoring, and Market Learning

SM-DPF separates pricing reference formation into three auditable layers. Layer 1 is the public investment anchor. Medical information systems and clinical service infrastructure are financed primarily to support diagnosis, treatment, and hospital administration; exchange-listed medical data are a companion output created when selected records are governed, processed, and made fit for a permitted secondary-use product. The anchor therefore applies a social-average, cost-informed logic only to establish a defensible monetary scale when comparable transactions are sparse [26,27]. It is not an estimate of intrinsic data value, an audited product cost, an official tariff, a price floor, or a final transaction price. The λ_{sw} , λ_{hw} , and λ_{ops} terms are expert-elicited, category-specific functional-relevance (apportionment) ratios: each ratio independently includes the share of its own expenditure category judged to support data generation, structuring, storage, security, quality control, and secondary-use readiness. They are not portfolio weights and are not required to sum to one. These ratios were fixed before the transaction diagnostics and were not estimated from the 23 transaction prices:

$$P_{anchor,t} = (\lambda_{sw}E_{sw,t} + \lambda_{hw}E_{hw,t} + \lambda_{ops}E_{ops,t})/V_t \quad (1)$$

$$\Pi_{base,t} = P_{anchor,t} / c_{unit} \quad (2)$$

The values used in the 2024 construction are $E_{sw} = 547 \times 10^8$ CNY, $E_{hw} = 655 \times 10^8$ CNY, and $E_{ops} = 359 \times 10^8$ CNY. The Operations and Maintenance category follows the cost boundary used for information technology services [28], the service volume denominator is taken from the national health statistics boundary [29], and the category construction is informed by the hospital-informatization evidence boundary [30]. Experts assigned $\lambda_{sw} = 0.95$ because software functions are most directly connected to data capture, structuring, coding, interfaces, and access control; the excluded 5% recognizes non-data-specific functions. They assigned $\lambda_{hw} = 0.80$ because storage, computation, backup, and security support data readiness, but hardware is also shared with broader clinical operations, and $\lambda_{ops} = 0.80$ because continuity, quality assurance, audit, and recovery also support general operations. The resulting contributions are 519.65×10^8 CNY (39.05%), 524.00×10^8 CNY (39.37%), and 287.20×10^8 CNY (21.58%), producing $N_{2024} = 1330.85 \times 10^8$ CNY. With $V_{2024} = 101.5 \times 10^8$ person-visits, $P_{anchor,2024} = 1330.85/101.5 = 13.1118$ CNY/person-visit. Dividing by $c_{unit} = 1$ CNY/person-visit gives $\Pi_{base,2024} = 13.1118$, rounded to 13.1. The denominator supplies only a cold-start normalization scale; it does not redefine the traded unit.

Layer 2 is the locked structural scoring model. The scoring model maps observable product and delivery attributes into hospital base score H , application score A , intrinsic data attribute multiplier D , processing intensity multiplier K , and scenario coefficient γ_s . These terms generate y_0 , the cold-start structural pricing reference before any future market learning update.

Layer 3 is a proposed future market learning governance interface. No market learning algorithm, automated update, dynamic recalibration, or market-calibration coefficient was implemented or evaluated in the present study. The interface is retained only as a governance design specifying how later framework versions could use sufficiently comparable exchange-observed transactions prospectively while preserving historical anchor and scoring rule versions. Historical records remain replayable under the versions used at the time of listing.

Figure 1 presents the three-layer SM-DPF architecture and the versioned replay path from anchor construction to structural scoring and later market learning. The Layer 1 anchor construction is shown in Equations (1) and (2), and the locked structural equations are shown in Equations (3)–(8). Tables 1–3 separate the three layers so that anchor construction inputs, locked scoring parameters, and future market learning governance fields are not treated as one parameter set.

Table 1. Anchor construction and unit convention used in the proof-of-concept SM-DPF model.

Parameter Group	Parameter	Symbol	Value or Rule	Interpretation
Anchor construction	Software expenditure component	$E_{sw,2024}$	547×10^8 CNY	Model construction input before apportionment
	Hardware expenditure component	$E_{hw,2024}$	655×10^8 CNY	Model construction input before apportionment
	Operations and maintenance component	$E_{ops,2024}$	359×10^8 CNY	Model construction input before apportionment
	Data-readiness relevance/apportionment ratios	$\lambda_{sw}, \lambda_{hw}, \lambda_{ops}$	0.95, 0.80, 0.80	Expert-elicited category-specific inclusion ratios; independently applied and not required to sum to one
	Medical-visit denominator	V_{2024}	101.5×10^8 person-visits	Official 2024 service volume denominator

Table 1. Cont.

Parameter Group	Parameter	Symbol	Value or Rule	Interpretation
Anchor normalization	Anchor unit constant	c_{unit}	1 CNY/person-visit	Used only to make Π_{base} dimensionless
Output unit	Product output-unit constant	c_{prod}	1 CNY/application-specific person-data group	Assigns the traded reference unit in Equation (4)
Anchor construction	Apportioned numerator	N_{2024}	1330.85×10^8 CNY	519.65 + 524.00 + 287.20; contribution shares 39.05%, 39.37% and 21.58%
Public investment anchor	2024 dimensionless anchor ratio	$\Pi_{\text{base},t}$	13.1	13.1118 before rounding; not an official price

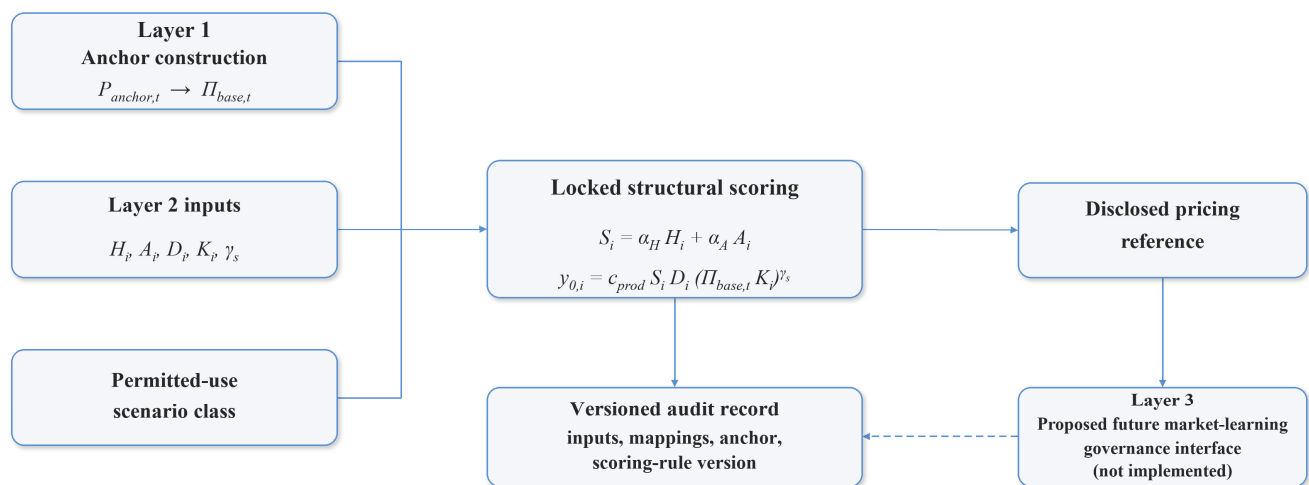


Figure 1. SM-DPF three-layer pricing reference governance chain and audit trail. Layer 3 is a proposed future governance interface and was not implemented or evaluated in this proof-of-concept.

Table 2. Locked structural scoring parameters used in the proof-of-concept.

Parameter Group	Parameter	Symbol	Value or Rule	Interpretation
Structural score	Hospital score weight	α_H	0.55	Weight on hospital source characteristics in Equation (3)
	Application score weight	α_A	0.45	Weight on application value in Equation (3)
Scenario response	Insurance claims	γ_{INS}	0.60	Moderately contractive insurance claims response used in the current row-level diagnostics
	Model pretraining	γ_{MPT}	1.00	Neutral model pretraining scenario response
	Pharmaceutical R&D	γ_{RD}	1.60	Expansionary research-and-development scenario response

Table 3. Market learning governance interface for future transaction-informed updates.

Governance Item	Recommended Handling	Current Proof-of-Concept Status
Calibration cell	Define by scenario, product type, data modality, processing depth, delivery conditions, and other comparability fields	Not estimated
Comparable transactions	Require sufficient exchange-observed records within the same calibration cell before update	Not satisfied
Calibration statistic	Use robust reference-to-transaction ratio summaries in future versions	Not estimated
Update direction	Apply prospectively to later references only	Not applied
Versioning	Preserve anchor version, scoring-rule version, and market learning version	Required
Historical replay	Keep historical references reproducible under the original versions used at listing	Required

3.3. Locked Scoring Model and Indicator Mappings

The model first combines hospital source and application value scores into a structural product score (Equation (3)). It then applies the intrinsic data attribute multiplier, a dimensionless public investment anchor ratio, processing-intensity multiplier, and scenario response to form the cold-start structural reference (Equation (4)):

$$S_i = \alpha_H H_i + \alpha_A A_i, \alpha_H = 0.55, \alpha_A = 0.45 \quad (3)$$

$$y_{0,i} = c_{\text{prod}} S_i D_i (\Pi_{\text{base},t} K_i)^{\gamma_s} \quad (4)$$

Here, S_i is the combined structural product score, D_i is the intrinsic data attribute multiplier, K_i is the processing-intensity multiplier, γ_s is the scenario-response coefficient, and $\Pi_{\text{base},t}$ is the dimensionless public investment anchor ratio. The constant $c_{\text{unit}} = 1$ CNY/person-visit is used only in Equation (2) to normalize the visit-based cold-start anchor. The distinct output unit constant $c_{\text{prod}} = 1$ CNY per application-specific person-data group appears in Equation (4), so $y_{0,i}$ has the unit of the traded reference rather than the anchor denominator. Both constants equal one numerically, so this dimensional clarification leaves every reported $y_{0,i}$ unchanged.

Hospital base score, application score, intrinsic data attribute multiplier, and processing intensity multiplier use the compact weighted forms in Equations (5)–(8):

$$H_i = \sum_{j=1}^{M_H} w_{H,j} h_{j,i}, \sum_{j=1}^{M_H} w_{H,j} = 1 \quad (5)$$

$$A_i = \sum_{j=1}^{M_A} w_{A,j} a_{j,i}, \sum_{j=1}^{M_A} w_{A,j} = 1 \quad (6)$$

$$D_i = \sum_{j=1}^{M_D} w_{D,j} d_{j,i}, \sum_{j=1}^{M_D} w_{D,j} = 1 \quad (7)$$

$$K_i = \sum_{j=1}^{M_K} w_{K,j} k_{j,i}, \sum_{j=1}^{M_K} w_{K,j} = 1 \quad (8)$$

In Equations (5)–(8), j indexes the indicators belonging to each component, M_H , M_A , M_D , and M_K denote the corresponding indicator counts, w denotes a non-negative component-specific weight, and $h_{j,i}$, $a_{j,i}$, $d_{j,i}$, and $k_{j,i}$ denote the mapped indicator values for

record i . The weights within each component sum to one. Disease-category interpretation follows the ICD-10 classification boundary [31]. Table 4 reports every weight and indicator, and Supplementary Table S1 reports the level-to-score mappings.

Table 4. Aggregation weights and scoring indicators used in Equations (5)–(8).

Component	Indicator	Symbol	Weight	Scoring Meaning
H	Hospital grade and scale	$w_{H,g}$	0.40	Institutional qualification, scale, reputation, and data authority
H	Informatization maturity	$w_{H,info}$	0.30	Data production, governance foundation, traceability, and standardization
H	Department capability	$w_{H,dept}$	0.30	Specialty capability and high-value clinical service capacity
A	Disease category	$w_{A,disease}$	0.20	International Classification of Diseases, 10th Revision (ICD-10)-style clinical grouping [31]
A	Scarcity	$w_{A,scarcity}$	0.30	Supply-demand scarcity related to disease, treatment protocol, or equipment substitutability
A	Outpatient/inpatient cost band	$w_{A,cost}$	0.20	Economic intensity associated with the treated condition
A	Physician seniority	$w_{A,physician}$	0.30	Professional authority and clinical quality of the data-generating source
D	Data quality	$w_{D,quality}$	0.20	Completeness, accuracy, and consistency of the data package
D	Content match	$w_{D,content}$	0.20	Match between supplied fields/content and the buyer's requirement list
D	Data modality	$w_{D,modality}$	0.30	Structured, multimodal, and unstructured data package differences
D	Scale match	$w_{D,scale}$	0.30	Match between supplied record scale and requested record scale
K	De-identification strength	$w_{K,deid}$	0.20	Compliance intensity and information-loss/processing burden
K	Processing degree	$w_{K,processing}$	0.40	Delivery upgrade, cleaning/annotation depth, and staff intensity
K	Number of content items	$w_{K,items}$	0.40	Field-list scale and delivery workload

The scenario coefficient is fixed by permitted use. The proof-of-concept uses insurance claims, model pretraining, and pharmaceutical R&D. Other candidate use scenarios are treated as outside the proof-of-concept claim boundary and are not used in the present analysis.

Table 3 reports the Layer 3 market learning governance interface. This interface is not a current proof-of-concept parameter set, and no market calibration coefficient is estimated in this study. It specifies the information that future framework versions should retain before transaction-informed updates are applied prospectively.

Table 4 reports the locked aggregation weights and indicator meanings used in Equations (5)–(8).

Complete level-by-level score mappings are provided in Supplementary Table S1. Supplementary Table S2 provides de-identified row-level replay fields for all 23 records, including scenario, H, A, $S_i = 0.55H_i + 0.45A_i$, D, K, scenario exponent γ_s , pricing reference output $y_{0,i}$, de-identified transaction price y_i , signed error, absolute error, and absolute percentage error. These fields support formula replay of $y_{0,i}$ under the disclosed rule chain. The audit record also identifies the scoring rule and anchor versions that should be retained for operational replay. Transaction parties, raw contracts, negotiation records, and identifiable institutional information remain undisclosed because they are exchange business records.

3.4. Expert Scoring and Parameter Fixation

The scoring mappings, structural weights, and anchor apportionment assumptions were fixed through a two-round expert consultation process following the reporting logic of CREDES [32]. The panel included 16 experts across medical supply, technical processing, market transaction, and academic/regulatory governance, with institutional input from the Data Price Committee of the China Price Association, medical data service companies, and the Institute for Internet Industry, Tsinghua University. Experts set initial indicator values and category-specific anchor inclusion ratios, adjusted them according to relative functional importance, and discussed scenario-specific discount and premium ranges. The anchor ratios and H/A and D/K mappings were version-locked independently of transaction price fitting. Scenario response values were then selected within the expert-defined ranges by iterative comparison with the same 23 proof-of-concept transactions.

In Round 1, experts reviewed initial indicator coefficients derived from the literature, policy context, and industry benchmarks and provided quantitative revisions and qualitative comments. In Round 2, group medians, means, and disputed items were returned to the experts, with attention to hospital source characteristics, specialty capability, disease scarcity, scenario differentiation, and processing intensity. The structural coefficients were locked after Round 2. The three scenario-response coefficients were subsequently chosen as initial discount, neutral, and premium settings by combining the expert-defined ranges with iterative inspection of the 23 transaction examples. No post-hoc adjustment was made to remove the retained RD-06 transaction-premium anomaly.

The 23 transaction prices were not used to estimate the H/A structural weights or the underlying D/K mappings. They did, however, inform the initial scenario-response coefficients 0.60, 1.00, and 1.60 and were then reused for the reported MAPE and rank correlation calculations. The manuscript therefore treats those calculations as in-sample diagnostics and does not describe them as independent validation, predictive accuracy, or evidence of market-wide generalizability.

Panel convergence was assessed using Kendall's coefficient of concordance and coefficient of variation. The local methods record reports Kendall's W values in the range 0.70–0.85 and coefficient of variation below 0.15 for core indicators. Item-block W and CV values are unavailable in the current materials. These diagnostics are reported conservatively as range-level expert-consensus diagnostics. They are insufficient for full Delphi validation or statistical validation of market price accuracy.

Table 5 reports the expert-panel composition used to fix the scoring mappings and weights.

Table 5. Expert panel composition.

Expert Quadrant	Seats	Background	Consultation Focus
Medical supply side	5	Tier-3 hospital information-center directors and medical record department directors	Hospital-based indicators and data-production issues
Technical processing side	2	CTOs/product directors from medical big-data companies	Cleaning, annotation, structuring, and processing intensity definitions
Market transaction side	4	Data exchange trading heads and insurance data procurement experts	Application potential, scarcity, market payment logic, and scenario differences
Academic/regulatory side	5	University professors, health commission officials, and pricing association experts	Scenario constraints, public/private value boundary, version governance, and auditability

3.5. Proof-of-Concept Records, Diagnostics, and Audit Trail

The proof-of-concept uses 23 de-identified transaction records. For each record, $y_{0,i}$ denotes the pricing reference output generated by the locked calculation and displayed to two decimals, while y_i denotes the de-identified transaction price or converted unit transaction price. Supplementary Table S2 consolidates the full row-level fields and errors. The same records informed initial scenario calibration and subsequent diagnostics, so all reported MAPE, median APE, and correlations are in-sample diagnostic summaries. They assess replayability, range fit, rank consistency, and anomaly visibility; they do not estimate out-of-sample transaction price performance.

The evaluation logic has three parts. First, executability checks whether the locked rule chain can generate a reference under the available inputs. Second, diagnostic consistency compares $y_{0,i}$ with de-identified transaction prices y_i . Third, scenario-level and record-level checks examine range fit, rank consistency, and anomaly cases. Comparator, error, rank correlation, and sensitivity calculations were implemented in Python 3.12.13 with NumPy 2.3.5. No public code repository was established for this revision; the deterministic equations, declared input ranges, and complete numerical inputs and outputs needed for direct replay are disclosed in the manuscript and Supplementary Tables S1–S7. Figure 2 summarizes the $y_{0,i}$ - y_i diagnostic comparison and should not be interpreted as external predictive validation.

For operational use, the audit record should include product identifier, scenario, H, A, D, K, scenario-coefficient version, anchor version, scoring-rule version, pricing reference output $y_{0,i}$, market learning status and version if used in a future implementation, calculation timestamp, and compliance note. A practical deployment may involve the exchange, data price committee functions, medical data service companies, research-institute consultation, and transaction partners in scenario classification, field checking, reference generation, negotiation support, and subsequent parameter review. Historical records should retain

the original anchor, scoring, and market learning versions so that older references remain replayable after later updates.

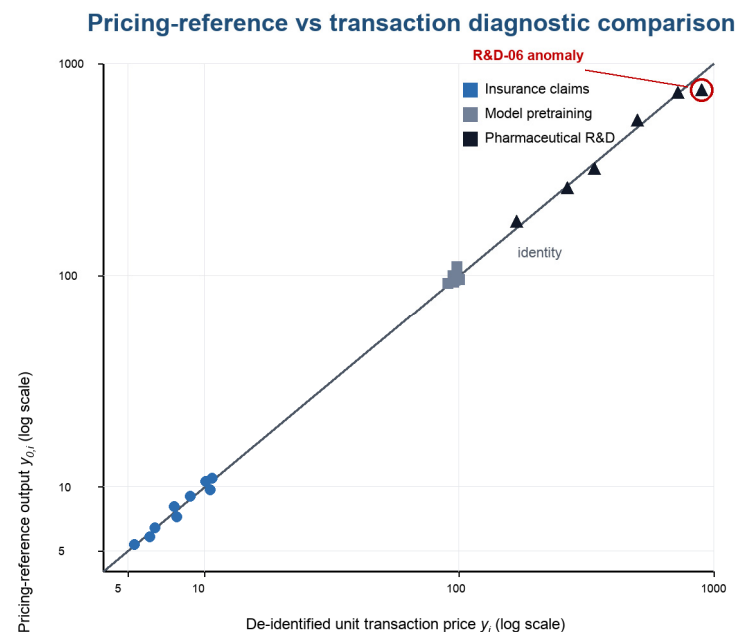


Figure 2. Diagnostic comparison between pricing reference output $y_{0,i}$ and de-identified transaction price y_i . The overall MAPE is 4.82%, the median APE is 4.55%, and the Spearman rho is 0.972; these values are descriptive in-sample diagnostics and do not function as out-of-sample prediction metrics.

4. Results

4.1. Scoring Model Disclosure and Record Status

The final scoring model is fully specified at the component level: hospital base score, application score, scenario exponent, intrinsic data attribute multiplier, and processing-intensity multiplier are all defined by locked mappings. The complete score mappings are provided in Supplementary Table S1, and the row-level replay records are provided in Supplementary Table S2. This disclosure supports two checks: whether the operational rule chain is inspectable, and whether $y_{0,i}$ is diagnostically consistent with de-identified transaction prices y_i in the early proof-of-concept sample.

The proof-of-concept disclosure supports row-level formula replay for all 23 records. The available fields include scenario, H, A, S, D, K, scenario exponent γ_s , pricing reference output $y_{0,i}$, de-identified transaction price y_i , signed error, absolute error, and absolute percentage error. The evidence therefore supports rule chain replay, diagnostic consistency checks, and anomaly diagnosis. Raw contracts, counterparties, and denominator materials for converted unit prices and negotiation details remain confidential exchange business records.

4.2. In-Sample Consistency of Disclosed PoC Outputs

The pricing reference outputs $y_{0,i}$ show diagnostic consistency with de-identified transaction prices y_i within the single-exchange proof-of-concept sample. These diagnostics are reported as range-fit, rank consistency, and anomaly checks. The model-pretraining subgroup is reported separately because $\rho = 0.381$ indicates limited rank discrimination in a narrow transaction price band. The diagnostic comparison should not be interpreted as external validation of a transaction price prediction model.

Table 6 reports the in-sample consistency diagnostics from the disclosed proof-of-concept records.

Table 6. In-sample consistency diagnostics from disclosed records.

Scenario	n	MAPE (%)	Median APE (%)	Q1 APE (%)	Q3 APE (%)	Spearman Rho ($y_{0,i}, y_i$)
Insurance claims	9	4.18	4.92	2.27	6.58	0.967
Model pretraining	8	3.94	2.32	1.86	4.84	0.381
Pharmaceutical R&D	6	6.94	6.71	4.00	7.62	1.000
Overall	23	4.82	4.55	2.10	6.77	0.972

These values should be interpreted as in-sample consistency under early proof-of-concept conditions. They do not show external predictive validity. The relatively low rank correlation in the model-pretraining group should not be overinterpreted because the group contains only eight records and a narrow transaction price range.

Across the 23 records, the overall MAPE was 4.82%, and the median APE was 4.55%. These are in-sample diagnostics because the initial scenario-response coefficients were informed by the same records. Insurance claims and model pretraining remained below 5% MAPE, while pharmaceutical R&D was 6.94% with RD-06 retained as a demand -premium anomaly. The model pretraining subgroup has only eight records, fixed $K = 1.80$, and a narrow y_i range of 90.60–100.50; its Spearman $\rho = 0.381$ therefore indicates limited rank discrimination rather than model failure or predictive validity.

The three scenarios occupy distinct empirical pricing reference bands. Insurance claims form the lowest band, model pretraining forms a narrow intermediate band, and pharmaceutical R&D forms the highest and most dispersed band. This pattern supports the scenario-conditioned structure of SM-DPF while remaining within the proof-of-concept boundary.

4.3. Scenario Diagnostics and Anomaly Response

The scenario diagnostics use the row-level proof-of-concept records. This keeps the figure and table aligned with the disclosed sample evidence and avoids implying a generic market-wide processing-intensity response surface.

Table 7 reports the scenario diagnostic summary from the 23 de-identified records.

Table 7. Scenario diagnostic summary from de-identified records.

Scenario	γ_s and K Range	$y_{0,i}$ and y_i Ranges	Diagnostic Interpretation
Insurance claims	$\gamma_s = 0.60$; $K = 0.40$ – 0.60	$y_{0,i} = 5.34$ – 11.00 ; $y_i = 5.30$ – 10.70	Good range fit; $\rho = 0.967$.
Model pretraining	$\gamma_s = 1.00$; $K = 1.80$	$y_{0,i} = 92.00$ – 110.26 ; $y_i = 90.60$ – 100.50	Narrow y_i band; $\rho = 0.381$, so rank discrimination is limited.
Pharmaceutical R&D	$\gamma_s = 1.60$; $K = 1.00$ – 1.80	$y_{0,i} = 180.00$ – 748.12 ; $y_i = 168.30$ – 900.20	R&D-06 retained as urgent-demand/scarcity/willingness-to-pay premium anomaly.

Figure 3 presents signed relative diagnostic residuals, defined as $(y_{0,i} - y_i)/y_i \times 100\%$, across the de-identified records and highlights the R&D-06 demand-premium anomaly.

The scenario diagnostics are based on the disclosed row-level records. Insurance claims use $\gamma_s = 0.60$ with $K = 0.40$ – 0.60 ; model pretraining uses $\gamma_s = 1.00$ with K fixed at 1.80; pharmaceutical R&D uses $\gamma_s = 1.60$ with $K = 1.00$ – 1.80 . R&D-06 has $y_{0,i} = 748.12$ and $y_i = 900.20$, corresponding to a signed relative residual of -16.89% . This gap is retained as

an anomaly case because the transaction context indicates urgent buyer demand, relative scarcity of the data product, and stronger willingness to pay. It is therefore interpreted as a demand-side transaction premium beyond the current rule chain variables, with no hidden recalibration of the model output.

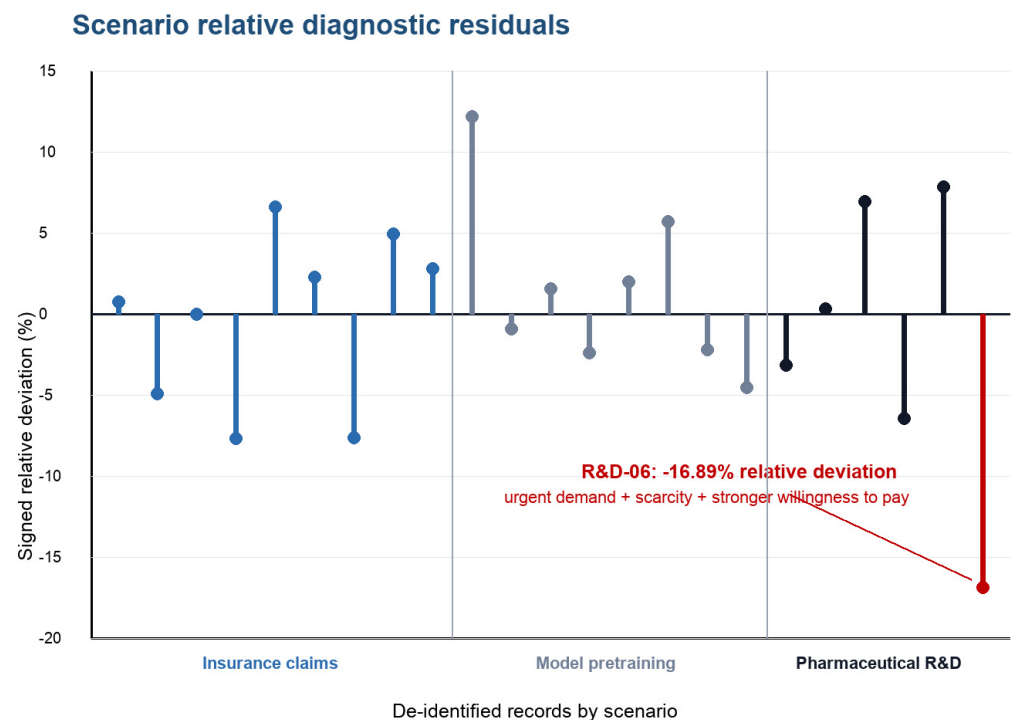


Figure 3. Relative diagnostic residuals and R&D-06 transaction-premium anomaly.

4.4. In-Sample Comparator Diagnostics

Six reviewer-requested comparators were evaluated on the same 23 records. Overall and scenario medians provide constant baselines; the additive, multiplicative, scalar, and log-linear forms test alternative functional structures. All fitted alternatives and the SM-DPF diagnostic are in-sample. The comparison does not establish predictive superiority because γ_s in SM-DPF were informed by these records and the comparator parameters were estimated from them. Table 8 reports the in-sample diagnostic comparison with requested baselines and alternative forms.

Table 8. In-sample diagnostic comparison with requested baselines and alternative forms.

Comparator	Specification	Estimated from Current y_i ?	MAPE (%)	Median APE (%)	MAE	Mean Signed Error	Spearman ρ
SM-DPF	$y_{0,i} = c_{prod} \cdot S_i \cdot D_i \cdot (13.1 \cdot K_i)^{\gamma_s}$	Partly: γ_s informed by these records	4.82	4.55	11.73	−5.16	0.972
Overall median	$\hat{y}_i = \text{median}(y_i)$	Yes	465.56	80.97	136.17	−67.50	N/A
Scenario median	$\hat{y}_i = \text{median}(y_i \mid \text{scenario})$	Yes	23.84	15.99	60.07	−16.33	0.939
Additive	$\hat{y}_i = \beta_0 + \beta_1 S_i + \beta_2 D_i + \beta_3 K_i + \text{scenario effects}$	Yes	212.28	17.86	36.46	−0.00	0.951
Multiplicative	$\hat{y}_i = \beta_s \cdot S_i \cdot D_i \cdot K_i$	Yes	10.63	6.36	19.32	5.45	0.972
Scalar	$\hat{y}_i = \beta \cdot y_{0,i}$	Yes	7.23	6.64	13.81	4.05	0.972
Log-linear	$\ln(\hat{y}_i) = \beta_0 + \beta_1 \ln S_i + \beta_2 \ln D_i + \beta_3 \ln K_i + \text{scenario effects}$	Yes	6.86	4.21	16.03	−3.30	0.966

The current SM-DPF result should be read as a replay diagnostic for a versioned rule chain. The lower in-sample error does not constitute independent model-selection evidence, while the poor performance of several cross-scenario baselines mainly reflects the very different price scales across the three permitted-use scenarios.

4.5. Local and Global Sensitivity Diagnostics

The earlier $\pm 20\%$ checks are retained and correctly labelled as local one-at-a-time perturbations. Supplementary Table S5A reports anchor, K, and scenario-coefficient checks. Supplementary Table S5B separately perturbs each category-specific apportionment ratio; because a relevance ratio cannot exceed 1, the +20% software test is capped at $\lambda_{sw} = 1.00$. The resulting whole-anchor multipliers are 0.9219–1.0205 for λ_{sw} , 0.9213–1.0787 for λ_{hw} , and 0.9568–1.0432 for λ_{ops} . These are bounded construction diagnostics, not empirical confidence intervals. A variance-based Sobol analysis was added as a global structural sensitivity diagnostic. It jointly varies H, A, D, K, and the dimensionless anchor over the declared admissible ranges and holds γ_s fixed within each scenario. First-order and total-effect indices are estimated from 65,536 Monte Carlo base samples with bootstrap confidence intervals. The analysis describes sensitivity of the deterministic mapping within this parameter domain; it is not an empirical probability model for future transactions.

Table 9 reports the summary of variance-based global sensitivity results. The dominant factor shifts with the scenario response: H has the largest first-order contribution under insurance claims, whereas K becomes dominant for model pretraining and pharmaceutical R&D. Supplementary Tables S5 and S6 report the full local and global results, including all confidence intervals.

Table 9. Summary of variance-based global sensitivity results.

Scenario	Largest First-Order Effect	S_1	Largest Total Effect	S_T
Insurance claims	H	0.425	H	0.482
Model pretraining	K	0.339	K	0.417
Pharmaceutical R&D	K	0.471	K	0.596

4.6. Auditability and Version-Replay Demonstration

The framework's auditability comes from retention of the calculation chain, not merely the final reference. A replayable record should preserve the product identifier, scenario, input scores, pricing reference output $y_{0,i}$, scoring-rule version, anchor version, and future market learning version when used. If the anchor or market learning interface is updated after additional market evidence accumulates, historical references should remain reproducible under their original versions.

This versioning logic separates three sources of change: updates to the public investment anchor, updates to product scoring inputs, and updates to the market learning governance interface. Such separation supports governance review and dispute resolution because reviewers can identify which part of the rule chain changed.

5. Discussion

5.1. Principal Findings

This study presents SM-DPF as an early proof-of-concept framework for scenario-based pricing references for exchange-listed medical data products. The framework is designed for governance review and reference formation, with final transaction price determination left to negotiated exchange processes. The proof-of-concept supports three limited findings: the rule chain is executable, disclosed outputs show in-sample consistency

with transaction prices, and the scenario mechanism produces ordered response across processing depth.

The most important finding is not the MAPE value. It is the feasibility of turning pricing reference formation into a versioned and replayable governance chain. This distinction matters because medical data exchanges must explain how references are formed under compliance constraints and sparse comparable evidence.

5.2. Governance Interpretation of the Three-Layer Framework

The three-layer structure addresses different governance problems in medical data circulation. The public investment anchor supplies an initial monetary scale under cold-start conditions. This helps public-sector data holders and other suppliers document why an initial listing reference is defensible under audit and stewardship scrutiny. The locked scoring model creates product-level differentiation based on hospital, application, intrinsic data attribute, and processing intensity scores, which helps buyers see how scenario, data quality, and processing depth affect the reference. The market learning governance interface prevents the cold-start anchor from becoming a permanent untested price once transaction records accumulate, while keeping the exchange in the role of transparent reference formation and avoiding unilateral price determination.

The 2024 anchor convention is a version-bound, cost-informed cold-start construction. Medical data originate in systems funded primarily for diagnosis, treatment, and hospital administration; the exchange-listed product is a secondary-use companion output. The category-specific relevance ratios include only the portion judged to support data-product readiness, and the national visit denominator supplies only a social-average scaling basis. The resulting 13.1 is outside the categories of intrinsic value, audited product cost, observed market benchmark, regulatory floor, and fitted transaction price coefficient. It initializes a dimensionally consistent reference for one person's application-specific data group under a defined permitted use. As sufficiently comparable medical data transactions accumulate, later versions should give those comparable greater evidential weight and make the original anchor less decisive, while preserving earlier parameter versions for audit replay.

5.3. Relationship to Existing Valuation Approaches

SM-DPF complements established valuation approaches. Cost evidence supplies an auditable starting scale when comparable transactions are scarce, while quality, processing and permitted-use indicators represent nonlinear product differentiation. The anchor does not override buyer demand, the seller's asking position, the exchange's target reference, or bilateral negotiation; the observed transaction price is formed only after those parties negotiate. In a more mature market, sufficiently comparable transactions should carry greater weight in prospective parameter review, and the cold-start anchor should become less decisive. The present paper defines that governance transition but does not implement a market learning estimator. Data contribution allocation remains outside the proof-of-concept boundary.

The narrower contribution is to combine these elements into a rule chain suitable for exchange governance. The framework is best evaluated as a pricing reference formation method. Universal valuation theory remains outside the proof-of-concept boundary.

5.4. Expert-Informed Scoring, Reproducibility and Auditability

The scoring system contains expert-informed structural coefficients and transaction-informed initial scenario-response coefficients. This introduces judgement and small-sample calibration risk, but both sources are visible and version-bound. The H/A weights and D/K mappings were fixed through expert consultation; the same 23 transactions then informed γ_s and were reused for diagnostics. The revised analysis therefore labels every

error, correlation and comparator result as in-sample and avoids claiming independent validation.

The expert consultation process is reported with its evidential boundary. The manuscript identifies the participating organizations and expert quadrants, the initial scoring and importance-adjustment process, and the scenario discount/premium ranges. It also distinguishes the expert-fixed structural rules from the transaction-informed scenario coefficients. Supplementary Materials provide complete score mappings, row-level inputs and outputs, and the available range-level consensus information. Exact item-block W and CV values are not supported by retained source records and are therefore not invented.

5.5. Limitations and Future Work

This study has several limitations. First, the 23 records come from one exchange context, and the same records informed the initial scenario-response coefficients and the subsequent diagnostics. The MAPE, correlations, comparator results, and anomaly analysis are therefore in-sample only and cannot establish external validity or market-wide predictive accuracy. Second, raw contracts, counterparty identities, negotiation logs, and contract-to-unit denominator records remain confidential even though the disclosed row-level fields support formula replay. Third, the public investment anchor and its apportionment ratios are expert-elicited, category-specific functional-relevance assumptions rather than audited accounting allocations or official cost weights; their current values require prospective review as better cost evidence and market comparable transactions emerge. Fourth, the Sobol analysis is a structural sensitivity assessment under declared input ranges and independence assumptions, not an empirical uncertainty distribution. Fifth, no market learning algorithm or external updating rule was implemented.

Future work should evaluate the framework on independent time windows and multi-exchange data. A market learning procedure may then be implemented prospectively, with predefined calibration cells, minimum sample thresholds, robust update statistics, and preserved historical versions. Until such work is completed, Layer 3 remains a proposed governance interface rather than an operational learning module.

6. Conclusions

SM-DPF addresses how an exchange can form an explainable cold-start pricing reference for medical data products when comparable transactions are sparse, product definitions are scenario-dependent, and auditability matters. The original three-layer structure is retained: a public investment anchor, a locked structural score, and a proposed future market learning governance interface. The revised evidence shows that $y_{0,i}$ can be replayed from disclosed fields and compared with de-identified transaction prices y_i . Because the same 23 records informed initial γ_s calibration and the diagnostics, all MAPE, correlation, and comparator results are explicitly interpreted as in-sample. The framework remains a governance-oriented pricing reference method; it does not predict transaction prices, establish market-wide validity, or implement market learning.

Supplementary Materials: The following supporting information can be downloaded at: <https://www.mdpi.com/article/10.3390/sym18081263/s1>. The revised Supplementary File contains Table S1: Complete score mappings; Table S2: The consolidated 23-record row-level dataset with H, A, S, D, K, γ_s , $y_{0,i}$, y_i and all reported errors; Table S3: The full 13.1 anchor calculation and apportionment rationale; Table S4: Six comparator specifications and results; Table S5: Local and bounded coefficient sensitivity checks; Table S6: Sobol global structural sensitivity results; Table S7: Worked replays; and Table S8: The data-disclosure and implementation boundary.

Author Contributions: Conceptualization, J.W. and W.D.; methodology, J.W., W.D. and K.Z.; formal analysis, W.D.; investigation, B.Q.; resources, J.W. and B.Q.; data curation, W.D. and B.Q.; writing—original draft preparation, W.D.; writing—review and editing, J.W., K.Z. and B.Q.; visualization, W.D.; supervision, J.W. and K.Z.; project administration, J.W. All authors have read and agreed to the published version of the manuscript.

Funding: This research was funded by the National Key Research and Development Program of China, Grant No. 2023YFB3106500, and the National Key R&D Program of China Young Scientist Project, Grant No. 2023YFB2704500.

Institutional Review Board Statement: Not applicable. This study used transaction-descriptive metadata only and did not involve human-subject intervention, individual-level health records, personally identifiable information, or sensitive personal information.

Informed Consent Statement: Not applicable. No personal health information was accessed, stored, or analyzed.

Data Availability Statement: The revised Supplementary File discloses the complete de-identified row-level analytical dataset required to reproduce $y_{0,i}$ and all reported in-sample diagnostics. It contains transaction-descriptive metadata and de-identified unit transaction prices supplied by Beijing International Big Data Exchange and its partners. Raw contracts, counterparties, negotiation logs, and identifiable denominator records remain confidential exchange business materials.

Acknowledgments: During preparation of this revised manuscript, the authors used OpenAI Codex (GPT-5-based coding agent; accessed on 17 July 2026) for language editing, document formatting, cross-document consistency checks and code-assisted verification of disclosed calculations and tables. The tool did not originate the transaction records or determine the study's substantive claims. The authors reviewed and edited all outputs and take full responsibility for the content of the publication.

Conflicts of Interest: Junwei Wang and Wei Dai are affiliated with Beijing International Big Data Exchange Co., Ltd., and the proof-of-concept records were supplied by the exchange and its partners. This relationship is disclosed as a potential competing interest. Structural scoring rules were expert-informed, scenario-response coefficients were initially calibrated with the 23 records, and the same records were used for in-sample diagnostics. The manuscript makes no independent-validation or transaction price prediction claim.

References

1. European Parliament; Council of the European Union. Regulation (EU) 2025/327 of the European Parliament and of the Council of 11 February 2025 on the European Health Data Space and Amending Directive 2011/24/EU and Regulation (EU) 2024/2847. Off. J. Eur. Union 2025, 2025/327. 5 March 2025. Available online: <https://eur-lex.europa.eu/eli/reg/2025/327/oj> (accessed on 17 July 2026).
2. OECD. *Health Data Governance for the Digital Age: Implementing the OECD Recommendation on Health Data Governance*; OECD Publishing: Paris, France, 2022. [CrossRef]
3. Ganna, A.; Carracedo, A.; Christiansen, C.F.; Di Angelantonio, E.; Dykstra, P.A.; Dzhambov, A.M.; Eils, R.; Green, S.; Schneider, K.L.; Varga, T.V.; et al. The European Health Data Space can be a boost for research beyond borders. *Nat. Med.* **2024**, *30*, 3053–3056. [CrossRef] [PubMed]
4. Elhussein, A.; Baymuradov, U.; NYGC ALS Consortium; Phatnani, H.; Kwan, J.; Sareen, D.; Broach, J.R.; Simmons, Z.; Arcila-Londono, X.; Lee, E.B.; et al. A framework for sharing of clinical and genetic data for precision medicine applications. *Nat. Med.* **2024**, *30*, 3578–3589. [CrossRef] [PubMed]
5. Marelli, L.; Stevens, M.; Sharon, T.; Van Hoyweghen, I.; Boeckhout, M.; Colussi, I.; Degelsegger-Márquez, A.; El-Sayed, S.; Hoeyer, K.; van Kessel, R.; et al. The European health data space: Too big to succeed? *Health Policy* **2023**, *135*, 104861. [CrossRef] [PubMed]
6. Fröhlich, H.; Funck Hansen, A.; Hilvo, M.; Jansen, G.; Madan, S.; Moazemi, S.; Negreanu, S.; Satagopam, V.; Gribbon, P.; Muehlendyck, C. Reality check: The aspirations of the European Health Data Space amidst challenges in decentralized data analysis. *J. Med. Internet Res.* **2025**, *27*, e76491. [CrossRef] [PubMed]
7. Central Committee of the Communist Party of China; State Council. Opinions on Building Basic Data Systems and Better Leveraging Data as a Factor of Production, 2022. Available online: http://www.gov.cn/zhengce/2022-12/19/content_5732695.htm (accessed on 17 July 2026).

8. National Data Administration; Cyberspace Administration of China; Ministry of Science and Technology; Ministry of Industry and Information Technology; Ministry of Transport; Ministry of Agriculture and Rural Affairs; Ministry of Commerce; Ministry of Culture and Tourism; National Health Commission; Ministry of Emergency Management; et al. Notice on Issuing the “Data Elements ×” Three-Year Action Plan (2024–2026), Guo Shu Zheng Ce [2023] No. 11. 31 December 2023. Available online: https://www.cac.gov.cn/2024-01/05/c_1706119078060945.htm (accessed on 17 July 2026).
9. China Appraisal Society. Guidance on Data-Asset Valuation, CAS (2023) No. 17, 2023. Available online: https://www.gov.cn/zhengce/zhengceku/202309/content_6901752.htm (accessed on 17 July 2026).
10. Hao, J.; Deng, Z.; Li, J. The evolution of data pricing: From economics to computational intelligence. *Heliyon* **2023**, *9*, e20274. [CrossRef] [PubMed]
11. Jones, C.I.; Tonetti, C. Nonrivalry and the economics of data. *Am. Econ. Rev.* **2020**, *110*, 2819–2858. [CrossRef]
12. Acemoglu, D.; Makhdoumi, A.; Malekian, A.; Ozdaglar, A. Too much data: Prices and inefficiencies in data markets. *Am. Econ. J. Microecon.* **2022**, *14*, 218–256. [CrossRef]
13. Chen, L.; Koutris, P.; Kumar, A. Towards model-based pricing for machine learning in a data marketplace. In *Proceedings of the 2019 International Conference on Management of Data, Amsterdam, The Netherlands, 30 June–5 July 2019*; ACM: New York, NY, USA, 2019; pp. 1535–1552. [CrossRef]
14. Xu, A.; Zheng, Z.; Li, Q.; Wu, F.; Chen, G. VAP: Online data valuation and pricing for machine learning models in mobile health. *IEEE Trans. Mob. Comput.* **2024**, *23*, 5966–5983. [CrossRef]
15. Sun, Y.; Li, B.; Yang, K.; Bi, X.; Zhao, X. TiFLCS-MARP: Client selection and model pricing for federated learning in data markets. *Expert Syst. Appl.* **2024**, *245*, 123071. [CrossRef]
16. Ghorbani, A.; Zou, J. Data Shapley: Equitable valuation of data for machine learning. In *Proceedings of the 36th International Conference on Machine Learning, Long Beach, CA, USA, 9–15 June 2019*; ML Research Press: Singapore, 2019; Volume 97, pp. 2242–2251. Available online: <https://proceedings.mlr.press/v97/ghorbani19c.html> (accessed on 17 July 2026).
17. Jia, R.; Dao, D.; Wang, B.; Hubis, F.A.; Hynes, N.; Gürel, N.M.; Li, B.; Zhang, C.; Song, D.; Spanos, C.J. Towards efficient data valuation based on the Shapley value. In *Proceedings of the 22nd International Conference on Artificial Intelligence and Statistics, Naha, Japan, 16–18 April 2019*; ML Research Press: Singapore, 2019; Volume 89, pp. 1167–1176.
18. Christiansen, C.F.; Doupi, P.; Schutte, N.; Ivanković, D. Piloting an infrastructure for the secondary use of health data: Learnings from the HealthData@EU Pilot. *Eur. J. Public Health* **2025**, *35*, iii3–iii4. [CrossRef] [PubMed]
19. Athey, S.; Catalini, C.; Tucker, C.E. *The Digital Privacy Paradox: Small Money, Small Costs, Small Talk*; NBER Working Paper No. 23488; National Bureau of Economic Research: Cambridge, MA, USA, 2017. [CrossRef]
20. Acquisti, A.; Taylor, C.; Wagman, L. The economics of privacy. *J. Econ. Lit.* **2016**, *54*, 442–492. [CrossRef]
21. Mussa, M.; Rosen, S. Monopoly and product quality. *J. Econ. Theory* **1978**, *18*, 301–317. [CrossRef]
22. Weiskopf, N.G.; Weng, C. Methods and dimensions of electronic health record data quality assessment: Enabling reuse for clinical research. *J. Am. Med. Inform. Assoc.* **2013**, *20*, 144–151. [CrossRef] [PubMed]
23. Kahn, M.G.; Callahan, T.J.; Barnard, J.; Bauck, A.E.; Brown, J.; Davidson, B.N.; Estiri, H.; Goerg, C.; Holve, E.; Johnson, S.G.; et al. A harmonized data quality assessment terminology and framework for the secondary use of electronic health record data. *eGEMS* **2016**, *4*, 18. [CrossRef] [PubMed]
24. Schwabe, D.; Becker, K.; Seyferth, M.; Klauf, A.; Schaeffter, T. The METRIC-framework for assessing data quality for trustworthy AI in medicine: A systematic review. *npj Digit. Med.* **2024**, *7*, 203. [CrossRef] [PubMed]
25. Goranitis, I.; Hayeems, R.Z.; Smith, H.S.; Buchanan, J.; Weymann, D.; Regier, D.A.; Mackley, M.P.; Scott, R.H.; Hill, S.L.; Chung, B.H.Y.; et al. Determining the value of genomics in healthcare. *Nat. Med.* **2025**, *31*, 4022–4033. [CrossRef] [PubMed]
26. Samuelson, P.A. The pure theory of public expenditure. *Rev. Econ. Stat.* **1954**, *36*, 387–389. [CrossRef]
27. Mitchell, J.; Ker, D.; Leshner, M. *Measuring the Economic Value of Data*; OECD Going Digital Toolkit Notes, No. 20; OECD Publishing: Paris, France, 2021. [CrossRef]
28. GB/T 28827.7-2022; Information Technology Service—Operations and Maintenance—Part 7: Cost Measurement Specification, 2022. Standardization Administration of China: Beijing, China, 2022. Available online: <https://openstd.samr.gov.cn/bzgk/gb/newGbInfo?hcno=801CD6C35A5EE192798BA8B6A9515AC8> (accessed on 17 July 2026).
29. National Health Commission of the People’s Republic of China. 2024 Statistical Bulletin on China’s Health Development, 2025. Available online: <https://www.nhc.gov.cn/guihuaxxs/c100133/202512/f1c3a3c617484a27a1a26a468afbaeee.shtml> (accessed on 17 July 2026).
30. China Hospital Information Management Association. 2023–2024 Survey on Hospital Informatisation in China, 2025. Available online: <https://www.chima.org.cn/Html/News/Articles/17294.html> (accessed on 17 July 2026).

31. World Health Organization. ICD-10 Version: 2019—International Statistical Classification of Diseases and Related Health Problems, 10th Revision. Available online: <https://icd.who.int/browse10/2019/en> (accessed on 17 July 2026).
32. Jünger, S.; Payne, S.A.; Brine, J.; Radbruch, L.; Brearley, S.G. Guidance on Conducting and REporting DElphi Studies (CREDES) in palliative care: Recommendations based on a methodological systematic review. *Palliat. Med.* **2017**, *31*, 684–706. [[CrossRef](#)] [[PubMed](#)]

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.