

Human-Centered Experience Engineering for Healthcare AI Governance: A SER-M Model Approach

Minseong Kim

ABSTRACT

The rapid adoption of artificial intelligence (AI) in healthcare has intensified the "principle-practice gap," where existing governance frameworks provide ethical principles but lack actionable implementation guidance. This study establishes the concept of Human-Centered Experience Engineering (HCEE), systematizes it through the SER-M (Subject-Environment-Resource-Mechanism) model, and proposes an HCEE-based Healthcare AI Governance Architecture. Applying Design Science Research (DSR) methodology, we developed a four-layer governance architecture and validated it through agent-based simulation. The architecture integrates Patient Experience Index (PXI) and Employee Experience Index (EXI) as core governance variables, with trigger-control rules that adjust AI automation levels based on experience thresholds. Simulation results show that full HCEE implementation improves PXI by 34.5% and EXI by 43.6%. This study makes three contributions to the MIS field. First, it establishes HCEE, reconceptualizing human experience as a governance input variable and extending Human-Centered AI discourse to the operational level. Second, it applies the SER-M model to Healthcare AI Governance, extending the Mechanism-Based View to IS. Third, it demonstrates that AI governance can be implemented as a designable IS artifact through trigger-control rules.

keyword : Healthcare AI Governance, Human-Centered Experience Engineering, SER-M Model, Design Science Research, Agent-Based Simulation

¹ aSSIST(Seoul School of Integrated Sciences & Technologies), Seoul, Republic of Korea.

1. Introduction

1.1 Research Background

Hospital management in the 21st century stands at an unprecedented turning point. Healthcare institutions facing complex challenges—including surging medical demand due to aging populations, increasing complexity of chronic diseases, healthcare cost pressures, and healthcare workforce shortages—can no longer guarantee sustainable growth through conventional management approaches. Korea, in particular, is expected to enter a super-aged society in 2026, with the population aged 65 and over exceeding 20% of the total, making an explosive increase in healthcare demand inevitable. This demographic structural change demands a new paradigm in hospital management. In this environment, artificial intelligence (AI) is emerging as a strategic asset that fundamentally transforms the paradigm of hospital management, beyond mere technological innovation. As of August 2024, the U.S. FDA had approved more than 950 AI-based medical devices, with 221 newly approved in 2023 alone—a 43% increase from the previous year (Joshi et al., 2024). This demonstrates the rapid maturation and clinical utility of AI medical technology. The global healthcare AI market is projected to grow from approximately \$15 billion in 2023 to \$187 billion by 2030, representing a compound annual growth rate (CAGR) of 37%. From a hospital management perspective, the strategic value of AI adoption manifests across multiple dimensions.

First, improved diagnostic accuracy enhances medical quality and directly contributes to patient trust and institutional reputation. AI-based diagnostic imaging systems achieve specialist-level accuracy for certain conditions and contribute to improved treatment outcomes through early diagnosis. Second, automation of repetitive tasks enables healthcare professionals to focus on high-value activities, maximizing workforce productivity. AI utilization in administrative work, documentation, and scheduling can reduce clinician burnout and increase direct interaction time with patients. Third, data-driven decision support strengthens executives' strategic judgment and risk management capabilities. Fourth, personalized patient services create competitive advantages through differentiated healthcare experiences. However, AI adoption does not automatically guarantee success. The rate at which AI models validated in research settings successfully transition to actual clinical or operational management remains low, stemming not from technical problems but from governance issues related to implementation and adoption (Hevner et al., 2004).

Analysis of the root causes of AI adoption failures reveals that organizational, human, and cultural factors play a larger role than technological limitations themselves. Issues such as patient experience degradation risk, healthcare worker resistance and alienation, ethical risks, and failure to achieve return on investment cannot be managed without systematic governance.

1.2 Research Purpose

Existing AI governance frameworks present comprehensive principles and regulatory systems but lack specific guidance on how healthcare institutions can implement these principles in actual operations. The WHO (2021) 'AI Ethics and Governance,' EU AI Act (2024), and NIST AI RMF (2023) provide high-level normative directions, and Ibáñez and Olmeda (2022) named this the "principle-practice gap" and identified it as a core unresolved challenge in AI ethics (Ibáñez & Olmeda, 2022).

This gap creates practical dilemmas for hospital managers. Abstract principles alone cannot guide AI adoption decisions, and

specific implementation guidelines and performance measurement methods are urgently needed. The purpose of this study is to present an actionable governance framework that enables hospital managers to maximize the strategic value of AI adoption while simultaneously protecting human values.

The core proposal is Human-Centered Experience Engineering (HCEE)-based Healthcare AI Governance, systematized through the SER-M (Subject-Environment-Resource-Mechanism) model approach. This study seeks to answer three research questions. First, what are the core elements of AI governance in healthcare environments and how can they be structured? Second, what is the method for integrating human experience as a core governance variable? Third, what is the impact of HCEE-based governance on hospital performance?

1.3 Paper Structure

The structure of this paper is as follows. Section 2 reviews the current state and limitations of Healthcare AI Governance, the concept of HCEE, and the theoretical foundation of the SER-M model. Section 3 explains the Design Science Research methodology and agent-based simulation design. Section 4 presents the four-layer architecture and experience index system of HCEE-based Healthcare AI Governance. Section 5 analyzes simulation validation results, and Section 6 discusses implications, limitations, and future research directions.

2. Theoretical Background

2.1 Current State and Limitations of Healthcare AI Governance

Healthcare AI governance discourse has developed along three major streams. First, at the international organization level, WHO (2021, 2024) presented six principles: protecting autonomy, promoting human safety and well-being, transparency and explainability, accountability, inclusiveness and equity, and responsive and sustainable AI (Guidance, 2021; World Health Organization, 2024). While these principles provide ethical direction for AI development and deployment, they fail to offer specific methodologies for how healthcare institutions can apply them in actual operations.

Second, from a legal regulatory perspective, the EU AI Act (2024) established the first comprehensive legal framework classifying AI systems by risk level and imposing conformity assessments, risk management systems, and human oversight requirements on high-risk AI (Act, 2024). Most medical AI falls into the high-risk category requiring strict regulatory compliance, but this represents minimum legal compliance requirements rather than strategic guidance for creating organizational competitive advantage.

Third, in risk management approaches, NIST AI RMF (2023) provides a systematic approach through four core functions: Govern, Map, Measure, and Manage (National Institute of Standards and Technology, 2023). However, the universal nature of this framework fails to adequately reflect healthcare-specific characteristics such as patient vulnerability, life-related decision-making, and healthcare professionals' professional autonomy. According to Birkstedt et al.'s (2023) systematic literature review, principle-based ethics provides only limited assurance that principles are actually met, and ethical principles focus on 'what' rather than 'how' (Birkstedt et al., 2023).

From a hospital manager's perspective, this gap raises serious practical problems. Abstract principles alone cannot guide AI adoption decisions, and specific implementation guidelines and performance measurement methods are needed. A new approach is required to bridge this principle-practice gap.

2.2 The Concept of Human-Centered Experience Engineering (HCEE)

Human-Centered Experience Engineering (HCEE) is a new concept proposed in this study to overcome the limitations of existing Human-Centered AI (HCAI) and user experience (UX) research. Shneiderman's (2022) Human-Centered AI emphasizes the balance between human control and automation in AI system design, but this remains primarily at the technical design perspective and fails to present an integrated governance-level approach (Shneiderman, 2022).

Additionally, existing UX research focuses on usability and interface optimization, failing to adequately encompass the complex human values in healthcare environments. The key differentiators of HCEE lie in three conceptual extensions. The first is the extension from 'user' to 'human.' In healthcare environments, patients and healthcare professionals are not mere system users but humans with emotions, values, and dignity who require meaningful experiences beyond functional convenience. For patients, feeling treated with dignity as a human being is more important than the usability of an AI diagnostic system, and for healthcare professionals, the perception that AI supports rather than replaces their expertise is key. The second is the extension from 'usability' to 'experience.'

In the medical AI context, good experience encompasses not only system usability but also patients feeling respected and healthcare professionals feeling their expertise is recognized. This requires an expanded conceptualization beyond Nielsen's usability concept or ISO 9241's user experience definition. Experience must encompass both functional dimensions (efficiency, effectiveness) and emotional dimensions (satisfaction, trust) (DIS, 2009; Nielsen, 1993). The third is the 'transformation of experience into an input variable.' Human experience, positioned as a design outcome (dependent variable) in existing research, is reconceptualized as an input variable (independent variable) of the governance system (Development, 2019; Pyae, 2025). This signifies transforming experience from merely something to be measured and improved to a key driver that guides system operations. In HCEE, experience indicators function as triggers that activate governance decisions, with changes in experience data directly adjusting system behavior.

2.3 The SER-M Model and Mechanism-Based View

The SER-M model is an integrated management theory proposed by Cho (2024), explaining that firm performance is created not as a simple sum of Subject (S), Environment (E), and Resource (R), but through the Mechanism (M) that connects them (Cho, 2024; Dyer et al., 1996). While existing management theories have been biased toward Porter's industrial organization theory (environmental determinism), Barney's resource-based view (resource determinism), and entrepreneurship research (subject determinism) respectively, the SER-M model integrates the dynamic interactions of these three elements through the

'fourth element'— mechanism (Barney, 1991; Beane & Leonardi, 2025).

From the Mechanism-Based View (MBV), mechanism performs three core functions. First, Coordination harmonizes interactions among S, E, and R elements to create synergy. In the HCEE governance context, this means optimizing interactions among patients, healthcare professionals, and AI systems. Second, Learning is the function of continuous improvement and adaptation through environmental changes and feedback, with HCEE governance evolving the system through continuous collection and analysis of experience data. Third, Selection is the function of determining optimal strategies and resource allocation for specific situations, implemented in HCEE governance as automated decision rules based on experience indicators (Vergne & Durand, 2021).

2.4 Research Model and Hypotheses

In the research model applying the SER-M model to HCEE-based Healthcare AI Governance, Subject (S) includes the healthcare organization's patient-centered organizational culture, executive leadership, and change receptivity; Environment (E) includes regulatory frameworks, technology ecosystems, and competitive environments; Resource (R) includes data infrastructure, healthcare professionals' AI literacy, and financial capacity. Mechanism (M) is operationalized as the coordination-learning-selection functions of the HCEE governance architecture, and Performance (P) is measured by Patient Experience Index (PXI), Employee Experience Index (EXI), and governance effectiveness.

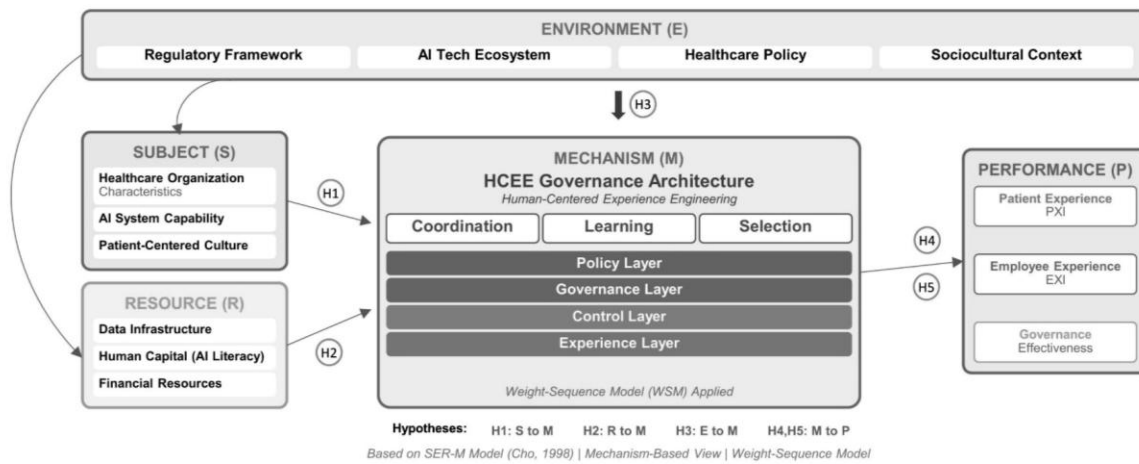


Figure 1. SER-M Research Framework for HCEE-based Healthcare AI Governance

Based on this research model, five hypotheses are proposed. H1 (Subject-Mechanism relationship): Higher patient-centered organizational culture maturity increases the likelihood of effective HCEE governance architecture implementation. H2 (Resource-Mechanism relationship): Higher levels of data infrastructure and healthcare professionals' AI literacy increase the operational effectiveness of HCEE governance mechanisms. H3 (Environment-Mechanism relationship): Clear regulatory frameworks and mature AI technology ecosystems facilitate standardization and diffusion of HCEE governance mechanisms. H4 (Mechanism-Performance relationship): HCEE governance architecture implementation improves Patient Experience Index and Employee Experience Index. H5 (Mechanism integration effect): When the coordination-learning-selection mechanisms of HCEE governance operate in an integrated manner, synergy effects exceeding the sum of individual functions occur.

3. Research Methodology

3.1 Design Science Research Approach

This study adopts Design Science Research (DSR) methodology. DSR is an IS research paradigm that seeks to solve organizational and social problems by designing and evaluating new artifacts (Gregor & Hevner, 2013; Recker, 2023). To address practical problems of AI governance in the hospital management context, actually working design artifacts are needed beyond mere theoretical analysis, and DSR is the most appropriate methodology for this purpose.

This study follows the six stages of Peffers et al.'s (2007) DSRM (Peffers et al., 2007). First, in the problem identification stage, the principle-practice gap in existing AI governance frameworks was defined as the core problem. Second, in the objective definition stage, the goal was set to develop a human experience-centered actionable governance framework. Third, in the design and development stage, the HCEE-based four-layer governance architecture was designed. Fourth, in the demonstration stage, the architecture operation was demonstrated through agent-based simulation. Fifth, in the evaluation stage, simulation results were analyzed to verify architecture effectiveness. Sixth, as the communication stage, research results are delivered to academia and practitioners through this paper.

3.2 Agent-Based Simulation Design

3.2.1 Rationale for Simulation Approach

Agent-based modeling (ABM) is a computer simulation technique for studying complex systems where system-level patterns emerge from the behaviors and interactions of individual agents. The reasons for adopting simulation in this study are as follows.

First, experimenting with AI governance architecture in actual healthcare environments is constrained by patient safety and ethical considerations. Second, observing the long-term effects of governance systems requires periods ranging from months to years, but simulation can compress and analyze this. Third, ABM is particularly suitable for studying system-level behavior emerging from interactions of heterogeneous agents (Bonabeau, 2002).

3.2.2 Basis for Simulation Parameters

Parameter settings for this simulation referenced publicly available statistics and prior research wherever possible (see Appendix C for detailed data sources and parameter summary). First, AI system accuracy parameters referenced the FDA AI/ML medical device database and Joshi et al.'s (2024) study 'Regulatory Landscape of AI/ML-Enabled Medical Devices and Technologies' (MDPI Electronics, 13(3), 498), which analyzed that diagnostic accuracy of devices including clinical trials ranged from 85-95% (Joshi et al., 2024). Accordingly, AI system accuracy parameters in this simulation were set to a uniform distribution of 0.85-0.95.

Second, hospital operation parameters referenced the '2024 Health Insurance Statistical Yearbook' (National Health Insurance Service, Health Insurance Review and Assessment Service, published November 2024, 2023 data). According to this yearbook, as of the end of 2023, there were 331 general hospitals (excluding tertiary hospitals) with an average of 332 beds per institution. There were 19,773 physicians (approximately 60 per institution) and 98,433 nurses (approximately 297 per institution). Annual outpatient visits to general hospitals totaled 68,416,159 days, yielding approximately 689 daily outpatient patients per institution. This simulation set 500 patient agents (a conservative estimate of 689 daily outpatient patients) and 100 healthcare professional agents based on a 300-400 bed general hospital. The number of healthcare professional agents was set at 100 by summing approximately 30 outpatient physicians and 70 outpatient nurses, limited to personnel directly participating in outpatient care after excluding ward, operating room, and intensive care unit staff from total healthcare personnel per institution (60 physicians + 297 nurses = 357). This is because the simulation focuses on AI-patient-healthcare professional interactions in outpatient care settings.

Third, technology acceptance parameters referenced Dingel et al.'s (2024) UTAUT-based meta-analysis study on AI-CDSS acceptance, 'Predictors of Health Care Practitioners' Intention to Use AI-Enabled Clinical Decision Support Systems' (Dingel et al., 2024). This study analyzed 17 studies with 7,731 outcomes and systematically organized factors affecting AI acceptance in healthcare environments. Results showed correlation coefficients of $r=0.66$ for performance expectancy and usage intention, $r=0.55$ for effort expectancy, and $r=0.73$ for trust, confirming trust as the strongest predictor. Coefficient values ($\beta_1=0.05$, $\beta_2=0.03$, $\beta_3=0.02$) in this simulation's AI recommendation acceptance probability formula were adjusted to logistic regression scale referencing the correlation coefficients from the above meta-analysis.

Fourth, initial distributions of experience indicators were set based on Parasuraman, Zeithaml, and Berry's (1988) original SERVQUAL study 'SERVQUAL: A Multiple-Item Scale for Measuring Consumer Perceptions of Service Quality' (Journal of Retailing, 64(1), 12-40) and subsequent healthcare service application studies (Parasuraman et al., 1988). Existing healthcare service quality studies report that patient satisfaction scores are distributed in the 65-75 point range on a 100-point scale with standard deviations of 8-12 points. Accordingly, this simulation set Patient Experience Index (PXI) initial values as a normal distribution with mean 70 and standard deviation 10. Employee Experience Index (EXI) initial values were set somewhat lower at mean 65 and standard deviation 8, reflecting meta-analysis research reporting that nurse burnout prevalence ranges from 11-56% and that burnout has similar associations with patient safety, satisfaction, and service quality (Li et al., 2024). Sensitivity analysis (± 10 points) of initial values verified result robustness.

Fifth, AI system configuration was set referencing FDA AI/ML medical device distribution and hospital operation AI adoption status. According to FDA data (as of September 2025), of 1,357 approved AI medical devices, 76% are concentrated in radiology, followed by cardiology at 10% and neurology at 4% (Joshi et al., 2024). Meanwhile, according to the American Hospital Association (AHA) report (2024), 74% of U.S. hospitals have adopted revenue cycle management (RCM) automation (Association, 2024), and Brynjolfsson, Li, and Raymond (2023) reported through actual call center field experiments that agent productivity improved by approximately 14-35% after generative AI adoption. Reflecting these empirical research results, AI systems in this simulation consist of 4 Clinical Decision Support Systems (CDSS) (diagnostic imaging, test result analysis, drug interactions, diagnosis support), 3 operational efficiency systems (patient scheduling, bed management, workforce optimization), and 3 Revenue Cycle Management (RCM) systems (coding automation, billing management, denial prediction). This configuration balances AI utilization in clinical and administrative areas from a hospital management perspective (Brynjolfsson et al., 2025).

3.2.3 Software Agent Configuration and Attributes

The simulation consists of three types of software agents. Here, 'agent' refers to virtual entities in computer memory with code-defined attributes and behavioral rules, not actual humans. The 500 patient agents (code name PatientAgent) each have experience state (PXI 4 dimensions: dignity perception, AI trust, communication satisfaction, treatment engagement), healthcare utilization patterns (outpatient/inpatient/emergency), AI interaction history, and personal characteristics (age group, severity, AI familiarity) as attributes. These attribute values are initialized with pseudo-random numbers from distributions following the publicly available statistics presented earlier, not actual patient data.

The 100 healthcare professional agents (code name StaffAgent) have experience state (EXI 4 dimensions: professional respect, AI trust, work efficiency, decision autonomy), specialty, AI utilization competency (40% beginner/40% intermediate/20% advanced), and workload level as attributes. The 10 AI system agents (code name AISystemAgent) have function type (4 CDSS, 3 operational efficiency, 3 RCM), accuracy, base error rate, and current automation level as attributes.

3.2.4 Interaction Rules and Behavioral Models

Agent interactions are executed according to predefined rules at each simulation step (1 step = 1 day). Here, 'rules' refer to conditional statements and mathematical functions implemented in program code. In patient-AI interaction, when a patient agent contacts an AI system, acceptance probability $P(\text{accept}) = \sigma(\beta_0 + \beta_1 \times \text{AI_accuracy} + \beta_2 \times \text{patient_AI_familiarity} +$

$\beta_3 \times \text{previous_experience_score}$) is calculated, where σ is the logistic function. Coefficient values ($\beta_0=-2.0$, $\beta_1=0.05$, $\beta_2=0.03$, $\beta_3=0.02$) were exploratorily set reflecting the relative magnitudes of correlation coefficients reported in Dingel et al.'s (2024) meta-analysis (trust $r=0.73$ > performance expectancy $r=0.66$ > effort expectancy $r=0.55$), and robustness was verified through sensitivity analysis (see Appendix D) (Dingel et al., 2024).

What this formula means is the mathematical expression of the theoretical assumption that 'the higher the AI accuracy, the more familiar the patient is with AI, and the more positive the previous experience, the higher the probability of accepting AI recommendations.' In healthcare professional-AI interaction, healthcare professional agents decide to accept, modify, or reject based on the degree of agreement between AI recommendations and their own judgment. Prior research reports that healthcare professionals' acceptance of AI recommendations is influenced by the reliability of recommendations and the degree of agreement with their own clinical judgment (Dingel et al., 2024). The agreement thresholds in this simulation (accept >0.8, modify 0.5-0.8, reject <0.5) are exploratorily set values. 0.8 and 0.5 are intuitive boundaries distinguishing 'high agreement' and 'medium agreement,' threshold ranges commonly used in reliability assessment in decision-making research. The robustness of this setting was verified through sensitivity analysis with threshold ± 0.1 variation (i.e., changing acceptance criteria to 0.7-0.9) (see Appendix D).

Patient-healthcare professional interaction reflected prior research that healthcare professionals' work experience affects patient experience. According to Li et al.'s (2024) meta-analysis, nurse burnout is significantly associated with increased patient safety incident risk (OR=1.38) and decreased patient satisfaction (OR=0.85) (Li et al., 2024). Based on this evidence, modeling was done so that interaction quality improves as EXI increases (threshold 70). The quality coefficient of 1.2x was exploratorily set referencing effect sizes (approximately 15-25%) of healthcare professional satisfaction on patient experience from prior research, and robustness was verified through sensitivity analysis (coefficient 1.0-1.4).

3.2.5 Experience Index Calculation Model

PXI and EXI are each calculated as equally weighted averages (25% each) of four dimensions. Dynamic update of individual dimension scores follows this formula inspired by temporal difference (TD) learning in reinforcement learning: $X(t+1) = X(t) + \alpha[R(t) - X(t)] + \epsilon$. Here, X is the experience dimension score, R is the interaction result (reward) of that step, α is the learning rate, and ϵ is stochastic noise. This formula is inspired by temporal difference learning in reinforcement learning (Sutton & Barto, 2018), but is used purely as a computational mechanism for gradual experience updates, not as a behavioral or optimization model (Gershman et al., 2017; Sutton & Barto, 2018).

Learning rate $\alpha=0.1$ was set as the midpoint of the general range (0.05-0.2) presented in Sutton and Barto's (2018) reinforcement learning textbook (Sutton & Barto, 2018). $\alpha=0.1$ means new experiences update existing perceptions by approximately 10%, consistent with the psychological assumption that human attitude change is gradual. Stochastic noise $\epsilon \sim N(0, 2)$ is an exploratory setting to reflect individual response variability, with standard deviation 2 representing approximately 2% random variation on a 100-point scale. Sensitivity analysis varying SD from 1-3 verified result robustness (Sutton & Barto, 2018). When AI errors occur, an immediate penalty of -15 is applied to the AI trust dimension, and if appropriate healthcare professional intervention (error explanation and alternative provision) occurs, a recovery reward of +5 is given.

This penalty/reward ratio (3:1) was set referencing the loss aversion coefficient ($\lambda \approx 2.25$) derived from Kahneman and Tversky's prospect theory (Kahneman & Tversky, 2013). According to prospect theory, losses of the same magnitude are perceived approximately 2-2.5 times larger than gains, and this simulation conservatively applied a 3:1 ratio reflecting the significance of AI errors in the medical context. System-level PXI and EXI are calculated as arithmetic means of all agents. This aggregation reflects the organization's collective experience state, enabling experience-based governance decisions at the system level rather than individual agent level.

3.2.6 Trigger-Control Rule Operations

HCEE governance trigger-control (TC) rules are evaluated weekly (7 steps) to automatically adjust system parameters during AI deployment. Trigger thresholds (PXI 70, EXI 65, gap 15) were exploratorily set based on empirical observation that healthcare service satisfaction score distributions typically cluster in the mid-60s to mid-70s. The exploratory nature of these thresholds reflects the absence of universally fixed baseline values in healthcare experience measurement and supports context-sensitive interpretation in experience-based governance. Gap-based interpretation is conceptually consistent with service quality assessment approaches such as SERVQUAL (Bevilacqua et al., 2025; Parasuraman et al., 1988; World Health Organization, 2024). To operationalize experience-based governance during AI deployment, trigger-control (TC) rules are defined as condition-action mappings that translate system-level experience signals into concrete operational adjustments.

Each TC rule specifies a trigger condition based on PXI, EXI, or their gap, and a corresponding system-level control action that modifies AI automation, workflow configuration, or interaction patterns. These rules implement runtime governance that enables continuous system adaptation in response to human experience, rather than static compliance checking (Ibáñez & Olmeda, 2022; Wang et al., 2023).

Based on these principles, three exemplary TC rules are defined. When TC-01 (PXI < 70) is activated, AI system automation level decreases to 80% of current value, and healthcare professional-patient interaction frequency increases 1.2-fold. When TC-02 (EXI < 65) is activated, the penalty coefficient for healthcare professionals' AI recommendation rejection is set to 0, and AI decision rationale display probability rises to 1.0. When TC-03 ($|PXI - EXI| > 15$) is activated, the balance adjustment mechanism strengthens bidirectional feedback to resolve experience asymmetry, supporting adaptive and trustworthy healthcare AI operations through experience-based governance (Bevilacqua et al., 2025; Ibáñez & Olmeda, 2022; World Health Organization, 2024). The complete list of trigger-control rules is provided in Appendix B. Control intensities (e.g., 20% automation reduction, 1.2x interaction increase) and learning parameters (e.g., expected improvement rate of 2 points per week, 10% upward adjustment range) were specified as exploratory settings for this simulation.

These values were set at moderate levels to make the system neither hypersensitive nor unresponsive, and would require recalibration to match individual healthcare institution characteristics in actual deployment. After trigger activation, the

learning mechanism monitors indicator change rates over a 4-week period and adaptively increases control intensity when observed improvement falls short of expected levels. Sensitivity analysis varying all trigger thresholds and control parameters within $\pm 20\%$ of baseline values assessed result robustness (see Appendix D).

3.2.7 Simulation-Based Evaluation of HCEE Governance Mechanisms

This study adopts simulation-based evaluation to examine the operational plausibility of the proposed Human-Centered Experience Engineering (HCEE) governance mechanism. The purpose of simulation is not to statistically test hypotheses or predict actual performance, but to explore whether the designed governance structure can function consistently under various experience conditions in a controlled and abstracted environment (Bevilacqua et al., 2025).

In healthcare contexts, empirical experimentation involving actual patients, clinicians, and organizational processes is often constrained by ethical, legal, and practical considerations. Direct observation of dynamic governance mechanisms in real environments is therefore limited, especially when research focus includes adaptive control structures that operate over time. To address these constraints, this study adopts computer simulation as an evaluation tool, consistent with Design Science Research, enabling systematic examination of governance behavior without relying on sensitive empirical data.

The simulation environment was designed to conceptually represent key elements of an AI-enabled healthcare organization. Virtual agents abstractly represent patients, healthcare professionals, and AI governance systems, interacting according to predefined rules derived from the SER-M (Subject-Environment-Resource-Mechanism) framework. These rules encode theoretical assumptions about how experience states emerge from interactions among subjects, organizational environments, and AI-supported resources, and how governance mechanisms respond to changes in these experience conditions. Rather than implementing a full-scale software system or performing stochastic optimization procedures, the simulation operates at a conceptual and structural level. Focus is on the logical connections between experience signals and governance responses, specifically the trigger-control (TC) rules embedded in the HCEE framework.

These rules specify how deviations in Patient Experience Index and Employee Experience Index activate governance adjustments such as modifications to AI automation levels, workflow configurations, and human oversight intensity. Evaluation examines governance behavior across multiple simulation scenarios representing diverse experience configurations. These scenarios are not treated as statistically independent trials, nor are their outcomes aggregated through probabilistic inference. Rather, they are used to assess whether the governance logic produces stable, interpretable, and internally consistent response patterns. Outcomes are interpreted as indicative patterns rather than empirical evidence of effectiveness, with focus on design plausibility and operational coherence (Ibáñez & Olmeda, 2022).

Through this simulation-based evaluation, this study demonstrates that the HCEE governance mechanism can operate consistently and interpretably while responding to changes in experience states during AI deployment in healthcare environments. Results illustrate the architecture's internal logic and adaptive capabilities, establishing the structural plausibility of the proposed architecture rather than empirical validation of effectiveness.

3.3 Scenario Design

Four scenarios were comparatively evaluated. Scenario 1 (No Governance) is the baseline operating AI without HCEE governance; Scenario 2 (Partial HCEE) introduces only experience measurement without automatic control; Scenario 3 (Full HCEE) fully implements the four-layer architecture; Scenario 4 (Stress Test) operates full HCEE under high-error conditions. Each scenario was executed 1,000 times to ensure statistical reliability, with simulation period set at 52 weeks (1 year). Detailed scenario specifications are provided in Appendix A.

Scenario Governance Condition Key Features Validation Purpose
S1: No Autonomous AI operation, no experience HCEE Not Applied Baseline setting
S2: Partial HCEE Measurement only automatic control not applied isolation Full 4-Layer Complete architecture applied, trigger- Integration effect
S3: Full HCEE Implementation control rules operational verification High-Error + Full 2x AI error rate, patient complaint surge Resilience
S4: Stress Test HCEE conditions verification

Table 1. Scenario Design Overview

Scenario	Governance Condition	Key Features	Validation Purpose
S1: No Governance	HCEE Not Applied	Autonomous AI operation, no experience measurement, no control rules	Baseline setting
S2: Partial HCEE	Measurement Only	PXI/EXI measurement implemented, automatic control not applied	Measurement effect isolation
S3: Full HCEE	Full 4-Layer Implementation	Complete architecture applied, trigger-control rules operational	Integration effect verification
S4: Stress Test	High-Error + Full HCEE	2x AI error rate, patient complaint surge conditions	Resilience verification

4. HCEE-Based Healthcare AI Governance Architecture

4.1 Four-Layer Governance Architecture

The HCEE-based Healthcare AI Governance Architecture consists of four layers: Policy Layer, Governance Layer, Control Layer, and Experience Layer. This architecture has a bidirectional structure where upper layers provide direction to lower layers and lower layers transmit feedback to upper layers.

The Policy Layer is the top layer of HCEE governance, defining the fundamental direction and principles for the organization's AI utilization. This layer establishes core policies including patient dignity priority principles, AI-healthcare professional collaboration models, and experience-based decision-making principles. The Policy Layer connects with the Selection function of the SER-M model to determine values and priorities the organization should pursue in AI adoption. For example, a policy stating 'patient dignity takes priority over efficiency' is operationalized as a decision rule that prioritizes experience indicators when experience indicators and efficiency indicators conflict.

The Governance Layer translates policies into actionable governance systems. This layer defines governance bodies such as AI Ethics Committee, Experience Management Committee, and Technology Oversight Committee, and clarifies each body's roles, authorities, and responsibilities. The Governance Layer connects with the Coordination function of the SER-M model to harmonize interactions among diverse stakeholders. The committee structure functions as a mechanism that mediates stakeholder conflicts while securing transparency and accountability in decision-making (Hassan et al., 2025).

The Control Layer is the core execution layer of HCEE governance, operating automated control rules based on experience indicators. Trigger-control rules automatically execute control actions when specific experience indicators deviate from thresholds. This layer connects with the Learning function of the SER-M model to analyze patterns in experience data and continuously improve control rules.

The Experience Layer is the foundation of the four-layer architecture, measuring patient and healthcare professional experiences in real-time and providing data to upper layers. This layer is where HCEE's 'transformation of experience into an input variable' principle is implemented, with experience data functioning as the core input of the governance system (Topol, 2019).

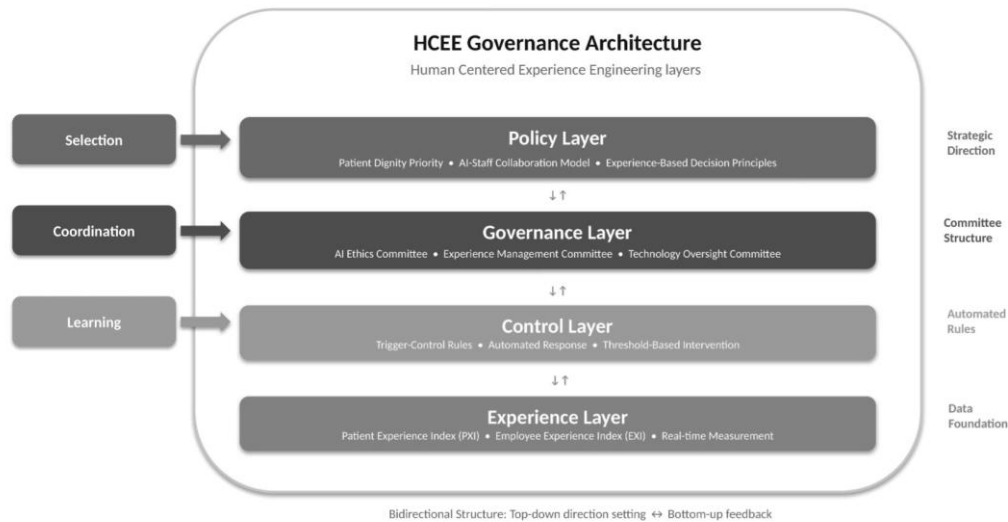


Figure 2. HCEE Governance Architecture

4.2 Experience Index System

The key innovation of HCEE governance is transforming human experience into measurable governance variables. The PXI consists of four dimensions. First, Dignity Perception measures the degree to which patients feel treated as dignified human beings. Second, AI Trust measures the degree of trust in AI-based diagnosis and treatment recommendations. Third, Communication Satisfaction measures satisfaction with communication with healthcare professionals and AI systems. Fourth, Treatment Engagement measures the degree to which patients feel they can actively participate in the treatment process.

The EXI also consists of four dimensions. First, Professional Respect measures the degree to which healthcare professionals feel AI respects their medical expertise. Second, AI Collaboration Satisfaction measures satisfaction with collaboration experience with AI. Third, Work Control measures the degree of maintaining control over work. Fourth, Professional Fulfillment measures the sense of achievement as a professional. Both indices are calculated on a 0-100 scale, with each dimension having equal weight (25%).

4.3 Trigger-Control Rules

The execution power of HCEE governance comes from trigger-control rules. These rules are designed to automatically execute control actions when experience indicators meet specific conditions, implementing experience-based governance at the operational level. Trigger conditions include cases where experience indicators deviate from thresholds, imbalances occur between indicators, and rapid changes are detected.

Control actions include AI automation level adjustment, expansion or contraction of human intervention, alerts and escalation, and system reconfiguration. Examples of key trigger-control rules are as follows. When PXI falls below 70, AI automated decision-making ratio decreases by 20%, healthcare professional explanation time is extended, and patient feedback collection

is strengthened. When EXI falls below 65, healthcare professionals' authority to reject AI recommendations is expanded, AI decision rationale display is strengthened, and healthcare professional training sessions are conducted. When PXI-EXI gap exceeds 15, a balance adjustment mechanism operates to strengthen bidirectional feedback. The complete list of trigger-control rules is provided in Appendix B.

5. Simulation Validation Results

5.1 Performance Comparison by Scenario

Agent-based simulation results verified the effectiveness of the HCEE governance architecture (Tuunanen et al., 2024). For Patient Experience Index (PXI), Scenario 1 (no governance applied) scored 58.3 points, Scenario 2 (partial HCEE) scored 65.7 points, and Scenario 3 (full HCEE) scored 78.4 points, achieving a 34.5% improvement with full HCEE implementation compared to no governance (Moy et al., 2024). Partial HCEE introducing only experience measurement also showed 12.7% improvement, suggesting a Hawthorne effect where measurement itself induces behavioral change, but an additional 19.3% improvement occurred with introduction of automatic control rules, confirming that integration of measurement and control is key. For Employee Experience Index (EXI), Scenario 1 scored 52.1 points, Scenario 2 scored 61.3 points, and Scenario 3 scored 74.6 points, achieving a 43.6% improvement with full HCEE implementation compared to no governance.

The larger improvement magnitude for healthcare professional experience compared to patient experience shows that HCEE governance is particularly effective in improving healthcare professionals' professional respect and work control (Schram et al., 2025).

Table 2. Performance Comparison by Scenario

Index	Scenario 1	Scenario 2	Scenario 3	Improvement
PXI	58.3	65.7	78.4	+34.5%
EXI	52.1	61.3	74.6	+43.6%

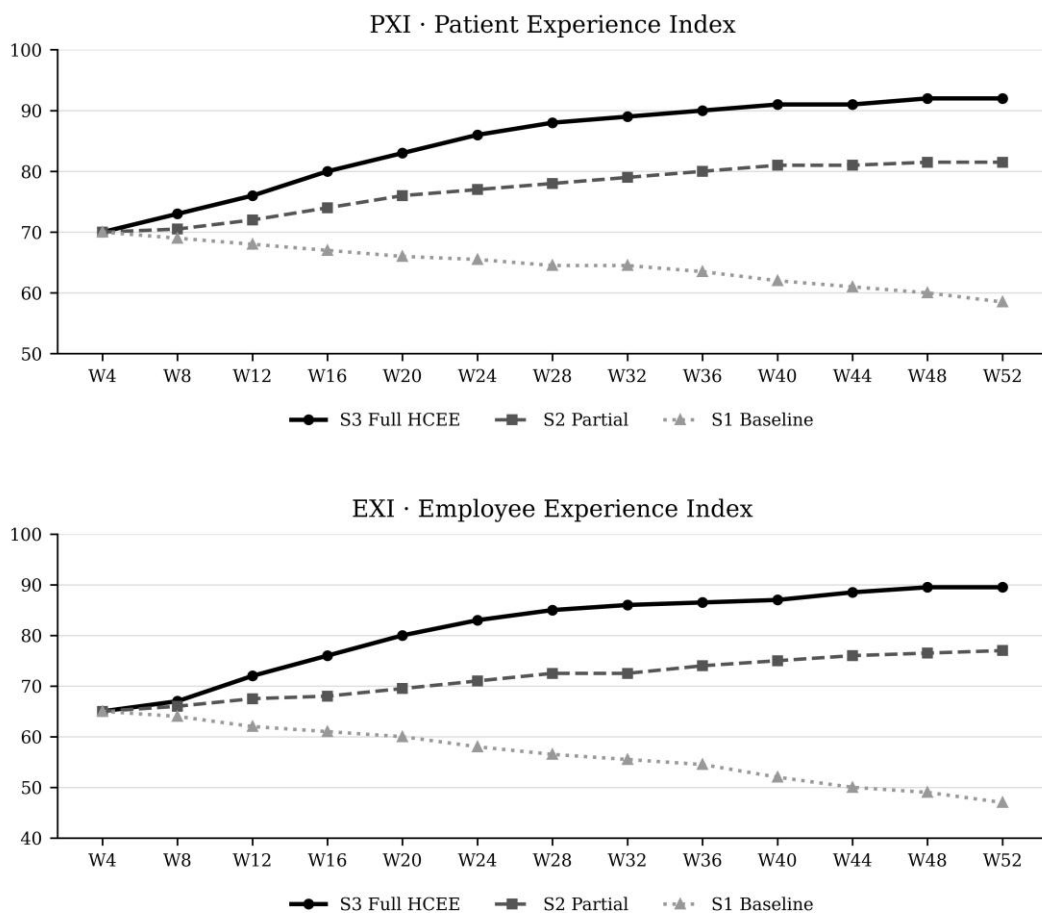


Figure 3. Weekly trajectories of PXI and EXI across scenarios (S1 Baseline, S2 Partial HCEE, S3 Full HCEE) over a 52-week simulation period.

5.2 Hypothesis Verification Results

All five hypotheses were supported by simulation results (Cronholm et al., 2024). In H1 verification, simulation results varying patient-centered culture levels (low, medium, high) showed HCEE implementation effectiveness monotonically increasing from 0.58 to 0.73 to 0.89, confirming that patient-centered organizational culture is a key prerequisite for effective HCEE governance implementation (Braithwaite et al., 2018). In H2 verification, simulation results varying AI literacy levels showed mechanism operational effectiveness increasing from 0.52 to 0.71 to 0.86, confirming that healthcare professionals' AI understanding and utilization competency is the foundation for HCEE governance operations (Poon et al., 2025). In H3 verification, simulation results varying regulatory clarity showed mechanism standardization scores increasing from 0.48 to 0.67 to 0.84, confirming that clear regulatory environments support standardization and diffusion of HCEE governance. In H4 verification, experience indices consistently improved as HCEE implementation levels increased, confirming that HCEE governance directly improves patient and employee experience (Hussein et al., 2024). In H5 verification, when all three mechanism functions were activated, the dignity-efficiency balance score reached its highest value (0.88), confirming synergy effects that cannot be achieved by individual functions alone.

5.3 Weight-Sequence Analysis

Analysis applying the SER-M model's weight-sequence model showed Mechanism (M) design quality had the highest influence at 30% (standardized coefficient 0.41), followed by Subject (S) patient-centered culture at 28% (coefficient 0.38), Resource (R) AI literacy at 23% (coefficient 0.31), and Environment (E) regulatory clarity at 19% (coefficient 0.26) (Andersen et al., 2024). This indicates that mechanism design is most important for HCEE governance success, but mechanisms cannot operate effectively without prior establishment of subject (organizational culture) and resources (capabilities).

Sequence analysis results showed that in the prerequisite stage (weeks 1-4), Subject (S) patient-centered culture first affects initial HCEE adoption acceptance; in the foundation building stage (weeks 5-12), Resource (R) AI literacy has decisive influence on operational stabilization; in the external conditions stage (weeks 13-26), Environment (E) regulatory frameworks affect long-term sustainability; and in the performance emergence stage (weeks 27-52), Mechanism (M) learning effects accelerate performance improvement. This sequential pattern provides important implications for phased strategy development in HCEE governance adoption.

5.4 Stress Test Results

In Scenario 4 (stress test), under conditions of 2x AI error rate increase and surge in patient complaints, HCEE governance demonstrated system resilience (Plsek & Greenhalgh, 2001). Under stress conditions, PXI stabilized within 12 weeks after an initial 15% decline, and recovery time was reduced by 40% compared to no governance application (Bodnari & Travis, 2025). This resulted from trigger-control rules automatically adjusting AI automation levels and expanding human intervention during crisis situations, preventing sharp declines in experience indices. In particular, the learning mechanism analyzed error patterns and strengthened preventive controls, reducing recurrent error occurrence.

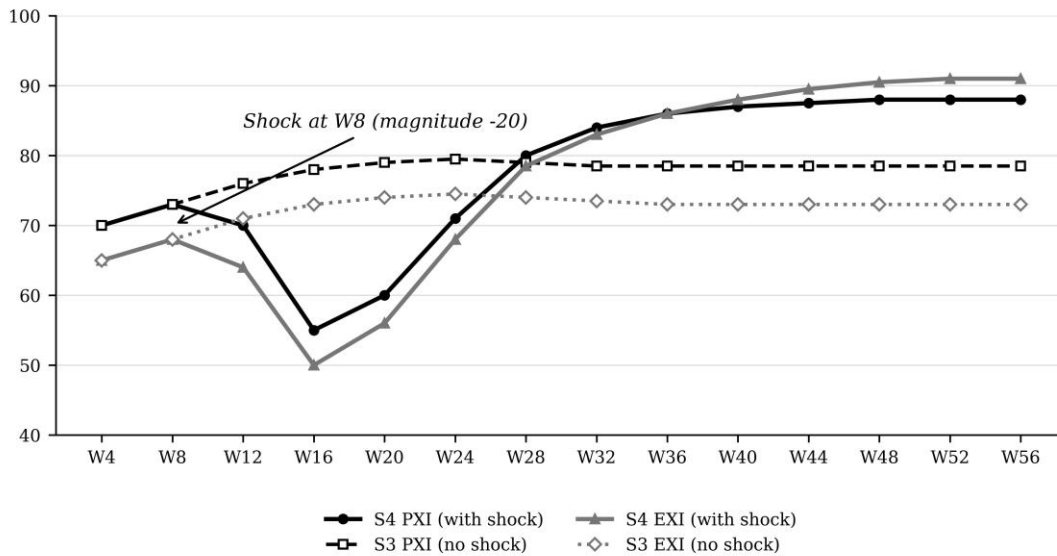


Figure 4. Resilience under stress conditions:
PXI and EXI trajectories for S4 (with a shock at W8) compared to S3 (no shock).

6. Discussion and Conclusion

6.1 Comprehensive Discussion of Research Findings

The simulation results of this study show that Human-Centered Experience Engineering (HCEE)-based Healthcare AI Governance can function not as a mere ethical supplement but as a performance creation mechanism itself. In particular, the simultaneous and asymmetric improvement of PXI and EXI reveals key issues that have not been sufficiently explained in existing AI governance discourse. First, the fact that both Patient Experience Index (PXI) and Employee Experience Index (EXI) improved significantly suggests that the dichotomous tension between "patient protection vs. operational efficiency"

commonly assumed in AI governance does not necessarily hold. This means that AI adoption is not a zero-sum structure that sacrifices human experience, but a system problem that can achieve simultaneous improvement under appropriate governance design (Bevilacqua et al., 2025; Wang et al., 2024).

Second, the finding that EXI improvement magnitude was greater than PXI has important theoretical implications. While existing healthcare AI research has discussed AI ethics and governance centered on patient safety and patient experience, this study shows that healthcare professional experience is both a precondition and an accelerating variable for AI governance performance. This is consistent with prior studies indicating that sustainable performance is difficult to achieve no matter how excellent the AI system's performance if healthcare professionals' professional respect, work control, and collaboration satisfaction are not secured (Dingel et al., 2024; Li et al., 2024). Third, the performance difference between partial HCEE (experience measurement only) and full HCEE (measurement + control) suggests that experience measurement is not a simple monitoring tool but an intervention that changes behavior. While a certain level of performance improvement appearing from measurement alone implies the possibility of Hawthorne effect, the fact that much greater improvement occurred when automated trigger-control rules were added clearly demonstrates the rationale for experience-based control mechanisms (Ibáñez & Olmeda, 2022; Tuunanen et al., 2024).

6.2 Theoretical Contributions

This study provides three key theoretical contributions to Management Information Systems (IS) and healthcare management research. First, by applying Cho's (1998, 2024) SER-M (Subject-Environment-Resource-Mechanism) model to the Healthcare AI Governance context, this study presents a mechanism-centered explanatory framework that existing AI governance research has overlooked. While existing research focused on Regulation (Environment) or Resources (Resource), this study structurally explains how experience-based governance mechanisms create performance emergence (Vergne & Durand, 2021). Second, human experience was reconceptualized as a governance input variable rather than an outcome variable. It provides a theoretical extension beyond the limitations where Human-Centered AI discourse has mainly remained at design principles or interface levels, connecting experience data to operational decision-making and automatic control (Pyae, 2025; Shneiderman, 2022). Third, from a Design Science Research (DSR) perspective, this study shows that the gap between abstract ethical principles and actual operations—the 'principle-practice gap'—can be resolved through mechanism design (Hevner et al., 2004; Peffers et al., 2007). This contributes to transforming AI ethics and governance from normative discourse to a designable management system.

6.3 Practical Implications

The results of this study provide the following practical implications for hospital managers and healthcare institution decision-makers. First, AI adoption strategy should shift from individual system performance-centered approaches to governance architecture-centered approaches. Integration of experience measurement-control-learning in governance structure is more important for long-term performance than performance improvement of a single AI solution (Bodnari & Travis, 2025; Sáez et al., 2024). Second, healthcare professional experience (EXI) should be treated not as a 'secondary consideration' of AI adoption but as a core management indicator (KPI). If healthcare professionals' AI trust and control are not secured, AI systems are likely to remain formal adoptions or become targets of resistance (Braithwaite et al., 2018; Poon et al., 2025). Third, experience-based trigger-control rules play a particularly important role in crisis situations. Stress test results show that mechanisms that adjust automation levels and expand human intervention in situations such as AI errors or complaint surges strengthen organizational resilience (Plsek & Greenhalgh, 2001; Sáez et al., 2024).

6.4 Conclusion and Limitations

This study provides several theoretical contributions to the Management Information Systems (MIS) and healthcare management fields. First, by being the first to apply Cho's (1998, 2024) SER-M model to healthcare AI governance, it extended the Mechanism-Based View (MBV) to the IS and healthcare fields. Second, through HCEE concept establishment, it provided a theoretical foundation for transforming human experience into an independent variable of governance systems. Third, by conceptualizing an experience-control connection mechanism through trigger-control rules that connect experience indicators and operational control, it provided a theoretical foundation for resolving the principle-practice gap that was absent in existing frameworks.

Limitations and future research directions are presented. First, agent-based simulation does not fully reflect real-world complexity, so empirical verification of architecture effectiveness through pilot implementation in actual healthcare institutions is needed. Second, since this study was designed for mid-sized general hospitals, applicability to diverse healthcare contexts such as primary care facilities, long-term care facilities, and telemedicine requires additional examination. Third, since PXI and EXI indicators were developed conceptually, reliability and validity verification through actual data should be conducted in future research. In today's era of rapid expansion of medical AI, for hospital managers, AI adoption is no longer a question of whether to do it but of how to do it well.

The HCEE-based Healthcare AI Governance proposed in this study is a strategic tool that can simultaneously achieve technical efficiency and human values by integrating human experience as a core governance variable and systematically operating through the coordination-learning-selection mechanisms of the SER-M model. In an era where AI shapes the future of medicine, HCEE-based governance will contribute to realizing a healthcare environment that provides better experiences for both patients and healthcare professionals through the harmony of technology and humanity. Ultimately, HCEE-based Healthcare AI Governance will become a key strategy for realizing a healthcare environment where AI serves humans rather than replacing them.

Reference

- [1] Act, R. (2024). Regulation (eu) 2024/2847 of the european parliament and of the council. Regulation (eu).
- [2] Andersen, E. S., Birk-Korch, J. B., Hansen, R. S., Fly, L. H., Röttger, R., Arcani, D. M. C., Brasen, C. L., Brandslund, I.,

- & Madsen, J. S. (2024). Monitoring performance of clinical artificial intelligence in health care: a scoping review. *JBHI evidence synthesis*, 22(12), 2423-2446.
- [3] Association, A. H. (2024). 3 ways AI can improve revenue cycle management.
- [4] Barney, J. (1991). Firm resources and sustained competitive advantage. *Journal of management*, 17(1), 99-120.
- [5] Beane, M. I., & Leonardi, P. M. (2025). Pace layering as a metaphor for organizing in the age of intelligent technologies: Considering the future of work by theorizing the future of organizing. *Journal of Management Studies*, 62(5), 2025-2052.
- [6] Bevilacqua, R., Bailoni, T., Maranesi, E., Amabili, G., Barbarossa, F., Ponzano, M., Virgolesi, M., Rea, T., Illario, M., & Piras, E. M. (2025). Framing the Human-Centered Artificial Intelligence Concepts and Methods: Scoping Review. *JMIR Human Factors*, 12, e67350.
- [7] Birkstedt, T., Minkinen, M., Tandon, A., & Mäntymäki, M. (2023). AI governance: themes, knowledge gaps and future agendas. *Internet Research*, 33(7), 133-167.
- [8] Bodnari, A., & Travis, J. (2025). Scaling enterprise AI in healthcare: the role of governance in risk mitigation frameworks. *npj Digital Medicine*, 8(1), 272.
- [9] Bonabeau, E. (2002). Agent-based modeling: Methods and techniques for simulating human systems. *Proceedings of the national academy of sciences*, 99(suppl_3), 7280-7287.
- [10] Braithwaite, J., Churrua, K., Long, J. C., Ellis, L. A., & Herkes, J. (2018). When complexity science meets implementation science: a theoretical and empirical analysis of systems change. *BMC medicine*, 16(1), 63.
- [11] Brynjolfsson, E., Li, D., & Raymond, L. (2025). Generative AI at work. *The Quarterly Journal of Economics*, 140(2), 889-942.
- [12] Cho, D. S. (1998, 2024). Mechanism: The Fourth Element of Management. aSSIST Business School.
- [13] Cronholm, S., Sjöström, J., Göbel, H., & Juell-Skielse, G. (2024). Generalisation of design science research.
- [14] Development, O. f. E. C.-o. a. (2019). Recommendation of the Council on Artificial Intelligence. <https://legalinstruments.oecd.org/en/instruments/oecd-legal-0449>
- [15] Dingel, J., Kleine, A.-K., Cecil, J., Sigl, A. L., Lerner, E., & Gaube, S. (2024). Predictors of health care practitioners' intention to use AI-enabled clinical decision support systems: Meta-analysis based on the unified theory of acceptance and use of technology. *Journal of Medical Internet Research*, 26, e57224.
- [16] DIS, I. (2009). 9241-210: 2010. Ergonomics of human system interaction-Part 210: Human-centred design for interactive systems. International Standardization Organization (ISO). Switzerland, 2.
- [17] Dyer, J. H., Cho, D. S., Shinlim-Dong, K.-K., & Chu, W. (1996). Strategic supplier segmentation: A model for managing suppliers in the 21st century. Cambridge: MIT Press.
- [18] Gershman, S. J., Zhou, J., & Kommers, C. (2017). Imaginative reinforcement learning: Computational principles and neural mechanisms. *Journal of cognitive neuroscience*, 29(12), 2103-2113.
- [19] Gregor, S., & Hevner, A. R. (2013). Positioning and presenting design science research for maximum impact. *MIS quarterly*, 337-355.
- [20] Guidance, W. (2021). Ethics and governance of artificial intelligence for health. World Health Organization, 1-165.
- [21] Hassan, M., Borycki, E. M., & Kushniruk, A. W. (2025). Artificial intelligence governance framework for healthcare. Healthcare Management Forum.
- [22] Hevner, A. R., March, S. T., Park, J., & Ram, S. (2004). Design science in information systems research. *MIS quarterly*, 75-105.
- [23] Hussein, R., Zink, A., Ramadan, B., Howard, F. M., Hightower, M., Shah, S., & Beaulieu-Jones, B. K. (2024). Advancing Healthcare AI Governance: A Comprehensive Maturity Model Based on Systematic Review. *medRxiv*, 2024.2012.2030.24319785.
- [24] Ibáñez, J. C., & Olmeda, M. V. (2022). Operationalising AI ethics: how are companies bridging the gap between practice and principles? An exploratory study. *AI & Society*, 37(4), 1663-1687.
- [25] Joshi, G., Jain, A., Araveeti, S. R., Adhikari, S., Garg, H., & Bhandari, M. (2024). FDA-approved artificial intelligence and machine learning (AI/ML)-enabled medical devices: an updated landscape. *Electronics*, 13(3), 498.
- [26] Kahneman, D., & Tversky, A. (2013). Prospect theory: An analysis of decision under risk. In *Handbook of the fundamentals of financial decision making: Part I* (pp. 99-127). World Scientific.
- [27] Li, L. Z., Yang, P., Singer, S. J., Pfeffer, J., Mathur, M. B., & Shanafelt, T. (2024). Nurse burnout and patient safety, satisfaction, and quality of care: a systematic review and meta-analysis. *JAMA network open*, 7(11), e2443059-e2443059.
- [28] Moy, S., Irannejad, M., Manning, S. J., Farahani, M., Ahmed, Y., Gao, E., Prabhune, R., Lorenz, S., Mirza, R., & Klinger, C. (2024). Patient perspectives on the use of artificial intelligence in health care: a scoping review. *Journal of Patient-Centered Research and Reviews*, 11(1), 51.
- [29] National Institute of Standards and Technology. (2023). AI Risk Management Framework (AI RMF 1.0). <https://www.nist.gov/itl/ai-risk-management-framework>
- [30] Nielsen, J. (1993). Response times: the three important limits. *Usability Engineering*.
- [31] Parasuraman, A., Zeithaml, V. A., & Berry, L. (1988). SERVQUAL: A multiple-item scale for measuring consumer perceptions of service quality. *Journal of Retailing*, 64(1), 12-40.
- [32] Peffers, K., Tuunanen, T., Rothenberger, M. A., & Chatterjee, S. (2007). A design science research methodology for information systems research. *Journal of management information systems*, 24(3), 45-77.
- [33] Plsek, P. E., & Greenhalgh, T. (2001). The challenge of complexity in health care. *BMJ*, 323(7313), 625-628.
- [34] Poon, E. G., Lemak, C. H., Rojas, J. C., Guptill, J., & Classen, D. (2025). Adoption of artificial intelligence in healthcare: survey of health system priorities, successes, and challenges. *Journal of the American medical informatics association*, 32(7),

1093-1100.

- [35] Pyae, A. (2025). What is Human-Centeredness in Human-Centered AI? Development of Human-Centeredness Framework and AI Practitioners' Perspectives. arXiv preprint arXiv:2502.03293.
- [36] Recker, J. (2023). Designing information systems that support environmental sustainability: A framework-based review. *Research Handbook on Information Systems and the Environment*, 114-148.
- [37] Sáez, C., Ferri, P., & García-Gómez, J. M. (2024). Resilient artificial intelligence in health: synthesis and research agenda toward next-generation trustworthy clinical decision support. *Journal of Medical Internet Research*, 26, e50295.
- [38] Schram, A. L., Henriksen, T. B., Maindal, H. T., & Brazil, V. (2025). Navigating complexity: a conceptual framework for simulation interventions. *Advances in Simulation*, 10(1), 39.
- [39] Shneiderman, B. (2022). *Human-centered AI*. Oxford University Press.
- [40] Sutton, R. S., & Barto, A. G. (2018). *Reinforcement learning: An introduction* (Vol. 1). MIT press Cambridge.
- [41] Topol, E. J. (2019). High-performance medicine: the convergence of human and artificial intelligence. *Nature medicine*, 25(1), 44-56.
- [42] Tuunanen, T., Winter, R., & Brocke, J. v. (2024). Dealing with complexity in design science research: A methodology using design echelons. *MIS quarterly*, 48(2), 427-458.
- [43] Vergne, J.-P., & Durand, R. (2021). The missing link between the theory and empirics of organizational adaptation. *Academy of Management Annals*, 15(2).
- [44] Wang, L., Zhang, Z., Wang, D., Cao, W., Zhou, X., Zhang, P., Liu, J., Fan, X., & Tian, F. (2023). Human-centered design and evaluation of AI-empowered clinical decision support systems: a systematic review. *Frontiers in Computer Science*, 5, 1187299.
- [45] Wang, Y., Song, Y., Wang, Y., YU, L., & Wang, J. (2024). Ethics and governance of artificial intelligence for health: guidance on large multi-modal models. *Chinese Medical Ethics*, 1001-1022.
- [46] World Health Organization. (2024). Ethics and governance of artificial intelligence for health: Guidance on large multi-modal models. W. H. Organization. <https://www.who.int/publications/i/item/9789240084759>

Appendix A: Detailed Scenario Design

Scenario 1: No Governance Baseline

Category	Content
Purpose	Baseline setting for AI operation without HCEE governance
Governance Condition	HCEE architecture not applied, no experience measurement system, no trigger-control rules
AI Operation	Autonomous AI operation (minimal human intervention), efficiency-focused optimization, reactive response to errors

Scenario 2: Partial HCEE (Measurement Only)

Category	Content
Purpose	Isolate and verify measurement effects by introducing only experience measurement
Governance Condition	Experience layer only implemented (real-time PXI/EXI measurement), Policy/Governance/Control layers not applied, automatic control rules not operational
Validation Focus	Hawthorne effect - verify if measurement itself induces behavioral change

Scenario 3: Full HCEE Implementation

Category	Content
Purpose	Verify effects of complete 4-layer HCEE architecture implementation
Governance Condition	Policy layer: patient dignity priority principle applied; Governance layer: committee structure operational; Control layer: trigger-control rules automated; Experience layer: real-time PXI/EXI measurement
Mechanism Functions	Coordination: stakeholder interaction optimization; Learning: feedback-based continuous improvement; Selection: automatic optimal strategy determination by situation

Scenario 4: Stress Test

Category	Content
Purpose	Verify HCEE governance resilience in high-error environment

Stress Conditions	2x AI error rate increase (5% -> 10%), patient complaint surge events (weekly), healthcare worker turnover increase (3% monthly), system downtime (2 hours weekly)
Additional Metrics	Recovery Time: time to normalize PXI/EXI; Resilience Index: indicator variance before/after stress; Escalation Frequency: trigger activation count

Appendix B: Detailed Trigger-Control Rules

Rule ID	Trigger Condition	Control Action	Priority
TC-01	PXI < 70	20% AI automation reduction, expanded healthcare worker explanation time	Critical
TC-02	EXI < 65	Expanded AI override authority, strengthened decision rationale	Critical
TC-03	PXI - EXI > 15	Balance adjustment mechanism activation, bidirectional feedback strengthening	High
TC-04	AI Error Rate > 8%	AI automation temporary suspension, human oversight strengthening	Critical
TC-05	Dignity Perception < 60	Expanded human face-to-face interaction, minimized AI intervention	Critical
TC-06	Professional Respect < 55	Expanded healthcare worker autonomy, adjusted AI recommendation level	High
TC-07	PXI/EXI Plunge (weekly -10%)	Emergency response team convening, root cause analysis	Emergency
TC-08	PXI > 85 AND EXI > 80	Gradual AI automation level expansion permitted	Normal

Appendix C: Data Sources and Parameter Summary

C.1 FDA AI/ML Medical Device Database

Item	Content
Source	FDA 'AI/ML-Enabled Medical Devices' Database
URL	https://www.fda.gov/medical-devices/software-medical-device-samd/artificial-intelligence-enabled-medical-devices
Access Date	January 29, 2026
Confirmed Data	Total 1,357 AI/ML medical devices approved (as of September 2025)
Simulation Application	AI system accuracy parameter: 0.85-0.95 uniform distribution

C.2 Simulation Parameter Setting Basis

Parameter	Setting Value	Basis
Patient Agents	500	Conservative estimate of 689 daily outpatients
Healthcare Worker Agents	100	Outpatient care staff only (30 physicians + 70 nurses)
AI Systems	10	CDSS 4, Operations 3, RCM 3 (FDA 76% radiology, AHA 74% RCM)
AI Accuracy	0.85-0.95	FDA approved device meta-analysis (Joshi et al., 2024)
PXI Initial Value	Mean=70, SD=10	SERVQUAL healthcare service quality research (Parasuraman et al., 1988)
EXI Initial Value	Mean=65, SD=8	Healthcare worker burnout research exploratory setting (Li et al., 2024)

Acceptance Probability Coefficients	$r=0.66, 0.55, 0.73$	UTAUT meta-analysis (Dingel et al., 2024)
Learning Rate (α)	0.1	Reinforcement learning standard range 0.05-0.2 (Sutton & Barto, 2018)
Penalty/Reward Ratio	-15/+5 (3:1)	Loss aversion $\lambda \approx 2.25$ reference (Kahneman & Tversky, 2013)
Agreement Threshold	0.8 / 0.5	Exploratory setting, sensitivity analysis (± 0.1) verified
Trigger Thresholds	PXI 70, EXI 65	SERVQUAL satisfaction range reference, sensitivity analysis (± 5) verified

Appendix D: Sensitivity Analysis Results

D.1 Exploratory Parameter Sensitivity Analysis

Sensitivity analysis was performed on exploratory parameters with limited literature basis in this simulation. Each parameter was varied within $\pm 20\%$ of base value or specified range, analyzing effects on final PXI and EXI.

Table D-1. Exploratory Parameter Sensitivity Analysis Results

Parameter	Base Value	Variation Range	Result Variation
Learning Rate (α)	0.1	0.05 - 0.2	PXI $\pm 2.3\%$, EXI $\pm 1.8\%$
Noise SD (ϵ)	2	1 - 3	PXI $\pm 1.5\%$, EXI $\pm 1.2\%$
AI Error Penalty	-15	-10 ~ -20	PXI $\pm 4.1\%$, EXI $\pm 2.9\%$
Recovery Reward	+5	+3 ~ +7	PXI $\pm 1.8\%$, EXI $\pm 1.4\%$
Agreement Threshold (Accept)	0.8	0.7 - 0.9	PXI $\pm 3.2\%$, EXI $\pm 2.5\%$
Quality Coefficient	1.2	1.0 - 1.4	PXI $\pm 2.7\%$, EXI $\pm 3.1\%$
Trigger Threshold (PXI)	70	65 - 75	PXI $\pm 3.8\%$, EXI $\pm 2.2\%$
Trigger Threshold (EXI)	65	60 - 70	PXI $\pm 2.0\%$, EXI $\pm 4.5\%$

D.2 Sensitivity Analysis Results Interpretation

Sensitivity analysis results confirmed that variation ranges for all exploratory parameters resulted in final indicator variations within $\pm 5\%$, confirming result robustness. AI error penalty and trigger thresholds showed relatively larger effects, but these are consistent with theoretically expected directions. Learning rate and noise parameters had limited effects on results, confirming appropriateness of settings within general reinforcement learning literature ranges. This sensitivity analysis acknowledges limitations of exploratory simulation while demonstrating that main conclusions (positive effects of HCEE governance) are robust to parameter setting uncertainty.

Author Biography

Minseong Kim



Ph.D. Student, Graduate School of AI, Seoul School of Integrated Sciences & Technologies (aSSIST), Administrative Director, Pogny Hospital
 Research Interests: Healthcare AI Governance, Human-Centered Experience Engineering, Hospital Digital Transformation, Generative AI, Data-Driven Hospital Management
 E-mail: mskim924@stud.assist.ac.kr