

Practical Applicability of an AI-Based Automated Classification Model in Elderly Care Institutions

Jee Won Moon , Tae yeon Oh

ABSTRACT

This study evaluates the practical applicability of an artificial intelligence (AI)-based automated text classification model in elderly care institutions. While prior research has primarily focused on improving algorithmic performance, limited empirical attention has been given to whether automated outputs demonstrate sufficient alignment with existing manual classification practices in real operational environments. This study reanalyzes a categorized dataset of 1,224 elderly care records to examine predictive agreement from an organizational applicability perspective. A logistic regression classifier with hyperparameter optimization was implemented, and performance was assessed using accuracy, precision, recall, and F1-score metrics. The optimized model achieved an accuracy of 0.790 and a weighted F1-score of 0.780. Category-level analysis indicated stable performance across primary operational classes, although variation in recall was observed in specific categories. The findings suggest that AI-assisted classification may function as a structured support mechanism within institutional documentation workflows. This study contributes empirical evidence regarding feasibility of practical implementation without asserting direct organizational performance transformation.

Keywords: AI text classification, elderly care records, predictive agreement, organizational applicability, logistic regression

1. INTRODUCTION

The digital transformation of healthcare and welfare institutions has accelerated the accumulation of operational data in textual form [1][2]. Elderly care institutions, in particular, generate large volumes of narrative records describing residents' physical conditions, cognitive changes, emotional expressions, and behavioral observations. These records serve multiple administrative purposes, including regulatory reporting, internal monitoring, and quality management.

Despite the growing volume of documentation, classification and structuring of narrative records remain largely dependent on manual processes. Staff members review textual entries and assign them to predefined operational categories. While human interpretation allows contextual nuance, manual classification requires substantial labor resources and may result in variability in category assignment. As the quantity of documentation increases, maintaining consistency and efficiency in classification becomes progressively more challenging.

Artificial intelligence (AI)-based text classification techniques have been introduced as a potential solution to address these limitations [3][4]. Machine learning models can be trained on labeled datasets to predict predefined categories for new textual inputs. Prior research in

healthcare and welfare domains has demonstrated that supervised learning models can achieve acceptable predictive performance under controlled experimental settings [3][4]. However, most studies have concentrated on improving algorithmic accuracy or comparing model architectures, rather than examining whether automated outputs align sufficiently with existing institutional classification standards.

For AI systems to be practically integrated into institutional workflows, predictive performance alone is not sufficient. Automated outputs must demonstrate stable agreement with manually assigned labels, particularly across core operational categories. If predictive alignment is unstable, the introduction of automated systems may increase corrective workload instead of improving efficiency. Therefore, assessing predictive agreement within real operational contexts is a necessary preliminary step before broader organizational application.

This study evaluates the practical applicability of an AI-based automated classification model using categorized elderly care records. Rather than emphasizing model innovation, the analysis focuses on whether automated predictions demonstrate sufficient alignment with manual classification practices to support structured institutional use. By adopting an organizational applicability perspective, this research seeks to provide empirical evidence regarding the feasibility of AI-

assisted record classification in elderly care institutions [5][6].

2. THEORETICAL BACKGROUND

2.1 AI-Based Text Classification in Healthcare Documentation

Text classification is a core task in natural language processing (NLP), widely applied to organize unstructured textual data into predefined categories. In healthcare and welfare contexts, machine learning-based classification models have been utilized to analyze nursing records, medical reports, and patient documentation [3][7]. Supervised learning approaches enable models to learn patterns from labeled datasets and generate category predictions for new textual inputs.

Among various classification algorithms, logistic regression remains frequently adopted in structured institutional settings due to its interpretability and computational efficiency [3][8]. While more complex neural architectures may offer performance improvements under large-scale datasets, relatively simpler models are often preferred in organizational environments where transparency and stability are important considerations [4][9].

Previous research has primarily focused on enhancing predictive performance through feature engineering, model comparison, and parameter optimization [3][10]. These studies contribute to methodological advancement; however, most evaluations are conducted under experimental conditions without explicit consideration of institutional classification practices embedded in daily workflows.

In operational settings such as elderly care institutions, record categorization functions not merely as a technical process but as part of routine administrative procedures. Therefore, the evaluation of automated classification models should consider their compatibility with existing classification standards and practical interpretability.

2.2 Predictive Agreement as a Criterion for Organizational Applicability

Organizational applicability of AI systems depends on more than technical accuracy. For automated classification to be integrated into institutional workflows, model outputs must demonstrate consistent alignment with manually assigned labels across core operational

categories. Predictive agreement refers to the degree of correspondence between automated predictions and human-generated reference labels. Standard performance metrics—precision, recall, and F1-score—are commonly employed to quantify such alignment [3][4].

Precision measures the proportion of correctly predicted instances among predicted cases within a category. Recall represents the proportion of correctly identified instances among actual cases. The F1-score balances these two measures.

In organizational contexts, balanced performance across categories may be more meaningful than marginal improvements in overall accuracy. Stable predictive patterns reduce the likelihood of systematic bias toward particular classes and enhance confidence in institutional application.

2.3 Operational Considerations and Error Implications

Automated classification errors may influence administrative monitoring and documentation processes. Although this study does not evaluate downstream operational outcomes, predictive instability across categories may affect the reliability of structured record management systems.

Prior research on AI adoption in healthcare and clinical decision environments emphasizes the importance of error awareness and monitoring structures in technology implementation [8][9]. While elderly care record classification does not constitute high-risk decision-making in a strict sense, inaccurate categorization may introduce reporting inconsistencies. Therefore, predictive stability serves as a foundational requirement for cautious institutional consideration.

From this perspective, the evaluation of AI-based classification models should move beyond purely technical optimization and examine their operational feasibility within structured documentation environments.

3. RESEARCH DESIGN AND DATA

3.1 Research Scope and Analytical Perspective

This study adopts an empirical validation approach to examine predictive agreement between automated and manual classification outcomes. The analytical objective differs from algorithm

development or performance maximization. Instead, the focus lies on assessing whether the observed performance level supports consideration of structured institutional application.

The dataset utilized in this study was originally collected and structured as part of the author’s prior master’s thesis research [6]. However, the analytical framework of the present study differs substantially in scope and objective. While the previous research concentrated on model construction and optimization strategies, this study reanalyzes the dataset from an organizational applicability perspective, emphasizing predictive alignment rather than technical advancement.

3.2 Data Description

The dataset consists of categorized textual records generated within an elderly care institution. The records were originally documented in narrative form and subsequently assigned to predefined operational categories. These categories reflect routine documentation practices within the institution.

A total of 1,224 records were included in the analysis. Among these, 871 records were manually labeled by trained personnel, and 353 records were incorporated through controlled data augmentation procedures designed to preserve semantic consistency [11].

(Table 1) Dataset Composition

Category	Count
Manual labeling	871
Data augmentation	353
Total	1,224

All records were anonymized prior to analysis in compliance with healthcare data protection standards [12]. No personally identifiable information was retained in the analytical dataset [2][13].

3.3 Classification Categories

The predefined operational categories represent recurring documentation themes observed in elderly care records. These include physical condition descriptions, cognitive changes, emotional states,

behavioral observations, and miscellaneous entries [14].

The classification scheme reflects practical administrative usage rather than theoretical taxonomy. Therefore, the evaluation of automated classification focuses on its ability to reproduce existing institutional categorization patterns rather than redefine category structures [15].

3.4 Data Preprocessing

Prior to model training, textual data underwent standard preprocessing procedures. These included token normalization, removal of non-informative characters, and vectorization using term-frequency-based feature representation. Class imbalance was addressed through weighting strategies within the classifier configuration.

The dataset was partitioned into training and evaluation subsets using stratified sampling to preserve category distribution proportions. Performance evaluation was conducted on the held-out evaluation set.

4. MODEL CONFIGURATION AND EXPERIMENTAL PROCEDURE

4.1 Model Selection Rationale

This study employs logistic regression as the primary classification algorithm. Although more complex neural network architectures have been widely adopted in natural language processing tasks, logistic regression remains a practical and interpretable baseline model in institutional environments [7]. The primary objective of this research is not to maximize predictive performance through architectural complexity, but to assess whether a relatively stable and transparent model can achieve sufficient predictive agreement for structured organizational application.

Logistic regression estimates the probability that a given textual input belongs to a predefined category based on feature weights. The model predicts class membership using the sigmoid function:

$$P(y = 1|x) = \frac{1}{1 + e^{-(\beta_0 + \beta_1 x_1 + \dots + \beta_n x_n)}}$$

where x_i represents feature values derived from textual vectorization and β_i denotes learned coefficients. In multi-class classification settings, a one-vs-rest (OvR) strategy is applied.

The interpretability of coefficient weights allows examination of feature influence patterns, which is advantageous in institutional contexts where transparency is desirable.

4.2 Hyperparameter Optimization

To improve model stability, hyperparameter tuning was conducted. The regularization strength parameter C was adjusted to balance bias and variance trade-offs. Class imbalance was addressed using a balanced class weighting scheme.

(Table 2) Hyperparameter Configuration

Parameter	Value
Regularization (C)	10
Solver	liblinear
Class Weight	Balanced
Multi-class Strategy	One-vs-Rest

The selected configuration demonstrated improved predictive consistency compared to baseline settings.

4.3 Experimental Design

The dataset was divided into training and evaluation subsets using stratified sampling to preserve category proportions. Model training was conducted on the training set, and evaluation was performed on a held-out test set.

Three experimental conditions were examined:

- 1) Manual-labeled dataset only
- 2) Manual + augmented dataset
- 3) Optimized configuration with hyperparameter tuning

This comparative design allows examination of how data augmentation and parameter tuning influence predictive agreement.

4.4 Evaluation Metrics

Model performance was evaluated using accuracy, precision, recall, and F1-score [3][4].

Accuracy measures overall correct predictions:

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN}$$

Precision and recall provide category-specific performance insight:

$$\text{Precision} = \frac{TP}{TP + FP}$$

$$\text{Recall} = \frac{TP}{TP + FN}$$

The F1-score is calculated as:

$$F1 = 2 \cdot \frac{\text{Precision} \cdot \text{Recall}}{\text{Precision} + \text{Recall}}$$

In organizational contexts, recall may carry particular importance in categories where omission errors are undesirable. Therefore, both macro-average and weighted-average F1-scores were examined to evaluate performance balance across categories.

4.5 Analytical Focus

The analytical emphasis of this study lies in predictive agreement rather than incremental accuracy improvement. Therefore, performance stability across categories is interpreted as a primary indicator of organizational feasibility. Variability in recall rates is analyzed to identify potential operational sensitivity in specific documentation categories.

5. RESULTS

5.1 Overall Performance Comparison

The predictive performance of the automated classification model was

evaluated under three experimental conditions. Table 3 summarizes the overall performance metrics.

(Table 3) Overall Model Performance Comparison

Condition	Accuracy	Precision	Recall	F1-score
Manual Only	0.78	0.76	0.74	0.75
Manual + Augmented	0.82	0.80	0.79	0.79
Multi-classStrategy	0.86	0.85	0.84	0.84

The optimized configuration demonstrated the highest level of predictive consistency across all evaluation metrics. The increase in F1-score from 0.75 (manual-only baseline) to 0.84 (optimized configuration) indicates that improvements occurred simultaneously in both precision and recall, rather than being driven by isolated metric inflation.

The inclusion of augmented data produced moderate gains, particularly in recall performance. However, the most substantial enhancement was observed following hyperparameter calibration. This suggests that model stability in institutional settings is sensitive not only to dataset size but also to parameter configuration.

Importantly, overall accuracy alone does not sufficiently capture organizational suitability. In structured documentation systems, balanced metric behavior provides a more meaningful evaluative basis than aggregate performance alone [3][4].

5.2 Category-Level Performance Analysis

To assess predictive agreement at a more granular level, category-specific precision, recall, and F1-scores were analyzed under the optimized configuration.

(Table 4) Category-Level F1-score (Optimized Model)

Category	Precision	Recall	F1-score
Physical Condition	0.88	0.85	0.86
Cognitive Change	0.83	0.81	0.82

Emotional State	0.84	0.83	0.83
Behavioral Observation	0.86	0.84	0.85
Miscellaneous	0.79	0.76	0.77

The highest predictive alignment was observed in the Physical Condition category (F1 = 0.86), followed closely by Behavioral Observation (F1 = 0.85). These results suggest that records describing observable physical or behavioral conditions may exhibit relatively consistent linguistic structures, facilitating model generalization.

The Miscellaneous category showed comparatively lower performance (F1 = 0.77). This outcome likely reflects semantic heterogeneity within the category, as broadly defined operational classes may incorporate diverse narrative patterns. The reduced performance does not indicate systemic failure but highlights the structural limitations of loosely bounded categories.

Crucially, no category exhibited extreme imbalance between precision and recall. The absence of sustained recall suppression across operational domains indicates that the model does not systematically under-identify specific documentation types.

5.3 Stability Across Experimental Conditions

Comparative analysis across experimental conditions reveals that data augmentation primarily contributed to recall improvement in categories with relatively smaller representation. This effect reduced classification bias by enhancing minority-category detection.

More notably, the optimized configuration reduced inter-category recall variance from 0.11 in the baseline condition to 0.07 after calibration. This reduction reflects improved distributional balance rather than isolated metric gain.

In institutional environments, such variance reduction is operationally significant. Lower inter-category volatility decreases the likelihood of systematic misclassification within particular documentation domains and supports administrative reliability.

Thus, predictive stability emerges not merely as a statistical

characteristic but as an organizationally relevant property.

5.4 Interpretation from an Organizational Perspective

From an institutional standpoint, the observed performance patterns indicate that the automated classification model is capable of reproducing manual categorization structures with substantial consistency.

While the model does not replicate professional judgment in its entirety, F1-scores exceeding 0.80 across core operational categories suggest that the system can function as a structured assistive mechanism within documentation workflows.

Importantly, the findings do not imply automation of final decision authority. Rather, they support a hybrid configuration in which automated classification provides preliminary categorization, followed by human verification where necessary.

In this context, predictive agreement at the observed level may reduce repetitive sorting workload while preserving documentation coherence and accountability structures.

6. DISCUSSION

6.1 Reframing Performance Beyond Metric Maximization

The empirical findings of this study call for a reframing of how performance should be interpreted within institutional environments.

Accuracy, precision, recall, and F1-score are conventionally treated as primary indicators of model quality in text classification research. However, when the objective shifts from technical optimization to organizational deployment, the meaning of these metrics necessarily changes.

In elderly care institutions, classification systems are embedded within structured documentation workflows that directly influence monitoring, reporting, and administrative coordination. Under such conditions, performance cannot be evaluated solely in terms of maximal aggregate accuracy. What matters more fundamentally is whether predictive behavior remains proportionate, balanced, and

structurally stable across operational categories.

A model may achieve high overall accuracy while simultaneously under-identifying specific documentation domains. Such imbalance, even if statistically subtle, can introduce distortion into reporting structures and administrative interpretation. Therefore, institutional adequacy requires an evaluative lens that extends beyond peak metric performance.

This study adopts that broader interpretive position.

6.2 Predictive Stability as an Institutional Threshold

The results indicate that predictive stability offers a more appropriate benchmark for deployment consideration than accuracy maximization alone.

Predictive stability, as defined in this study, entails:

- 1) Balanced precision and recall across predefined operational categories.
- 2) Controlled inter-category variance in recall performance.
- 3) Absence of systematic suppression in any documentation domain.

The optimized model satisfied these criteria. F1-scores exceeded 0.80 across major categories, and recall variance was reduced following parameter calibration. No category demonstrated persistent under-identification.

These empirical properties suggest that the model maintains structural alignment with existing manual classification patterns rather than distorting them. Stability in this sense reflects proportional behavioral consistency across operational domains.

In institutional documentation systems where categorization affects monitoring and compliance processes, such proportional consistency constitutes a meaningful reliability threshold.

Accordingly, this study advances predictive stability as a primary evaluative benchmark for structured institutional deployment of AI-based classification systems.

6.3 Clarifying the Boundaries of Interpretation

It is essential to delimit the scope of this conclusion.

The findings do not imply replacement of professional judgment, nor do they assert that logistic regression represents a universally optimal architecture. The contribution of this study lies not in technical superiority but in evaluative repositioning.

Within a predefined institutional classification framework, the model demonstrates sufficient predictive alignment to function as a structured assistive layer. This suggests feasibility for preliminary categorization within a human-supervised workflow, where automated outputs are subject to oversight rather than autonomous authority.

The system, therefore, is conceptualized not as a decision-maker but as a stabilizing administrative component embedded within existing accountability structures.

6.4 From Computational Evaluation to Governance Judgment

The broader significance of this study lies in connecting computational evaluation with institutional governance judgment.

AI-based classification models are often evaluated as isolated technical artifacts. However, institutional integration is ultimately a governance decision. It requires assessing whether model behavior aligns with documentation norms, reporting structures, and accountability expectations.

By formalizing predictive stability as an operational threshold, this study provides a decision-oriented interpretive framework. The evaluative question shifts from:

"Which model achieves the highest accuracy?"

to:

"Does the model maintain balanced and reliable behavior across operational domains under real documentation constraints?"

This reframing does not diminish technical rigor. Instead, it repositions performance metrics within the institutional realities of service organizations.

The contribution of the study, therefore, lies not in competing for incremental metric gains, but in clarifying how predictive behavior should be judged when organizational integration is under consideration.

7. LIMITATIONS

7.1 Contextual Delimitation of the Empirical Setting

The empirical analysis presented in this study is situated within the documentation system of a single elderly care institution. Such documentation structures are not neutral technical artifacts; they are shaped by organizational routines, staffing patterns, regulatory expectations, and local operational cultures. Linguistic conventions, category boundaries, and narrative styles emerge from these institutional configurations.

Accordingly, the predictive stability observed in this study should be understood as stability within a specific organizational configuration rather than as a universal property of the model itself. The findings demonstrate how the stability criterion operates under a concrete documentation regime.

This contextual delimitation does not weaken the conceptual contribution of the study. Rather, it clarifies the empirical grounding of the proposed stability framework. By explicitly specifying the scope of application, the study avoids unwarranted generalization while preserving the structural validity of its evaluative perspective.

Future research may extend this investigation by incorporating multi-institutional datasets in order to examine whether predictive stability operates consistently across heterogeneous documentation environments.

7.2 Algorithmic Scope and Architectural Generalizability

This study intentionally adopted logistic regression as a transparent and structurally stable modeling approach. The objective was not to compete in algorithmic sophistication, but to assess whether predictive stability could be established within an interpretable and institutionally viable configuration.

Recent advances in transformer-based and deep learning architectures have demonstrated notable gains in raw classification accuracy. However, superior metric performance under experimental conditions does not automatically translate into institutional suitability. Organizational deployment requires consideration of interpretability, behavioral consistency, and variance control across operational domains.

Because cross-architecture benchmarking was not conducted, the generalizability of the proposed stability criterion across model families remains empirically open. Whether predictive stability persists under highly expressive neural architectures constitutes an important direction for subsequent investigation.

Thus, this limitation is not merely technical but conceptual. It concerns the portability of stability as a deployment benchmark beyond a single modeling paradigm.

7.3 Absence of Longitudinal Institutional Implementation

The present research evaluates predictive alignment between automated and manual classification outcomes under controlled analytical conditions. It does not include longitudinal observation of implemented workflows within the institution.

While predictive stability may serve as a precondition for deployment, the dynamic interaction between automated classification outputs and human verification practices was not empirically examined. Changes in documentation efficiency, redistribution of workload, and behavioral adaptation among staff remain outside the scope of this study.

Future research may therefore investigate how stability-informed deployment decisions translate into measurable administrative outcomes over time. Such longitudinal validation would allow the proposed evaluative framework to be assessed not only at the metric level but also at the organizational process level.

7.4 Formalization of Stability as an Operational Threshold

This study conceptualizes predictive stability as an operational reliability threshold but does not codify it into a fixed quantitative cutoff. The interpretation relies on balanced performance patterns and reduced inter-category variance rather than on a predetermined

numerical boundary.

Acceptable stability levels may vary depending on documentation criticality, institutional risk tolerance, or regulatory sensitivity. Therefore, the explicit calibration of deployment thresholds remains an open theoretical and practical task.

This limitation signals a direction for conceptual consolidation rather than a deficiency in empirical design. Further refinement of stability-based criteria may translate the proposed framework into more concrete operational guidelines.

8. CONCLUSION

This study examined the practical applicability of an AI-based automated classification model within the operational context of an elderly care institution. Rather than pursuing incremental gains in predictive accuracy or architectural sophistication, the analysis focused on whether predictive stability could serve as a meaningful benchmark for institutional deployment decisions.

The findings indicate that predictive stability—understood as balanced inter-category performance and consistent alignment with manual categorization—constitutes a more operationally relevant criterion than peak metric performance alone. In structured documentation environments, the reliability of category behavior across routine records may hold greater institutional significance than isolated improvements in aggregate accuracy.

By repositioning performance interpretation from competitive optimization toward stability-oriented assessment, this study advances a deployment-sensitive analytical perspective. AI-based classification is framed not as an autonomous performance artifact, but as a system embedded within organizational routines where interpretability, variance control, and behavioral consistency are central considerations.

The study does not assert universal generalizability nor claim immediate organizational transformation. Instead, it demonstrates that predictive stability can be empirically observed and analytically articulated within a defined institutional configuration. This articulation provides a principled basis for deployment decisions grounded in measured reliability rather than technological enthusiasm.

In this respect, the contribution of the study lies in clarifying what constitutes sufficient empirical evidence for organizational consideration of AI systems. Within structured administrative environments, stability—rather than maximal accuracy—emerges as the central evaluative axis for responsible integration.

References

- [1] S. Cheresharov, G. Dragomirov, G. Gustinov, S. Hadzhikoleva, and K. Yotov, "Transforming Nursing Home Care: An Integrated Approach Using Sensors, AI, and Monitoring Technologies," *Computer Science and Interdisciplinary Research Journal*, Vol. 1, No. 1, 2024.
- [2] T. M. Song, "Trends and Utilization of Health and Welfare Big Data in Korea," *Health and Welfare Policy Forum*, Vol. 23, No. 3, 2018.
- [3] H. S. Kim, D. J. Ahn, J. H. Lim, and H. Y. Lee, "A Case Study on Automatic Classification of Record Text Using Machine Learning," *Journal of the Korean Society for Information Management*, Vol. 34, No. 4, 2017.
- [4] H. Sakai and S. S. Lam, "Large Language Models for Health Care Text Classification: Systematic Review," *JMIR AI*, Vol. 5, e79202, 2026.
- [5] H. J. Park and H. S. Lim, "A Study on the Case Analysis and Introduction of NLP: Analysis Framework and Implications," *Journal of IT Services*, Vol. 21, No. 2, 2022.
- [6] J. W. Moon, "Design and Performance Evaluation of a Text-Based Machine Learning Model for Automatic Classification of Records of Changes in the Status of Elderly Long-Term Care Recipients," *Master's Thesis, aSSIST University, Seoul*, 2025.
- [7] D. Y. Hong and J. S. Huh, "Research Trends in Record Management Using Unstructured Text Data Analysis," *Journal of Korean Society for Archives and Records Management*, Vol. 23, No. 4, 2023.
- [8] D. H. Lee et al., "Clinical Decision Support System (CDSS) Technology Trends," *Electronics and Telecommunications Trends*, Vol. 31, No. 3, 2016.
- [9] M. Katebi, M. Bahreini, R. Bagherzadeh, and S. Pouladi, "Artificial Intelligence and Nursing Management: Opportunities, Challenges, and Ethical Considerations—A Scoping Review," *Journal of Nursing Management*, 2025.
- [10] H. Ryuno, T. Mukaihata, T. Takemura, C. Greiner, and Y. Yamaguchi, "Application of Artificial Intelligence to Electronic Health Record Data in Long-Term Care Facilities: A Scoping Review Protocol," *BMJ Open*, Vol. 15, e098091, 2025.
- [11] J. W. Gong, G. Y. Kim, Y. S. Kim, and B. D. Oh, "Development of ChatGPT-Based Medical Text Expansion Tool for Synthetic Text Generation," *Proceedings of the Korea Society of Computer and Information Conference*, 2023.
- [12] S. Y. Shin, "Anonymization of Healthcare Data for Privacy Protection," *Communications of the Korean Institute of Information Scientists and Engineers*, Vol. 36, No. 6, 2018.
- [13] J. H. Choi, J. H. Seong, M. H. Kim, and H. C. Kwon, "Redefining Entity Types and Generating Training Data for Personal Information De-identification," *Journal of KIISE*, Vol. 50, No. 2, 2023.
- [14] H. R. Seo, "Development and Utilization of Artificial Intelligence Technology Based on Electronic Medical Records for Improving Hospital Processes," *Master's Thesis, University of Ulsan*, 2024.
- [15] W. S. Kang et al., "Disability Classification Design Based on Specific Behavior Record Data," *Proceedings of the Korea HCI Conference*, 2012.

Author Biography



Jeewon Moon

aSSIST University / Ph. D. Candidate

She is a multidisciplinary expert with experience in clinical nursing, medical review, and long-term care management.

Her research centers on "Silent Risk" and AI governance (VG-HITL) to enhance ethical accountability and decision-making in high-risk care systems.

E-mail : mjw301130@gmail.com



Tae-yeon Oh

aSSIST University/ Assistant Professor/ Corresponding author

Seoul National University, Ph. D. in Sport Management

Sogang University, MA in Economics Areas of Interest: data analytics using AI and explainable artificial intelligence (XAI).

E-mail : tyoh@assist.ac.kr