



Applying Explainable Artificial Intelligence Principles to Interface Design: Improving User Trust, Understandability, and Usability in a Chicken Weight Monitoring System

Maritza Angelina Az Zahra¹, Divi Galih Prasetyo Putri¹, Margareta Hardiyanti²

¹ Universitas Gadjah Mada, Sleman, Indonesia

² Kyoto University, Kyoto, Japan

Abstract: Artificial Intelligence (AI) has been increasingly adopted in smart farming to support monitoring and decision-making processes. However, many AI-based systems still operate as black boxes, making their outputs difficult for end users to understand and potentially reducing user trust. Although Explainable Artificial Intelligence (XAI) has been proposed to improve transparency, studies integrating XAI principles into interface design and evaluating their effects on user experience remain limited, particularly in smart farming contexts. Rather than proposing a new XAI algorithm, this study operationalizes human-centered XAI principles into interface design and empirically evaluates their influence on user trust, understandability, and usability in an existing AI-based chicken weight monitoring system. A concurrent embedded mixed-methods approach with a within-subject and counterbalanced design was conducted involving 16 participants. The redesigned interface incorporated human-centered XAI principles and was evaluated using the Trust in Automation Scale (TiAS), an understandability questionnaire, and the System Usability Scale (SUS). The results showed statistically significant improvements across all evaluated aspects ($p < 0.001$). Trust increased from 50.52 to 73.82, understandability from 54.25 to 81.88, and usability from 43.28 to 74.38, with large effect sizes observed in all measurements. Qualitative findings indicated that clearer and contextual explanations improved users' interpretation of system outputs. These findings suggest that integrating XAI principles into interface design can support more transparent and understandable interaction in AI-based monitoring systems.

Keywords: Explainable Artificial Intelligence (XAI), chicken weight monitoring, user trust, understandability, usability

Article History:

Received: 19 June 2026

Accepted: 7 July 2026

Published: 7 July 2026

Corresponding Author: Maritza Angelina A Zahra, Email: maritzaangel04@gmail.com

DOI: 10.65917/aisa.v2i1.72

1 Introduction

Artificial Intelligence (AI) has increasingly been adopted across smart farming applications to support monitoring, prediction, and data-driven decision-making. AI enables systems to process large amounts of operational data and generate insights that assist users in managing agricultural activities more efficiently. In livestock production, AI-based technologies have been utilized to support monitoring production, AI-based technologies have been utilized to support monitoring activities and improve operational effectiveness through automated analysis and continuous observation processes [1], [2], [3]. In poultry farming, monitoring systems supported by AI can assist users in observing production conditions and interpreting indicators that support daily operational decisions [4].

As AI becomes more integrated into operational environments, interaction between users and AI-generated outputs becomes increasingly important. However, many AI systems still operate as black-box systems, where users receive recommendations or predictions without understanding how those outputs are produced [5]. This limitation creates challenges, particularly for users without technical knowledge of AI mechanisms. Previous studies indicate that insufficient understanding of system behavior may increase uncertainty during interaction and reduce user confidence when making decisions based on system outputs [6]. To address these challenges, Explainable Artificial Intelligence (XAI) has emerged as an approach that improves transparency by communicating how and why AI systems generate predictions [7], [8], [9]. Previous studies explain that effective explanations should consider user needs, explanation presentation, and contextual relevance [10]. Human-centered XAI emphasizes explanations that match users' goals and level of expertise, rather than exposing technical model details [11]. Studies have also shown that explainability can contribute to building more calibrated trust, reducing cognitive burden, and improving users' ability to understand AI-supported decisions [12], [13].

Although explainability has received increasing attention, most existing studies focus on explanation methods, technical transparency, or conceptual XAI frameworks [10], [14], [15]. Recent work suggests that explainability should also be considered an interface design problem because the effectiveness of explanations depends on how information is presented to users [11], [16], [17], [18]. However, empirical studies implementing XAI principles in real interface redesign and evaluating their impact on user trust, understandability, and usability remain limited, particularly in smart farming applications [19], [20].

This study addresses these limitations by redesigning an existing AI-based chicken weight monitoring system interface and empirically evaluating their effects on user trust, understandability, and usability. Rather than proposing a new XAI algorithm, this study operationalizes human-centered XAI principles into interface components and evaluates their contribution to user experience in a smart farming application.

The following research questions guide this study:

- **RQ1:** How does the implementation of XAI elements in interface design influence user trust toward the application?
- **RQ2:** How does the implementation of XAI elements in interface design improve users' understandability in interpreting information generated by the application?
- **RQ3:** How does redesigning the interface by integrating XAI principles affect the usability of the application?



To address these questions, the contributions of this study are as follows:

- Implementing human-centered XAI principles into user interface redesign for an AI-based monitoring system.
- Evaluating trust, understandability, and usability simultaneously to provide a broader assessment of user experience in XAI applications.
- Extending the application of XAI into the smart farming domain through empirical evaluation.

The remainder of this paper is organized as follows. **Section 2** presents related work covering XAI, human-centered design, and previous studies on XAI interface implementation. **Section 3** describes the research methodology, including system context, interface redesign process, participants, instruments, and data analysis procedures. **Section 4** presents the results and discussion of the interface redesign and evaluation outcomes. Finally, **Section 5** concludes the study and presents limitations and future research directions.

2 Related Work

XAI has developed from a model-centered perspective toward a broader human-centered approach that considers how explanations are presented, interpreted, and utilized during interaction with AI systems. Existing studies have explored various dimensions of explainability, including transparency, trust formation, user understanding, and interface design. However, findings across previous work indicate that explainability remains fragmented across different evaluation perspectives and implementation approaches. This section synthesizes prior studies and positions the current research within existing literature.

2.1 Explainability, Transparency, and Human-Centered XAI

Early XAI studies primarily focused on increasing transparency by exposing internal reasoning processes of AI models and enabling users to understand how outputs are generated [14]. Existing approaches include interpretable models, post-hoc explanations, feature-based explanations, and surrogate models intended to improve interpretability and accountability. However, subsequent studies argued that explainability should not be viewed solely as a technical property of algorithms. Explainable systems involve not only computational mechanisms but also interfaces and interaction processes through which users access and interpret information [11]. Therefore, explanation effectiveness depends not only on explanation quality but also on whether explanations are communicated in forms aligned with user needs [10].

Recent literature further shifts attention toward Human-Centered Explainable AI (HC-XAI), which emphasizes explanation design based on users' goals, context, and cognitive processes. Human-centered perspectives suggest that explanations become more useful when delivered using concise language, contextual information, and adaptive presentation mechanisms rather than exposing excessive technical details. In addition, transparency has been positioned as a foundational requirement for enabling users to develop realistic expectations and interpret AI-supported decisions appropriately [21].

Across these studies, explainability is increasingly understood as a communication and interaction challenge rather than solely a model interpretation problem. This perspective becomes relevant for systems where users rely on interfaces as the primary medium for understanding AI outputs.

2.2 Trust and Human-AI Interaction

Trust has become one of the central evaluation dimensions in XAI research because explainability ultimately aims to support effective collaboration between users and AI systems [19]. Previous findings demonstrate that explanations can improve users' confidence and decision performance by helping users interpret system behaviour more appropriately. In this context, explanations contribute to calibrated trust, where users neither overtrust nor undertrust AI-generated outputs.

Nevertheless, prior studies report that trust does not always increase proportionally with explanation availability. Human-centered XAI studies indicate that interaction design, explanation accessibility, and contextual relevance strongly influence how users perceive explanations and whether explanations are considered useful [20]. Systematic evaluation studies also reveal that trust alone may not sufficiently represent explanation effectiveness because users may trust systems without fully understanding the reasoning behind recommendations [15]. Therefore, explainability evaluation increasingly emphasizes combining trust assessment with complementary dimensions that reflect actual interpretation and comprehension.

2.3 Understandability and User Experience Evaluation

Studies show that explanation quality is influenced not only by informational content but also by how explanations are structured and presented to users [10]. Explanation strategies that utilize contextual information, simplified presentation, and visual support tend to improve comprehension and reduce cognitive burden.

Systematic reviews further indicate that there is still no universally accepted framework for evaluating XAI because each application context prioritizes different dimensions and explanation objectives [15]. Commonly evaluated dimensions include understandability, trustworthiness, usefulness, transparency, and controllability. In addition, recent human-centered evaluation studies suggest that explanation effectiveness should be measured beyond interpretation accuracy and include broader user experience outcomes such as usability and interaction quality [16]. Users may understand explanations but still experience difficulties using the system effectively.

2.4 Interface Design for Explainable Systems

As explainability research matures, increasing attention has been directed toward interface design as a mechanism for operationalizing XAI principles [16], [17], [18].

Previous studies emphasize that interfaces play a central role in determining how users access explanations, interpret outputs, and develop expectations regarding system behaviour [16]. Rather than treating explanations as static textual descriptions, interface explainability encourages explanations to become integrated components of interaction. Research on GUI design patterns for XAI proposes principles, including trustworthiness, causality, informativeness, interactivity, and accessibility, to improve explanation delivery [18]. These principles aim to make explanations easier to navigate and more actionable during user interaction.



Further studies focusing on UI prioritization in XAI suggest that trust and understandability represent two of the most influential factors affecting user experience [17]. Findings indicate that users generally prefer explanations that are adaptive and contextual rather than excessive and technically complex.

Although previous studies provide valuable design guidance, several limitations remain. Existing research often focuses on explanation mechanisms without implementing interface redesign directly, evaluates trust independently without measuring understandability, or remains limited to conceptual recommendations. Furthermore, empirical implementation within smart farming applications remains relatively underexplored.

2.5 AI-based Smart Farming Applications

Recent studies have demonstrated the growing adoption of AI in smart farming, particularly in Precision Livestock Farming (PLF), to support automated monitoring of poultry health, welfare, and production. A review by [22] explains the integration of sensors, computer vision, machine learning, and IoT for continuous poultry monitoring in PLF. Similarly, [23] identifies computer vision as a key technology for real-time poultry behavior, welfare, and health monitoring. In addition, [24] summarizes the use of computer vision and deep learning for chicken weight estimation, disease detection, and behavior analysis. [25] highlights the application of AI for livestock monitoring and decision support across various production systems. Furthermore, [26] emphasizes the role of computer vision in enabling automated monitoring for modern agriculture and livestock production.

Although existing studies have advanced AI-based monitoring for smart farming, they mainly focus on model development and monitoring performance rather than on how AI-generated information is communicated through user interfaces. Therefore, this study operationalizes human-centered XAI principles at the interface level and empirically evaluates their effects on user trust, understandability, and usability in an existing AI-based chicken weight monitoring system.

2.6 Research Positioning

Based on the reviewed literature, this study positions itself at the intersection of XAI, interface redesign, and user experience evaluation within smart farming systems.

Table 1: Comparison of Previous Studies

Aspect	Previous Studies	This Study
Explainability principles	✓	✓
Human-centered approach	✓	✓
Direct implementation into UI	Limited	✓
Trust evaluation	✓	✓
Understandability evaluation	Limited	✓
Simultaneous usability evaluation	Rare	✓
Smart farming application	Limited	✓

Unlike previous studies that predominantly emphasize explanation techniques, conceptual frameworks, or isolated evaluation dimensions, this study directly implements human-centered

XAI principles into the redesign of an operational monitoring interface and evaluates their influence on trust, understandability, and usability simultaneously. By applying this approach within a chicken weight monitoring system, this study extends empirical evidence regarding explainable interface design in smart farming contexts.

3 Methods

3.1 Research Design

This study employed a concurrent embedded mixed-methods design to evaluate the implementation of XAI principles in redesigning the interface of an AI-based chicken weight-monitoring application. Concurrent embedded mixed methods integrates quantitative and qualitative approaches within a single research framework, while assigning different analytical priorities to each method, where one method serves as the primary source of findings and the other acts as complementary support [27].

In this study, quantitative evaluation was used as the primary method to measure changes in user trust, understandability, and usability, whereas qualitative findings were embedded to provide contextual interpretation regarding participant experiences and interaction behaviour. This combination was selected to produce not only measurable evidence of interface performance but also a deeper understanding of how explainable interface elements influence user perception. The integration of numerical and narrative findings allows for a more comprehensive interpretation of the observed outcomes [28].

This study adopted a within-subject experimental design, meaning that each participant evaluated both interface conditions: the existing interface and the modified interface implementing XAI principles. Within-subject design enables repeated measurement on the same participants across multiple treatment conditions, allowing direct comparison while controlling variability between individuals [29]. Since each participant acts as their own control, differences in responses are more likely to reflect the effect of interface modifications rather than differences in participant characteristics.

To reduce sequence effects and minimize learning bias, the evaluation procedure implemented a counterbalanced design. Counterbalanced design controls ordering, learning, and fatigue effects by varying the sequence of treatment exposure across participant groups so that each condition receives an equal opportunity during testing [30]. Through this arrangement, comparison between the existing and modified interfaces becomes more objective and less affected by internal procedural bias. Participants were randomly assigned to two counterbalanced groups. Group A (n=8) evaluated the existing interface first, followed by the modified interface, whereas Group B (n=8) completed the evaluation in the opposite order. Participant codes ending with “_A” and “_B” indicate the assigned counterbalancing sequence.

The overall research procedure is presented in Figure 1.

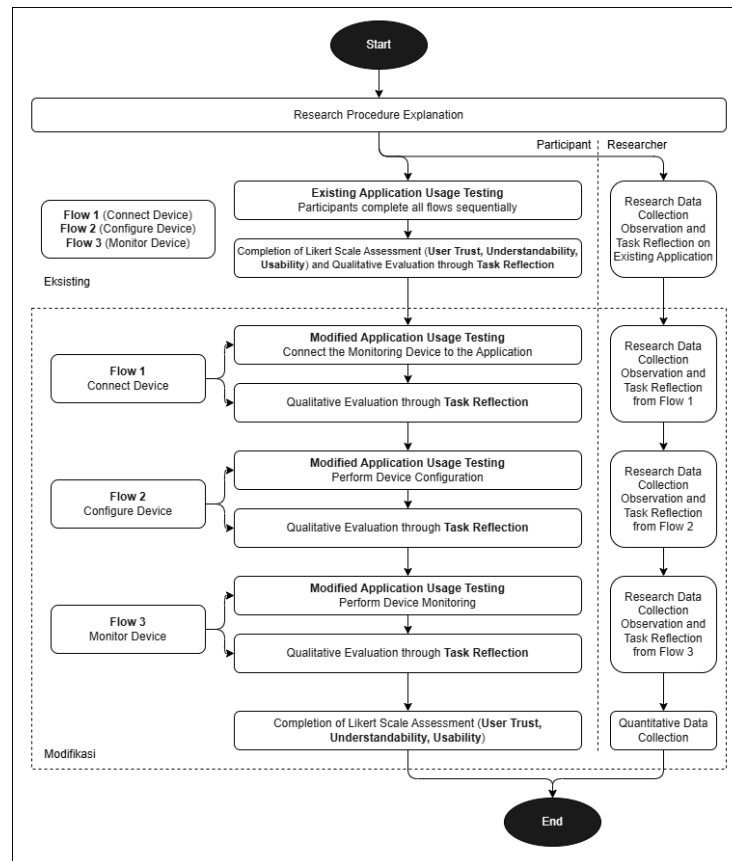


Figure 1: Research Testing Flow

As illustrated in Figure 1, the study began with an explanation of the research objectives, evaluation procedures, and participant instructions to ensure a consistent understanding before interaction started. Participants were then asked to evaluate the existing interface by completing all predefined operational tasks sequentially. The interaction consisted of three operational flows representing the main activities supported by the application: device connection, device configuration, and monitoring activities.

During the first flow, participants connected monitoring devices to the application and verified whether device communication was successfully established. The second flow required participants to configure operational parameters according to the monitoring requirements. In the final flow, participants observed monitoring information generated by the system, interpreted the presented output, and completed monitoring-related activities through the interface.

After completing all interaction flows using the existing interface, participants completed quantitative evaluation instruments measuring trust, understandability, and usability. Simultaneously, researchers conducted observational recording to identify interaction patterns, confusion points, interpretation difficulties, and user responses toward the existing interface.

Participants then proceeded to evaluate the modified interface, which incorporated human-centered XAI principles. The same operational flows were maintained to preserve consistency and ensure that observed differences originated from interface redesign rather than task variation.

Unlike the valuation of the existing interface, participants completed task reflection activities after each interaction flow during the modified interface evaluation. These reflections were intended to capture immediate qualitative responses regarding explanation clarity, confidence during interaction, interpretation processes, and perceived usefulness of the redesigned interface components.

After all interaction tasks were completed, participants filled in post-evaluation questionnaires covering trust, understandability, and usability dimensions. Quantitative findings were subsequently analyzed statistically, while qualitative findings obtained through observation and task reflection were analyzed thematically. The final interpretation integrated both forms of evidence to provide a more comprehensive understanding of how XAI-based interface redesign influenced user experience.

3.2 System Context

This study was conducted using an existing AI-based chicken weight monitoring system developed for poultry farming operations. The system employs computer vision to estimate chicken weight from captured images and had already been implemented prior to this study. The underlying AI model, prediction mechanism, and application functionality were not modified. Instead, this study focused exclusively on redesigning the user interface by integrating human-centered XAI principles while preserving the existing AI model and operational workflows.

3.3 XAI-Based Interface Redesign

The interface redesign in this study was guided by a literature-driven synthesis of XAI approaches related to explanation presentation in user interfaces. Rather than directly adopting individual design recommendations from previous studies, this study synthesized recurring explanation characteristics identified across prior literature to establish a conceptual foundation for redesigning the existing interface.

Through this synthesis process, several approaches to presenting explanations in AI systems were identified with the common objective of making AI-generated information more accessible to non-technical users and supporting decision-making activities. The resulting synthesis produced six principal aspects of XAI that describe how explanations should be presented within the interface so that users can understand system outputs without requiring knowledge of technical implementation details.

These principles were not treated as independent interface features but as design directions that informed the modification of the existing interface. Therefore, the identified principles were used throughout the study to guide redesign decisions, construct evaluation instruments, and organize redesign outcomes into meaningful categories.

- **Comprehensible Explanation:** explanations should be concise, understandable for non-technical users, and presented at a higher level of abstraction without exposing technical details of AI algorithms or creating information overload [31], [32], [33], [34], [35], [36]. In the redesign process, this principle was implemented by simplifying explanatory content, improving information hierarchy, rewriting interface texts into clear language, and introducing supporting explanation elements to reduce ambiguity and facilitate interpretation during monitoring activities.
- **Progressive and Context-Aware Explanation:** explanations should be delivered gradually and displayed only when relevant to the user's interaction context so that additional information remains available without increasing interface complexity [35], [37], [38], [39], [40]. To



implement this principle, explanatory content was reorganized into multiple levels of detail and displayed progressively according to user needs, allowing users to access more detailed information only when necessary.

- **Reasoning and Causal Explanation:** focuses on communicating reasons and contributing factors behind AI-generated outputs so users can understand relationships between observed conditions and system results [41], [42], [43], [44]. In the redesigned interface, explanatory elements were added to provide contextual reasoning and help users understand why certain monitoring outcomes appeared without exposing underlying computational processes.
- **Actionable and Goal-Oriented Explanation:** explanations should not only provide information but also support users in determining possible actions to improve or optimize outcomes generated by the system [31], [37], [45], [46], [47]. The redesign implemented this principle by introducing recommendation-oriented explanations and supportive guidance to assist users in making operational decisions based on monitoring results.
- **Trust Calibration and Bias Awareness:** helping users understand when system outputs can be trusted and what limitations or contextual conditions may influence output reliability [35], [44], [46], [48]. To implement this principle, redesign decisions introduced explanatory cues that communicate confidence-related information and encourage users to develop balanced expectations toward AI-generated results.
- **Mental Model Alignment:** aims to align users' understanding with how the AI system operates through explanation structure and interface visualization [35], [41], [43], [44], [49]. In the redesign process, this principle guided the restructuring of information hierarchy, interaction flow, and explanation sequencing to better match user expectations and improve the interpretation of monitoring information.

After all redesign decisions were implemented, the modified interface was evaluated using the same operational scenarios and procedures as the existing interface to ensure that observed differences originated from explainability improvements rather than functionality changes.

3.4 Participants

This study involved 16 participants recruited using a purposive sampling strategy to ensure alignment between participant characteristics and the operational context represented by the evaluated system. Participant recruitment emphasized contextual relevance rather than broad demographic representation because the objective of this study was to evaluate user trust, understandability, and usability when interacting with AI-generated information in a poultry monitoring environment. Therefore, participants were selected based on their involvement in poultry farming activities and familiarity with operational monitoring processes.

To participate in this study, individuals were required to satisfy the following inclusion criteria:

- be involved in poultry farming activities as farmers or operational workers (ABK);
- possess prior experience with monitoring-related operational activities;
- be capable of completing all evaluation scenarios independently; and
- provide consent to participate throughout the study procedure.

Because this study employed a within-subject experimental design, all participants completed the evaluation under both interface conditions, namely the existing interface and the modified interface implementing XAI principles. Each participant experienced identical operational scenarios following the predefined evaluation sequence. To maintain participant confidentiality and comply with ethical considerations, demographic information is reported in an aggregated form.

Table 2: Participant Demographics

Code	Gender	Age	Role	Length of Employment	Last Education	Familiar with AI
P01_A	M	37	Farmer	>3 years	Senior High School	✓
P02_B	M	20	ABK	<1 year	Junior High School	✓
P03_B	M	52	ABK	<1 year	Senior High School	✓
P04_A	M	40	Farmer	>3 years	Senior High School	✓
P05_B	F	38	Farmer	>3 years	Diploma/Bachelor's Degree	✓
P06_A	F	46	ABK	1-3 years	Senior High School	✓
P07_B	M	42	ABK	1-3 years	Senior High School	✓
P08_A	M	42	Farmer	>3 years	Senior High School	✓
P09_B	M	27	Farmer	>3 years	Diploma/Bachelor's Degree	✓
P10_A	M	24	ABK	1-3 years	Senior High School	✓
P11_B	M	23	ABK	<1 year	Senior High School	✓
P12_A	M	49	Farmer	>3 years	Senior High School	✓
P13_A	M	42	ABK	>3 years	Senior High School	✓
P14_B	M	52	ABK	1-3 years	Senior High School	✓
P15_B	M	55	ABK	<1 year	Senior High School	✓
P16_A	M	31	Farmer	>3 years	Diploma/Bachelor's Degree	✓

Based on the demographic profile, participants were predominantly male and ranged from 20 to 55 years old. Participant roles consisted of both farmers and farm operators (ABK), with most participants having more than one year of working experience. In terms of educational background, the majority completed senior high school education, while several participants held a diploma or bachelor's-level qualifications. All participants reported prior familiarity with AI-related technologies and therefore were able to interact with the evaluated system without requiring additional technical training. Familiarity with AI was self-reported through a demographic questionnaire administered before the evaluation. Participants indicated whether they had previous experience using AI-based applications in their daily activities or work. No formal assessment of AI expertise was conducted.

3.5 Instruments

To evaluate the impact of implementing XAI principles in interface redesign, this study employed a combination of quantitative and qualitative evaluation instruments. Quantitative evaluation focused on three user experience dimensions that are frequently discussed in human-centered XAI research, namely user trust, understandability, and usability, while qualitative evaluation was supported through task reflection to capture contextual experiences during interaction. All quantitative instruments used a five-point Likert scale, ranging from 1 (strongly disagree) to 5 (strongly agree). The same evaluation procedure was applied to both interface conditions to enable direct comparison of participant responses after interacting with the existing and modified interfaces.



3.4.1 User Trust

User trust was evaluated using the Trust in Automation Scale (TiAS). TiAS is a questionnaire instrument used to measure user trust toward automated systems through the evaluation of a set of trust-related statements [50]. The instrument captures both negative aspects, including perceptions of danger, suspicion, and distrust, as well as positive aspects such as reliability, integrity, security, and confidence toward system outputs. This approach reflects that trust is not a single construct but emerges from the combination of perceived system quality and perceived risk during interaction.

The trust questionnaire consisted of 12 evaluation statements adapted from the TiAS instrument and used the explanation object as the target of evaluation. The instrument included statements related to: trust toward explanations, perceived reliability, confidence in generated outputs, integrity of explanation, perceived safety and dependability, and suspicion and perceived harmful outcomes. Examples of evaluation items include “I can trust the explanation”, “The explanation is reliable”, and “I am confident in the explanation”.

Participant responses were aggregated into an overall trust score and converted to percentages to facilitate comparison between the existing and modified interfaces. Higher scores indicate stronger user trust toward the evaluated interface.

3.4.2 Understandability

Understandability was evaluated using a custom questionnaire developed specifically for this study. Unlike trust and usability, this study did not adopt a standardized instrument because existing literature indicates that understandability in XAI remains context-dependent and strongly influenced by how explanations are presented within user interaction. Therefore, the understandability instrument was developed through a synthesis of previous studies on explanation quality, interpretability, explanation effectiveness, transparency, and human-centered XAI evaluation and resulted in 10 evaluation statements measuring users’ understanding of AI-generated explanations during system interaction.

Examples of questionnaire statements include:

- “I generally understand how the system in this application works based on the explanations provided.” [37]
- “I understand the main factors or information that influence the results displayed by the system.” [51]
- “The explanations presented help me interpret the results provided by the system.” [52]

These statements were intended to evaluate not only whether participants could read the displayed information but also whether they could correctly interpret explanations and apply them during decision making. Responses from all items were aggregated and converted into percentage scores to support comparison between interface conditions.

Since the instrument was developed specifically for this study context, formal external validation outside the evaluated environment had not been conducted before implementation and is therefore a limitation of this study.

3.4.3 Usability

Usability was evaluated using the System Usability Scale (SUS). SUS is a standardized questionnaire-based evaluation instrument used to measure overall system usability through ten statements with five response options ranging from strongly disagree to strongly agree. SUS was selected because it enables efficient assessment of perceived usability while capturing user perceptions regarding ease of use, efficiency, and comfort during interaction [53]. The questionnaire consisted of 10 standard SUS statements, including items such as “I think that I would like to use this system frequently”, “I found the system unnecessarily complex”, and “I thought the system was easy to use”.

Participant responses were processed using the standard SUS scoring procedure to generate usability scores ranging from 0 to 100, where higher scores indicate better usability performance. Finally, usability scores were interpreted according to the SUS grading criteria used in this study.

3.4.4 Task Reflection

To complement quantitative findings, this study incorporated task reflection as a qualitative evaluation instrument. Task reflection was conducted after participants completed each interaction scenario to explore users’ global mental models regarding the system, including how they understood system functionality, interpreted presented explanations, and perceived their influence on trust and decision making. These aspects represent subjective dimensions of explainability evaluation that are often difficult to capture through structured questionnaires alone. Therefore, qualitative findings were used to support the interpretation of quantitative results.

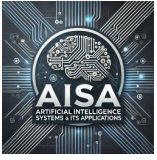
Participants were encouraged to reflect on their interaction experiences and provide feedback regarding difficulties encountered, interpretation processes, perceived usefulness of explanations, and overall impressions during system use. One example of the reflection questions used in this study was: “After using this application, to what extent do you understand how the system generates the displayed information or results?” This question was designed to evaluate users’ alignment of mental models with system behaviour. The reflection questions were developed and grouped based on the same XAI aspects used during the redesign and quantitative evaluation stages to support systematic analysis of user needs and interaction experiences.

3.6 Data Analysis

Data analysis was conducted using both quantitative and qualitative approaches to obtain a comprehensive understanding of the effects of implementing XAI principles in interface redesign. The integration of both approaches aligns with the concurrent embedded mixed methods design adopted in this study.

3.5.1 Quantitative Analysis

Quantitative analysis was performed to examine differences in participant perceptions between the existing interface and the redesigned interface implementing XAI principles. The analysis was conducted using questionnaire data collected from the dimensions of user trust, understandability, and usability. Because this study employed a within-subject design, statistical analysis was performed using score differences obtained from the same participants under both interface conditions. The quantitative analysis procedure consisted of three stages:



Normality Test

Normality testing was conducted as an initial step to determine the appropriate statistical method for paired comparison analysis. The tested data consisted of score differences between the existing and redesigned interfaces. This study applied the Shapiro-Wilk test, which is commonly recommended for small to medium sample sizes, particularly when the number of participants is fewer than 50 [54]. The Shapiro-Wilk statistic was calculated using Equation (1):

$$SW = \frac{(\sum_{i=1}^n w_i x_{(i)})^2}{\sum_{i=1}^n (x_i - \bar{x})^2} \quad (1)$$

where SW represents the Shapiro-Wilk statistic, w_i represents the Shapiro-Wilk weight coefficient, $x_{(i)}$ represents ordered observations, x_i represents individual observations, $\bar{x}$ represents the sample mean, and n represents the number of observations.

Normality decisions were determined using the significance value (p-value) with $\alpha = 0.05$. Data were considered normally distributed when $p > 0.05$, whereas $p \leq 0.05$ indicated a non-normal distribution.

Hypothesis Test

Hypothesis testing was conducted to determine whether statistically significant differences existed between the existing and redesigned interfaces. Because the evaluation involved two related interface conditions, this study employed the paired-samples t-test for normally distributed data. The paired-samples t-test is an inferential statistical method used to examine whether significant differences exist between two dependent observations [55]. The paired-samples t-test statistic was calculated using Equation (2):

$$t = \frac{\bar{d}}{s_d / \sqrt{n}} \quad (2)$$

where t represents the test statistic, $\bar{d}$ represents the mean difference between paired observations, s_d represents the standard deviation of score differences, and n represents the number of paired observations.

The hypotheses were defined as:

$$H_0: \mu_{existing} = \mu_{modified}$$

$$H_1: \mu_{existing} \neq \mu_{modified}$$

Decisions were based on two-tailed significance values with $\alpha = 0.05$. If normality assumptions were not satisfied, analysis was continued using the Wilcoxon Signed-Rank Test as a nonparametric alternative.

Effect Size Analysis

Effect size analysis was conducted to measure the magnitude of interface redesign effects beyond statistical significance. Effect size provides additional interpretation regarding the strength of treatment effects and supports practical interpretation of findings, especially in studies with limited sample sizes [56]. For parametric analysis, effect size was calculated using Cohen's d :

$$d = \frac{\bar{X}_1 - \bar{X}_2}{S_p} \quad (3)$$

where the pooled standard deviation was calculated using Equation (4):

$$S_p = \sqrt{\frac{(n_1 - 1)s_1^2 + (n_2 - 1)s_2^2}{n_1 + n_2 - 2}} \quad (4)$$

where d represents Cohen's effect size, $\bar{X}_1$ and $\bar{X}_2$ represents group means, S_p represents the pooled standard deviation, and n and s_1^2 s_2^2 represent the sample size and variance for each condition.

The interpretation of effect size followed the conventional threshold proposed by Cohen, where a Cohen's d value of approximately 0.20 represents a small effect, approximately 0.50 represents a medium effect, and values of 0.80 or greater indicate a large effect.

3.5.2 Qualitative Analysis

Qualitative analysis was conducted to obtain a deeper understanding of user experiences during interaction with the application, particularly regarding how participants interpreted system explanations, utilized presented information, and connected explanations to decision-making processes. Qualitative data were collected through task reflection administered after participants completed evaluation tasks. The study employed thematic analysis, a qualitative analysis method used to identify recurring patterns of meaning (themes) within textual responses.

The analysis procedure consisted of four stages: familiarization with participant responses to obtain an overall understanding of the collected qualitative data, coding relevant statements according to the objectives of the task reflection evaluation, grouping similar codes into recurring themes and sub-themes to identify common response patterns, and interpreting the resulting themes to explain the quantitative findings. The coding was conducted by the first author using an iterative thematic analysis process. Initial codes were generated from participant responses and were repeatedly reviewed and refined until stable themes and sub-themes were established. Because coding was performed by a single researcher, formal inter-coder agreement was not assessed, which is acknowledged as a limitation of this study.

To support transparent reporting, thematic findings were summarized in a structured table containing the identified theme, sub-theme, theme description, and the number and percentage of respondents mentioning each theme. The thematic results were subsequently integrated with the quantitative findings to explain how explicability influenced trust, understanding, and perceived usability.

4 Results and Discussion

4.1 Interface Design Result

The redesign process produced a modified interface for the AI-based chicken weight monitoring application by integrating six XAI principles synthesized from previous studies. These principles were translated into interface elements supporting transparency, interpretation, and decision-making based on user needs and evaluation findings from the existing interface.

4.1.1. Comprehensible Explanation



Figure 3: Implementation of XAI for Comprehensible Explanation

In the existing interface, the QR scanning page provided limited guidance and required users to infer the required actions independently, resulting in differences in user understanding and confusion during the initial interaction. To address this issue, the redesigned interface implemented the Comprehensible Explanation principle by introducing brief instructional guidance directly within the QR scanning page. The explanation was presented using non-technical and natural language to improve users' understanding when interacting with the system from the beginning of the process [32], [44]. The information focused only on essential steps required to complete QR pairing successfully to maintain concise communication and avoid unnecessary cognitive burden [31]. As illustrated in Figure 2, the explanation was embedded directly into the interface to provide contextual support at the moment users needed the information, allowing users to understand system procedures more easily without requiring prior technical knowledge [37].

4.1.2. Progressive and Context-Aware Explanation

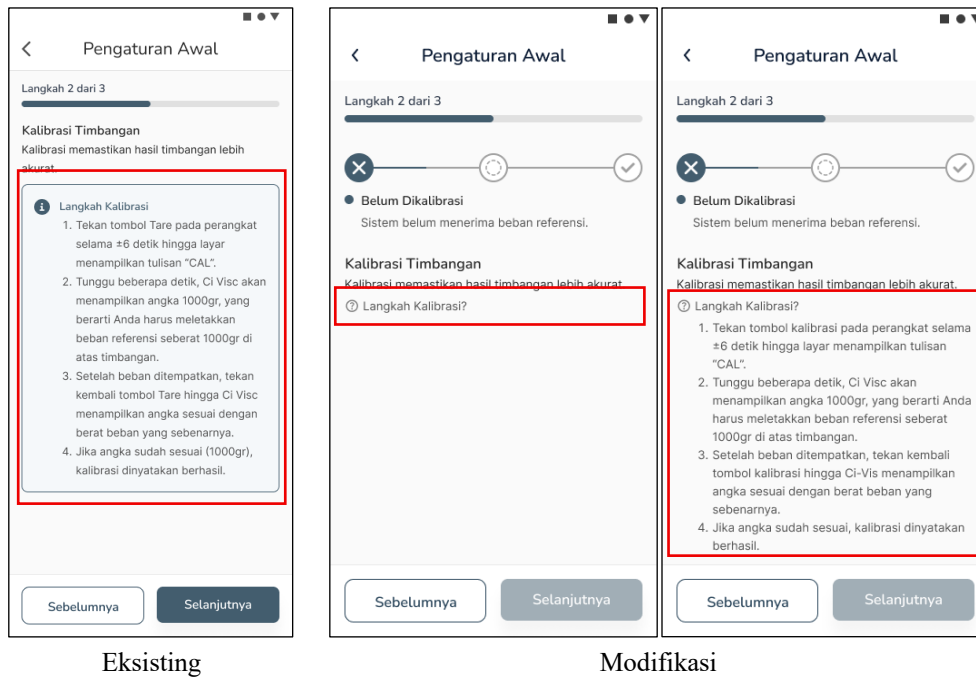


Figure 4: Implementation of XAI for Progressive and Context-Aware Explanation

In the existing interface, calibration instructions remained continuously visible regardless of user needs, reducing flexibility during interaction and potentially displaying unnecessary information for experienced users. To improve this condition, the redesigned interface implemented the Progressive and Context-Aware Explanation principle by allowing calibration guidance to be expanded or collapsed on demand. As shown in Figure 3, concise information was initially presented in a collapsed state and revealed more detailed calibration steps only when users requested additional guidance. This approach enables users to access explanations according to their level of familiarity while keeping the interface concise and contextually relevant [37], [38].

4.1.3. Reasoning and Causal Explanation

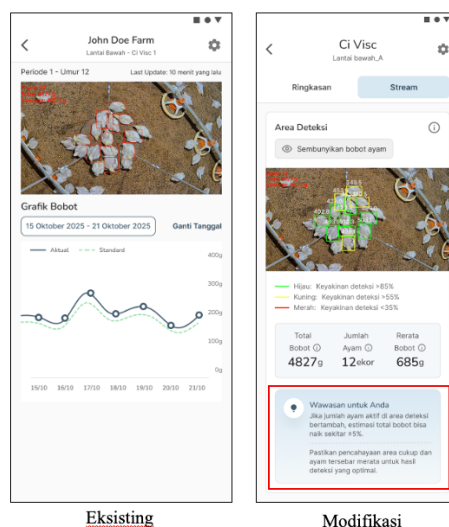


Figure 5: Implementation of XAI for Reasoning and Causal Explanation

The existing interface presented chicken weight monitoring results without additional explanations describing the factors behind the observed changes. As a result, users could view the output but received limited support in understanding the reasons underlying the displayed condition. To address this limitation, the redesigned interface implemented the Reasoning and Causal Explanation principle through the addition of an “Insights for You” section that provides concise explanatory insight directly within the monitoring interface. As illustrated in Figure 4, the explanation highlighted key factors contributing to weight changes and presented cause-and-effect relationships to support user interpretation of the results [42]. The insight was also designed to provide contextual understanding and reduce the risk of misinterpreting unusual or significant changes in monitoring data [43]. Beyond explaining system outputs, the explanation encouraged users to better understand potential actions or considerations based on the presented information [46].

4.1.4. Actionable and Goal-Oriented Explanation

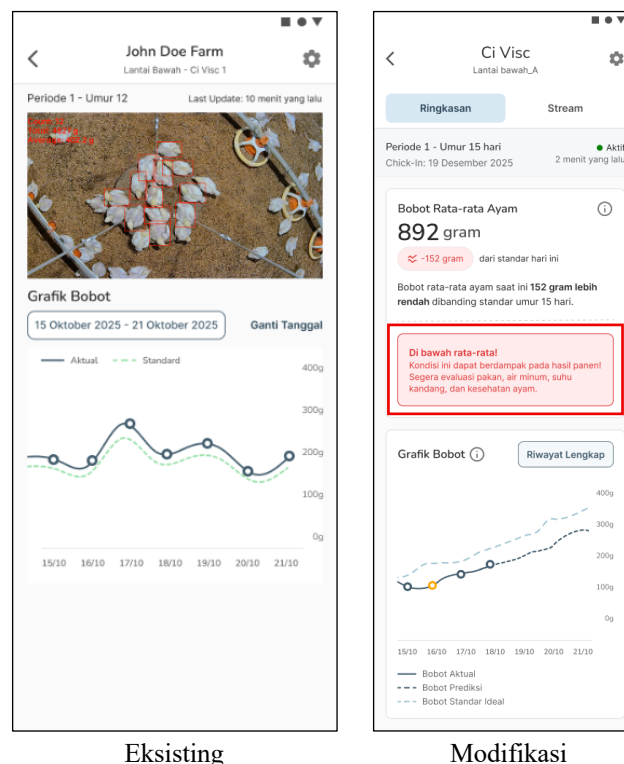


Figure 6: Implementation of XAI for Actionable and Goal-Oriented Explanation

The existing monitoring interface presented chicken weight results without indicating what actions users should take based on the displayed condition. As a result, users needed to interpret the results independently and determine appropriate responses without additional guidance. To address this limitation, the redesigned interface implemented the Actionable and Goal-Oriented Explanation principle by embedding action-oriented recommendations directly within the

summary card displayed at the beginning of the monitoring page. As illustrated in Figure 5, messages such as warning indicators translate monitoring outcomes into suggested follow-up actions, allowing users to immediately recognize conditions that require attention. The recommendations were designed to align with the detected context and guide users toward relevant decisions rather than presenting monitoring data alone [37], [45], [46].

4.1.5. Trust Calibration and Bias Awareness

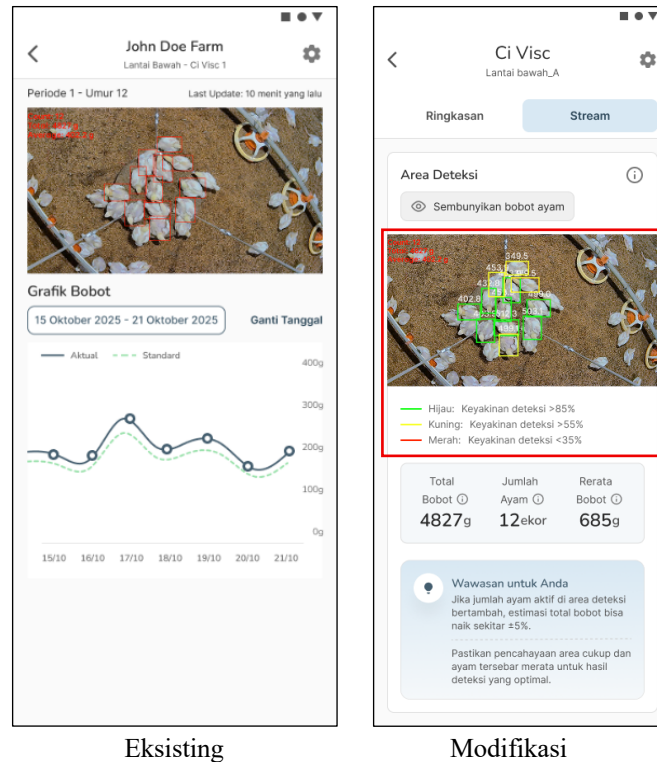


Figure 7: Implementation of XAI for Trust Calibration and Bias Awareness

The existing monitoring interface displayed detection results without communicating the level of confidence associated with the generated output. This limited users' ability to judge whether the displayed results should be accepted directly or interpreted with additional consideration. To support more calibrated trust, the redesigned interface introduced a confidence score within the monitoring page as an indicator of the system's confidence in the detected results. As shown in Figure X, this information provided users with additional context regarding output quality and encouraged more informed interpretation of the displayed values/ presenting confidence information also increased transparency by explicitly communicating uncertainty and reducing the tendency to rely on system outputs without evaluation. In addition, confidence values enabled users to observe variations across monitoring conditions and identify situations that may require further attention before making decisions.

4.1.6. Mental Model Alignment

Unlike the existing interface, which primarily emphasized presenting monitoring results, the redesigned system introduced a workflow explanation to provide users with a broader understanding of how the application operates from input to output. The explanation outlined the main processing stages and was designed to help users build a mental model of the system



by understanding the relationship between input data, system processing, and generated results [41]. To improve accessibility, the workflow was communicated using structured and non-technical language that simplified complex processes into understandable concepts [49]. This implementation aimed to reduce black-box perception, increase transparency, and support users in developing more informed interpretations of system outputs rather than relying solely on final results [35], [43], [49].

4.2 Quantitative Evaluation

This study quantitatively evaluated the impact of implementing XAI principles on three user experience dimensions: user trust, understandability, and usability. Evaluation was conducted by comparing participant responses between the existing interface and the modified interfaces using a within-subject design. Statistical analysis was performed using paired comparison methods to determine whether the observed differences represented meaningful improvements. Prior to hypothesis testing, normality testing was conducted using the Shapiro-Wilk test, and the results indicated that the evaluated dimensions satisfied the normality assumption, allowing further analysis using the paired-samples t-test.

Table 3: Quantitative Evaluations Results

Aspect	Version	Mean ($\bar{x}$)	Data Analysis			
			Normality	Hypothesis	Effect Size	Interpretation
User Trust	Existing	50.52	Normal	< 0.001	-3.916	Large effect
	Modified	73.82	Normal			
Understandability	Existing	54.25	Normal	< 0.001	-5.235	Large effect
	Modified	81.88	Normal			
Usability	Existing	43.28	Normal	< 0.001	-4.194	Large effect
	Modified	74.38	Normal			

As shown in Table 3, the implementation of XAI-based interface redesign produced improvements across all evaluated dimensions.

For user trust, the average score increased from 50.52 in the existing interface to 73.82 after redesign implementation. Statistical testing showed a significant difference between conditions ($p < 0.001$), indicating rejection of the null hypothesis. The calculated effect size (Cohen's $d = -3.916$) indicated a large practical effect, suggesting that the observed improvement was not only statistically significant but also meaningful in terms of user perception.

A similar pattern as observed for understandability, which increased from 54.25 to 81.88. Hypothesis testing also demonstrated a statistically significant difference ($p < 0.001$) with a large effect size (Cohen's $d = -5.235$). Among the evaluated dimensions, understandability showed the largest magnitude of change, suggesting that integrating the explanation element directly into the interface substantially improved users' ability to interpret and understand AI-generated outputs.

For usability, the average score improved from 43.28 under the existing interface to 74.38 after redesign. Statistical analysis again showed a significant difference ($p < 0.001$) accompanied by a large effect size (Cohen's $d = -4.194$). These findings indicate that the redesigned interface not only

improved comprehension of explanations but also enhanced overall interaction quality and perceived ease of use.

Across all dimensions, the modified interface consistently achieved higher mean scores and large practical effects. These results suggest that integrating explainability principles into interface design contributed positively to users' interaction experience beyond improving transparency alone. The findings indicate that explanation mechanisms embedded within interface elements may support users in developing greater confidence, a stronger understanding of system outputs, and more effective interaction during monitoring activities. However, these improvements should not be attributed solely to XAI principles, as clearer instructions, improved information hierarchy, and a more structured interface layout may also have contributed to the observed improvements.

4.3 Qualitative Findings

The qualitative data were analyzed using thematic analysis to identify recurring patterns across participant responses and group them into broader themes. The coding process followed the same conceptual structure used during redesign and evaluation, namely the six XAI principles implemented in the interface. This approach allowed the qualitative findings to support the interpretation of quantitative outcomes directly.

Table 4: Summary of Quantitative Findings

Interface Version	Theme	Key Findings	Respondents (%)
Existing	Limited explanation quality	Explanations were insufficient and lacked contextual relevance	100%
	Limited support for understanding	Users interpreted outputs independently due to missing reasoning	100%
	Low decision support	Outputs were rarely used directly for operational decisions	100%
	Limited transparency and trust	Users relied on manual verification	100%
	Improved explanation clarity	Information became easier to interpret	100%
Modified	Better contextual understanding	Explanations became more relevant to the interaction context	87.5%
	Improved reasoning and mental model	Users better understood the relationships between conditions and outputs	100%
	Increased trust and decision support	Confidence toward system outputs improved	100%
	Remaining improvement opportunities	Users requested deeper explanations and additional features	93.75%

As shown in Table 4, the thematic analysis revealed substantial differences in how participants experienced the existing and modified interfaces. For the existing interface, participants consistently reported that explanations did not provide sufficient support for understanding system outputs. Information was perceived as incomplete and insufficiently contextualized, requiring users to interpret results independently. One participant stated, “*Menurutku masih kurang jelas, soale penjelasannya dikit.*” (P13), while another explained that “*Ini ada kemungkinan alatnya nimbang kotoran gak ya?*” (P01). As a result, participants continued relying more on field observation and prior experience than on system recommendations, which reduced trust and limited understanding of AI-generated outputs.



Following implementation of the XAI-based redesign, participants reported improvements in explanation clarity, interpretation, and interaction experience. Explanations were perceived as more complete and easier to connect with operational conditions. One participant stated, “*Jauh lebih lengkap ini penjelasannya.*” (P01), while another highlighted that “*Bagus ini ada histori datanya, kalau yang sebelumnya cuman terbatas 7 hari kan ya.*” (P12). Participants also reported greater confidence in interpreting results and using system outputs during monitoring activities. However, several participants noted that more detailed information and additional supporting features would further improve the system experience.

4.4 Discussion

This study shows that the integration of XAI into the interface design of the chicken weight monitoring system leads to consistent improvements in user trust, understandability, and usability. To answer RQ1, the findings demonstrate that integrating XAI principles into the interface design significantly improved user trust compared with the existing interface. The redesigned interface provided more transparent and contextual explanations through instructional guidance, causal insights, confidence indicators, and workflow visualization, enabling participants to better understand how monitoring results were generated. As a result, users developed greater confidence in interpreting system outputs and were better able to validate the presented information. These findings are supported by the qualitative results, which indicate that participants perceived the explanations as clearer and more informative, thereby reducing the black-box perception commonly associated with AI systems.

To answer RQ2, the redesigned interface also substantially improved users’ understandability of AI-generated information. The integration of comprehensible explanations, progressive information disclosure, reasoning-based insights, and workflow explanations enabled participants to develop a clearer mental model of how the system transformed input data into monitoring results. Rather than only observing the final outputs, users were able to understand the relationship between monitoring data, system processing, and generated results. The qualitative findings further support this improvement, showing that participants experienced fewer difficulties in interpreting system outputs and considered the redesigned explanations easier to follow.

Regarding RQ3, the findings indicate that redesigning the interface by integrating XAI principles also improved overall usability. Participants completed monitoring tasks more efficiently because explanatory elements were embedded directly into the interaction flow, reducing unnecessary interpretation and supporting faster decision-making. However, these improvements should not be attributed solely to the implementation of XAI principles. General interface design enhancements, including clearer instructions, improved information hierarchy, and a more organized visual presentation, may also have contributed to the observed improvements in usability and overall interaction quality.

Overall, the findings indicate that user trust, understandability, and usability are closely interconnected. Improvements in explanation clarity strengthened participants’ understanding of system behavior, which subsequently increased confidence in AI-generated outputs and enhanced overall interaction quality. Rather than proposing a new XAI algorithm, this study demonstrated how human-centered XAI principles can be operationalized into interface design and empirically evaluated within an existing AI-based chicken weight monitoring system. By translating explainability principles into concrete interface elements and evaluating their effects in a smart farming context, this study extends the application of human-centered XAI beyond algorithm development toward interface-level implementation. These findings are consistent

with previous studies [10] which highlight the role of AI explanations in improving understandability and trustworthiness, as well as studies [15] and [20] which positions trust, understandability, and usability as interrelated dimensions in human-centered XAI evaluation. Furthermore, the result aligns with [19], which demonstrates that XAI can support calibrated trust through explanations, and [17], which emphasizes that trust and understandability are key determinants of user experience quality.

5 Conclusion

This study evaluated the impact of integrating XAI principles into the interface design of a chicken weight monitoring system. The findings show that the modified interface consistently improves user trust, understandability, and usability compared to the existing interface. Quantitative results indicate significant differences across all dimensions with large effect sizes, where user trust increased from 50.52 to 73.82, understandability from 54.25 to 81.88, and usability from 43.28 to 74.38. These results are strengthened by qualitative findings that show clear explanations and better contextual information help users interpret system outputs more easily and increase their confidence in using the system. These improvements are driven by the integration of multiple XAI components, including comprehensible explanations, contextual and progressive information, causal reasoning, actionable recommendations, confidence indicators, and workflow explanations. These elements work together within the interface to reduce uncertainty, improve understanding of system behavior, and support decision-making during monitoring activities.

The results confirm that explanation design functions as an integrated mechanism that shapes overall user experience in AI-based systems. However, although the redesigned interface integrated XAI principles, the observed improvements should not be attributed solely to XAI. They may also reflect the influence of broader interface redesign decisions, such as clearer instructions, improved information organization, and enhanced visual presentation, which collectively contributed to the overall interaction quality.

Although the results show consistent improvements, this study has limitations that should be considered. The evaluation involved a limited number of participants, which may affect the generalizability of the findings. In addition, the instrument used to measure understandability was adapted from previous studies and has not undergone full external validation. Future research is recommended to involve larger and more diverse participant groups and to explore adaptive explanation approaches that adjust the level of detail based on user needs and context, in order to further improve clarity, reduce cognitive load, and enhance usability in real-world applications.

Acknowledgments

The authors would like to thank the participants involved in this study.

Funding Information

The authors declare that no funding was received for this study.

Conflict of Interest Statement

The authors declare no conflicts of interest.



Ethical Approval

This study involved human participants for usability evaluation. Participation was voluntary and informed consent was obtained.

Data Availability

Data are available upon reasonable request from the corresponding author.

References

- [1] Yohanes Bowo Widodo, "Integrating Artificial Intelligence for Smart and Adaptive Information Systems," *jisem*, vol. 10, no. 9s, pp. 751–757, Feb. 2025, doi: 10.52783/jisem.v10i9s.1301.
- [2] A. Zsombok and I. Zsombok, "Revolutionizing Predictive Maintenance: How AI-Driven Solutions Enhance Efficiency and Reduce Costs Across Industries," *IJSR*, vol. 12, no. 12, pp. 1168–1171, Dec. 2023, doi: 10.21275/SR231214124830.
- [3] A. Yaseer and H. Chen, "A Review of Sensors and Machine Learning in Animal Farming," in *2021 IEEE 11th Annual International Conference on CYBER Technology in Automation, Control, and Intelligent Systems (CYBER)*, Jiaxing, China: IEEE, Jul. 2021, pp. 747–752. doi: 10.1109/CYBER53097.2021.9588295.
- [4] R. Sasirekha, K. R. S. G. A. Mohamed, and U. Iroda, "Smart Poultry House Monitoring System Using IoT," *E3S Web Conf.*, vol. 399, p. 04055, 2023, doi: 10.1051/e3sconf/202339904055.
- [5] B. Y. Lim, A. K. Dey, and D. Avrahami, "Why and why not explanations improve the intelligibility of context-aware intelligent systems," in *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems*, Boston MA USA: ACM, Apr. 2009, pp. 2119–2128. doi: 10.1145/1518701.1519023.
- [6] J. Zhou, S. Z. Arshad, S. Luo, and F. Chen, "Effects of Uncertainty and Cognitive Load on User Trust in Predictive Decision Making," in *Human-Computer Interaction – INTERACT 2017*, vol. 10516, R. Bernhaupt, G. Dalvi, A. Joshi, D. K. Balkrishan, J. O'Neill, and M. Winckler, Eds., in *Lecture Notes in Computer Science*, vol. 10516, Cham: Springer International Publishing, 2017, pp. 23–39. doi: 10.1007/978-3-319-68059-0_2.
- [7] U. B. Khakurel and D. B. Rawat, "Evaluating explainable artificial intelligence (XAI): algorithmic explanations for transparency and trustworthiness of ML algorithms and AI systems," in *Artificial Intelligence and Machine Learning for Multi-Domain Operations Applications IV*, T. Pham, L. Solomon, and M. E. Hohil, Eds., Orlando, United States: SPIE, Jun. 2022, p. 7. doi: 10.1117/12.2620598.
- [8] P. Liu, L. Wang, and J. Li, "Unlocking the Potential of Explainable Artificial Intelligence in Remote Sensing Big Data," *Remote Sensing*, vol. 15, no. 23, p. 5448, Nov. 2023, doi: 10.3390/rs15235448.
- [9] P. J. Phillips, C. A. Hahn, P. C. Fontana, D. A. Broniatowski, and M. A. Przybocki, "Four Principles of Explainable Artificial Intelligence," Aug. 17, 2020. doi: 10.6028/NIST.IR.8312-draft.
- [10] S. Laato, M. Tiainen, A. K. M. Najmul Islam, and M. Mäntymäki, "How to explain AI systems to end users: a systematic literature review and research agenda," *INTR*, vol. 32, no. 7, pp. 1–31, Dec. 2022, doi: 10.1108/INTR-08-2021-0600.

- [11] S. Mohseni, N. Zarei, and E. D. Ragan, "A Multidisciplinary Survey and Framework for Design and Evaluation of Explainable AI Systems," *ACM Trans. Interact. Intell. Syst.*, vol. 11, no. 3–4, pp. 1–45, Dec. 2021, doi: 10.1145/3387166.
- [12] L.-V. Herm, "Impact Of Explainable AI On Cognitive Load: Insights From An Empirical Study," 2023, *arXiv*. doi: 10.48550/ARXIV.2304.08861.
- [13] A. Hudon, T. Demazure, A. Karran, P.-M. Léger, and S. Sénécal, "Explainable Artificial Intelligence (XAI): How the Visualization of AI Predictions Affects User Cognitive Load and Confidence," in *Information Systems and Neuroscience*, vol. 52, F. D. Davis, R. Riedl, J. Vom Brocke, P.-M. Léger, A. B. Randolph, and G. Müller-Putz, Eds., in *Lecture Notes in Information Systems and Organisation*, vol. 52., Cham: Springer International Publishing, 2021, pp. 237–246. doi: 10.1007/978-3-030-88900-5_27.
- [14] N. Burkart and M. F. Huber, "A Survey on the Explainability of Supervised Machine Learning," *jair*, vol. 70, pp. 245–317, Jan. 2021, doi: 10.1613/jair.1.12228.
- [15] J. Kim, H. Maathuis, and D. Sent, "Human-centered evaluation of explainable AI applications: a systematic review," *Front. Artif. Intell.*, vol. 7, p. 1456486, Oct. 2024, doi: 10.3389/frai.2024.1456486.
- [16] A. Mangold, J. Zietz, S. Weinhold, and S. Pannasch, "On the Design and Evaluation of Human-centered Explainable AI Systems: A Systematic Review and Taxonomy," 2025, *arXiv*. doi: 10.48550/ARXIV.2510.12201.
- [17] D. Lei, Y. He, and J. Zeng, "What Is the Focus of XAI in UI Design? Prioritizing UI Design Principles for Enhancing XAI User Experience," in *Artificial Intelligence in HCI*, vol. 14734, H. Degen and S. Ntoa, Eds., in *Lecture Notes in Computer Science*, vol. 14734., Cham: Springer Nature Switzerland, 2024, pp. 219–237. doi: 10.1007/978-3-031-60606-9_13.
- [18] S. Brdnik, "GUI Design Patterns for Improving the HCI in Explainable Artificial Intelligence," in *28th International Conference on Intelligent User Interfaces*, Sydney NSW Australia: ACM, Mar. 2023, pp. 240–242. doi: 10.1145/3581754.3584114.
- [19] B. Leichtmann, C. Humer, A. Hinterreiter, M. Streit, and M. Mara, "Effects of Explainable Artificial Intelligence on trust and human behavior in a high-risk decision task," *Computers in Human Behavior*, vol. 139, p. 107539, Feb. 2023, doi: 10.1016/j.chb.2022.107539.
- [20] Y. Rong *et al.*, "Towards Human-Centered Explainable AI: A Survey of User Studies for Model Explanations," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 46, no. 4, pp. 2104–2122, Apr. 2024, doi: 10.1109/TPAMI.2023.3331846.
- [21] N. Balasubramaniam, M. Kauppinen, A. Rannisto, K. Hiekkänen, and S. Kujala, "Transparency and explainability of AI systems: From ethical guidelines to requirements," *Information and Software Technology*, vol. 159, p. 107197, Jul. 2023, doi: 10.1016/j.infsof.2023.107197.
- [22] N. Li, Z. Ren, D. Li, and L. Zeng, "Review: Automated techniques for monitoring the behaviour and welfare of broilers and laying hens: towards the goal of precision livestock farming," *Animal*, vol. 14, no. 3, pp. 617–625, 2020, doi: 10.1017/S1751731119002155.
- [23] B. Paneru *et al.*, "Computer vision models for precision poultry farming: A narrative review of behavioral and welfare monitoring studies," *Poultry Science*, vol. 105, no. 7, p. 106887, Jul. 2026, doi: 10.1016/j.psj.2026.106887.
- [24] E. Supriyanto, R. R. Isnanto, and S. Hadi Purnomo, "Computer Vision in Chicken Monitoring System Using Machine Learning: A General Review," *E3S Web Conf.*, vol. 448, p. 02014, 2023, doi: 10.1051/e3sconf/202344802014.
- [25] W. A. Nurtrisha, L. Ramadani, R. Y. Fa'rifah, F. Hamami, and N. I. Utama, "Artificial Intelligence for Precision Livestock Farming: A Systematic Review of Applications, Models, and Evaluation Metrics," *JUSIFO: J. Sistem Inf.*, vol. 11, no. 2, pp. 121–132, Dec. 2025, doi: 10.19109/jusifo.v11i2.31179.
- [26] B. Petrovic, V. Tunguz, and P. Bartos, "Application of computer vision in livestock and crop production—A review," *Comput. Artif. Intell.*, vol. 1, no. 1, p. 360, Nov. 2023, doi: 10.59400/cai.v1i1.360.
- [27] D. S. Azhari, Z. Afif, M. Kustati, and N. Sepriyanti, "Penelitian Mixed Method Research untuk Disertasi," *Innovative: Journal of Social Science Research*, vol. 3, no. 2, pp. 8010–8025, 2023.



- [28] Sugiyono, *Metode Penelitian Kuantitatif, Kualitatif, dan R&D*, Cetakan ke-19. Bandung: Alfabeta, CV, 2013.
- [29] A. K. Montoya, "Selecting a Within- or Between-Subject Design for Mediation: Validity, Causality, and Statistical Power," *Multivariate Behavioral Research*, vol. 58, no. 3, pp. 616–636, May 2023, doi: 10.1080/00273171.2022.2077287.
- [30] H. T. Gebremariam and D. M. Asgede, "Effects of students' self-reflection on improving essay writing achievement among Ethiopian undergraduate students: a counterbalanced design," *Asian. J. Second. Foreign. Lang. Educ.*, vol. 8, no. 1, p. 30, Oct. 2023, doi: 10.1186/s40862-023-00203-7.
- [31] J. Zhang, "User Interface Design Based on Human-Centered Explainable AI Methods," Master's Thesis, Aalto University, 2022.
- [32] W. Samek, T. Wiegand, and K.-R. Müller, "Explainable Artificial Intelligence: Understanding, Visualizing and Interpreting Deep Learning Models," 2017, *arXiv*. doi: 10.48550/ARXIV.1708.08296.
- [33] J. Tullio, A. K. Dey, J. Chalecki, and J. Fogarty, "How it works: a field study of non-technical users interacting with an intelligent system," in *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems*, San Jose California USA: ACM, Apr. 2007, pp. 31–40. doi: 10.1145/1240624.1240630.
- [34] T. Miller, P. Howe, and L. Sonenberg, "Explainable AI: Beware of Inmates Running the Asylum Or: How I Learnt to Stop Worrying and Love the Social and Behavioural Sciences," 2017, *arXiv*. doi: 10.48550/ARXIV.1712.00547.
- [35] V. Bellotti and K. Edwards, "Intelligibility and Accountability: Human Considerations in Context-Aware Systems," *Human-Computer Interaction*, vol. 16, no. 2–4, pp. 193–212, Dec. 2001, doi: 10.1207/S15327051HCI16234_05.
- [36] J. L. Herlocker, J. A. Konstan, and J. Riedl, "Explaining collaborative filtering recommendations," in *Proceedings of the 2000 ACM conference on Computer supported cooperative work*, Philadelphia Pennsylvania USA: ACM, Dec. 2000, pp. 241–250. doi: 10.1145/358916.358995.
- [37] Q. V. Liao, D. Gruen, and S. Miller, "Questioning the AI: Informing Design Practices for Explainable AI User Experiences," 2020, doi: 10.48550/ARXIV.2001.02478.
- [38] F. Doshi-Velez and B. Kim, "Towards A Rigorous Science of Interpretable Machine Learning," 2017, *arXiv*. doi: 10.48550/ARXIV.1702.08608.
- [39] J. D. Moore and C. L. Paris, "Requirements for an expert system explanation facility," *Computational Intelligence*, vol. 7, no. 4, pp. 367–370, Nov. 1991, doi: 10.1111/j.1467-8640.1991.tb00409.x.
- [40] Z. C. Lipton, "The Mythos of Model Interpretability," 2016, *arXiv*. doi: 10.48550/ARXIV.1606.03490.
- [41] R. R. Hoffman, S. T. Mueller, G. Klein, and J. Litman, "Metrics for Explainable AI: Challenges and Prospects," 2018, *arXiv*. doi: 10.48550/ARXIV.1812.04608.
- [42] J. Y. Halpern and J. Pearl, "Causes and Explanations: A Structural-Model Approach. Part II: Explanations," 2002, *arXiv*. doi: 10.48550/ARXIV.CS/0208034.
- [43] A. Nguyen, A. Dosovitskiy, J. Yosinski, T. Brox, and J. Clune, "Synthesizing the preferred inputs for neurons in neural networks via deep generator networks," 2016, *arXiv*. doi: 10.48550/ARXIV.1605.09304.
- [44] D. Gunning, E. Vorm, J. Y. Wang, and M. Turek, "DARPA 's explainable AI (XAI) program: A retrospective," *Applied AI Letters*, vol. 2, no. 4, p. e61, Dec. 2021, doi: 10.1002/ail2.61.
- [45] A. Adadi and M. Berrada, "Peeking Inside the Black-Box: A Survey on Explainable Artificial Intelligence (XAI)," *IEEE Access*, vol. 6, pp. 52138–52160, 2018, doi: 10.1109/ACCESS.2018.2870052.

- [46] S. Wachter, B. Mittelstadt, and C. Russell, “Counterfactual Explanations without Opening the Black Box: Automated Decisions and the GDPR,” 2017, doi: 10.48550/ARXIV.1711.00399.
- [47] C. T. Wolf, “Explainability scenarios: towards scenario-based XAI design,” in *Proceedings of the 24th International Conference on Intelligent User Interfaces*, Marina del Ray California: ACM, Mar. 2019, pp. 252–257. doi: 10.1145/3301275.3302317.
- [48] T. Kulesza, S. Stumpf, M. Burnett, S. Yang, I. Kwan, and W.-K. Wong, “Too much, too little, or just right? Ways explanations impact end users’ mental models,” in *2013 IEEE Symposium on Visual Languages and Human Centric Computing*, San Jose, CA, USA: IEEE, Sep. 2013, pp. 3–10. doi: 10.1109/VLHCC.2013.6645235.
- [49] M. Eiband, H. Schneider, M. Bilandzic, J. Fazekas-Con, M. Haug, and H. Hussmann, “Bringing Transparency Design into Practice,” in *23rd International Conference on Intelligent User Interfaces*, Tokyo Japan: ACM, Mar. 2018, pp. 211–223. doi: 10.1145/3172944.3172961.
- [50] S. Brdnik, V. Podgorelec, and B. Šumak, “Assessing Perceived Trust and Satisfaction with Multiple Explanation Techniques in XAI-Enhanced Learning Analytics,” *Electronics*, vol. 12, no. 12, p. 2594, Jun. 2023, doi: 10.3390/electronics12122594.
- [51] D. S. Weld and G. Bansal, “The challenge of crafting intelligible intelligence,” *Commun. ACM*, vol. 62, no. 6, pp. 70–79, May 2019, doi: 10.1145/3282486.
- [52] J. Dieber and S. Kirrane, “A novel model usability evaluation framework (MUSE) for explainable artificial intelligence,” *Information Fusion*, vol. 81, pp. 143–153, May 2022, doi: 10.1016/j.inffus.2021.11.017.
- [53] D. Saputra, E. Ardiyan Syah, and F. Darnis, “Usability Testing on the Simponik Website using the System Usability Scale (SUS),” *Sinkron*, vol. 7, no. 4, pp. 2584–2592, Nov. 2022, doi: 10.33395/sinkron.v7i4.11916.
- [54] M. Isnaini, M. W. Afgani, A. Haqqi, and I. Azhari, “Teknik Analisis Data Uji Normalitas,” 2025.
- [55] B. Darma, *Statistika penelitian menggunakan SPSS (Uji validitas, uji reliabilitas, regresi linier sederhana, regresi linier berganda, uji t, uji F, R2)*. 2021.
- [56] H. Hardiana, P. Putriyani, S. Djafar, N. Nurdin, and P. P. Alam, “Meta Analisis Lembar Kerja Digital Berbasis Geogebra Terhadap Kemampuan Penalaran,” *JRHotsPM*, vol. 5, no. 2, Jun. 2025, doi: 10.51574/kognitif.v5i2.3006.