

Explainability Practices and Challenges from the Experience of a Brazilian Industrial Research and Innovation Company

Lívia Mancine
Universidade Federal de Goiás
Instituto Federal Goiano
Goiânia, GO, Brazil
livia.mancine@ifgoiano.edu.br

Renata Braga
Universidade Federal de Goiás
Goiânia, GO, Brazil
renatadbraga@ufg.br

Davi Viana
Universidade Federal do Maranhão
São Luís, MA, Brazil
davi.viana@ufma.br

Marcos Kalinowski
Pontifícia Universidade Católica do
Rio de Janeiro
Rio de Janeiro, RJ, Brazil
kalinowski@inf.puc-rio.br

Renato F. Bulcão-Neto
Universidade Federal de Goiás
Goiânia, GO, Brazil
rbulcao@ufg.br

Abstract

Explainability plays a key role in the development of Artificial Intelligence (AI)-enabled systems. Although the literature recognizes the potential of Requirements Engineering (RE) to achieve explainability, its practical adoption remains limited, often lacking real-world applicability and systematic methods to embed explainability from the early stages of AI-enabled system development. This study investigates how AI project leaders perceive explainability and the challenges they face when addressing it as a software requirement in a Brazilian Research, Development, and Innovation (RD&I) unit. We performed semi-structured interviews, which were analyzed through Grounded Theory's open and axial coding. Our findings reveal that explainability is perceived as a transversal quality factor closely related to trust and transparency. Moreover, participants report cultural and knowledge barriers and emphasize the need to adapt RE processes to better integrate explainability. This study offers qualitative and exploratory insights into how explainability is currently understood and applied in practice, revealing gaps in RE processes and highlighting the need to advance RE practices focused on explainability in AI-enabled systems.

Keywords

Requirements Engineering, Explainability, Artificial Intelligence, Industry, Experimentation, Interview, Open and Axial Coding

1 Introduction

Artificial Intelligence (AI)-enabled systems are increasingly adopted to enhance products, optimize processes, and support decision-making. However, their adoption introduces new challenges for Requirements Engineering (RE), particularly regarding the need to ensure that system decisions are understandable and trustworthy.

AI pipelines are capable of producing accurate predictions; however, they often fail to incorporate two crucial phases: understanding and explanation. Understanding involves analyzing the problem domain, data, and model behavior, including training and quality assurance [30]. Explanation, in turn, is essential throughout the AI lifecycle, especially when models are deployed in real-world applications. Together, these phases contribute to model interpretability, trustworthiness, and effective deployment [7].

Explainability refers to the capability of a system to provide intelligible and contextually meaningful explanations regarding its functioning and decision-making processes. It is considered a non-functional requirement (NFR) with a substantial impact on the quality of AI-enabled systems [29]. The need for explainability may be driven by factors such as transparency, user acceptance, and regulatory compliance [6, 9]. While particularly critical in high-risk scenarios, explainability is relevant across any system that incorporates intelligent components, supporting stakeholders in understanding, justifying, and auditing automated decisions [23].

The importance of explainability is further reinforced by international guidelines and regulatory initiatives. In Brazil, Proposed Law No. 2.338/2023 and the Brazilian Artificial Intelligence Plan (PBIA) explicitly require AI systems to be explainable, auditable, and transparent [20, 25]. These demands highlight the need to address explainability as a first-class concern during system development.

Despite this relevance, existing studies indicate that the incorporation of explainability in RE remains limited and often ad hoc [13]. There is a lack of systematic approaches to address explainability from early development stages, as well as uncertainties regarding what should be explained, when, and to whom [17].

To investigate this gap, we conducted a qualitative exploratory study with professionals leading AI projects in an applied research and innovation environment in Brazil, aiming to understand how they perceive explainability and the challenges they face when treating it as a software requirement. Data were collected through semi-structured interviews with 13 participants and analyzed using qualitative coding procedures based by Grounded Theory (GT), specifically open and axial coding.

The main contribution of this paper is an empirical characterization of how explainability is currently understood, elicited, specified, validated, and managed as a requirement in real AI-enabled projects, revealing concrete gaps in RE practice and deriving implications for the systematic integration of explainability into AI development processes. Our findings indicate that, although explainability is widely recognized as important, it is still addressed in an implicit and reactive manner rather than being systematically incorporated throughout RE activities. This highlights a gap between the recognized importance of explainability and its effective

operationalization in practice, reinforcing the need to integrate explainability from the early stages of RE.

The results show that explainability is perceived as a transversal quality attribute closely associated with trust and transparency. However, its treatment in practice is still fragmented, with challenges related to stakeholder understanding, cultural aspects, and the absence of structured RE support.

The remainder of this paper is organized as follows. Section 2 presents the background. Section 3 reviews the related work. Section 4 describes the research method. Section 5 presents the results and findings. Sections 6 and 7 discuss the findings and the threats to validity, respectively. Finally, Section 8 concludes the paper.

2 Background

Explainability as a requirement to be addressed by RE in software systems has been discussed in previous work. Chazette et al. [15] define a system as explainable with respect to an aspect X if a corpus of information (the explanation) is provided to an addressee A in such a way that A is able to understand X . From this perspective, explainability can be understood as a software quality attribute that enables stakeholders to comprehend specific aspects of a system's behavior or decisions. This quality has gained increasing relevance in the field of RE due to the growing complexity of modern software systems, particularly AI-enabled systems [5].

Research on explainability has focused on Explainable AI (XAI) [4, 27]. In this context, explanations typically concentrate on clarifying AI models and algorithms from a technical perspective. As a result, XAI explanations are predominantly developer-oriented and do not adequately address the information needs of other stakeholders, such as domain experts, clients, and end users. This limitation highlights the gap between model-level interpretability and stakeholders' understanding at the system level. Beyond AI-specific contexts, explainability has also gained prominence within RE [10, 19, 29].

Within RE, explainability is commonly regarded as a cross-cutting NFR closely related to trust, transparency, user understanding, and user acceptance. It is often characterized as a “means to an end,” rather than an end in itself [6, 15]. From this perspective, explainability should be considered user-centered, as explanations need to be tailored to different stakeholder profiles and their specific information needs. Consequently, a system may need to incorporate multiple types of explanations to effectively support diverse users throughout their interaction with the system.

Despite its recognized importance, previous studies have identified persistent gaps in the integration and maturity of explainability within the development of AI-enabled systems [17, 19]. In particular, the tertiary study by Martins et al. [19] synthesized the existing literature and identified research gaps that informed the research problem addressed in this work. Based on these findings, the present study adopts an empirical and qualitative approach to investigate how practitioners perceive and address explainability throughout RE activities in real AI projects. RE can support the systematic integration of explainability by extending the notion of explainability beyond model interpretability and systematically addressing stakeholders' understanding needs at the system level.

Regarding RE practices, traditional methods are still not well-defined for AI projects [9, 31]. Establishing effective RE practices in

this context is challenging, primarily due to the lack of professionals dedicated to formal RE activities [1] and the lack of techniques and tools tailored to data-driven projects, as research tends to focus on using AI to support RE rather than exploring how RE can enhance the development of AI systems [3]. It is therefore not surprising that recent studies report that practitioners consider RE to be the most complex phase in AI-enabled system projects [21].

In addition to technical challenges, the rapid growth of AI has prompted global regulatory initiatives such as the EU AI Act¹, OECD², and UNESCO ethical guidelines³, Brazil's Bill 2338/2023 [25], and the Brazilian Plan of AI (PBI) [20]. ISO/IEC 42001 complements these efforts by providing a management framework for responsible, auditable AI. These initiatives emphasize transparency and understandable explanations, particularly in high-risk systems, reinforcing the role of RE in operationalizing explainability as a critical requirement for compliance, ethics, and stakeholder trust.

3 Related Work

Explainability in AI-enabled systems has been investigated from multiple perspectives, including user needs, stakeholder alignment, and practitioner viewpoints. Previous work emphasizes tailoring explanations to different audiences.

To address communication gaps between technical teams and users, Hoffman et al. [12] introduced the Stakeholder Playbook, which categorizes explanation needs by audience and purpose to help developers align explanations with stakeholder goals. Similarly, Lakkaraju et al. [14] found that physicians and policymakers prefer interactive, natural-language explanations, reinforcing the need for context-aware explanation strategies.

From a practitioner perspective, Habiba et al. [30] report that AI professionals recognize explainability as essential for achieving transparency, improving models, and mitigating bias. However, their findings also indicate challenges in operationalizing explainability across different projects, particularly due to the absence of standardized practices.

In the context of RE and quality attributes, Habibullah et al. [10] interviewed experts from academia and industry to assess the state of practice regarding NFRs in AI-enabled systems, highlighting explainability as a central concern in RE. Similarly, Haindl et al. [11] analyzed quality attributes such as trust, explainability, and auditability, emphasizing their role in human-AI collaboration. Likewise, Vogelsang and Borg [32] identified an industry tension where data scientists prioritize model performance over stakeholder understanding, reinforcing the need to treat explainability as a first-class requirement.

Despite these advances, unclear expectations and divergent stakeholder perspectives persist. Stakeholder-centered approaches remain crucial to bridge gaps between model-generated explanations and human understanding [29], while current interpretability methods often fall short of user needs [2].

While previous studies have investigated explainability needs, stakeholder perspectives, and quality attributes, there is still limited empirical evidence on how explainability is actually incorporated

¹<https://artificialintelligenceact.eu/the-act/>

²<https://www.oecd.org/en/topics/ai-principles.html>

³<https://unesdoc.unesco.org/ark:/48223/pf0000381137>

into RE activities in real AI projects, especially in Brazilian industrial contexts. This work contributes by providing qualitative and exploratory insights into how explainability is perceived and addressed in practice, highlighting gaps in current RE activities and informing opportunities for a more systematic integration of explainability into AI-enabled systems.

4 Research Method

The following goal is formulated according to the Goal–Question–Metric (GQM) template [26].

Examine current RE practices in AI projects **for the purpose of** understanding **with respect to** practices, challenges, and practitioners’ perceptions of explainability as a software requirement **from the point of view of** professionals **in the context of** AI-enabled systems.

We investigate the following research questions:

- **RQ1)** What RE practices do professionals use to incorporate explainability in their AI projects? This question aims to identify how professionals use Requirements Engineering to address explainability in AI projects.
- **RQ2)** What challenges do practitioners face when defining and managing explainability requirements? This question is motivated by the need to understand the challenges practitioners face when defining and managing explainability as a software requirement in AI projects.
- **RQ3)** Which quality attributes can be achieved through explainability? Answering this research question will help us explore which quality attributes practitioners associate with explainability, reinforcing its role as a means to an end rather than an isolated requirement.
- **RQ4)** What trade-offs between explainability and other quality attributes need to be considered when integrating explainability into AI projects? The answer to this question helps examine how professionals perceive quality attributes that may conflict with explainability when integrating it into AI projects.

4.1 Data Collection

The study was conducted at the Center of Excellence in Artificial Intelligence (CEIA), affiliated with the Institute of Informatics (INF) at the Federal University of Goiás (UFG), Brazil. CEIA is a research, development, and innovation (RD&I) center that conducts applied AI projects in collaboration with industry and public-sector partners. This research was approved by the Research Ethics Committee of the Federal University of Goiás. All participants provided informed consent prior to the interviews.

The participants were project coordinators, researchers, and AI professionals involved in multiple projects across domains such as healthcare, manufacturing, energy, CRM, e-commerce, public administration, automotive, security, education, logistics and transportation, and compliance. Their experience provided a broad view of how explainability is discussed and implemented in real-world AI projects.

Table 1: Criteria for Considering Explainability as a Requirement

ID	Criteria	Context
1	Critical nature of the system	Systems that make critical decisions for individuals, such as medical diagnosis, credit approval, social benefits authorization, or autonomous vehicles (without human intervention).
2	Business opportunity associated – not necessarily critical	Explainability is a need or expectation of a specific stakeholder.
3	Data bias potential	The data may contain biased patterns, replicating recurring unethical behaviors.
4	Need for compliance with legal requirements related to privacy	The project must comply with legal and regulatory requirements that demand explanations for automated decisions, such as Brazil’s LGPD – Art. 20.
5	Requirement for expert validation	Models that require prior validation to perform specific actions (e.g., medical or legal decision support).
6	Auditability requirement – not necessarily critical	Models that need to be auditable to identify and justify their decisions, such as in internal investigations or bias analysis.

Participants were selected based on publicly available data from the RD&I unit, targeting individuals directly involved in leading AI projects. We considered only projects initiated from 2022 onward to ensure the inclusion of recent initiatives, reflecting current practices and challenges related to explainability in AI-enabled systems.

The dataset comprised 85 AI projects, including project titles, coordinators’ names, email addresses, and short descriptions. Additionally, two researchers with experience in RE established the following criteria to identify projects with a potential demand for explainability, as presented in Table 1.

All 85 projects met at least one of these criteria. Therefore, invitations were sent via email to 16 researchers responsible for these projects, of whom 13 agreed to participate in the study. It is worth noting that each researcher could be responsible for one or more projects. The participant selection followed a non-random, purposeful sampling strategy targeting projects developed by the RD&I unit that were likely to require explainability.

The invitation email included a concise description of the interview’s objectives and listed the projects under the respondent’s coordination. This procedure allowed each coordinator to select the project with which he/she was most familiar to discuss aspects related to explainability, resulting in 13 interviews corresponding to 13 distinct projects, and ensuring a contextualized and specific discussion during the interview. To accommodate participants’ availability, we offered two participation modes: semi-structured interviews (remote/synchronous) or semi-structured interviews (in-person). All participants answered the same set of questions.

Participation was voluntary and non-remunerated. All interviews were conducted remotely, recorded via a web conferencing platform, and automatically transcribed using online transcription tools. Participants were informed in advance and explicitly authorized the recording and transcription of their interviews. All Personally Identifiable Information (PII) was anonymized prior to analysis.

4.2 Interview Design

The semi-structured interview contained six main sections:

- (1) *Informed consent request.* It describes the purpose of the study, participant responsibilities, confidentiality rules, data

protection measures, the voluntary nature of participation, and the estimated duration of the interview.

- (2) *Preliminary concepts knowledge.* We clarified the concepts related to explainability from the RE perspective. By providing a definition, we aimed to establish a common understanding to support the interview responses and to contextualize explainability within each participant’s project.
- (3) *Participants’ profile.* To validate the experience of the interviewees regarding both the development of projects using AI and the practice of RE.
- (4) *Characterization of AI projects.* We analyzed the project’s application domain, the AI tasks involved, and the presence of RE experts in the project team.
- (5) *RE practice for explainability.* We investigated how projects considered the explainability requirements and how it was addressed from the RE standpoint. To guide this investigation, we used SWEBOK [22] as a reference, which defines the main RE activities: elicitation, analysis, specification, validation, and management. For each of these activities, we also asked how processes related to explainability were conducted throughout the project.
- (6) *Challenges and needs related to explainability.* We asked participants about the main difficulties in implementing explainability, as well as which quality attributes can be achieved through it and the associated trade-offs.

All the questions and response options rely on a previous work [17] in which we identify current gaps and RE practices on explainability in AI-enabled systems.

The interviews were conducted individually. Most of the questions were open-ended, and not all questions were answered; if the participant had not implemented explainability in the project, specific RE-related questions were not applicable. This approach enabled the collection of rich qualitative insights, as participants elaborated on the practices they typically implement, the challenges encountered, and their overall perspective on explainability. The interview data are available in the online repository [18].

4.3 Interview Study Protocol

We conducted an initial pilot interview with one participant to validate the interview guide and ensure question clarity. The pilot participant had a strong academic and professional background in AI and Software Engineering, with experience in applied research and familiarity with RE practices. As no refinements were necessary, we proceeded directly to the data collection phase. Individual semi-structured interviews were then conducted between February and April 2025 to address the research questions. Each session lasted approximately 50–60 minutes, totaling around 650 minutes of recordings.

The interviews were conducted remotely using the *Google Meet* platform⁴. Recordings were performed with *OBS Studio*⁵, while automatic transcription was supported by *Read.ai*⁶ and *Tactiq.io*⁷. After each interview, transcriptions were manually reviewed to

correct potential errors, and all Personally Identifiable Information (PII) was anonymized prior to the analysis.

4.4 Data Analysis

The data from the semi-structured interviews were analyzed through a qualitative approach complemented by descriptive quantitative summaries [16]. The qualitative analysis aimed to examine the 13 interviews conducted in order to answer the predefined research questions within the GQM paradigm.

For qualitative analysis, we used coding procedures based on GT, specifically open and axial coding, as a systematic strategy to organize and analyze the interview data. Our objective was not to apply GT in its entirety or to develop an emergent theory. Nevertheless, GT coding procedures can also be applied to gain a deeper understanding of a phenomenon while allowing prior research questions and existing literature to inform the analysis [28]. The coding process followed a predominantly inductive approach based on the interview data. A set of preliminary assumptions derived from the literature was used exclusively to support the study design and the development of the interview protocol. These assumptions were not conceived as testable hypotheses and were not used to guide the coding process, generate analytical categories, or test causal relationships. Additional details on the interview protocol are available in the supplementary material [18].

All analyses were conducted using the *Atlas.ti* software⁸. One author performed open coding of the transcripts for each quotation, while another author manually reviewed the generated codes. Both authors collaboratively discussed and decided which codes should be retained or excluded. The analysis comprised a total of 185 quotations.

The open coding stage involved a detailed reading of the transcripts to identify relevant information and associate codes with interview quotations. These codes represented phenomena of interest observed in the data. Next, axial coding was carried out to group the identified codes according to their properties, forming broader categories. All details regarding the coding procedure and the generated codes are available in an online repository [18].

To ensure adequate coverage of the topics investigated, the interview protocol was designed based on findings from the literature and on RE activities described in the SWEBOK framework [22]. This helped ensure that the interview questions systematically covered the main RE activities, namely elicitation, analysis, specification, validation, and management. In addition, the protocol included questions addressing practitioners’ perceptions of the importance of explainability and its potential trade-offs with other quality attributes, such as performance and security.

5 Results and Findings

This section presents the results of the semi-structured interview study organized according to the research questions (RQs). We begin by presenting an overview of the participants’ profile and the AI projects included in the study. The subsequent subsections present the findings associated with each research question.

To complement the qualitative findings, we also report descriptive quantitative summaries obtained from closed-ended questions.

⁴<https://meet.google.com>

⁵<https://obsproject.com/pt-br>

⁶<https://www.read.ai/pt>

⁷<https://tactiq.io/pt-br>

⁸<https://atlasti.com/>

Table 2: Characterization of AI Projects Coordinated by Interviewees

ID	RE Exp. (yrs)	Application Domain	AI Tasks	Analyst
P01	0	Manufacturing	Predictive analysis, classification, regression	No
P02	<1	Healthcare	Anomaly detection, computer vision	No
P03	1–5	Energy, Manufacturing	Prescriptive analysis, recommendation	No
P04	1–5	CRM, e-commerce	Classification, recommendation	No
P05	1–5	Manufacturing	Optimization and manufacturing processes	Yes
P06	6–10	Public administration	Predictive analysis, NLP, image recognition	Yes
P07	>10	Healthcare	Discovery, descriptive, prescriptive analysis	Yes
P08	>10	Healthcare	Predictive analysis, classification	Yes
P09	>10	Automotive	Predictive analysis, anomaly detection, recommendation	Yes
P10	>10	Compliance	NLP, recommendation, image recognition	Yes
P11	>10	Automotive, Security	Computer vision	No
P12	>10	Education, Healthcare	Classification, NLP, recommendation	No
P13	-	Logistics and Transportation	Robotic navigation, route optimization, computer vision	No

5.1 Participants’ and AI Projects’ Profile

The participants’ experience varied across different levels of familiarity with RE practices. Among the 13 interviewees, 11 reported previous experience with RE practices, one participant reported no prior experience in RE, and one participant did not answer this question. All interviewees had at least two years of experience in AI innovation projects.

Although most participants had previous experience with RE, not all projects involved a dedicated requirements analyst (see Table 2). Approximately half of the projects did not include professionals in this role, revealing a practical gap and emphasizing the need for systematic mechanisms to handle explainability as a requirement.

We observed a significant diversity of application domains, with a particular emphasis on critical areas such as healthcare and safety. These contexts demand a higher degree of trust and transparency, attributes directly associated with explainability. Finally, multiple AI tasks were identified in the 13 projects including predictive analysis, classification, recommendation, and computer vision. This suggests the need for a systematic and standardized process to effectively deal with explainability across different AI tasks. Table 2 summarizes these findings.

5.2 RE Practices and Explainability (RQ1)

We investigated RE activities and the concern with explainability requirements in the analyzed projects. Our analysis showed that six projects explicitly addressed explainability, and the following subsections report the findings according to the RE activities.

Figure 1 illustrates the relationship between the presence of an RE analyst, the level of concern for explainability, and the factors influencing these decisions. Specifically, it shows how projects are distributed across three categories of concern (explicit, implicit, and none), and how these categories are associated with motivations such as compliance requirements, client requests, lack of knowledge, and resource constraints.

Overall, Figure 1 reveals that explicit concern for explainability is more frequently associated with the presence of an RE analyst, whereas projects without this role tend to exhibit implicit or no concern. In cases of implicit concern (e.g., P01 and P02), explainability was not initially treated as a structured requirement, but rather emerged during project execution as a need to justify system decisions.

This behavior is reflected in participants’ accounts. P01 described a scenario in which explanations were required to support reliability estimation, but were not formally treated as requirements: *“We needed to estimate reliability, which is essentially a probability of failure. To do this, we used statistical models, and later AI models to predict future behavior. However, this was treated as part of the modeling process, not as an explicit requirement.”* Similarly, P02 highlighted that explainability was considered implicitly during decision-making: *“I believe we address explainability implicitly. Many of the decisions we make depend on the results we obtain, so there is a concern with explaining them, but without explicitly treating it as a requirement.”*

In contrast, participants who reported no concern for explainability indicated that, in their view, because the project was an AI-enabled system, limited attention was given to requirements in general. This perspective is illustrated by P03, who stated: *“Since our main goal is to deliver AI solutions, we pay little attention to requirements.”* This statement suggests a conceptual framing in which AI systems are treated primarily as model-centric solutions rather than as software systems, which may contribute to the limited consideration of RE practices, including those related to explainability.

For projects with explicit concern (e.g., P09 and P10), the primary drivers were compliance requirements and explicit client requests. These findings suggest that the presence of an RE analyst contributes to a more systematic treatment of explainability, while its absence is associated with weaker or more reactive consideration of this requirement.

This interpretation is illustrated by participants’ accounts. P09 emphasized that explainability requirements were directly driven by stakeholder demands: *“It was a requirement from the client, who needed to meet their specific needs based on the data we were providing.”* Similarly, P10 associated the need for explainability with compliance-related constraints, highlighting that system decisions must align with legal requirements: *“It is related to compliance, as decisions need to be verified according to legal criteria.”*

These observations are consistent with P06, suggesting that although explainability is recognized as important, it may not be systematically prioritized and is often addressed reactively. In some cases, explainability requirements emerge implicitly throughout project execution as a need to justify system decisions.

Table 3 summarizes the projects in which explainability was considered, highlights each project’s application domain, the main reasons for addressing explainability, and the RE techniques employed. The results are organized according to RE activities [22].

Elicitation. The projects that addressed explainability conducted the elicitation phase using techniques such as interviews, document analysis, prototyping, observation, brainstorming, scenarios, reverse engineering, mind maps, and questionnaires, with sources including clients, end users, technical documents, artifact reuse,

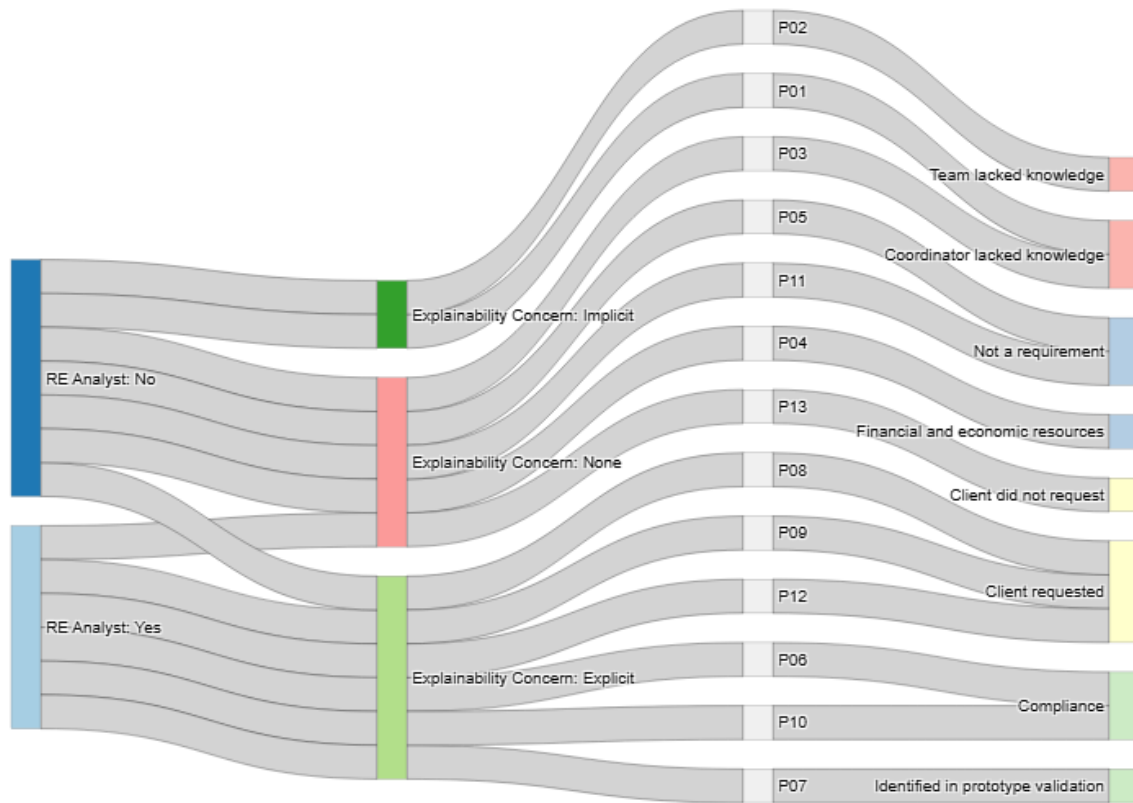


Figure 1: AI projects requirements analyst and explainability concern (and related reasons).

Table 3: Overview of Projects and RE Practices Related to Explainability

ID	Application Domain	Reason for Explainability	Elicitation Technique	Elicitation Source	Analysis Technique	Specification Format	Validation Technique
P06	Compliance	Compliance requirement	Document analysis, interviews, laws and regulations	Client, technical documentation, end user	Prioritization	NL document (justification)	-
P07	Healthcare	Identified during prototype validation with the client	Prototyping	Client	-	-	Prototype
P08	Healthcare	Client request	Document analysis, interviews, prototyping	Client, artifact reuse	Use case modeling; clustering technique	User stories	Prototype
P09	Automotive	Client request	Scenarios, interviews, observation	Scientific papers, technical documentation	-	User interface mockup	Prototype
P10	Compliance	Compliance requirement	Document analysis, laws and regulations	Technical documentation	-	Initial project stage	-
P12	Education, Healthcare	Client request	Brainstorming, scenarios, reverse engineering, interviews, mind maps, observation, prototyping, questionnaires	Scientific papers, technical documentation	-	-	-

and scientific literature. Projects with implicit or no focus on explainability did not report formal techniques.

Participants' accounts illustrate the diversity of elicitation practices adopted in projects with explicit concern for explainability. P09 reported the use of scenarios, observation, and interviews to explore how explainability could be better addressed in the proposed solution: "Scenarios were presented, and based on these scenarios we

discussed how the solution could be better proposed. We also used observation and interviews." Similarly, P06 emphasized the importance of interviews combined with technical and regulatory analysis: "The elicitation was conducted through interviews. There is also a certain understanding of the problem context involved. I would highlight interviews together with the analysis of technical documentation, laws, and regulations."

In particular, P12 reported the use of multiple elicitation techniques and emphasized that the need for explainability was explicitly driven by client demands, especially in sensitive domains such as healthcare. As noted by P12: *“There must be reasons behind the outcomes, because sometimes AI can provide results that are even more relevant than those of specialists. This is a major concern in healthcare, especially in mental health, where knowledge is still limited. We do not fully understand the brain, and many aspects remain unresolved even at the state of the art. Therefore, explainability becomes a fundamental requirement.”* This statement illustrates how domain complexity and uncertainty reinforce the importance of explainability as a key requirement during elicitation.

These observations indicate that elicitation practices related to explainability are present in some projects, particularly when external demands, domain complexity, or regulatory concerns require greater justification of system decisions.

Analysis. Among the six projects with explicit attention to explainability, four did not apply formal analysis techniques; the remaining employed clustering, prioritization, or use-case modeling methods. All coordinators were asked about potential alignments or conflicts between explainability requirements and other quality attributes. Attributes most frequently perceived as associated were trust, transparency, ethics, traceability, performance, and usability, while conflicts were mainly reported with performance and privacy.

The lack of structured analysis is reflected in participants’ accounts. P08 reported that explainability was modeled as a conventional requirement: *“We tried to model the explainability requirement as a regular requirement, which resulted in an additional use case related to predicting certain outcomes.”*

However, as reported by P08, the requirement was not further detailed, making it unclear whether such practices were sufficient to adequately capture explainability needs.

Additionally, participants’ accounts further reinforce the relationship between explainability and quality attributes. P07 emphasized that explainability is closely related to NFR, particularly usability, highlighting its contribution to user understanding and satisfaction: *“Explainability is strongly related to quality attributes, especially usability, as it helps ensure that the system is easy to understand and supports users in locating and interpreting information, ultimately contributing to overall product quality and user satisfaction.”* Similarly, P13 noted that explainability tends to complement other requirements rather than conflict with them: *“I believe it complements other requirements well; it does not interfere with them, as this is a trend we have observed.”*

These observations indicate that explainability is primarily perceived as complementary to other quality attributes, although trade-offs may arise in specific contexts, particularly regarding performance and privacy.

Participants’ accounts illustrate these trade-offs. P10 highlighted concerns related to privacy, emphasizing the need to make system behavior transparent without exposing sensitive information: *“You need to make clear how the system works without exposing sensitive data, so privacy becomes a concern.”* In addition, P06 pointed out a potential tension with performance, noting that explainability may affect model effectiveness or perception: *“It is not enough for the*

model to be good; it also needs to appear reliable. This may create a trade-off with model performance.”

One participant (P11) highlighted a conceptual conflict related to the nature of AI models: *“These are essentially non-deterministic algorithms that learn patterns approximately, which may hinder the generation of precise and complete explanations.”*

Overall, these findings suggest that, although trade-offs related to explainability are recognized, they are not systematically addressed in current RE practices. Furthermore, the limited adoption of formal analysis techniques indicates that explainability is rarely analyzed in a structured manner.

Specification. The documentation of explainability requirements was limited: only three projects explicitly specified them through user stories, user interface screens, or natural language (NL) descriptions used to justify the adoption of specific algorithms. The remaining projects did not formalize explainability requirements.

P06 reported documenting the rationale behind model choices: *“We documented why we chose to initially test tree-based models to achieve greater specificity. This was described in a document of fewer than 20 pages explaining this decision.”* This reflects an effort to operationalize explainability through the justification of technical choices, although not necessarily formalized as structured explainability requirements.

P08 indicated that explainability was treated at the same level as other requirements and specified using user stories: *“We treated this requirement at the same level as others. For example, we used a user story such as: ‘As a medical professional, I would like the model to provide explanations for its predictions.’”*

This suggests that established RE practices may be used to support the specification of explainability requirements. However, the level of detail remains relatively limited, and there is limited evidence to assess whether the structure of these practices is sufficient to adequately capture explainability needs. These observations suggest that the specification of explainability requirements remains limited, particularly due to the absence of evidence regarding the adequacy and validation of the adopted approaches.

Validation. Three projects reported validation activities, where prototypes were used as a means to validate explainability requirements through stakeholder interactions. Participants described validation as an iterative process involving feedback and refinement. For instance, P07 noted that validation played a key role in structuring explainability requirements: *“We first identified the requirements, then prototyped and validated them. It was during validation that the explainability requirement became more structured.”* Similarly, P09 highlighted the importance of stakeholder feedback: *“We presented the solution and gathered their feedback.”*

Coordinators were also asked which stakeholders should be involved in explainability validation. Domain experts were the most frequently mentioned (nine quotes), followed by business owners and business analysts (six each), AI engineers and end users (three each), requirements engineers (two quotes), and data scientists and project coordinators (one each).

Participants emphasized the importance of involving multiple roles to ensure alignment between technical solutions and stakeholder needs. For example, P07 highlighted the relevance of domain experts, engineers, and analysts to connect system behavior with

real-world demands: “Domain experts, engineers, and analysts are essential to ensure that the solution aligns with the needs and allows revisiting models or decision criteria.” Similarly, P04 stressed the role of domain experts, end users, and requirements engineers in bridging technical and business perspectives: “Domain experts, end users, and requirements engineers are important, as they help connect the solution with business needs.” These findings suggest that validation activities are present but limited.

Management. No project with explicit concern for explainability implemented formal management activities throughout the project lifecycle. This finding suggests that, without structured elicitation, analysis, and specification steps, managing explainability becomes unfeasible or is not perceived as necessary. The absence of formal management practices reinforces this finding.

5.3 What challenges do they face in this process? (RQ2)

Despite its perceived benefits, participants reported several challenges that hinder the effective incorporation of explainability into AI projects. One of the challenges reported by participants was the lack of clear explainability objectives, reflecting uncertainty about what should be delivered as explainability and how it should be operationalized. As noted by P07: “I think there is also a lack of understanding about what should actually be delivered as explainable. In fact, it is not only a lack of understanding, but also difficulty in recognizing the need itself.”

This lack of clarity was often compounded by practical constraints, including tight project deadlines, team-related issues, budget limitations, and limitations in existing RE processes. P04 summarized these constraints as a “triangle” involving “deadline, team, and cost.” Similarly, P06 emphasized the challenges associated with financial resources, infrastructure, and the absence of structured RE processes to support explainability implementation.

In addition, many of the reported difficulties are closely related to early RE activities, such as elicitation, analysis, and negotiation of requirements. Participants highlighted the importance of RE processes in improving the quality of explainability-related outcomes. As emphasized by P07: “The Requirements Engineering process is extremely important. One of the difficulties is that when you do not apply it, it becomes difficult to achieve quality in the delivered solution.” Treating explainability as an explicit concern of RE can help mitigate these challenges by clarifying stakeholder needs, aligning expectations, and systematically managing trade-offs throughout the project lifecycle. Figure 2-B presents the main factors identified by practitioners as difficulties in its implementation.

5.4 Which quality attributes can be achieved through explainability? (RQ3)

Based on data from 13 interviewees, the quality attributes most frequently associated with explainability were trust and transparency, followed by comprehensibility and user satisfaction. As noted by P09: “I think that, in a way, it contributes a little to each one [quality attributes], you know? And in the end, it results in satisfaction. You convince the user. Comprehensibility comes from having more clarity. Security is about being able to trust what is being done. Transparency

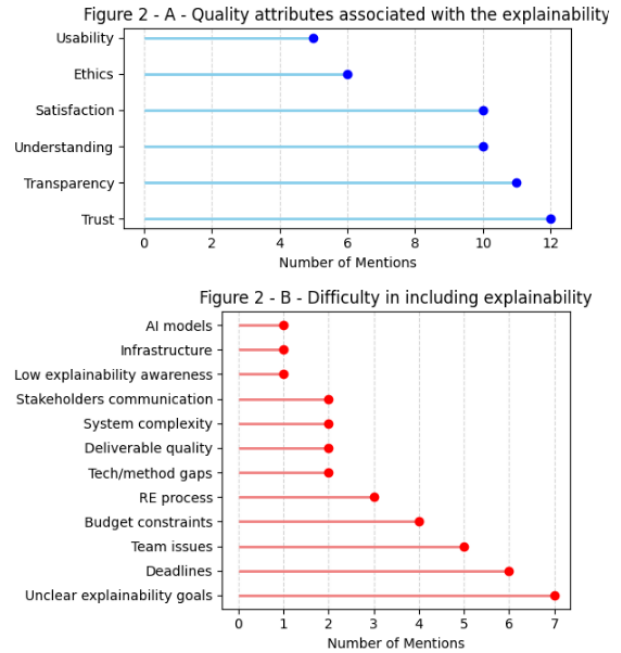


Figure 2: Explainability-related quality attributes and challenges in AI projects.

is understanding what is being done. Comprehensibility is closely related to satisfaction.” Participants also associated explainability with security, ethical compliance, usability, and AI acceptance. These attributes reinforce the view of explainability as a cross-cutting quality that primarily serves as a means of achieving user-centered goals, such as increasing users’ confidence in the system and supporting their understanding of system behavior and decisions. Rather than being perceived as an end in itself, explainability is valued for its contribution to broader quality concerns that directly affect user experience and acceptance. Figure 2-A illustrates the quality attributes associated with explainability.

5.5 What trade-offs between explainability and other quality attributes need to be considered when integrating explainability into AI projects? (RQ4)

Although participants did not report direct conflicts between explainability and other quality attributes, they highlighted concerns regarding unnecessary data exposure, particularly in relation to security. Furthermore, when asked whether explainability is a cultural issue, nine out of the thirteen interviewees agreed, framing explainability as a socio-technical concern rather than a purely technical one. From a GT perspective, this perception was associated with the codes stakeholder awareness, domain influence, and expectations about AI capabilities. P13 explicitly described explainability as a cultural issue, highlighting unrealistic expectations regarding AI capabilities: “People believe that AI will solve every problem in their lives and knows everything about everything, but that is not exactly what AI does.”

Similarly, P10 associated this cultural issue with the lack of formalization of explainability-related requirements in Software Engineering practices: *“I believe this is still a cultural issue, but technically there is also a lack of formalization. Today, explainability and transparency are not properly classified as NFRs. We are still using traditional software documentation approaches to document intelligent systems, even though these systems involve different requirements.”*

These observations reinforce the need to evolve RE practices to better support AI-enabled systems. They also highlight the role of RE in making culturally influenced and implicit concerns about explainability visible and systematically addressed as software requirements, emphasizing the need to address explainability beyond purely technical solutions.

6 Discussion

Our findings suggest that explainability is recognized as relevant. However, it remains weakly institutionalized in RE practice. In some projects, it is triggered by external demands (clients’ requests or compliance activities) rather than being addressed proactively as part of RE activities. This indicates a gap between the conceptual importance of explainability in the literature and its proper operationalization in real AI projects. In this section, we discuss the results according to the RQ.

RQ1 - What RE practices do professionals use to incorporate explainability in their AI projects?

Our findings indicate that practitioners employ a variety of elicitation techniques, such as interviews, document analysis, prototyping, scenarios, and observation, to elicit explainability-related requirements. However, these practices are not consistently followed by RE analysis, specification, validation, and management activities. This suggests that explainability is considered during elicitation but is rarely refined, documented, validated, and managed throughout the RE lifecycle. As a consequence, explainability tends to remain an implicit consideration rather than a systematically documented and maintained software requirement. This lack of continuity across RE activities may hinder the refinement, documentation, traceability, validation, and long-term evolution of explainability requirements, making their treatment dependent on project-specific decisions rather than on a systematic RE process. These findings reinforce the need to integrate explainability across the different RE activities from the early stages of project conception. They also highlight opportunities for RE approaches that provide systematic support for eliciting, analyzing, specifying, validating, and managing explainability requirements throughout the software development lifecycle.

RQ2 - What challenges do they face in this process? Our findings indicate that practitioners mainly struggle to understand what should be delivered as explainability, revealing both conceptual and practical barriers. A key challenge is the lack of clarity regarding the expected explainability outcomes, suggesting that explainability is still not sufficiently defined or operationalized within RE practices. This lack of conceptual clarity is further reinforced by limitations in the current state of the art, indicating that existing tools and methodologies do not yet provide sufficient support to address explainability demands. In addition, communication with stakeholders emerged as another relevant challenge, which may

be aggravated by the absence of RE processes that explicitly incorporate explainability from the early stages of development. These findings suggest that addressing explainability requires more than technical solutions. It also requires structured RE practices that provide practical guidance for incorporating explainability throughout the RE process, fostering a shared understanding among development teams, users, and stakeholders.

RQ3 - Which quality attributes can be achieved through explainability? Our findings indicate that explainability is perceived as an enabler of multiple quality attributes rather than as an isolated characteristic of AI systems. Participants associated explainability with increased trust, transparency, user satisfaction, and understanding, while also highlighting its contribution to usability, security, ethical compliance, and AI acceptance. These findings suggest that the value of explainability extends beyond making AI models interpretable. Instead, explainability contributes to improving how users understand, interact with, and trust AI systems, demonstrating its influence on both technical and human-centered quality concerns. Taken together, these results reinforce the view of explainability as a user-centered software requirement whose benefits extend across multiple quality attributes. This highlights the importance of considering explainability from the early stages of RE to ensure that these quality attributes are systematically addressed throughout the development process.

RQ4 - What trade-offs between explainability and other quality attributes need to be considered when integrating explainability into AI projects? The results show that participants did not report direct conflicts between explainability and other quality attributes, such as performance and security. However, they indicated that explainability should consider risks related to unnecessary data exposure, particularly in relation to security concerns. These findings suggest that the main challenge is not the existence of direct trade-offs between explainability and other quality attributes, but rather ensuring that explainability is implemented without compromising sensitive information. Beyond technical considerations, participants’ accounts also suggest that organizational and cultural aspects constitute relevant barrier to the effective integration of explainability in AI projects. In this sense, culture does not emerge as a traditional quality trade-off, but rather as a socio-technical constraint that influences how explainability is interpreted, valued, and operationalized alongside other quality attributes. Taken together, these findings reinforce that integrating explainability into AI projects requires considering both technical and organizational aspects, highlighting the importance of incorporating explainability throughout the RE process rather than addressing it only during implementation.

6.1 Practical Implications and Recommendations

Based on the previous discussion, we summarize practical implications and recommendations for explainability-enabled AI projects:

Explainability from the project’s early stages: It should be treated as an NFR from the design stage and planned alongside attributes such as trust, transparency, satisfaction, and usability.

Integrate explainability systematically into RE phases: Explainability is often limited to elicitation, without analysis or

documentation. It is recommended to adopt a complete RE process—including analysis, specification, and validation—to ensure traceability and maintenance throughout the system lifecycle.

Conceptual clarity and practical guidelines: The unclear definition of explainability underscores the need for guidelines, examples, and standards that translate the concept into verifiable and implementable requirements.

External factors and business value: The inclusion of explainability is often motivated by customer or compliance requirements. Organizations should also recognize it as a competitive advantage and value-added attribute.

Stakeholder communication: Explainability should be explicitly approached as a user-centered concern, aligned with users' and stakeholders' profiles, knowledge levels, and expectations. Explanations must be intelligible, relevant, and meaningful to their intended audience, rather than generic or purely technical. Engaging stakeholders early supports a shared understanding of explainability goals and ensures that explanations effectively address users' needs and contexts of use.

Quality and security: When assessing explainability, data protection and privacy must be ensured so that transparency does not compromise security or expose sensitive information.

7 Threats to Validity

We discuss the threats to validity and limitations of our study [24].

Construct validity: To strengthen construct validity, the interview protocol was grounded in RE activities defined in SWEBOK [22], ensuring alignment between the interview questions, the investigated constructs, and the research objectives. The interview guide was also reviewed by the research team prior to data collection. Nevertheless, differences in participants' interpretation of explainability may have influenced the responses.

Internal validity: To mitigate this threat, interviews were conducted with experienced AI professionals, and participants were encouraged to report their actual practices regarding explainability requirements. We also minimized social desirability bias [8] by emphasizing that there were no right or wrong answers. However, participants' perceptions and organizational context may still have influenced the findings.

External validity: A notable threat is the limited transferability of the results, as the study was conducted within a single RD&I organization. Although we sought to mitigate this limitation by including participants with different professional roles and experience, as well as projects from diverse application domains (Manufacturing, Healthcare, Energy, CRM, E-commerce, Public Administration, Automotive and Security, Education, Logistics and Transportation, and Compliance), the findings may not transfer to organizations with different levels of maturity, domains, or development practices. Therefore, the results should be interpreted as context-dependent.

Reliability: To improve reliability, data analysis followed systematic open and axial coding procedures. Open coding was performed by one author and reviewed by another, while axial coding was performed collaboratively. Weekly meetings and brainstorming sessions were held to discuss coding decisions and consolidate the findings. Although these measures reduced researcher bias, some degree of interpretation remains inherent to qualitative analysis.

8 Conclusions and Future Work

In this study, we conducted semi-structured interviews with 13 AI professionals to investigate how they perceive and address explainability in their projects. The main contribution of this study is an empirical characterization of how explainability is addressed throughout RE activities in real AI-enabled projects. The results indicate that, although explainability is widely recognized as important, its treatment as a requirement remains limited. In projects that consider explainability, elicitation is more consistently performed, while other RE activities (e.g., analysis and specification) are often overlooked, even when a requirements analyst is involved. This reveals a gap between the recognized importance of explainability and its effective operationalization throughout the RE process. Implementing explainability also faces several challenges, including a lack of understanding of what should be delivered, limitations of existing methods and tools, and communication issues among stakeholders.

Despite these challenges, the findings suggest that explainability should be systematically considered from the early stages of AI-enabled system development. Treating it only after model implementation may limit its potential to support design decisions and stakeholder communication. As a cross-cutting quality concern closely related to trust and transparency, explainability should also be balanced with data protection and privacy requirements.

These findings, together with prior studies [17, 19], motivated the development of an RE framework to systematically incorporate explainability throughout RE activities. The current version supports key activities from the Inception and Elaboration phases, including the definition of the explainability vision, elicitation, and specification through user stories. Preliminary validation using illustrative scenarios provides initial evidence of its feasibility and practical applicability. As future work, the framework will be extended to additional RE phases and evaluated through real-world case studies to assess its applicability and effectiveness in supporting explainability in AI-enabled systems.

Artifact Availability

The supplementary materials supporting this study are publicly available at <https://doi.org/10.5281/zenodo.21415990>. The repository contains the interview protocol, study design, GQM model, project descriptions, qualitative codes, complementary results, and the accepted paper.

Acknowledgments

The authors thank the Center of Excellence in Artificial Intelligence (CEIA) and the Institute of Informatics at the Federal University of Goiás (INF/UFG) for supporting this study and facilitating access to the professionals who participated in the interviews. This study was financed in part by the CAPES (Finance Code 001). The third author would like to thank FAPEMA, CNPq (Grants 311533/2025-6 and 448761/2025-4), and FAPESP (IARA 20/09835-1 and EcoSustain 23/00811-0). The fourth author thanks CNPq (Grants 312275/2023-4 and 420191/2025-9), FAPERJ (Grant E-26/204.256/2024), and Instituto Kunumi.

References

- [1] Antonio Pedro Santos Alves et al. 2023. Status Quo and Problems of Requirements Engineering for Machine Learning: Results from an International Survey. In *Product-Focused Software Process Improvement: 24th International Conference, PROFES 2023*. Springer-Verlag, 159–174. doi:10.1007/978-3-031-49266-2_11
- [2] Diogo V Carvalho et al. 2019. Machine learning interpretability: A survey on methods and metrics. *Electronics* 8, 8 (2019), 832. doi:10.3390/electronics8080832
- [3] Fabiano Dalpiaz and Nan Niu. 2020. Requirements Engineering in the Days of Artificial Intelligence. *IEEE Software* 37, 4 (2020), 7–10. doi:10.1109/MS.2020.2986047
- [4] Arun Das and Paul Rad. 2020. Opportunities and challenges in explainable artificial intelligence (xai): A survey. *arXiv preprint arXiv:2006.11371* (2020). doi:10.48550/arXiv.2006.11371
- [5] Hannah Deters et al. 2024. How explainable is your system? towards a quality model for explainability. In *International Working Conference on Requirements Engineering: Foundation for Software Quality*. Springer, 3–19. doi:10.1007/978-3-031-57327-9_1
- [6] Hannah Deters et al. 2025. Exploring the means to measure explainability: Metrics, heuristics and questionnaires. *Information and Software Technology* (2025), 107682. doi:10.1016/j.infsof.2025.107682
- [7] Rudresh Dwivedi et al. 2023. Explainable AI (XAI): Core ideas, techniques, and solutions. *ACM computing surveys* 55, 9 (2023), 1–33.
- [8] Adrian Furnham. 1986. Response bias, social desirability and dissimulation. *Personality and individual differences* 7, 3 (1986), 385–400. doi:10.1016/0191-8869(86)90014-0
- [9] Gökem Giray. 2021. A software engineering perspective on engineering machine learning systems: State of the art and challenges. *Journal of Systems and Software* 180 (2021), 111031. doi:10.1016/j.jss.2021.111031
- [10] Khan Mohammad Habibullah et al. 2023. Non-functional requirements for machine learning: Understanding current use and challenges among practitioners. *Requirements Engineering* 28, 2 (2023), 283–316.
- [11] Philipp Haindl et al. 2022. Quality characteristics of a software platform for human-ai teaming in smart manufacturing. In *International Conference on the Quality of Information and Communications Technology*. Springer, 3–17.
- [12] Robert R Hoffman et al. 2023. Explainable AI: roles and stakeholders, desirments and challenges. *Frontiers in Computer Science* 5 (2023), 1117848. doi:10.3389/fcomp.2023.1117848
- [13] Faiza Khan Khattak et al. 2024. MLHops: Machine learning health operations. *IEEE Access* (2024). doi:10.1109/ACCESS.2024.3521279
- [14] Himabindu Lakkaraju et al. [n.d.]. Rethinking explainability as a dialogue: a practitioner’s perspective (2022). doi:10.48550/arXiv.2202.01875
- [15] Wasja Brunotte Larissa Chazette and Timo Speith. 2021. Exploring explainability: a definition, a model, and a knowledge catalogue. In *2021 IEEE 29th international requirements engineering conference (RE)*. IEEE, 197–208. doi:10.1109/RE51729.2021.00025
- [16] Paulo Malcher et al. 2023. Investigating open innovation practices to support requirements management in software ecosystems. In *International Conference on Software Business*. Springer Nature Switzerland Cham, 35–50.
- [17] Livia Mancine et al. 2024. Estado da Arte sobre Engenharia de Requisitos e Explicabilidade em Sistemas Baseados em Aprendizado de Máquina. In *Extended Proceedings of the 30th Brazilian Symposium on Multimedia and Web Systems*. 143–158. doi:10.5753/webmedia_estendido.2024.243944 (in Portuguese).
- [18] Livia Mancine et al. 2025. Explainability Requirement Practices and Challenges from a Brazilian Industrial Research and Innovation Company. Zenodo. doi:10.5281/zenodo.21415990
- [19] Mariana Crisostomo Martins et al. 2025. Research Synthesis in Requirements Engineering for Machine Learning-Based AI Systems: A Tertiary Study. *JSERD* 13, 2 (2025), 13–129. doi:10.5753/jsrd.2025.4892
- [20] Ministério da Ciência, Tecnologia e Inovação (MCTI) and Centro de Gestão e Estudos Estratégicos (CGEE). 2025. *IA para o Bem de Todos: Plano Brasileiro de Inteligência Artificial*. Technical Report. MCTI / CGEE, Brasília, DF. Retrieved Jun 01, 2025 from https://www.cgee.org.br/documents/10195/11009772/CGEE_PBIA.PDF
- [21] Nadia Nahar et al. 2022. Collaboration challenges in building ML-enabled systems: communication, documentation, engineering, and process. In *Proceedings of the 44th International Conference on Software Engineering (ICSE ’22)*. 413–425. doi:10.1145/3510003.3510209
- [22] Joanna Isabelle Olszewska. 2024. *Guide to the Software Engineering Body of Knowledge v4.0*. Vol. 4.0. IEEE Computer Society. Retrieved Oct 10, 2025 from <https://www.computer.org/education/bodies-of-knowledge/software-engineering/v4>
- [23] Avi Rosenfeld and Ariella Richardson. 2019. Explainability in human-agent systems. *Autonomous agents and multi-agent systems* 33, 6 (2019), 673–705.
- [24] Per Runeson and Martin Höst. 2009. Guidelines for conducting and reporting case study research in software engineering. *Empirical software engineering* 14, 2 (2009), 131–164. doi:10.1007/s10664-008-9102-8
- [25] Senado Federal do Brasil. 2023. PL N° 2338, de 2023. Retrieved April 20, 2026 from <https://www25.senado.leg.br/web/atividade/materias/-/materia/157233>
- [26] Forrest Shull, Janice Singer, and Dag IK Sjøberg. 2008. *Guide to advanced empirical software engineering*. Vol. 93. Springer.
- [27] Ilia Stepin et al. 2021. A Survey of Contrastive and Counterfactual Explanation Generation Methods for Explainable Artificial Intelligence. *IEEE Access* 9 (2021), 11974–12001. doi:10.1109/ACCESS.2021.3051315
- [28] Anselm L. Strauss and Juliet Corbin. 1994. *Grounded Theory Methodology: An Overview*. In *Handbook of Qualitative Research*, Norman K. Denzin and Yvonna S. Lincoln (Eds.). Sage Publications, Inc., Thousand Oaks, CA, 273–285.
- [29] Umm-E-Habiba. 2023. Requirements Engineering for Explainable AI. In *2023 IEEE 31st International Requirements Engineering Conference (RE)*. 376–380. doi:10.1109/RE57278.2023.00058
- [30] Umm-E-Habiba and other. 2025. How do ML practitioners perceive explainability? an interview study of practices and challenges. *Empirical Software Engineering* 30, 1 (2025), 18. doi:10.1007/s10664-024-10565-2
- [31] Hugo Villamizar et al. 2021. Requirements engineering for machine learning: A systematic mapping study. In *2021 47th Euromicro Conference on Software Engineering and Advanced Applications (SEAA)*. IEEE, 29–36. doi:10.1109/SEAA53835.2021.00013
- [32] Andreas Vogelsang and Markus Borg. 2019. Requirements engineering for machine learning: Perspectives from data scientists. In *2019 IEEE 27th International Requirements Engineering Conference Workshops (REW)*. IEEE, 245–251.