

Simultaneous pursuit of accountability for regulatory compliance, financial benefits, and societal impacts in artificial intelligence (AI) projects

Gloria J. Miller, DBA
0000-0003-2603-0980
maxmetrics, Heidelberg
Email: g.j.m@ieee.org

Abstract—This exploratory study, grounded in agency theory, employs quantitative analyses to investigate the simultaneous pursuit of accountability for regulatory compliance, financial benefits, and societal impacts within artificial intelligence (AI) projects. An agent-principal matrix was developed, synthesizing knowledge from the AI stakeholder model into 11 accountability indicators. These indicators establish a standard of responsibility among project actors for regulatory compliance, ethical practices, and financial benefits. Using quantitative methods, we analyzed survey data on accountability and defined the scope of AI systems under development. We identified two clusters of AI systems—autonomous and non-autonomous—based on seven features. We then examined how these two types of systems, as well as the importance of sustainability and fairness, impact the promotion of accountability. Results indicate that accountabilities shift based on the scope of the AI system and the project role. Regulatory compliance, financial benefits, and societal impacts are not mutually exclusive project goals and coexist. The findings quantify subjective and theoretical speculation about accountabilities within AI projects. Additionally, the study contributes empirical data to the literature on AI, ethics, and project management.

Index Terms—Accountability, Artificial intelligence, AI, Project management, Ethics, Agency Theory

I. INTRODUCTION

THE launch of ChatGPT in 2023 dramatically increased interest, enthusiasm, and concerns regarding artificial intelligence (AI); ChatGPT is a generative AI chatbot based on a large language model that produces conversational outputs. Recognizing the potential impact of AI systems on the public and the environment, the European Commission and numerous countries have proposed frameworks, laws, regulations, and rules to responsibly manage these impacts, emphasizing interdependent values, principles, and actions [1, 2, 3]. At the same time, organizations prioritize achieving the target financial and business benefits from their AI investments [4, 5, 6].

This exploratory study investigated the relationship between accountability and the simultaneous pursuit of regulatory compliance, financial benefits, and societal impacts in AI projects. An AI project refers to a temporary organization that develops AI systems before deploying them for productive use

or putting them into the market [7]. In the public's opinion, accountability is, on average, perceived as one of the most important ethical principles before fairness, security, privacy, and accuracy [8]. The present study addresses the following research question: *Who within an AI project is accountable for regulatory compliance, financial benefits, and societal impacts?*

The study employed a questionnaire-based survey and quantitative analysis to measure AI project teams' expectations and perceptions regarding accountability. Its theoretical framework is grounded in agency theory and builds upon an existing AI stakeholder accountability model [7] that maps project success factors and AI ethical principles. These factors establish an accountability standard governing the relationship between project actors, regulations, societal impacts, and financial benefits. A quantitative analysis of survey data then identifies what individuals involved in AI projects are held accountable for.

This paper begins with a literature review, followed by a description of the study's theoretical background, hypothesis development, and research methodology, including data collection and analysis. The findings are then presented and discussed, alongside the study's contributions, implications, limitations, and conclusions.

II. LITERATURE REVIEW

A. AI Projects

AI refers to a wide range of efforts to replicate complex human capabilities, such as language use, vision, and autonomous action, using computational models. This encompasses various system categories, including chatbots, humanoid robots, behavioral software, and predictive systems [7]. AI systems function as artificial agents, capable of making decisions and acting with little or no human involvement. The complex data and models in AI systems are often referred to as "black-boxes" because the end-users, and sometimes the developers, are unable to explain or interpret the model decisions [9]. Consequently, concerns have been raised regarding civil and criminal liabilities, data protection, security,

trustworthiness, and transparency related to moral decision-making without human involvement [9, 10, 8].

An AI project is a temporary organization dedicated to developing AI systems before their deployment for productive use or commercialization [7]. Project teams comprise the natural or legal people who are best placed to understand the system's design and its potential impacts on individuals and society [7, 11]. Heaton et al. [12] reinforces the pivotal role of the team, arguing that developers should design transparent, explainable AI systems. However, AI projects often involve tensions between financial targets, ethics, internal controls, and compliance [13].

Specifically, ethical and community-oriented AI projects should engage stakeholders, such as end-user groups or affected communities, as participatory collaborators in the development of AI systems [12, 11]. However, participatory development is time-consuming and often faces systemic barriers, including a lack of support and resources for brokering relationships, as well as short project timelines. AI end-user groups and affected communities are vastly underrepresented in AI research relative to developers, designers, and organizational leaders [14].

B. AI Societal Impacts

Within AI projects, project actors, including project sponsors, managers, and team members, function as moral agents [15]. Their in-project decisions can significantly impact the lives of individuals and their liberty, human rights, or civil rights [12]. The actions of project actors can also result in environmental, financial, reputational, and political harm. AI systems are artificial agents that embody a moral code, as the outcomes of their decisions or recommendations can benefit or harm individuals or society. Consequently, project actors bear a moral responsibility to apply fair, ethical, and transparent processes to the AI systems they develop Heaton et al. [12].

AI projects may also have societal and environmental impacts. Training large models consumes significant amounts of energy and water, and produces carbon emissions, resulting in high energy costs and negative environmental impacts [10, 16]. Moreover, data representativeness has power, value, and cost implications. Determining who is included in or excluded from AI datasets is valuable. Hence, data curation involves tradeoffs between costs, accuracy, quality, completeness, representativeness, and equality [16, 11].

Ryan and Stahl [10] described 11 ethical categories that users and developers should consider when fulfilling their moral responsibilities within AI projects: beneficence, dignity, freedom and autonomy, justice and fairness, non-maleficence, privacy, responsibility, solidarity, sustainability, transparency, and trust. While agencies and governing bodies, including UNESCO and the EU High-Level Expert Group, have established ethical principles that define an acceptable moral code for AI systems, these principles are non-binding for developers, users, and platforms [1]. Furthermore, the principles alone are limited in their impact on AI design and governance, as it can be challenging to translate concepts, theories, and values into practice [10, 17].

C. AI Regulations

Regional legislators have encoded ethical principles into laws and regulations relating to AI. For instance, the EU introduced the AI Act, China the Chinese Internet Information Service Algorithm Recommendation Management Regulations, and Korea the AI Basic Act [1, 2, 3]. Conversely, in May 2025, the United States (US) House of Representatives proposed legislation that would ban federal, state, and local governments from enforcing laws or regulating AI for a 10-year period [18]; this sweeping moratorium was unsuccessful, but its proposal demonstrates the political appetite for regulating AI at the federal level in the US. DePaula et al. [19] found that the 50 US states are addressing AI in an uneven, piecemeal fashion. “While several legislations regulate an aspect of AI, a common response was to create task forces and commissions to study AI, prohibit potentially discriminatory practices of AI, and provide funding to capture potential economic or educational benefits from AI” [19, p. 822]. Regardless, research by Hopkins and Booth [20] found that regulations are ineffective in changing the behaviors of technology professionals.

The EU AI Act defines an AI system by its capability features (such as levels of autonomy, adaptiveness, and inferential abilities). It classifies such systems according to the risks they pose to society and the environment [1, 9]. The Act assigns the roles of deployers and providers to those working in the field of AI system development and use. The deployer is the natural or legal person responsible for the system during its development. This role is responsible for completing risk and compliance assessments, as well as registration, before proceeding to the next stage. The provider is the natural or legal person accountable for the system's operations.

D. AI Financial Benefits

AI ‘projects aim to realize their target benefits and enhance organizational performance... “target” benefits are intentionally sought by project funders and strategically set before project commencement’ [5, p. 656]. AI systems utilize data and algorithms to either benefit an organization's value proposition or monetize data. Generative AI applications such as ChatGPT are improving the performance in tasks that were previously unaffected by automation technology; empirical studies show performance gains between 14% and 50% in several business contexts [4].

“Value propositions define an organization's position in a market, its core skills and how these are applied for the benefit of others” [6, p. 168]. Consequently, introducing AI systems presents a business and technical challenge that should yield target benefits and tactical goals, such as increased revenue or improved performance efficiency. Organizations expect projects to be executed efficiently, delivering results within a given timeline and budget. They must therefore balance meeting target objectives with consideration of societal and environmental impacts, as well as the need to act ethically and comply with regulations.

III. THEORETICAL FRAMEWORK AND HYPOTHESIS DEVELOPMENT

This section outlines the theoretical framework and the development of hypotheses.

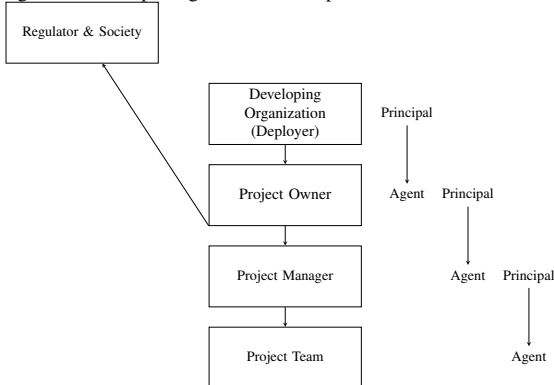
A. Theoretical Framework

Jensen and Meckling [21, p. 308] defined agency relationships as contracts “under which one or more persons (the principal(s)) engage another person (the agent) to perform some service on their behalf which involves delegating some decision-making authority to the agent.” Within the principal–agent relationship, agents are accountable to principals for their conduct. The agents have obligations to fulfill and should face consequences for their failure to perform. Moreover, agents must have the capacity to act intentionally, make choices, and exert influence with some degree of autonomy in the decision-making process [12].

Relating the roles in the EU AI Act to actors in the AI stakeholder accountability model [7] and agency theory, the deployer represents the developing organization, while the provider represents the operating organization. Within the deployer’s role, the project owner acts as an agent for the principal, the developing organization. Organizational strategic goals and policies flow downward from the principals to agents [14, 22].

When building AI systems, the project owner functions as the principal, with the project manager serving as the agent. Project management standards stipulate that the project manager is responsible for planning and managing the project, directing the performance of planned activities and overseeing the project’s technical, administrative, and organizational interfaces [23]. The team thus serves as an agent of the project manager. Fig. 1 illustrates the principal–agent relationships concerning the developing organization and the regulators and society.

Figure 1. Principal Agent Relationship



The AI stakeholder accountability model specifies and defines the complex relationships between actors and AI success factors, as well as the conduct and obligations that characterize them [7]. For this study, the model serves as the basis for defining accountability, i.e., the obligations the agent has to the principal; the success factors describe the conduct, events, acts, or deliverables in that relationship.

Table I synthesizes the knowledge from the AI stakeholder accountability model into a matrix that consolidates the 17 success groups into 11 accountability indicators, organized across the columns. The rows represent the agent–principal relationships and signify whether the project agent has an obligation to the principal for a given indicator. Success factors encompass ethical principles and activities mandated by AI regulations, representing deliverables, acts, or events that an AI project must include [7].

The synthesis identifies ethical and regulatory deliverables in the agent–principal relationship for AI projects, providing a basis for quantitatively evaluating AI accountabilities in practice. We selected the AI stakeholder accountability model as a theoretical basis for its alignment of project success with the AI ethical principles from Ryan and Stahl [10], regulatory and financial accountabilities, and project actors. Success factors identify the deliverables necessary for achieving project objectives. They can be used to establish an accountability standard that governs the relationship between project actors and stakeholders. Furthermore, the model determines what actions the project actors should be held accountable for and identifies the project owner as accountable to the regulators on behalf of the developing organization.

B. Hypothesis Development

Table I defines the obligations of project agents in achieving regulatory compliance, developing ethical AI systems, and delivering financial benefits using the 11 accountability indicators. Thus, in practice, accountabilities should vary by project role as defined in the matrix. This leads to the formulation of the following hypothesis:

Hypothesis 1: The (A) project manager, (B) owner, and (C) team will be responsible for their subset of the 11 accountability indicators in the agent–principal matrix.

In the context of AI projects, [7, p. 475] notes that “[f]irst and foremost, the scope established by the project owner is an essential artifact in designing AI systems.” The project scope is therefore an independent factor that determines the tasks and accountabilities of individual project actors. Similarly, the EU AI Act requires different obligations depending classification of AI system according to its use case and technical features [1, 9]. For example, AI use cases for social scoring or manipulation are prohibited. Furthermore, system features such as model type, adaptiveness after deployment, personal data use, or inferring input data outputs would affect the system classification, computing costs, and environmental impact [1, 9, 16]. Therefore, the following hypothesis is proposed:

Hypothesis 2: Depending on the system scope, there will be a significant difference in accountability for (A) regulatory compliance, (B) energy costs, and (C) sustainability.

Within AI projects, hidden information and actions create challenges for the agent–principal relationship. Hidden information refers to the principal’s difficulties in determining whether the agent has the necessary qualities and skills to perform a task in the principal’s interest [24]. Furthermore,

Table I
ACCOUNTABILITY BY AGENT-PRINCIPAL RELATIONSHIP AFTER [7]

Agent	Principal	Accountabilities										Sustainability SY
		Energy Cost EC	Ethic Practices EP	Financial Benefits FB	Intellect Property IP	Societal Impacts SI	Prj Efficiency PE	Privacy Protect PP	Product Quality PQ	Reg. Com- pliance RC	Transparency Usability STU	
Prj Mgr	Owner	X	X		X		X	X		X	X	
Owner	Organization	X		X	X		X		X	X		
	Regulator		X					X		X		
	Society		X			X		X	X		X	X
Team	Prj Mgr	X					X					
	Owner		X		X			X		X	X	

Abbreviations: Prj-Project, Mgr-Manager; Reg-Regulatory

some AI systems are black-boxes with hidden features caused by complex data and models [9]. Nevertheless, the AI systems are the result of actions by the project teams, influenced by their choices, biases, values, and intentions [12]. Hidden actions refer to activities that may be in the agent's interests but not the principal's. Agency theory explains the issue of moral hazards, where within the principal-agent relationship, individuals may pursue their personal interests even when they conflict with the team's objectives [21, 24, 25].

Hypothesis 3: For an autonomous system scope, agents are more likely to be accountable for (A) system transparency and usability, and (B) protecting intellectual property or financial gains.

The performance of project actors is based on what their organization considers important and what they will be held accountable for. Thus, if the organization considers the development of fair, understandable, and sustainable systems important, project agents will be accountable for tasks that deliver responsible AI systems [7, 8, 10].

Hypothesis 4: If the organization emphasizes the importance of creating fair and understandable algorithms: (A) the project manager is more likely to be responsible for implementing ethical practices; (B) the owner is more likely to be responsible for societal impacts; and (C) the team is more likely to be responsible for system transparency and understandability, product quality, and societal impacts.

Hypothesis 5: If the organization emphasizes the importance of environmental sustainability, the (A) project manager, (B) owner, and (C) team are more likely to be responsible for sustainability, energy cost, or both.

Organizations undertake projects to achieve target benefits and enhance organizational performance [4, 5, 6].

Hypothesis 6: (A) A majority of (a) owners and (b) teams will be involved in achieving financial benefits, intellectual property, or both. (B) A majority of (a) project managers and (b) owners will be accountable for the project's efficiency.

The AI stakeholder accountability model does not associate achieving financial benefits, including revenue generation, cost efficiency, or project efficiency, with any ethical principle. Furthermore, the project team must make several tradeoffs when building AI systems, including accuracy, energy cost, financial costs, project efficiency, and regulatory compliance, as well as transparency and intellectual property rights, and

legal safeguards and system flexibility [7, 13, 25].

Hypothesis 7: The following will be negatively correlated: (A) financial benefits and regulatory compliance responsibility, (B) financial benefits and privacy protections responsibility, and (C) intellectual property and system transparency and understandability.

IV. RESEARCH METHODOLOGY

This section describes the research methodology, including the data collection methods, measurement instrument, analysis methods, and reliability and validity measures. The unit of analysis for the study is the project.

A. Data Collection

This study collected data on accountability within an AI project using a web-based survey hosted on SurveyMonkey, a global survey platform that utilizes a proprietary panel of survey respondents as its audience. Data were collected in November 2023 and November 2024 using a SurveyMonkey audience panel for a US audience.

Respondents provided consent to participate. No personally identifiable data were collected, and no compensation was offered. Respondents who selected "None" as an AI project role were disqualified. The survey targeted respondents in consulting, information technology, management, and project management job roles.

For the 2023 survey, 228 invitations were sent to potential participants. Of these, 121 (47%) were disqualified based on their AI project role, providing 107 (53%) usable responses. SurveyMonkey reported an error rate of +/- 10%. For the 2024 survey, 151 invitations yielded 8 (5%) disqualifications and 143 (95%) usable responses, and an error rate of +/- 8%.

Of the combined 2023 and 2024 responses, 57 of the usable respondents had not completed at least three AI projects, and 29 entries had invalid data in the experience field. These were excluded, resulting in a final total of 164 valid responses.

B. Measurement Instrument

We employed a three-stage process of content, clarity, and expert review to develop the measurement instrument, which was in English. The expert review was performed by three professionals with experience in at least three AI projects in the last ten years. All stages resulted in adjustments to the text.

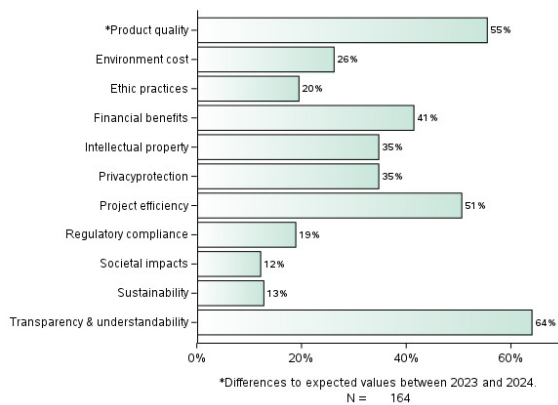


Figure 2. Accountability Indicator Demographics

This instrument defined algorithms, AI systems, and AI projects as follows:

“An algorithm is a defined, repeatable process and outcome based on data, processes, and assumptions. An Artificial Intelligence (AI) system is a computer system that incorporates algorithms that learn from data and support human decision-making or autonomously make decisions. An AI project is a temporary organization (from few months to many years) that develops one or more AI systems” (author’s definitions).

Respondents selected their *project role* from seven options and were grouped into composite role measures for the project owner, manager, and team. The roles and groups were taken from Miller [7].

Table I, which was also synthesized from Miller [7], was used as a bases for defining the measures for individual project *accountability* and *agent–principal relationship* accountabilities. For individual project *accountability*, respondents selected from 11 options or entered “other”; these options are shown in Table II under “Accountability Questions.” Each selection created a binary measure with a value of one if it was selected and zero if not. Fig. 2 visualizes the statistics for the responses to this question. For the *agent–principal relationship* accountability combinations, a dummy variable was created with an indicator for each agent-principal pair that had an expected accountability in Table I.

Respondents selected one or more of six characteristics to describe a specific *AI system scope* they developed as part of their project [1, 9]. Fig. 3 presents the statistical responses regarding AI system characteristics. The respondents were also asked “Of the 10 levels of automation... [i]n one of my last 3 AI projects, AI was at level...?” A table based on Sheridan et al. [26] provided 10 levels of automation to guide the participants’ responses, allowing them to select answers from 1 to 10. A dummy variable was created where levels 1 to 6 were considered human decisions and levels 7 to 10 were considered AI decisions. The responses to questions about *AI system scope* and the *level of automation* were used to classify the AI system as autonomous or non-autonomous systems using latent class analysis (LCA).

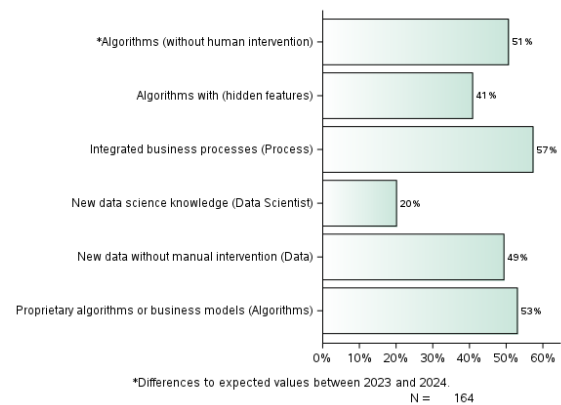


Figure 3. AI System Scope Demographics

For *experience*, respondents entered the number of AI projects they had completed in the last decade; those reporting fewer than three were not included in the analysis. To ascertain the importance of *fairness* and *sustainability*, two seven-point Likert scale questions were included, and binary measures were created for each.

Table II
ACCOUNTABILITY QUESTION

What aspects of project performance were you accountable for in the AI project referenced in question Q2? (check all that apply)			
	Response Options	Analysis Identifier	ID
AI system scope	Product quality (training data, algorithm, user interface, system configuration)	Product quality	PQ
	System transparency and understandability, usage controls, decision quality	Transparency & understandability	STU
Regulatory	Ethic practices, training, policies, oversight	Ethic practices	EP
	Regulatory and legal compliance	Regulatory compliance	RC
	Protecting data privacy and confidentiality	Privacy Protections	PP
Societal Impacts	Social impacts of system usage (civil and human rights protections, etc.)	Societal impacts	SI
	Environmental sustainability of system usage	Sustainability	SY
Benefits	New revenue generation, cost savings, license fees	Financial benefits	FB
	Project time, cost, and quality performance—iron triangle	Project efficiency	PE
	Environmental costs, energy costs environmental impacts	Energy cost	EC
	Intellectual property protections (trade secrets, proprietary algorithms, obtaining patents, etc.)	Intellectual property	IP

C. Data Analysis

The SAS Studio Release 3.8 (Enterprise Edition) was utilized for quantitative analysis, statistical tests, and bias checks. The accountability data were binomial and non-normally distributed. Thus, parametric measures were suitable for the analysis. The Wilcoxon test was used to compare the means of variables and determine the significance of the comparison. Kendall’s tau-b correlation coefficient was

used to evaluate the strength, direction, and significance of the relationship between measurement items. The correlation coefficient does not indicate which variable causes a change in another variable, nor does it imply causality. The Wilcoxon scores and correlations with a p -value of less than 0.05 were considered significant with a 95% confidence interval [27]. The rank shows the relevant importance of the indicator.

LCA clustering was employed to identify homogeneous clusters within the scope of the AI system. The dimensions were modeled in R version 4.0.2 using poLCA for the analysis. Seven items for the system scope and level of automation were specified as indicators in the model. LCA was executed with 50,000 iterations for 1–10 classes and 10 repetitions per model. The class model was selected by evaluating the Bayesian information criterion (BIC) fit statistics, log-likelihood, entropy, and the model's theoretical interpretability [28].

D. Reliability and Validity

This study relied on the quality of demographics and professional data from SurveyMonkey for sample construction. SurveyMonkey conducts periodic quality checks using surveys and panel calibration studies to help ensure the quality of open-ended responses and identify poor respondent behavior [29]. Moreover, Bentley et al. [30] conducted a comparative study on several audience panels and found the SurveyMonkey audience panel had an error margin that was “statistically indistinguishable” from more expensive panels. Thus, the sample construction was considered valid for the study.

The survey included a question to establish whether respondents had participated in both 2023 and 2024 to account for the effects in repeated surveys; 31% indicated they took both surveys. This repetition impacts the uniqueness of the responses and the effect size of the results. First, our unit of measure is the project. A person is likely to change projects between years; machine learning or AI projects are from a few weeks to several months [16]. Thus, the project characteristics may vary for the repeated responses, which is valuable to our analysis. Second, effect sizes must be evaluated based on the average number of repeated measures, as additional repeated measures provide less novel information [31]. Analysis of effect sizes was not relevant to our study; nevertheless, this factor must be considered, as it decreases the sample size, which impacts the generalizability of the study data.

For further sample validity, we used a screening question and performed quantitative checks on the survey responses to identify missing data, extreme responses (where the same response was recorded for all questions), and irregular responses to critical questions, including the entry of odd responses to open-ended questions regarding years of experience. Consequently, we removed 29 entries that were deemed invalid.

We conducted validity checks for common method bias and response bias. Response bias was assessed between survey years (2023 and 2024) to determine whether there were significant differences between the responses. Responses to the question on product quality in the accountabilities section and algorithm use without human intervention in the system scope section both exhibited significantly higher scores in 2024 than

Table III
DEMOGRAPHICS

Title	Label	N	Percent
Age	18 – 29	15	9.15
	30 – 44	84	51.2
	45 – 60	59	36.0
	> 60	6	3.66
Gender	Male	103	62.8
	Female	61	37.2
Experience	< 5 projects	47	28.7
	5 projects	41	25.0
	6 – 7 projects	14	8.54
	8 – 15 projects	30	18.3
	> 15 projects	32	19.5
Owner	Chief Executive Officer, Business Leader	31	18.9
	Chief Information Officer, IT Manager	53	32.3
Project Manager	Project Leader/Manager, Product Owner	40	24.4
Team	Data Scientist, Computer Vision, AI Specialist	14	8.54
	Data Engineer, Data Custodian, Curator	6	3.66
	Architect, Developer, Computer Scientist, IT	7	4.27
	Business User/ Analyst, Subject Matter Expert	13	7.93
Source	SurveyMonkey 2023	54	32.9
	SurveyMonkey 2024	110	67.1
AI system	Autonomous system	79	48.2
	Non—autonomous system	85	51.8
Automation	Level 1 – 6 Human Decisions	85	51.8
	Level 7 – 10 AI Decisions	79	48.2
Fairness	Not important	9	5.49
	Important	155	94.5
Sustainability	Not important	24	14.6
	Important	140	85.4

in 2023. These differences were managed in the analysis and reported in the results to avoid spurious interpretations.

Given that all variables were collected using the same survey instrument from a single source, Harman's single-factor test was performed to check for common method bias [32], using unrotated factor analysis to analyze the scope and accountability variables. A single factor explained 16.53% of the variance, which is well below the 50% heuristic for indicating bias. Thus, common method bias was not considered an issue.

We applied the recommendations presented by Weller et al. [28] to prepare the LCA clusters.

V. FINDINGS

This section describes the study's findings. Table IV provides a summary of the hypotheses and findings.

A. Demographics

Table III, Fig. 2, and Fig. 3 present the demographics of the respondents. Based on the AI system scope, less than half of the participants (42%) indicated that they had developed *autonomous* systems. Fig 4 shows the mean score for the features by AI system scope classification. The majority of respondents reported that *fairness* (95%) and sustainability (85%) were important to their organizations.

B. Agent-Principal Accountabilities

The data did not support the accountability expectations from the AI stakeholder accountability model. Table V displays the basic statistics for the ratio of accountabilities

Table IV
SUMMARY OF HYPOTHESES FINDINGS

ID	Hypothesis	Supported
1	The (A) project manager, (B) owner, and (C) team will be responsible for their subset of the 11 accountability indicators in the agent–principal matrix.	(A) No (B) Yes (C) Yes
2	Depending on the system scope, there will be a significant difference in accountability for (A) regulatory compliance, (B) energy costs, and (C) sustainability.	(A) No (B) Yes (C) Yes
3	For an autonomous system scope, agents are more likely to be accountable for (A) system transparency and usability, and (B) protecting intellectual property or (C) financial gains.	(A) No (B) No (C) Yes
4	If the organization emphasizes the importance of creating fair and understandable algorithms: (A) the project manager is more likely to be responsible for implementing ethical practices; (B) the owner is more likely to be responsible for societal impacts; and (C) the team is more likely to be responsible for system transparency and understandability, product quality, and societal impacts.	No No No
5	If the organization emphasizes the importance of environmental sustainability, the (A) project manager, (B) owner, and (C) team are more likely to be responsible for sustainability, energy cost, or both.	No
6	(A) A majority of (a) owners and (b) teams will be involved in achieving financial benefits, intellectual property, or both. (B) A majority of (a) project managers and (b) owners will be accountable for the project's efficiency.	No No
7	The following will be negatively correlated: (A) financial benefits and regulatory compliance responsibility, (B) financial benefits and privacy protections responsibility, and (C) intellectual property and system transparency and understandability.	No No Yes

Table V
AGENT-PRINCIPAL RELATIONSHIP STATISTICS

Agent	Principal	N	MIN	MEAN	MAX	STD
All	All Indicators	164	0.00	3.71	11.0	2.06
Owner	Organization	84	0.00	2.40	6.00	1.18
	Regulator	84	0.00	0.70	3.00	0.89
	Society	84	0.00	2.79	7.00	1.67
Prj Mgr	Owner	40	0.00	2.48	5.00	1.22
Team	Prj Mgr	40	0.00	0.65	2.00	0.74
	Owner	40	0.00	1.70	5.00	1.30

Legend: Roles: Prj-Project, Mgr-Manager
Statistics: N-Number, Std-Standard Deviation

in the agent–principal relationship. For example, the project manager and owner relationship is expected to include seven accountabilities, as shown in Table 1. However, the maximum was five, the mean was 2.48, and the standard deviation was 1.22; thus, no project manager claimed all the accountabilities suggested by the model. The project owner and teams fared better since the maximum was achieved; nevertheless, the means were relatively low. Table VI presents the percentage of accountability by respondent role. Therefore, hypotheses 1B and 1C are supported for the owner and team, but 1A is rejected for the project manager.

Notably, the respondents' experience level influenced several accountabilities. Respondents with 15 or more projects were accountable for intellectual property (IP) and product quality (PQ) more than other experience groups; those with fewer than five projects were responsible for ethical practices (EP), and those with fewer than six projects were more responsible for project efficiency (PE).

C. Impact of AI System Scope

We assessed accountability for each project role, where the AI system scope was the independent variable and the

Table VI
AGENT ACCOUNTABILITY STATISTICS

	Accountabilities										
	EC	EP	FB	IP	SI	PE	PP	PQ	RC	STU	SY
Agent	Percent										
Owner	13	9	21	18	7	29	17	32	10	35	8
PM	7	5	10	9	4	12	8	15	4	16	1
Team	6	5	10	9	1	10	10	9	5	13	4

Legend: Roles: PM-Project Manager;Accountabilities: EC-Energy Cost, EP-Ethic Practices, FB-Financial Benefits, IP-Intellectual Property, SI-Societal Impacts, PE-Project Efficiency, PP-Data Privacy, PQ-Product Quality, RC-Regulatory, STU-Traceability & Understandability, SY-Sustainability
Statistics: Owner=84, Prj Mgr=40, Team=40

accountabilities were the dependent variables. The scope had a significant influence on the accountabilities. There were significant differences in several indicators: financial benefits (FB) ($H=8.55$, $\rho=0.00$), energy cost (EC) ($H=13.28$, $\rho=0.00$), and sustainability (SY) ($H=5.19$, $\rho=0.02$). Consequently, agent–principal accountabilities for the owner–organization ($H=14.69$, $\rho=0.00$) and team–project manager ($H=15.26$, $\rho=0.00$) relationships were impacted. Hypothesis 2A is therefore rejected, and hypotheses 2B and 2C are supported. Table VII shows the Kruskal-Wallis mean rank statistics for each accountability and agent–principal pair by AI system scope.

System transparency and understandability (STU) ranked highest for both autonomous and non-autonomous scopes. The Kendall correlation presented in Table VIII revealed a negative, insignificant correlation between STU and the AI system scope ($\tau = -0.06$, $\rho > 0.05$). Thus, possessing hidden features did not affect transparency requirements or move the agent away from STU. Nevertheless, the EC and FB accountabilities were ranked significantly higher for the autonomous system scope. Intellectual property (IP) had no significant difference by scope. Therefore, hypothesis 3C is supported, while hypotheses 3A and 3B are rejected.

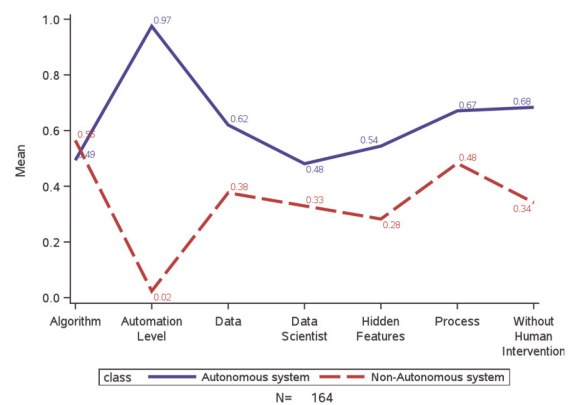


Figure 4. AI system scope classification by features for the survey population.

D. Regulatory Accountabilities

No significant difference was found in regulatory compliance (RC) accountability in terms of the AI system scope. Since many types of laws and regulations, such as those related to healthcare and disabilities, apply in the US regardless of

Table VII
ACCOUNTABILITY MEAN RANK STATISTICS BY AI SYSTEM SCOPE

			Autonomous system (N= 79)			Non-Autonomous system (N= 85)			Overall (N= 164)			Kruskal Wallis	
Category	Label	ID	Mean	Std	Rank	Mean	Std	Rank	Mean	Std	Rank	H	ρ Value
Agent-Principal													
Agent – Principal	Owner – Society		2.75	1.62	1	2.52	1.69	1	2.63	1.66	1	0.29	0.59
	Prj Mgr – Owner		2.71	1.53	2	2.28	1.44	2	2.49	1.49	2	2.33	0.13
	Owner – Organization		2.65	1.20	3	1.93	1.18	3	2.27	1.24	3	14.69	0.00
	Team – Owner		1.73	1.30	4	1.71	1.23	4	1.72	1.26	4	0.04	0.85
	Team – Prj Mgr		0.97	0.64	5	0.58	0.62	6	0.77	0.66	5	15.26	0.00
	Owner – Regulator		0.67	0.89	6	0.79	0.93	5	0.73	0.91	6	0.77	0.38
Accountabilities													
AI system scope	Transparency & Usability	STU	0.67	0.47	1	0.61	0.49	1	0.64	0.48	1	0.62	0.43
	Product Quality	PQ	0.57	0.50	3	0.54	0.50	2	0.55	0.50	2	0.13	0.72
Regulatory	Privacy Protection	PP	0.33	0.47	7	0.36	0.48	4	0.35	0.48	6	0.23	0.63
	Ethic Practice	EP	0.16	0.37	10	0.22	0.42	7	0.20	0.40	8	0.90	0.34
	Regulatory Compliance	RC	0.18	0.38	9	0.20	0.40	8	0.19	0.39	9	0.14	0.71
Societal Impacts	Sustainability	SY	0.19	0.39	8	0.07	0.26	11	0.13	0.34	10	5.19	0.02
	Social Impact	SI	0.15	0.36	11	0.09	0.29	10	0.12	0.33	11	1.27	0.26
Benefits	Project Efficiency	PE	0.58	0.50	2	0.44	0.50	3	0.51	0.50	3	3.52	0.06
	Financial Benefit	FB	0.53	0.50	4	0.31	0.46	6	0.41	0.49	4	8.55	0.00
	Intellectual Property	IP	0.39	0.49	6	0.31	0.46	6	0.35	0.48	6	1.34	0.25
	Energy Cost	EC	0.39	0.49	6	0.14	0.35	9	0.26	0.44	7	13.28	0.00

Abbreviations: N=Number observations, Std=standard deviation, KW=Kruskal-Wallis, ρ =significance, Prj=Project, Mgr=Manager

the technology, the results are not entirely unexpected. Approximately 10% of the owners, 4% of the project managers, and 5% of the team members claimed RC accountability. RC accountability was significantly correlated with FB ($\tau = .16$, $\rho < .05$), privacy protections (PP) ($\tau = .27$, $\rho < .0001$), ethical practices (EP) ($\tau = .27$, $\rho < .0001$), and societal impacts (SI) ($\tau = .20$, $\rho < .05$). In sum, both hypothesis 2A, positing that the scope will impact RC, and hypothesis 7A, proposing that there will be negative tradeoff between RC and FB, are rejected.

E. Societal Accountabilities

When controlling for interaction for the importance of sustainability or fairness, there was no significant difference in agent–principal accountabilities. First, all owners indicated that fairness was important. Second, there were no significant differences for any agents regarding the accountabilities of EP, STU, PQ, or SI when controlling for fairness or sustainability.

Both SI ($\tau = .22$, $\rho < .01$) and SY ($\tau = .16$, $\rho < .05$) were significantly correlated with FB. This could imply that organizations see their projects as benefiting society and considering sustainability goals. The AI stakeholder accountability model attributes these accountabilities exclusively to the owner. However, only a small percentage of owners (8%) claimed them, and the other roles also claimed these responsibilities. Hypotheses 4 and 5 are therefore rejected.

F. Financial Accountabilities

Using the accountabilities of FB or IP to represent the financial gains from the project, a minority of agents had these responsibilities: project owners (31%), project managers (15%), and team members (14%). Surprisingly, only 12% of project managers and 29% of owners claimed responsibility for PE, a standard metric of project management performance including budget, time, and quality. Thus, hypotheses 6A and 6B are rejected.

EC was significantly correlated with ($\tau = .29$, $\rho < .0001$) and impacted by ($SS=2.33$, $df(1)=13.05$, $\rho < .0001$) the AI

system scope, but it was not affected by the sustainability importance measure as expected.

G. Tradeoffs Between Accountabilities

The tradeoff between IP and STU ($\tau = -.01$, $\rho > .05$) was negative, but this was not significant. However, the tradeoffs anticipated by the AI stakeholder accountability model between FB and RC (.16, $\rho < .01$) and FB and PP (.01, $\rho > .05$) were not realized. Thus, hypotheses 7A and 7B are rejected, while hypothesis 7C is supported. Fig. 5 visualizes the Kendall Tau correlations that were significant with an effect size greater than .20.

VI. DISCUSSION AND CONCLUSIONS

This study employed the AI stakeholder accountability model, grounded in agency theory, to establish an accountability standard governing the relationships among project actors, regulation, ethical principles, and benefits. A quantitative analysis of survey data examined individuals' accountabilities in practice within AI projects. Additionally, we identified two clusters of AI systems—autonomous and non-autonomous—based on seven features. The *autonomous* cluster describes systems with a high level of machine-based decision-making that allow for limited human intervention, while the *non-autonomous* excludes these features.

This study examined how AI systems (both autonomous and non-autonomous) influence accountability, as well as the importance of sustainability and fairness in this context. As expected, we found that accountabilities shift based on the scope of the AI system; unexpectedly, we also determined that most expected accountabilities are not undertaken in practice.

The intersection between target benefits, societal impacts, and regulatory compliance depends on the project owner's role. Our findings suggest that the owner's accountability for regulatory compliance remains relatively unchanged, regardless of the scope of the AI system. However, their accountability towards the organization does change.

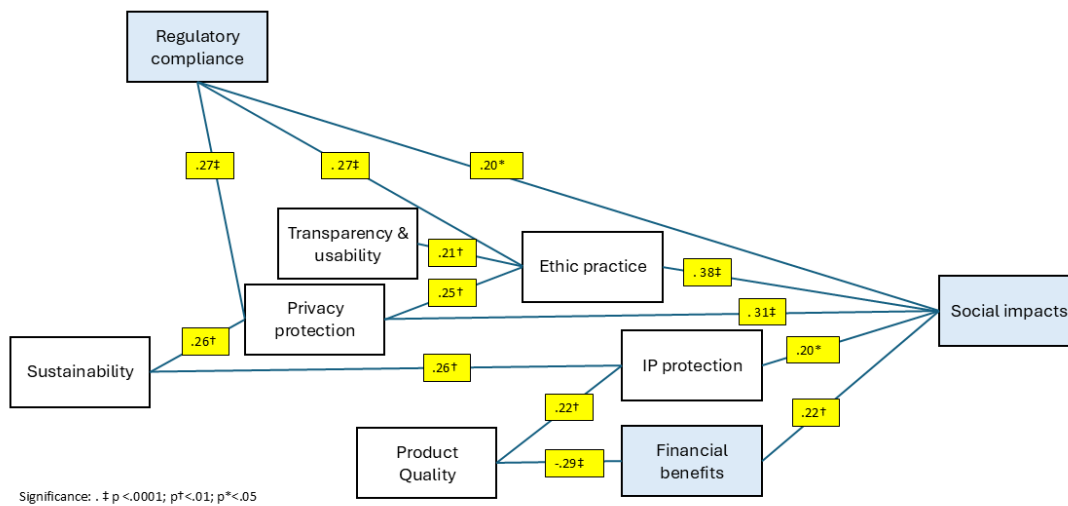


Figure 5. Kendall Tau Correlation Visualization

Most respondents practiced accountability regarding system transparency and understandability. This finding is consistent with Heaton et al. [12], who suggest that developers should design transparent and explainable AI systems to promote trust and mitigate any negative consequences that could arise from system use. Moreover, the transparency and understandability indicator is significantly correlated with ethical practice.

Ethical practice plays a central role in regulatory compliance and social impacts, and is further correlated with financial performance. This suggests that ethical practices may be the driving force behind the achievement of all other benefits. The analysis across experience levels further supports the argument for centering ethical practices. It showed that respondents with four or fewer projects indicated a greater responsibility for ethical practices than any other experience group. Taken together with the other indicators, this suggests that novices rely on documented practices (ethical practice) or traditional project performance measures (project efficiency) in their job performance. At the same time, more experienced individuals focus on creating content (intellectual property) and producing quality results (product quality).

We examined the fluidity of accountabilities, utilizing owner, project manager, and team roles from the sponsoring organization (i.e., the deployer, as defined by the EU AI Act). The accountabilities change based on the AI system's scope; if the project includes a scope typical of building autonomous decision-making tools, there is greater accountability, with a keen focus on the business impacts of product quality, intellectual property, and energy costs.

The lack of some expected tradeoffs may be attributed to the study's design. Wachnik [25] notes that many moral hazards in projects are due and the dangers linked to opportunistic supplier behaviour in a client-supplier relationship. The study lacked information on the respondents' organizational relationships, so it did not permit such analysis. Nevertheless, some tradeoffs between compliance and obtaining target benefits appear to be based on role assignments. Not all accountabilities

are observed by a single agent, which helps avoid conflicts between personal and organizational interests [21, 24, 25]. Preventing moral hazards should incentivize the optimization of project organizations, including their contractual structures, team size, and role assignments.

Despite the robustness of the AI stakeholder accountability model, individuals hold less than 50% of the expected accountabilities in practice. However, we anticipate that this will change over time. First, the monetary and environmental cost of developing new and tuning existing AI models becomes extremely expensive, quickly [16]. Second, the EU AI Act comes into full effect by 2027. In response to non-compliance, the EU can impose "administrative fines of up to 35 000 000 EUR...or up to 7% of its total worldwide annual turnover...whichever is higher" [1, p. 115]. Third, generative AI capabilities extended with digital tools or integrated with robotics are estimated to produce "around 0.4-0.9 p.p. [percentage point] contribution to annual labour productivity growth, when assuming a standard capital multiplier of 1.5" [4, p. 10]. Finally, there are strong incentives for firms to implement one or more AI systems for economic benefits or to address competitive pressures [6].

This study shows that for ethical AI development, it is necessary to establish a project role assignment matrix; this finding is consistent with Herrera [9] and Kuehnert et al. [14]. Here, we propose that the AI stakeholder accountability model evaluated in the current study could provide a baseline standard for assigning AI project responsibilities.

A. Theoretical Contributions

Our study's findings quantify subjective and theoretical speculation about who is accountable for the financial, regulatory, and societal impacts of AI projects. This study identifies the relationships among system features, automation levels, and accountability indicators in AI projects by empirically investigating the practices of experienced project participants. The results contribute to the literature on AI, ethics, and

Table VIII
KENDALL TAU B CORRELATION COEFFICIENTS

Kendall Tau b Correlation Coefficients											
	AIS	PQ	STU	PE	FB	EC	SY	IP	RC	PP	EP
PQ	-.03										
STU	-.06	-.03									
PE	-.15	.07	-.03								
FB	-.23†	-.29‡	.12	.04							
EC	-.29‡	-.05	-.04	-.02	.06						
SY	-.18*	.05	.06	.09	.16*	.10					
IP	-.09	.22†	-.01	.03	.06	.18*	.26†				
RC	.03	-.01	.10	-.02	.16*	.00	.14	.11			
PP	.04	.19*	.09	-.07	.01	.15	.26†	.17*	.27‡		
EP	.07	.07	.21†	.18*	.15	.16*	.18*	.09	.27‡	.25†	
SI	-.09	.11	.12	.14	.22†	.16*	.19*	.20*	.20*	.31‡	.38‡

Significance: ‡p<.0001, †p<.01, *p<.05; Abbreviations: AIS=AI System Scope, Exp=Experience, PQ=Product Quality, STU=System Transparency and Usability, PE=Project Efficiency, FB=Financial Benefits, EC=Energy Cost, SY=Sustainability, IP=Intellectual Property, RC=Regulatory Compliance, PP=Privacy Protections, EP=Ethic Practices, SI=Social Impacts

project management by aligning regulatory requirements with a project success model. Although research on ethical principles in AI is widely available, this research addresses the criticism that organizational leaders are underrepresented in AI research [14].

B. Practical Implications

Project sponsors should understand the interplay between various project assignments and define performance measures to avoid conflicts that may create moral hazards. Furthermore, the scope of the AI system should not be underestimated when defining a project organization. Thus, when staffing decisions require a tradeoff in role assignments, it is crucial to consider ethics, regulations, and performance efficiency.

C. Limitations and Future Research

One limitation of the present study was its data collection strategy. The survey was distributed to randomly selected individuals in the US, did not consider a single project, and relied on self-report measures. First, relying on a sample consisting exclusively of respondents from the US may limit the generalizability of the findings, especially given that regions such as Europe have more advanced AI regulations in place. Second, the study did not assess different people from the same project team. Thus, it limited the ability to investigate team assignments and moral hazards. Third, the survey relied on self-reports by the participants. Thus, bias may have been introduced by individual memories or recalls.

To mitigate the limitations of this study, future research could focus on a case study to quantify the engagement of separate roles, or employ strategies to assess a single project environment. Future studies should investigate AI performance measures at a granular level by project role. Finally, this analysis was conducted from the deployer's perspective and should also be conducted from the promoter's perspective. The system configuration (SC), user interface (UI), and usage control (UC) qualities were not included in the survey as standalone items and were therefore excluded from the analysis. This was an acceptable exclusion, as the AI stakeholder accountability model identified these qualities as important for

operations, which were not analyzed in this study. Although risk and quality management are important components of the EU AI Act, our survey did not include an independent question on these elements, so this aspect could not be independently analyzed. These are interesting avenues for future research to explore.

D. Conclusion

In conclusion, regulatory compliance, financial benefits, and societal impacts are not mutually exclusive project goals and coexist as shown in Fig. 5.

REFERENCES

- [1] European Union, "Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence and amending Regulations (EC) No 300/2008, (EU) No 167/2013, (EU) No 168/2013, (EU) 2018/858, (EU) 2018/1139 and (EU) 2019/2144 and Directives 2014/90/EU, (EU) 2016/797 and (EU) 2020/1828 (Artificial Intelligence Act)," 2024. [Online]. Available: <http://data.europa.eu/eli/reg/2024/1689/oj>
- [2] M. Theresia, "Newly enacted law sets basis for nat'l development of AI," Dec 27, 2024 2024.
- [3] M. Sloane and E. Wüllhorst, "A systematic review of regulatory strategies and transparency mandates in AI regulation in europe, the united states, and canada," *Data & Policy*, vol. 7, 2025. <https://dx.doi.org/10.1017/dap.2024.54>
- [4] F. Filippucci, P. Gal, and M. Schief, "Miracle or myth? assessing the macroeconomic productivity gains from artificial intelligence," OECD Publishing, Report, 2024.
- [5] O. Zwikael, Y.-Y. Chih, and J. R. Meredith, "Project benefit management: Setting effective target benefits," *International Journal of Project Management*, vol. 36, no. 4, pp. 650–658, 2018. <https://doi.org/10.1016/j.ijproman.2018.01.002>
- [6] C. F. Breidbach and P. Maglio, "Accountable algorithms? the ethical implications of data-driven business models," *Journal of Service Management*, vol. 31, no. 2, pp. 163–185, 2020. <https://doi.org/10.1108/JOSM-03-2019-0073>
- [7] G. J. Miller, "Stakeholder-accountability model for artificial intelligence projects," *Journal of Economics and Management*, vol. 44, no. 1, pp. 446–494, 2022. <https://doi.org/10.22367/jem.2022.44.18>
- [8] K. Kieslich, B. Keller, and C. Starke, "Artificial intelligence ethics by design. evaluating public perception on the importance of ethical design principles of artificial intelligence," *Big Data & Society*, vol. 9, p. 205395172210929, 2022. <https://dx.doi.org/10.1177/20539517221092956>
- [9] F. Herrera, "Attentiveness on criticisms and definition about explainable artificial intelligence," in *2024 19th Conference on Computer Science and Intelligence Systems (FedCSIS)*, 2024, Conference Proceedings, pp. 45–52. <https://dx.doi.org/10.15439/2024F0001>

- [10] M. Ryan and B. C. Stahl, "Artificial intelligence ethics guidelines for developers and users: clarifying their content and normative implications," *Journal of Information, Communication and Ethics in Society*, vol. 19, no. 1, p. 61–86, 2021. <https://dx.doi.org/10.1108/JICES-12-2019-0138>
- [11] L. H. Ajmani, N. A. Abdelkadir, and S. Chancellor, "Secondary stakeholders in AI: Fighting for, brokering, and navigating agency," p. 1095–1107, 2025. <https://dx.doi.org/10.1145/3715275.3732071>
- [12] D. Heaton, J. Clos, E. Nichele, and J. E. Fischer, "The social impact of decision-making algorithms: Reviewing the influence of agency, responsibility and accountability on trust and blame," p. Article 11, 2023. [10.1145/3597512.3599706](https://dx.doi.org/10.1145/3597512.3599706)
- [13] P. S. Scoleze Ferrer, G. D. A. Galvão, and M. M. de Carvalho, "Tensions between compliance, internal controls and ethics in the domain of project governance," *International Journal of Managing Projects in Business*, vol. 13, no. 4, p. 845–865, 2020. <https://dx.doi.org/10.1108/IJMPB-07-2019-0171>
- [14] B. Kuehnert, R. Kim, J. Forlizzi, and H. Heidari, "The "who", "what", and "how" of responsible AI governance: A systematic review and meta-analysis of (actor, stage)-specific tools," p. 2991–3005, 2025. <https://dx.doi.org/10.1145/3715275.3732191>
- [15] T. M. Jones, "Ethical decision making by individuals in organizations: An issue-contingent model," *Academy of Management Review*, vol. 16, no. 2, p. 366–395, 1991. <https://doi.org/10.5465/amr.1991.4278958>
- [16] E. Strubell, A. Ganesh, and A. McCallum, "Energy and policy considerations for deep learning in NLP," *arXiv preprint arXiv:1906.02243*, 2019. <https://doi.org/10.48550/arXiv.1906.02243>
- [17] B. Mittelstadt, "Principles alone cannot guarantee ethical AI," *Nature Machine Intelligence*, vol. 1, no. 11, p. 501–507, 2019. <https://doi.org/10.1038/s42256-019-0114-4>
- [18] 119th Congress (2024–2025), "Committee print: Providing for reconciliation pursuant to h. con. res. 14, the concurrent resolution on the budget for fiscal year 2025," 2025.
- [19] N. DePaula, L. Gao, S. Mellouli, L. F. Luna-Reyes, and T. M. Harrison, "Regulating the machine: An exploratory study of us state legislations addressing artificial intelligence, 2019–2023," p. 815–826, 2024. <https://dx.doi.org/10.1145/3657054.3657148>
- [20] A. Hopkins and S. Booth, "Machine learning practices outside big tech: How resource constraints challenge responsible development," in *AIES 2021: Proceedings of the 2021 AAAI/ACM Conference on AI, Ethics, and Society*. ACM, 2021, Conference Proceedings, p. 134–145. <https://dx.doi.org/10.1145/3461702.3462527>
- [21] M. C. Jensen and W. H. Meckling, "Theory of the firm: Managerial behavior, agency costs and ownership structure," *Journal of Financial Economics*, vol. 3, no. 4, pp. 305–360, 1976. [https://doi.org/10.1016/0304-405X\(76\)90026-X](https://doi.org/10.1016/0304-405X(76)90026-X)
- [22] R. Derakhshan, R. Turner, and M. Mancini, "Project governance and stakeholders: a literature review," *International Journal of Project Management*, vol. 37, no. 1, pp. 98–116, 2019. <https://doi.org/10.1016/j.ijproman.2018.10.007>
- [23] International Organization for Standardization, *Project, programme and portfolio management—Guidance on project management (ISO Standard No. 21502:2020-12)*. ISO, 2020.
- [24] T. Guggenberger, L. Lämmermann, N. Urbach, A. Walter, and P. Hofmann, "Task delegation from AI to humans: A principal-agent perspective," in *ICIS 2023 Convention, Hyderabad, India*, 2023, Conference Paper.
- [25] B. Wachnik, "Moral hazard in it project completion. a multiple case study analysis," pp. 1557–1562, 2015. <http://dx.doi.org/10.15439/2015F68>
- [26] T. Sheridan, W. Verplank, and T. L. Brooks, "Human and computer control of undersea teleoperators," 1978.
- [27] J. F. J. Hair, W. C. Black, B. J. Babin, and R. E. Anderson, *Multivariate Data Analysis Seventh Edition*. Pearson College Division, 2014.
- [28] B. E. Weller, N. K. Bowen, and S. J. Faubert, "Latent class analysis: A guide to best practice," *J. Black Psychol.*, vol. 46, no. 4, p. 287–311, 2020. <https://doi.org/10.1177/0095798420930932>
- [29] SurveyMonkey, "SurveyMonkey answers to the ESOMAR questions to help buyers of online samples," 2024. [Online]. Available: <https://www.surveymonkey.com/mp/legal/esomar-37/>
- [30] F. Bentley, K. O'Neill, K. Quehl, and D. Lottridge, "Exploring the quality, efficiency, and representative nature of responses across multiple survey panels," in *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems*, 2020, Conference Proceedings, pp. 1–12. <https://dx.doi.org/10.1145/3313831.3376671>
- [31] A. Gordon and K. Thorson, *Dealing with Repeated Measures: Design Decisions and Analytic Strategies for Over-Time Data**. Cambridge University Press, 2024, pp. 532–564. <https://dx.doi.org/10.1017/9781009170123.023>
- [32] P. M. Podsakoff, S. B. MacKenzie, J.-Y. Lee, and N. P. Podsakoff, "Common method biases in behavioral research: A critical review of the literature and recommended remedies," *J. Appl. Psychol.*, vol. 88, no. 5, pp. 879–903, 2003. <https://dx.doi.org/10.1037/0021-9010.88.5.879>