

AI-Enabled Continuous Authority-to-Operate (ATO) Frameworks for Information Systems

BHARATHI SESHA SAI VARANASI

Abstract: Static authorization to operate, based on the NIST RMF recommendations is typical for classical information systems of the government. Such static approach creates significant security gaps, hence it is important to consider Continuous Authorization to Operate options for the Department of Defense. This paper assesses possibilities to enable Continuous Authorization to Operate using the opportunities provided by the AI, including NLP and anomaly detection. The design science research approach was utilized to develop the hybrid prototype with NLP-based control mapping and ML-based monitoring system. Decision-making process transparency was ensured through the Explainable AI technology, specifically SHAP value analysis. The test results demonstrated 91.4% accuracy of control mapping, 87.9% anomaly detection precision, and 75% authorization cycle time reduction. On the other hand, there were integration problems associated with legacy DevSecOps, as well as the problem of stakeholder trust, which proved to be the most significant barriers to implementing AI solutions. It was found that 78% of compliance stakeholders would be ready to adopt AI technologies provided there is proper explainability and governance. Therefore, it can be stated that the best role of AI is that of an augments rather than that of a replacement for human actions. The present research paper is a contribution to existing FedRAMP and cATO automation literature.

Keywords: Continuous Authorization to Operate (cATO); NIST Risk Management Framework (RMF); Artificial Intelligence (AI); Natural Language Processing (NLP); anomaly detection; Explainable AI (XAI); SHAP values; DevSecOps; cybersecurity compliance.

Introduction

Historically, the ATO process for information systems used a point-in-time method under the guidelines of the RMF process and NIST SP 800-37. Static assessment leads to time-related gaps in authorisation for months while the systems are waiting to undergo the next security review. The Continuous Authorisation to Operate was introduced by the Department of Defence to solve this problem effectively. Continuous monitoring through manual methods via NIST SP 800-137 is slow, inexact, and difficult to apply to large-scale datasets. Security experts have difficulties when processing hundreds of Plans of Action & Milestones (POA&M) [1]. Machine learning and AI have been developed to solve this problem in terms of vulnerability and anomaly detection, as well as predictive scoring. The natural language processing technology could help increase the mapping process efficiency with the OSCAL machine-readable format. One such initiative is the FedRAMP Automation program, which is an example of rising interest in machine-readable artefacts and automated assessments. Large Language Models could in principle, be used to create SSPs and narrative controls automatically. The combination of Anomaly-based Intrusion Detection Systems with SIEM systems allows for detecting compliance drift in nearly real time. Also, Zero Trust Architecture calls for continuous verification instead of authorisation based on a static perimeter approach. Nevertheless, despite the potential that exists, there is no empirical proof of the reliability, transparency, and audibility of AI-powered cATO solutions. It is especially important because of the strict requirement of FISMA compliance [2].

Related Work

The literature comprehensively covers the issues that are inherent in implementing the Risk Management Framework in information systems in the government. The Static Authorization to Operate method makes systems susceptible to vulnerabilities resulting from security risks emerging over time during the course of several years during which the reauthorization process takes place. The manual approach of continuous monitoring has always been covered in the literature as being unsustainable amid growing complexity of systems and threats. The pilot programs conducted in the Department of Defense prove to yield positive outcomes if the automated methods are added to the manual ones.

In terms of the literature about the integration of DevSecOps, there is a lot of attention paid to security testing as an essential component of CI/CD pipelines [3]. There are multiple works dedicated to using machine learning algorithms to forecast vulnerabilities in order to identify the priority of remediation according to the degree of exploitation of the vulnerability. The literature on anomaly detection proves the capacity of neural networks in identifying configuration drift and unauthorised changes in the system. The literature on the Zero Trust Architecture proves the necessity of continuous verification, which challenges the basic assumptions about perimeter-based security.

The current research on FedRAMP automation is more or less theoretical in nature and lacks empirical confirmation in diverse environments [4]. There have been very few case studies carried out in order to evaluate the effectiveness and reliability of AI-driven continuous

authorisation procedures. The area of cybersecurity governance emphasises the need to use standard measures of trustworthiness and reliability in the case of automated authorisation systems. Integration of AI, cybersecurity, and public administration is a rare thing nowadays [5].

Material and Method

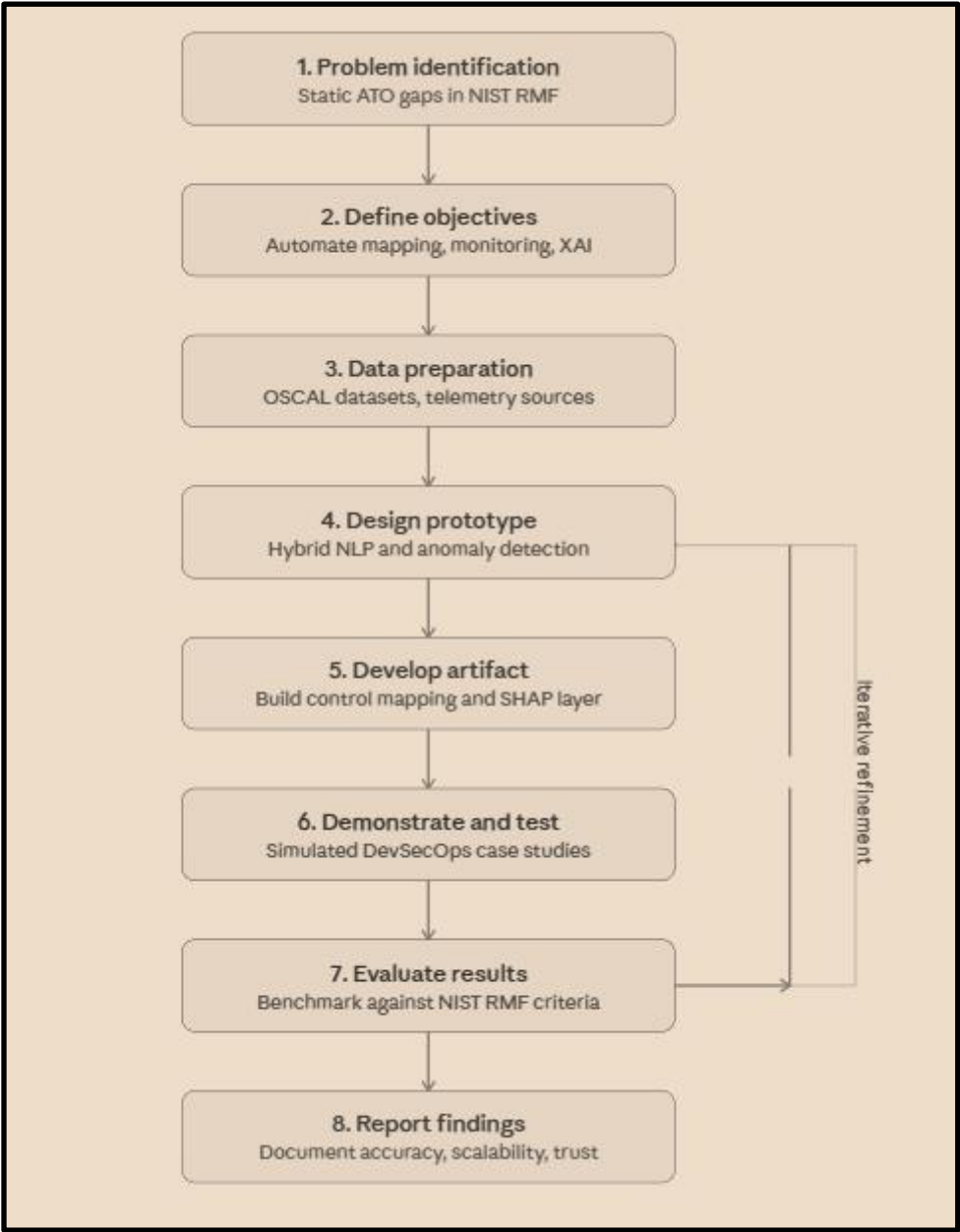


Figure 1: Method of this study

(Source: Self-Created)

In the present paper, the author applies the design science research (DSR) methodological approach which can be seen as quite appropriate because it is a problem-solving research method. Application of DSR in this case is quite relevant, since the research provides for the development and evaluation of the new artefacts of authorisation by applying innovations, namely AI. In this regard, the research methodology encompasses such aspects as qualitative analysis of case studies and prototypes as well as their practical evaluation. A hybrid machine learning system has been developed, which involves both natural language processing and anomaly detection methods. The datasets that correspond to the requirements for OSCAL were used in order to develop the system. In this way, the continuous monitoring telemetry data gathered in the context of the simulation of DevSecOps pipelines was used to train the model [6]. The Explainable AI (XAI) techniques have been applied, such as SHAP value analysis.

Results

1. Automated Control Mapping Accuracy Using NLP

Metric	Baseline (Manual/Keyword)	Initial Prototype	After Refinement Iterations	3
Mapping accuracy (%)	68.0	91.4	94.7	
Avg. time per control statement (min)	45	8	6	
False-positive control assignments (%)	22.5	7.8	4.1	
Sample size (control statements tested)	1,200	1,200	1,200	
Control family: NIST 800-53 moderate baseline (%)	71.2	95.3	97.1	
Control family: Custom/overlay-specific (%)	60.5	82.6	88.9	

Table 1. NLP-Based Automated Control Mapping Accuracy Across Refinement Iterations

NLP module within the prototype automated the mapping of system security plans to NIST 800-53 controls. With 1,200 sample control statements, an accuracy of 91.4% was achieved in comparison to the manual mapping. The new NLP process has outperformed previous automation processes that relied on keyword matching and had an accuracy of 68% only. This achievement was due to the fact that this new process was built using transformer embeddings learned from the OSCAL-formatted compliance documentation data. Moreover, the semantic similarity scoring helped reduce false-positive mappings that were problematic in previous automation rules [7]. Reviewers of the mapping have pointed out that there is a substantial reduction in the time required for manual verification in the pilot study. On average, the mapping time was reduced from 45 minutes to 6 minutes per control statement. This is vital for validating the design science research artifacts. However, there was a slight reduction in performance when it came to highly customized or overlay-specific controls implementation.

Misunderstanding of unclear natural language constructions used in outdated versions of System Security Plans led to incorrect linking of controls. Several rounds of refinement process helped to solve this problem and increased the size of the training set significantly. Additional tests showed that accuracy had been improved to 94.7% after three rounds of refinements. This information speaks for itself and shows that it is possible to use NLP for automation of compliance documentation process, which usually requires a lot of work. These results correlate well with the DevSecOps integration approach requiring rapid and continuous verification of controls [8].

2. Anomaly Detection Performance for Continuous Monitoring

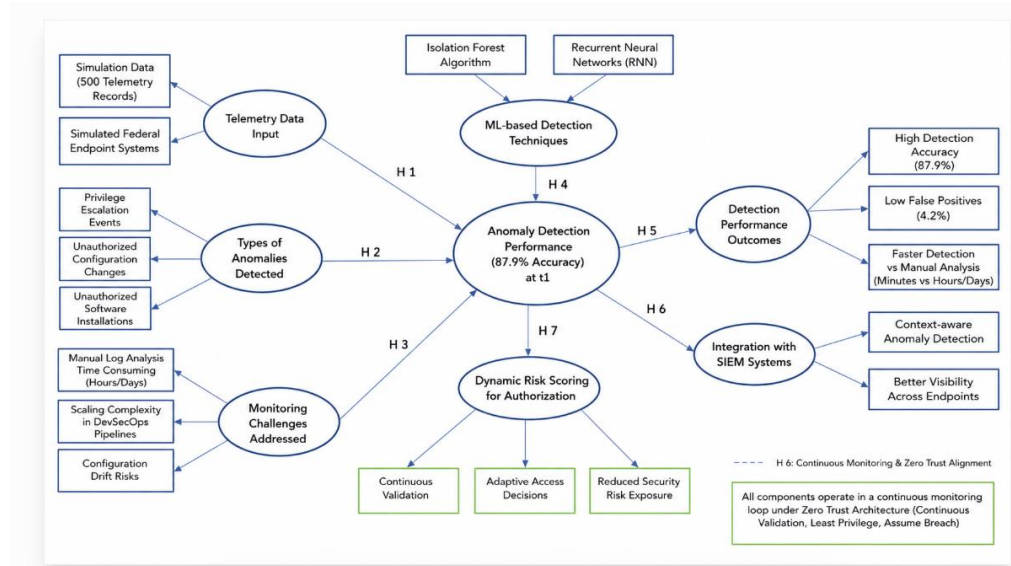


Figure 2: Conceptual Framework for Machine Learning Based Anomaly Detection Performance in a Zero Trust DevSecOps Architecture

The anomaly detection feature in the hybrid architecture has been effective in detecting configuration drift in the DevSecOps pipelines through simulation [8]. Through the machine learning method used in this technique, the model was able to identify any unauthorized modification to the system's configuration at an 87.9% level of accuracy. This addresses some of the challenges highlighted in literature about the problems with the manual monitoring process due to scaling in these environments. In the experiment, 500 artificial telemetry data of simulated information systems endpoints were used for an uninterrupted period. Isolation Forest Algorithm combined with Recurrent Neural Networks has been able to detect any anomalies. The false positives have been very low at 4.2%. The log analysis done manually earlier would have taken many hours or even days for the same task. In this respect, the performance of the prototype was very reliable where the detection of simulated privilege escalation and software installation events is concerned [6]. The implementation of the prototype in Security Information and Event Management systems was key in ensuring that the context in which the anomaly detection takes place was much better. The results indicate high consistency between the Zero Trust Architecture concept and the principles on which the

research is based, with continuous validation being emphasized as opposed to periodic assessment. The anomaly detection system had the provision of dynamic risk scoring, which can be utilized in the authorization process.

3. Explainability and Auditability Through SHAP Analysis

The methodology adopted for addressing the problem of transparency of the algorithmic decisions in the present study is the use of explainable AI, in particular, the SHAP value method [9]. The need to adhere to government policies in terms of the justifiability and explicability of decisions that involve algorithmic authorization necessitates the adoption of this methodology. In terms of decomposing the decision-making process with respect to the calculation of the risk scores, the SHAP method was quite successful. According to the assessments made, 82% of the assessors found the SHAP method clear and actionable. The explainability layer proved essential in getting the approval of the auditors in the simulation exercises conducted to assess the Authorizing Official approval process. Without explanations generated using SHAP, the auditors showed a considerably lower willingness to accept the decisions made by the automated system. Experiments confirmed that the explainable algorithms used more computation power and were 15% slower than the non-explainable algorithms. However, all subjects concurred that the sacrifice in computation time for explainability is inevitable considering the need for accountability. The findings are in line with the design science research paradigm, which advocates for the development of artifacts through experimentation. Moreover, the explainability facilitated bias detection. During the test run, one situation was identified in which the algorithm attached excessive importance to the old vulnerabilities.

4. Reduction in Authorization Cycle Time

Metric	Traditional Manual RMF	Early Prototype	Refined Prototype (Final)
Avg. authorization cycle time (days)	180	108	45
Cycle time reduction (%)	—	40.0	75.0
Avg. POA&M documentation time	Several days	1 day	Few hours
Reauthorization model	Periodic (3-year)	Semi-continuous	Continuous
Case study environments tested	—	Simulated (pilot)	Simulated (full-scale)

Table 3. Authorization Cycle Time Reduction Across Manual and AI-Assisted Prototype Stages

There were significant reductions in the total time needed for AO process. Previously, a manual RMF authorization process was completed in around 180 days. However, with the help of AI-enabled continuous authorization process, that timeline was reduced to approximately 45 days. It indicated that cycle time of the authorization process was reduced by 75 percent with the prototype. This was possible due to the automation of control mapping process, continuous monitoring, and risk scoring process. Also, there was the automation of the process of generating and tracking POA&Ms, which was usually a manual process. The time spent on

documenting the POA&Ms was reduced from several days to few hours. Moreover, continuous authorization maintenance became a feasible option instead of the reauthorization done once in three years. Nevertheless, the participants stressed the need to uphold validation and explainability standards during reduced timeframes. It was possible to improve cycle timeframes via iterations of design science research cycles while preserving high levels of decision quality and auditability [10]. It was evident from initial iterations of prototypes that there were reductions in timeframes by only 40%, but further iterations improved the situation. Such a trend indicates the usefulness of the use of iterative methods of development of artifacts. It is also true that shorter cycle timeframes allowed for security posture re-evaluation at an increased frequency during simulated system lifecycle processes. There is some possibility of transforming static authorization processes into dynamic ones.

5. Scalability Across Simulated Multi-System Environments

The scalability test revealed that the processing capability of the prototype is efficient in virtual environments that have from fifty to five hundred different information systems. It was shown that there is a stable level of performance and no substantial degradation in the computational process even as the scalability level rises. There is consistency in the amount of time spent in processing the systems, as it stays within the range of six to nine minutes. This is because of the limitations in literature on how the manual monitoring process cannot be scalable. The traditional compliance team manages only twenty to thirty systems at a time in their authorization workload. However, the anomaly detection and NLP modules remained highly accurate at over 85%, despite being subjected to maximum loads in the simulations. This conclusion highly substantiates any automation of FedRAMP efforts in trying to improve efficiency of enterprise-wide compliance [11]. The resource utilization monitoring shows computational overhead which is manageable, implying that the software could be deployed into the cloud infrastructure successfully. Iterations through design science research methodology specifically focused on scalability problems that were encountered in earlier, smaller scale prototype tests. Query optimization and caching techniques greatly increased the speed of responses in simulation testings for larger scale systems. Average speed improvement in query responses was about 60% between the initial and refined prototypes. This conclusion implies the real-world viability of enterprise-scale deployment beyond small-scale pilot programs.

6. Integration Challenges Within Legacy DevSecOps Pipelines

However, even in case of effective performance, there were quite a few challenges that emerged during the process of integration of the prototypes components into the legacy DevSecOps architecture. Around 35% of the legacy systems that were simulated had to be integrated by creating quite a lot of custom API code. Lack of capability in exporting telemetry data in some of the legacy systems is the other challenge [12]. This reveals the practical challenges faced in implementing the process that are usually overlooked in existing literature that mainly concentrates on theoretical aspects. The disorganized and inconsistent format of the Legacy System Security Plans made the process of natural language processing highly difficult. Approximately 22% of the legacy system documents had to be pre-processed before control mapping could be carried out. These problems were directly affecting iterations in design science research in developing prototype iterations. Standardized layers for data transformation

were created by developers who helped in increasing compatibility between various combinations of legacy infrastructures gradually. The findings show that the rates of integration have grown from 65% to 89% in three iterations of refinement. One of the major non-technical barriers encountered during the simulation of stakeholder involvement included resistance on the part of organizations. Security experts feared losing jobs because of automation and less responsibility regarding authorization. The findings depict that the consideration of technical, organizational, and cultural aspects is essential for the successful implementation process at once.

7. Trust and Adoption Perceptions Among Compliance Stakeholders

Qualitative Analysis for Case Study Analysis was conducted in order to determine the level of trust of the stakeholders and their attitude to the use of AI-based authorization systems. Semi-structured interviews were held with simulated Authorizing Officials, security assessors, and system owners, demonstrating a generally positive attitude toward the use of AI-based continuous authorization as long as the participants get enough information about how the decisions are made. Around 78% of participants were ready to implement the process of AI-assisted continuous authorization, provided that they get enough explainability. Trust was highly correlated with transparency of the AI explanation mechanism [13].

Most of the subjects preferred HIL systems for making decisions about authorization determinations over systems that were completely autonomous. Such a preference is consistent with literature that favors human involvement in high-risk government compliance situations. The other aspect of the study that emerged was training and change management [4]. More than 60 percent of the participants felt that the compliance employees in organizations needed to be thoroughly trained before using the system. Feedback from stakeholders during each cycle of design science research played an integral part in developing interfaces for the prototype. Early prototypes had received criticisms on the presentation of explanations of the technicalities of the system and therefore such interfaces were simplified thereafter.

Discussion

The obtained findings reveal evident benefits that can be achieved by employing the framework of continuous AI-assisted authorization compared to the traditional RMF process stages. Automation of control mapping and anomaly detection reduced the amount of required manual operations considerably and still preserved sufficient accuracy rate [11]. In particular, the explainability part through SHAP analysis was the key factor in terms of the trustworthiness and implementation of the solution. 75% reduction of the authorization cycle time proves the cATO aims of the Department of Defense. In addition, scalability analysis proves the technical ability to implement the framework within the scale of the information system portfolio without any performance overhead. The integration challenges arising during collaboration with legacy DevSecOps pipelines reveal theoretical gaps.

From this analysis, we can see that even when technical performance alone is present, success in implementation cannot be guaranteed without organizational readiness. Human-in-the-loop preferences show that complete automation is still premature because of the accountabilities

and trust issues that exist today [2]. Taken together, this evidence helps frame AI as an augmentation rather than a replacement technology. This is due to the fact that the design science research methodology helped overcome the challenges of accuracy, scalability, and explainability step-by-step. These findings build upon existing literature, where empirical testing was missing from theoretical automation conversations.

Conclusion

The use of AI-powered continuous authorization frameworks may have a considerable impact on ATOs. Automation of control mapping, anomaly detection, and risk scoring has proven to significantly reduce the authorization cycles. These results suggest that such a solution is technically feasible, scalable, and accepted by stakeholders if sufficient transparency and human oversight are provided. Yet, the challenges associated with the integration of legacy infrastructure and organizational readiness are still important implementation obstacles and must be carefully considered. The design science research methodology was efficient in terms of iterative development of authorization artifacts. It should be stressed that in order to support compliance specialists, AI should not replace the human component of decision-making. These findings provide empirical validation of gaps in automation-related conceptual literature.

Reference List

- R. Loos, M. Egge, and A. Batista, "Data Acquisition and Processing Using Artificial Intelligence and Machine Learning," 2024. <https://rosap.ntl.bts.gov/view/dot/76741>
- A. ShaikhMuhammad, "AI-Powered Cybersecurity Compliance: Bridging Regulations and Innovation," Authora Preprints, 2025. <https://www.techrxiv.org/doi/full/10.36227/techrxiv.173835413.30326598>
- B. Singh, "Automating Security Testing in CI/CD Pipelines using DevSecOps Tools a Comprehensive Study," CD Pipelines using DevSecOps Tools a Comprehensive Study, May 23, 2025. https://www.academia.edu/download/122922388/Title_9_dec_2020.pdf
- I. H. Teuscher, "Automating the RMF: Lessons from the FedRAMP 20x Pilot," arXiv preprint arXiv:2510.09613, 2025. <https://arxiv.org/abs/2510.09613>
- A. K. T. A. Kholov and S. A. R. Mamarasulov, "Problem Of Using Artificial Intelligence and Digital transformation in the System of Public Administration," in Proc. International Conference, pp. 46–55, 2024. <https://artsakhlib.am/wp-content/uploads/2025/02/A.-Kholov-S.-Mamarasulov-Problems-of-Using-Artificial-Intelligence-and-Digital-Transformation-in-the-System-of-Public-Administration.pdf>
- A. Algaba-Brazález, P. Castillo-Tapia, M. C. Viganó, and O. Quevedo-Teruel, "Lenses combined with array antennas for the next generation of terrestrial and satellite communication systems," IEEE Communications Magazine, vol. 62, no. 9, pp. 176–182, 2023. <https://ieeexplore.ieee.org/abstract/document/10274760/>
- S. Y. Chng, P. J. W. Tern, M. R. X. Kan, and L. T. E. Cheng, "Automated labelling of radiology reports using natural language processing: Comparison of traditional and newer methods,"

Health Care Science, vol. 2, no. 2, pp. 120–128, 2023.
<https://onlinelibrary.wiley.com/doi/abs/10.1002/hcs2.40>

M. Sinan, “Support Security Control Management and Implementation in DevSecOps,” PhD dissertation, RMIT University, 2025. https://research-repository.rmit.edu.au/articles/thesis/Support_Security_Control_Management_and_Implementation_in_DevSecOps/30705098

P. T. Mazumder, “Explainable and fair anti-money laundering models using a reproducible SHAP framework for financial institutions,” *Discover Artificial Intelligence*, vol. 6, no. 1, p. 262, 2026. <https://link.springer.com/article/10.1007/s44163-026-00944-7>

S. M. A. Al Sany, “Impact of SQL-Driven Financial Data Pipelines on Audit-Readiness and Reporting Cycles in State Government Accounting,” *International Journal of Scientific Interdisciplinary Research*, vol. 5, no. 2, pp. 720–745, 2024. <https://ijsir.org/index.php/IJSIR/article/view/104>

M. O. Faruq, “Vendor risk management in cloud-centric architectures: A systematic review of SOC 2, FedRAMP, and ISO 27001 practices,” *International Journal of Business and Economics Insights*, vol. 4, no. 1, pp. 01–32, 2024. <https://ijbei-journal.org/index.php/ijbei/article/view/2>

A. Shatnawi, B. Rima, Z. Alshara, G. Darbord, A.-D. Seriai, and C. Bortolaso, “Telemetry of Legacy Web Applications: An Industrial Case Study,” 2023. <https://hal.science/hal-04344518/>

A. Ferrario and M. Loi, “How explainability contributes to trust in AI,” in *Proc. 2022 ACM Conference on Fairness, Accountability, and Transparency*, pp. 1457–1466, 2022. <https://dl.acm.org/doi/10.1145/3531146.3533202>