

Insider Threats in Cybersecurity: Modeling Patterns with MITRE Profiles through Technical, Organizational, and Individual (TOI) Factors

Ling Leung¹, Suzanne Kieffer¹, and Esma Aïmeur²

¹ Université catholique de Louvain, Louvain-la-Neuve, Belgium
{ling.leung,suzanne.kieffer}@uclouvain.be

² Université de Montréal, Montréal, Canada
aimeur@iro.umontreal.ca

Abstract. Insider threats arise from the interplay of technical, organizational, and individual conditions, yet most analytical frameworks address these dimensions separately. This paper presents an exploratory, data-driven approach that combines structured document analysis with a locally deployed large language model in order to identify explicit and implicit technical, organizational, and individual (TOI) factors in textual descriptions of insider-related events. The TOI coding structure is derived from the systematic review by Nassir et al. A hybrid reference corpus of sourced and synthetic indicators is generated under a controlled prompting protocol and verified by hand before being used to calibrate a separate detection model, so that generation and detection remain methodologically independent. Recurring configurations of factors are then compared with the insider profiles defined by MITRE (malicious, negligent, mistaken, and outsmarted) to examine whether specific factor combinations align with distinct insider types. The study is in progress; this paper reports the analytical framework and the corpus construction procedure.

Keywords: Insider threat · MITRE taxonomy · Socio-technical factors · Large language models · Cybersecurity.

1 Introduction

Concerns about insider threats in the literature are not new, and there is an impressive body of knowledge in this broad field. Unlike an external attacker, an insider holds legitimate access to the organization’s facilities and is familiar with its critical assets, which makes malicious activity easier to carry out and to conceal [1]. Detecting insider attacks is correspondingly harder than detecting externally originated ones, since the insider is an authorized user and the boundary between legitimate and malicious activity is difficult to draw [1].

Indeed, insider risk now ranks among the costliest categories of security incident: organizations reported an average annual cost of USD 19.5 million in 2025,

up from USD 17.4 million the previous year, with negligence the single largest cost driver [20].

Yet, despite the prominent role of human in insider risk, cybersecurity research has historically devoted considerably more attention to technological safeguards than to the people interacting with them. This imbalance is noteworthy because human actions remain implicated in a substantial proportion of security incidents [21].

2 Background

The concept of an insider lacks an universally accepted definition within the cybersecurity literature [16]. Homoliak et al. [13] note that most definitions rely on static descriptors that emphasize the following attributes: access, trust, knowledge, and awareness of security policies. Generally, an insider is defined as a trusted individual [4] — typically a current or former employee, contractor, or business partner [7] — who possesses legitimate access to an organization’s networks, systems, or sensitive data [7], has knowledge of its security system [19, 10, 3], and is aware of the organization’s policies, procedures, technology, and vulnerabilities [7].

Building on this understanding, the concept of an insider threat refers to a trusted individual who possesses authorized access to an organization’s systems, data, or infrastructure and misuses this privilege [11], or whose access results in misuse, in ways that compromise the confidentiality, integrity, or availability of information assets [8]. Such misuse may be intentional, for example through data theft, fraud, or sabotage, or unintentional, as in the case of errors, negligence, or policy violations.

According to Reveraert and Sauer [24], definitions of insider threats can be grouped into two conceptual perspectives: the harm-oriented and the privilege-oriented. Both make intentionality central, but they locate it differently: the harm-oriented perspective in the intention to cause harm, the privilege-oriented perspective in the intention to misuse privileged access.

- The harm-oriented perspective defines insider threats primarily by an individual’s intent to cause harm to the organization. In narrow interpretations, only malicious insiders are considered—typically described as individuals with special access or knowledge who intend to negatively impact confidentiality, integrity, or availability [8]. Broader definitions also classify all harmful incidents involving insiders as insider. As Reveraert and Sauer [24] highlight, these varying scopes create conceptual ambiguity, with definitions ranging from overly restrictive to overly inclusive depending on how intentionality is framed.
- The privilege-oriented perspective instead focuses on the insider’s intention to misuse authorized access to, or knowledge of, organizational assets, irrespective of whether harm is the intended outcome. Reveraert and Sauer favor this perspective, but observe that existing privilege-oriented definitions

do not draw an explicit line between intentional and unintentional misuse. They therefore introduce a further distinction: unintentional misuse of privilege constitutes an *insider hazard*, whereas intentional misuse constitutes an *insider threat*. The notion of “misuse” is nevertheless often interpreted subjectively and presumes the existence of clear organizational rules or policies, although such definitions may be missing, vague, or inconsistently enforced in practice. Beyond organizational policy, social and ethical factors further complicate interpretation, often making it difficult to distinguish acceptable insider behaviors from unacceptable ones [23].

The definitions and perspectives discussed above clarify the conceptual foundations of insider threats. However, they remain largely theoretical and provide limited practical guidance for identifying or explaining insider incidents in real-world contexts. To address this limitation, researchers have sought to operationalize these concepts through taxonomies that classify insider profiles and contributing factors, thereby linking theoretical constructs with observable characteristics.

3 Related Work

3.1 Historical Evolution of Insider Threat Research

Research on insider threats has progressed from a narrow technical focus toward a multidisciplinary perspective integrating human, organizational, and technical domains.

From the 1990s to the early 2000s, research predominantly emphasized technical detection methods such as rule-based and anomaly-based intrusion detection systems, access controls, and log analysis [2, 9, 13]; this technical line of work continues today in user-behavior analytics [25]. During this period, insider threats were largely conceptualized as technical adversaries whose actions could be detected through system monitoring and privilege restriction. However, recurring incidents exposed the limitations of purely technology-centered approaches [15] and highlighted the neglect of human and contextual influences that were often overlooked or deprioritized [21].

Research into the “weakest link” phenomenon reinforced this awareness. Numerous studies have shown that humans often constitute the most fragile component of security systems, as their actions can be easily exploited and are sometimes neglected in security design and evaluation. This vulnerability has been linked to limited security awareness, exposure to social engineering, human error, and inconsistent compliance with security policies. Drawing on Kuhn’s notion of paradigm shifts [17], Yan et al. [27] describe a broader reorientation of cybersecurity research toward the capabilities and vulnerabilities of ordinary users, rather than an exclusive emphasis on technological robustness.

As Hunker and Probst [14] observe, this recognition led to the emergence of two main schools of thought: a technical school, encompassing system-based defenses such as access control, intrusion detection, and monitoring mechanisms;

and a sociological-organizational school, which examined the motivational, psychological, and cultural factors influencing insider behavior. This conceptual division reflected the growing awareness that technical controls alone were insufficient to address the complexity of insider risks.

Subsequent research revealed significant overlap between these perspectives, prompting their convergence into socio-technical approaches that analyze human, organizational, and technological factors as interdependent components of a unified system. This convergence marked a decisive turning point in the field, shifting the focus from viewing insiders as purely technical actors to understanding them as individuals shaped by organizational dynamics and personal motivations [14].

3.2 Technical, Organizational, and Individual Factors

As discussed in the previous subsection, restricting insider-threat research to technical indicators limits understanding of early warning signs and underlying causes [12]. Insider threats arise from the interplay of multiple factors rather than from isolated intentions. To integrate these factors into analysis, it is necessary to determine how they can be observed and how their interactions can be systematically examined.

Building upon this need for multidimensional understanding, Nassir et al. [18] conducted a systematic review entitled *Revealing the Multi-Perspective Factors Behind Insider Threats in Cybersecurity*, which provides one of the most comprehensive syntheses of insider-threat determinants to date. Following the PRISMA protocol and applying qualitative content analysis to 32 studies (2014–2023), the authors derived an empirical classification encompassing technical, organizational, and human dimensions. Throughout this paper we refer to these three dimensions as technical, organizational, and individual (TOI) factors; the “individual” dimension corresponds to what Nassir et al. term human factors.

Human factors were grouped under four themes—inadequate security training and awareness, mental health issues, personal problems, and negative personality traits—each associated with sub-themes such as stress, disgruntlement, addiction, or carelessness. Technical factors comprised three themes: lack of resources, inadequate monitoring systems, and weak security evaluation and validation. Organizational factors comprised four: organizational practice issues, insufficient risk management, poor management systems, and workplace stressors. This classification highlights how vulnerabilities emerge from the convergence of behavioral, technical, and organizational deficiencies, and it distinguishes itself by its methodological rigor, empirical scope, and balanced integration of socio-technical dimensions. It is therefore adopted in the present work as a robust and comprehensive analytical structure for examining insider-threat incidents in a multidimensional manner.

3.3 Taxonomies of Insider Threats

To extend this multidimensional analysis, several taxonomies have been developed to translate these insights into structured representations of insider behaviors, intentions, and profiles. These taxonomies help connect theoretical constructs with real-world observations and support more effective detection, prevention, and mitigation strategies. Among them are the following.

1. **The CERT taxonomy.** Developed by the CERT Insider Threat Center at Carnegie Mellon University, this classification rests on an empirical database of insider cybercrime built since 2001, comprising more than seven hundred U.S. cases at the time of the reference guide [5]. It divides malicious activity into three crime types—information technology sabotage, theft of intellectual property, and fraud—to which national security espionage is added as a fourth category in CERT’s case-analysis work [5, 7]. Each category represents recurring behavioral patterns and organizational conditions associated with insider misconduct, and is used as a framework for incident analysis and mitigation. Unintentional insider threats are examined separately in CERT’s research [6, 12]; they involve accidental disclosure of information, mishandling of sensitive data, loss of physical or digital assets, and social-engineering compromise such as phishing or execution of malicious code.
2. **The taxonomy of Homoliak et al.** [13]. Homoliak et al. introduced a structural and orthogonal taxonomy organized around the 5W1H analytical framework: *who* (the type or actor profile involved), *what* (the action or activity undertaken by the insider), *when* (temporal aspects including timing and duration), *where* (the physical or logical environment in which the incident occurs), *why* (the insider’s motivation or driving cause), and *how* (the method or technique used to carry out the action). This taxonomy integrates technical and behavioral dimensions, supporting systematic forensic analysis and structured data collection. It is designed to compare and analyze existing insider threat classifications.
3. **The MITRE taxonomy** [17]. MITRE’s human-focused typology refines the distinction between malicious and non-malicious insiders by subdividing the latter into three types: negligent (familiar with the rules but careless), mistaken (genuine errors arising because security is not the employee’s primary task and because demands increase while resources decrease), and outsmarted (manipulated by adversaries exploiting so-called “human zero-day” vulnerabilities). MITRE stresses that the distinction has operational consequences: responding to a mistaken or outsmarted employee as if they were negligent—through remedial training or escalation to a supervisor—can itself increase risk. This human-centered typology captures both intentional and unintentional behaviors.³
4. **The VISTA taxonomy** [22]. The VISTA (InclusiVe InSider Threat tAxonomy) model offers a holistic, multidimensional view of insider threats. It

³ MITRE refers to these categories as insider *types*; for consistency with the other taxonomies reviewed here, we use *profiles* throughout.

classifies behavior along four dimensions: knowledge (unaware vs. aware), compliance (compliant vs. non-compliant), volition (deliberate vs. incidental harm), and goal orientation (for self, society, ideals, or malice). From these dimensions, nineteen insider types are derived and grouped into seven categories: untrained, fallible, disempowered, whistleblower, misbehavior, ideology, and malicious. Each category is linked to specific mitigation strategies, integrating technical, organizational, environmental, and motivational factors.

The reviewed taxonomies reflect the progressive broadening outlined in Section 3.1. The CERT taxonomy established an empirical foundation by categorizing malicious insider incidents into recurring behavioral patterns, but excluded unintentional and socio-technical dimensions. Homoliak et al. expanded the analytical scope through the 5W1H framework, enabling systematic comparison of technical and behavioral aspects across incidents. The MITRE taxonomy enhanced practical applicability by combining theoretical insight with operational use, redefining insider types according to the employee’s awareness and degree of agency, and distinguishing negligent, mistaken, and outsmarted insiders to support real-incident analysis.

VISTA [22] further advanced the field by offering a detailed classification of insider types and linking each category to specific mitigation strategies.

MITRE’s classification cuts across both perspectives identified by Reveraert and Sauer [24]. Malicious insiders exemplify deliberate intent to harm, whereas negligent, mistaken, and outsmarted insiders create risk without such intent and, in Reveraert and Sauer’s terms, fall largely on the side of insider hazards rather than insider threats. The outsmarted category further extends this continuum by encompassing situations in which employees unintentionally create risk through unforeseen circumstances—what MITRE refers to as human zero-day vulnerabilities. This distinction matters because each category entails distinct risks and different forms of harm: organizations can effectively mitigate insider incidents only when they understand the specific type of threat they face and apply appropriately tailored countermeasures.

The present study adopts the MITRE taxonomy for its balance between conceptual inclusiveness and operational clarity. Its well-defined behavioral categories provide a consistent framework for mapping observed actions to insider types and associated intentions, thereby supporting the analysis of real-context incidents.

4 Methodology

This study adopts a data-driven exploratory approach that combines structured document selection, manual content analysis, and pattern-oriented comparative analysis supported by a locally deployed large language model (LLM). The objective is to train the model to recognize explicit and implicit factors associated with insider threats and to examine recurring combinations of these factors that may align with known profiles from the MITRE taxonomy.

The analytical framework is derived from Nassir et al. [18], whose review classifies insider-threat determinants into technical, organizational, and individual (TOI) categories. This taxonomy provides the conceptual basis and initial coding structure for factor identification. No predefined insider typology is applied at this stage, in order to preserve an open, exploratory orientation.

4.1 LLM-Assisted Analytical Procedure

The LLM supports a three-step analytical procedure. First, it is used to detect and classify explicit and implicit factors within textual descriptions of insider-related incidents according to the TOI factor categories. Second, it identifies recurring patterns, defined as co-occurring configurations of TOI factors that repeatedly appear across different incidents. Third, it examines whether certain patterns are associated with, or appear to contribute to, the escalation from insider status to insider-threat behavior, and whether specific configurations of factors correspond to one or more insider profiles defined in the MITRE taxonomy. The mapping between TOI factor configurations and MITRE profiles was conducted manually by the researcher, according to the behavioral definitions and distinguishing characteristics outlined in the MITRE taxonomy. This procedure does not assume predefined causal relationships but instead relies on data-driven pattern discovery, in which the model assists in extracting structured information from textual data to identify the recurrence and escalation of factors contributing to insider-threat behavior.

4.2 Corpus Preparation and Model Calibration

To operationalize the TOI framework while maintaining methodological rigor, a two-stage procedure was implemented, combining prompt-based corpus generation, context engineering, and human validation.

Stage 1 — Prompt-based corpus generation through context engineering. The generation process was implemented using a controlled prompting protocol, referred to here as the *TOI Scenario Generator*, in which each prompt instantiates a single TOI sub-theme and constrains the model to produce indicator texts conforming to a fixed output schema. A dedicated LLM was employed exclusively for corpus construction to ensure a clear separation between generation and analytical modeling. For each TOI sub-theme defined in Tables 2–4 of Nassir et al. [18], the model received as structured input the corresponding `factor_type`, `theme`, `subtheme`, and `definitional_description_and_exemplar_quote`. These structured elements were embedded within the model’s context window (context engineering) to constrain reasoning within realistic and literature-grounded boundaries. The model was instructed to generate four `indicator_text` entries per sub-theme: two explicit (`explicit: true`, `implicit: false`) and two implicit (`explicit: false`, `implicit: true`), each describing a concise, realistic manifestation of the factor. When available, the

LLM produced *sourced indicators* grounded in verifiable materials such as peer-reviewed publications, CERT or ENISA reports, or industry analyses (e.g., the Verizon DBIR). When no eligible material was found, it generated *synthetic indicators* labeled `grounding_mode: synthetic`. All generations followed a predefined JSON schema specifying `explicitness`, `grounding_mode`, and `citation_metadata`. This procedure yielded the hybrid reference corpus used for subsequent validation and analytical calibration.

Stage 2 — Human validation and factor-level verification. Each generated incident was systematically reviewed by the researcher to verify the correctness of explicit and implicit factor identification, the accuracy of labeling, and the coherence with the TOI framework. Validation ensured that sourced indicators remained faithful to their original references and that synthetic indicators accurately represented their designated sub-themes. Only validated examples meeting all schema and quality criteria were retained. The resulting corpus—comprising verified sourced and synthetic incidents—was then used to calibrate a separate local LLM dedicated to TOI factor detection. Maintaining this strict model separation between generation and detection prevented circular reasoning and reinforced methodological independence. Systematic human verification further preserved interpretive consistency and analytical rigor throughout the process.

4.3 Analytical Application and Comparative Interpretation

Once the corpus was finalized and the model calibrated, the model was applied to an independent application corpus composed of previously unprocessed documents describing insider-related events. The analysis relied on context-engineered queries that structured the model’s reasoning process and standardized factor extraction across all documents. The model assisted in identifying explicit and implicit TOI factors from these texts, which then served as the analytical foundation for pattern identification and interpretation.

Following factor extraction, a comparative, pattern-oriented analysis was conducted to identify recurrent configurations of technical, organizational, and individual factors across incidents. The analysis examined which combinations of factors consistently co-occurred before or during the emergence of insider-threat behavior.

At the interpretive stage, the identified patterns were compared with the MITRE insider taxonomy (malicious, negligent, mistaken, outsmarted) to examine potential conceptual and behavioral overlaps. Model outputs were validated against human annotations using standard performance metrics—precision (accuracy of detected factors), recall (coverage of relevant factors), and the F_1 score (harmonic mean of both)—to evaluate the reliability of factor identification and ensure consistency between automated and manual analyses.

5 Results and Discussion

Although results are still in progress, the proposed analytical approach is expected to reveal recurrent patterns linking TOI factors to specific insider profiles. Such patterns may appear as co-occurring conditions (e.g., weak organizational culture combined with technical monitoring gaps) or as recurring trajectories from individual drivers to harmful outcomes. At the same time, the analysis may also reveal variability, with weak or inconsistent regularities, suggesting that insider incidents resist stable patterns. Future work should explore how this systemic, integrative perspective can help organizations move beyond fragmented mitigation strategies toward more context-specific and dynamic defenses.

References

1. Al-Mhiqani, M.N., Ahmad, R., Abidin, Z.Z., Abdulkareem, K.H., Mohammed, M.A., Gupta, D., Shankar, K.: A new intelligent multilayer framework for insider threat detection. *Computers and Electrical Engineering* **97**, 107597 (2022). <https://doi.org/10.1016/j.compeleceng.2021.107597>
2. Anderson, J.P.: Computer security threat monitoring and surveillance. Tech. rep., James P. Anderson Company, Fort Washington, PA (1980)
3. Bishop, M.: Position: Insider is relative. In: *Proceedings of the Workshop on New Security Paradigms*. pp. 77–78. ACM (2005)
4. Brackney, R.C., Anderson, R.H. (eds.): *Understanding the Insider Threat: Proceedings of a March 2004 Workshop*. No. CF-196-ARDA in Conference Proceedings, RAND Corporation, Santa Monica, CA (2004), https://www.rand.org/content/dam/rand/pubs/conf/proc/proceedings/2005/RAND_CF196.pdf
5. Cappelli, D., Moore, A.P., Trzeciak, R.F.: *The CERT Guide to Insider Threats: How to Prevent, Detect, and Respond to Information Technology Crimes (Theft, Sabotage, Fraud)*. SEI Series in Software Engineering, Addison-Wesley Professional, Upper Saddle River, NJ (2012)
6. CERT Insider Threat Team: *Unintentional insider threats: A foundational study*. Tech. Rep. CMU/SEI-2013-TN-022, Software Engineering Institute, Carnegie Mellon University, Pittsburgh, PA (2013)
7. Collins, M.L., Theis, M.C., Trzeciak, R.F., Strozer, J.R., Clark, J.W., Costa, D.L., Cassidy, T., Albrethsen, M.J., Moore, A.P.: *Common sense guide to mitigating insider threats, fifth edition*. Technical Note CMU/SEI-2016-TR-015, Software Engineering Institute, Carnegie Mellon University, Pittsburgh, PA (2016), <https://www.odni.gov/files/NCSC/documents/nittf/20180209-CERT-Common-Sense-Guide-Fifth-Edition.pdf>
8. Collins, M.L., Theis, M.C., Trzeciak, R.F., Strozer, J.R., Clark, J.W., Costa, D.L., et al.: *Common sense guide to prevention and detection of insider threats*. Tech. rep., CERT, Software Engineering Institute, Carnegie Mellon University, Pittsburgh, PA (2016)
9. Denning, D.E.: An intrusion-detection model. *IEEE Transactions on Software Engineering* **SE-13**(2), 222–232 (1987). <https://doi.org/10.1109/TSE.1987.232894>
10. Garfinkel, R., Gopal, R., Goes, P.: Privacy protection of binary confidential data against deterministic, stochastic, and insider threat. *Management Science* **48**(6), 749–764 (2002)

11. Greitzer, F.L., Frincke, D.A., Zabriskie, M.: Social/ethical issues in predictive insider threat monitoring. In: *Information Assurance and Security Ethics in Complex Systems: Interdisciplinary Perspectives*, pp. 132–161. IGI Global, Hershey, PA (2010)
12. Greitzer, F.L., Strozer, J., Cohen, S., Bergey, J., Cowley, J., Moore, A., Mundie, D.: Unintentional insider threat: Contributing factors, observables, and mitigation strategies. In: *2014 47th Hawaii International Conference on System Sciences (HICSS)*. pp. 2025–2034. IEEE (2014). <https://doi.org/10.1109/HICSS.2014.256>
13. Homoliak, I., Toffalini, F., Guarnizo, J., Elovici, Y., Ochoa, M.: Insight into insiders and it: A survey of insider threat taxonomies, analysis, modeling, and countermeasures. *ACM Computing Surveys* **52**(2), 1–40 (2019). <https://doi.org/10.1145/3303771>
14. Hunker, J., Probst, C.W.: Insiders and insider threats: An overview of definitions and mitigation techniques. *Journal of Wireless Mobile Networks, Ubiquitous Computing, and Dependable Applications* **2**(1), 4–27 (2011). <https://doi.org/10.22667/JOWUA.2011.03.31.004>
15. Jeong, J.J., Mihelcic, J., Oliver, G., Rudolph, C.: Towards an improved understanding of human factors in cybersecurity. In: *2019 IEEE 5th International Conference on Collaboration and Internet Computing (CIC)*. IEEE (2019). <https://doi.org/10.1109/CIC48465.2019.00047>
16. Kont, M., Pihelgas, M., Wojtkowiak, J., Trinberg, L., Osula, A.M.: Insider threat detection study. Tech. rep., NATO Cooperative Cyber Defence Centre of Excellence, Tallinn, Estonia (2015), <https://ccdcoe.org/uploads/2018/10/InsiderThreatStudyCDCOE.pdf>
17. MITRE Insider Threat Research & Solutions: MITRE’s human-focused insider threat types. Online resource, <https://insiderthreat.mitre.org/insider-types/> (2022), public Release Case Number 22-1798; last updated 5 June 2022; accessed 5 November 2025
18. Nassir, N.F.M., Abdul Rauf, U.F., Zainol, Z., Abdul Ghani, K.: Revealing the multi-perspective factors behind insider threats in cybersecurity. *Journal of Media and Information Warfare* **17**(2), 65–82 (Oct 2024), <https://jmiw.uitm.edu.my/images/Journal/Vol17No2/Revealing.pdf>
19. Pfleeger, S.L., Predd, J.B., Hunker, J., Bulford, C.: Insiders behaving badly: Addressing bad actors and their actions. *IEEE Transactions on Information Forensics and Security* **5**(1), 169–179 (2010)
20. Ponemon Institute: 2026 cost of insider risks global report. Tech. rep., Ponemon Institute, sponsored by DTEX Systems (2026), <https://ponemon.dtex.ai/>
21. Rahman, T., Rohan, R., Pal, D., Kanthamanon, P.: Human factors in cybersecurity: A scoping review. In: *Proceedings of the 12th International Conference on Advances in Information Technology (IAIT2021)*. Association for Computing Machinery, Bangkok, Thailand (Jun 2021). <https://doi.org/10.1145/3468784.3468789>
22. Renaud, K., Warkentin, M., Pogrebna, G., van der Schyff, K.: VISTA: An inclusive insider threat taxonomy, with mitigation strategies. *Information & Management* **61**(1), 103877 (2024). <https://doi.org/10.1016/j.im.2023.103877>
23. Reveraert, M.: Exploring Insider Threat Awareness and Mitigation: More than the Devil in Disguise. Ph.D. thesis, University of Antwerp, Antwerp, Belgium (2023)
24. Reveraert, M., Sauer, T.: Redefining insider threats: A distinction between insider hazards and insider threats. *Security Journal* **34**(4), 755–775 (2021). <https://doi.org/10.1057/s41284-020-00259-x>
25. Singh, M., Mehtre, B.M., Sangeetha, S.: User behavior based insider threat detection using a multi fuzzy classifier. *Multimedia Tools and Applications* **81**(16), 23757–23778 (2022). <https://doi.org/10.1007/s11042-022-12173-y>