



Securing Artificial Intelligence (SAI); Guide to Cyber Security for AI Models and Systems

Reference

DTR/SAI-0012

Keywords

artificial intelligence, security

ETSI

650 Route des Lucioles
F-06921 Sophia Antipolis Cedex - FRANCE

Tel.: +33 4 92 94 42 00 Fax: +33 4 93 65 47 16

Siret N° 348 623 562 00017 - APE 7112B
Association à but non lucratif enregistrée à la
Sous-Préfecture de Grasse (06) N° w061004871

Important notice

The present document can be downloaded from the
[ETSI Search & Browse Standards](#) application.

The present document may be made available in electronic versions and/or in print. The content of any electronic and/or print versions of the present document shall not be modified without the prior written authorization of ETSI. In case of any existing or perceived difference in contents between such versions and/or in print, the prevailing version of an ETSI deliverable is the one made publicly available in PDF format on [ETSI deliver](#) repository.

Users should be aware that the present document may be revised or have its status changed,
this information is available in the [Milestones listing](#).

If you find errors in the present document, please send your comments to
the relevant service listed under [Committee Support Staff](#).

If you find a security vulnerability in the present document, please report it through our
[Coordinated Vulnerability Disclosure \(CVD\)](#) program.

Notice of disclaimer & limitation of liability

The information provided in the present deliverable is directed solely to professionals who have the appropriate degree of experience to understand and interpret its content in accordance with generally accepted engineering or other professional standard and applicable regulations.

No recommendation as to products and services or vendors is made or should be implied.

No representation or warranty is made that this deliverable is technically accurate or sufficient or conforms to any law and/or governmental rule and/or regulation and further, no representation or warranty is made of merchantability or fitness for any particular purpose or against infringement of intellectual property rights.

In no event shall ETSI be held liable for loss of profits or any other incidental or consequential damages.

Any software contained in this deliverable is provided "AS IS" with no warranties, express or implied, including but not limited to, the warranties of merchantability, fitness for a particular purpose and non-infringement of intellectual property rights and ETSI shall not be held liable in any event for any damages whatsoever (including, without limitation, damages for loss of profits, business interruption, loss of information, or any other pecuniary loss) arising out of or related to the use of or inability to use the software.

Copyright Notification

No part may be reproduced or utilized in any form or by any means, electronic or mechanical, including photocopying and microfilm except as authorized by written permission of ETSI.

The content of the PDF version shall not be modified without the written authorization of ETSI.

The copyright and the foregoing restriction extend to reproduction in all media.

© ETSI 2025.
All rights reserved.

Contents

Intellectual Property Rights	5
Foreword.....	5
Modal verbs terminology.....	5
Introduction	5
1 Scope	6
2 References	6
2.1 Normative references	6
2.2 Informative references.....	6
3 Definition of terms, symbols and abbreviations.....	11
3.1 Terms.....	11
3.2 Symbols.....	13
3.3 Abbreviations	13
4 How to use the present document.....	14
4.1 Purpose	14
4.2 Relationship to ETSI TS 104 223.....	15
5 Guidance on implementation.....	15
6 Examples to meet AI Security Provisions	17
6.1 Principle 1: Raise awareness of AI security threats and risks	17
6.1.1 Provision 5.1.1-1.....	17
6.1.2 Provision 5.1.1-1.1.....	17
6.1.3 Provision 5.1.1-2.....	19
6.1.4 Provision 5.1.1-2.1.....	19
6.1.5 Provision 5.1.1-2.2.....	20
6.2 Principle 2: Design the AI System for Security as well as Functionality and Performance	21
6.2.1 Provision 5.1.2-1.....	21
6.2.2 Provision 5.1.2-1.1.....	23
6.2.3 Provision 5.1.2-2.....	23
6.2.4 Provision 5.1.2-3.....	24
6.2.5 Provision 5.1.2-4.....	25
6.2.6 Provision 5.1.2-5.....	26
6.2.7 Provision 5.1.2-5.1.....	26
6.2.8 Provision 5.1.2-6.....	27
6.2.9 Provision 5.1.2-7.....	28
6.3 Principle 3: Evaluate the threats and manage the risks to the AI system.....	28
6.3.1 Provision 5.1.3-1.....	28
6.3.2 Provision 5.1.3-1.1.....	29
6.3.3 Provision 5.1.3-1.2.....	30
6.3.4 Provision 5.1.3-1.3.....	31
6.3.5 Provision 5.1.3-2.....	32
6.3.6 Provision 5.1.3-3.....	32
6.3.7 Provision 5.1.3-4.....	33
6.4 Principle 4: Enable human responsibility for AI systems.....	33
6.4.1 Provision 5.1.4-1.....	33
6.4.2 Provision 5.1.4-2.....	35
6.4.3 Provision 5.1.4-3.....	35
6.4.4 Provision 5.1.4-4.....	36
6.4.5 Provision 5.1.4-5.....	36
6.5 Principle 5: Identify, track, and protect assets.....	37
6.5.1 Provision 5.2.1-1.....	37
6.5.2 Provision 5.2.1-2.....	38
6.5.3 Provision 5.2.1-3.....	39
6.5.4 Provision 5.2.1-3.1.....	39

6.5.5	Provision 5.2.1-4.....	40
6.5.6	Provision 5.2.1-4.1.....	41
6.5.7	Provision 5.2.1-4.2.....	42
6.6	Principle 6: Secure the infrastructure	42
6.6.1	Provision 5.2.2-1.....	42
6.6.2	Provision 5.2.2-2.....	43
6.6.3	Provision 5.2.2-3.....	44
6.6.4	Provision 5.2.2-4.....	45
6.6.5	Provision 5.2.2-5.....	45
6.6.6	Provision 5.2.2-6.....	46
6.7	Principle 7: Secure the supply chain.....	47
6.7.1	Provision 5.2.3-1.....	47
6.7.2	Provision 5.2.3-2.....	48
6.7.3	Provision 5.2.3-2.1.....	49
6.7.4	Provision 5.2.3-2.2.....	49
6.7.5	Provision 5.2.3-3.....	50
6.7.6	Provision 5.2.3-4.....	51
6.8	Principle 8: Document Data, Models, and Prompts	51
6.8.1	Provision 5.2.4-1.....	51
6.8.2	Provision 5.2.4-1.1.....	52
6.8.3	Provision 5.2.4-1.2.....	52
6.8.4	Provision 5.2.4-2.....	53
6.8.5	Provision 5.2.4-2.1.....	54
6.8.6	Provision 5.2.4-3.....	54
6.9	Principle 9: Conduct appropriate testing and evaluation	55
6.9.1	Provision 5.2.5-1.....	55
6.9.2	Provision 5.2.5-2.....	56
6.9.3	Provision 5.2.5-2.1.....	57
6.9.4	Provision 5.2.5-3.....	57
6.9.5	Provision 5.2.5-4.....	58
6.9.6	Provision 5.2.5-4.1.....	59
6.10	Principle 10: Communication and processes associated with End-users and Affected Entities	59
6.10.1	Provision 5.3.1-1.....	59
6.10.2	Provision 5.3.1-2.....	60
6.10.3	Provision 5.3.1-2.1.....	61
6.10.4	Provision 5.3.1-2.2.....	61
6.10.5	Provision 5.3.1-3.....	62
6.11	Principle 11: Maintain Regular Security Updates, Patches, and Mitigations	63
6.11.1	Provision 5.4.1-1.....	63
6.11.2	Provision 5.4.1-1.1.....	64
6.11.3	Provision 5.4.1-2.....	64
6.11.4	Provision 5.4.1-3.....	65
6.12	Principle 12: Monitor the system's behaviour	65
6.12.1	Provision 5.4.2-1.....	65
6.12.2	Provision 5.4.2-2.....	67
6.12.3	Provision 5.4.2-3.....	67
6.12.4	Provision 5.4.2-4.....	68
6.13	Principle 13: Ensure proper data and model disposal.....	69
6.13.1	Provision 5.5.1-1.....	69
6.13.2	Provision 5.5.1-2.....	70
Annex A:	Mapping from design and organization principles to SAI	71
Annex B:	Bibliography	73
History		74

Intellectual Property Rights

Essential patents

IPRs essential or potentially essential to normative deliverables may have been declared to ETSI. The declarations pertaining to these essential IPRs, if any, are publicly available for **ETSI members and non-members**, and can be found in ETSI SR 000 314: *"Intellectual Property Rights (IPRs); Essential, or potentially Essential, IPRs notified to ETSI in respect of ETSI standards"*, which is available from the ETSI Secretariat. Latest updates are available on the [ETSI IPR online database](#).

Pursuant to the ETSI Directives including the ETSI IPR Policy, no investigation regarding the essentiality of IPRs, including IPR searches, has been carried out by ETSI. No guarantee can be given as to the existence of other IPRs not referenced in ETSI SR 000 314 (or the updates on the ETSI Web server) which are, or may be, or may become, essential to the present document.

Trademarks

The present document may include trademarks and/or tradenames which are asserted and/or registered by their owners. ETSI claims no ownership of these except for any which are indicated as being the property of ETSI, and conveys no right to use or reproduce any trademark and/or tradename. Mention of those trademarks in the present document does not constitute an endorsement by ETSI of products, services or organizations associated with those trademarks.

DECT™, **PLUGTESTS™**, **UMTS™** and the ETSI logo are trademarks of ETSI registered for the benefit of its Members. **3GPP™**, **LTE™** and **5G™** logo are trademarks of ETSI registered for the benefit of its Members and of the 3GPP Organizational Partners. **oneM2M™** logo is a trademark of ETSI registered for the benefit of its Members and of the oneM2M Partners. **GSM®** and the GSM logo are trademarks registered and owned by the GSM Association.

Foreword

This Technical Report (TR) has been produced by ETSI Technical Committee Securing Artificial Intelligence (SAI).

Modal verbs terminology

In the present document "**should**", "**should not**", "**may**", "**need not**", "**will**", "**will not**", "**can**" and "**cannot**" are to be interpreted as described in clause 3.2 of the [ETSI Drafting Rules](#) (Verbal forms for the expression of provisions).

"**must**" and "**must not**" are **NOT** allowed in ETSI deliverables except when used in direct citation.

Introduction

The growing deployment and technological advancements of Artificial Intelligence (AI) has further reiterated the need for tailored security requirements for AI systems. The present document will guide stakeholders across the AI supply chain on implementations of ETSI TS 104 223 [i.1] by providing non-exhaustive scenarios as well as examples of practical solutions to meet these provisions.

1 Scope

The present document gives guidance to help stakeholders in the AI supply chain in meeting the cyber security provisions defined for AI models and systems in ETSI TS 104 223 [i.1]. These stakeholders could include a diverse range of entities, including large enterprises and government departments, independent developers, Small and Medium Enterprises (SMEs), charities, local authorities and other non-profit organizations. The present document will also be useful for stakeholders planning to purchase AI services. Additionally, the present document has been designed to support the future development of AI cyber security standards, including specifications that could inform future assurance and certification programmes. Where relevant, the present document signposts supporting specifications and international frameworks.

2 References

2.1 Normative references

Normative references are not applicable in the present document.

2.2 Informative references

References are either specific (identified by date of publication and/or edition number or version number) or non-specific. For specific references, only the cited version applies. For non-specific references, the latest version of the referenced document (including any amendments) applies.

NOTE: While any hyperlinks included in this clause were valid at the time of publication, ETSI cannot guarantee their long-term validity.

The following referenced documents may be useful in implementing an ETSI deliverable or add to the reader's understanding, but are not required for conformance to the present document.

- [i.1] ETSI TS 104 223: "Securing Artificial Intelligence (SAI); Baseline Cyber Security Requirements for AI Models and Systems".
- [i.2] [ETSI TR 104 222](#): "Securing Artificial Intelligence; Mitigation Strategy Report".
- [i.3] [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).
- [i.4] [ISO/IEC 22989:2022](#): "Information technology — Artificial intelligence — Artificial intelligence concepts and terminology".
- [i.5] NCSC: "[Machine learning principles](#)".
- [i.6] NCSC: "[Guidelines for Secure AI system development](#)".
- [i.7] CISA: "[Joint Cybersecurity Information - Deploying AI Systems Securely](#)".
- [i.8] MITRE: "[ATLAS Framework](#)".
- [i.9] [NIST AI 100-2 E2023](#): "Adversarial Machine Learning: A Taxonomy and Terminology of Attacks and Mitigations".
- [i.10] OWASP®: "[AI Exchange](#)".
- [i.11] OWASP®: "[Top 10 for LLM Applications](#)".
- [i.12] ETSI TR 104 032: "Securing Artificial Intelligence (SAI); Traceability of AI Models".
- [i.13] ETSI TR 104 225: "Securing Artificial Intelligence TC (SAI); Privacy aspects of AI/ML systems".

- [i.14] ETSI TR 104 066: "Securing Artificial Intelligence (SAI); Security Testing of AI".
- [i.15] ISO/IEC 42001:2023: "Information technology — Artificial intelligence — Management system".
- [i.16] ISO/IEC 25059:2023: "Software engineering — Systems and software Quality Requirements and Evaluation (SQuaRE) — Quality model for AI system)".
- [i.17] NIST: "[AI Risk Management Framework](#)".
- [i.18] International Scientific Report on the Safety of Advanced AI: "[Interim Report](#)".
- [i.19] ICO: "[Data Protection Audit Framework](#)".
- [i.20] ICO: "[Generative AI: eight questions that developers and users need to ask](#)".
- [i.21] OWASP: "[LLM Applications Cybersecurity and Governance Checklist](#)".
- [i.22] ICO: "[How should we assess security and data minimisation in AI?](#)".
- [i.23] AI Village Defcon 2024 Report: "[Generative AI Red Teaming Challenge: Transparency Report](#)".
- [i.24] OWASP®: "[Threat Modeling Cheat Sheet](#)".
- [i.25] ICO: "[Newsletters](#)".
- [i.26] NCSC: "[News on AI](#)".
- [i.27] MITRE: "[Slack Channel](#)".
- [i.28] OWASP®: "[Slack Invite](#)".
- [i.29] NIST: "[NIST Secure Software Development Framework for Generative AI and for Dual Use Foundation Models Virtual Workshop](#)".
- [i.30] OWASP®: "[OWASP Secure Coding Practices-Quick Reference Guide](#)".
- [i.31] NCSC: "[Secure Development and Deployment Guidance](#)".
- [i.32] ICO: "[Do we need to consult the ICO?](#)".
- [i.33] NCSC: "[Risk management](#)".
- [i.34] ICO: "[What is the impact of Article 22 of the UK GDPR on fairness?](#)".
- [i.35] ICO: "[AI and data protection risk toolkit](#)".
- [i.36] [PCI® DSS](#).
- [i.37] [FCA Standards](#).
- [i.38] EU Artificial Intelligence Act: "[The AI Act Explorer](#)".
- [i.39] NCSC: "[Risk management - Threat Modelling](#)".
- [i.40] NIST: "[NIST AI RMF Playbook](#)".
- [i.41] NIST: "[NIST AI RMF Crosswalk Documents](#)".
- [i.42] ICO: "[What are the accountability and governance implications of AI?](#)".
- [i.43] European AI Alliance: "[Implementing AI Governance: from Framework to Practice](#)".
- [i.44] OWASP®: "[New OWASP AI Security Center of Excellence \(CoE\) Guide](#)".
- [i.45] [MITRE ATT&CK®](#).
- [i.46] CSA: "[Guidelines and Companion Guide on Securing AI Systems](#)".
- [i.47] NCSC: "[Secure Design Principles](#)".

- [i.48] MLOps: "[MLOps Principles](#)".
- [i.49] NIST: "[Cybersecurity Supply Chain Risk Management](#)".
- [i.50] [NIST SP 800-161 Rev. 1](#): "Cybersecurity Supply Chain Risk Management Practices for Systems and Organizations".
- [i.51] NCSC: "[Supply Chain Security Guidance](#)".
- [i.52] ICO: "[UK GDPR Guidance and Resources](#)".
- [i.53] NCSC: "[Protecting bulk personal data](#)".
- [i.54] European Commission: "[GDPR Compliance Guidelines by EU Commission](#)".
- [i.55] [ISO/IEC 27001:2022](#): "Information security, cybersecurity and privacy protection — Information security management systems — Requirements".
- [i.56] ICO: "[Reporting Processes](#)".
- [i.57] ICO: "[Breach identification, assessment and logging](#)".
- [i.58] NCSC: "[Cyber Security Toolkit for Boards: Developing a positive cyber security culture](#)".
- [i.59] NCSC: "[Responding to a cyber incident - a guide for CEOs](#)".
- [i.60] NCSC: "[Cloud Security Guidance - Using a cloud platform securely - Apply access controls](#)".
- [i.61] NCSC: "[Zero trust architecture design principles](#)".
- [i.62] OWASP®: "[Threat Modeling Process](#)".
- [i.63] [NIST SP 800-154](#): "Guide to Data-Centric System Threat Modeling".
- [i.64] NCSC: "[Introduction to Logging for Security Purposes](#)".
- [i.65] OWASP®: "[Top 10 for APIs - API4:2019 Lack of Resources & Rate Limiting](#)".
- [i.66] OWASP®: "[Threat modeling in practice](#)".
- [i.67] OWASP®: "[AI Top 10 API Security Risks - 2023](#)".
- [i.68] Yi Dong, Ronghui Mu, et al.: "[Building Guardrails for Large Language Models](#)".
- [i.69] OWASP®: "[Threat Modeling Playbook](#)".
- [i.70] NCSC: "[Risk Management - Cybersecurity Risk Management Framework](#)".
- [i.71] [ISO 9001](#): "What does it mean in the supply chain?".
- [i.72] CISA: "[Software Bill of Materials \(SBOM\)](#)".
- [i.73] World Economic Forum: "[Adopting AI Responsibly: Guidelines for Procurement of AI Solutions by the Private Sector: Insight Report](#)".
- [i.74] MITRE: "[AI Risk Database](#)".
- [i.75] [NIST SP 800-137](#): "Information Security Continuous Monitoring (ISCM) for Federal Information Systems and Organizations".
- [i.76] NCSC: "[Early Warning](#)".
- [i.77] NCSC: "[Building a Security Operations Centre \(SOC\) - Threat Intelligence](#)".
- [i.78] ICO: "[Audit Framework toolkit on AI - Human review](#)".
- [i.79] [NISTIR 8312](#): "Four Principles of Explainable Artificial Intelligence".
- [i.80] ICO: "[Explaining Decisions made with AI](#)".

- [i.81] G. Detommaso, M. Bertran, R. Fogliato, A. Roth: "[Multicalibration for Confidence Scoring in LLMs](#)".
- [i.82] H. Luo, L. Specia: "[From Understanding to Utilization: A Survey on Explainability for Large Language Models](#)".
- [i.83] ICO: "[Audits](#)".
- [i.84] ICO: "[A Guide to ICO Audit - Artificial Intelligence \(AI\) Audits](#)".
- [i.85] OWASP®: "[LLM and Generative AI Security Solutions Landscape](#)".
- [i.86] CISA: "[CISA, JCDC, Government and Industry Partners Conduct AI Tabletop Exercise](#)".
- [i.87] OWASP®: "[Guide for Preparing and Responding to Deepfake Events](#)".
- [i.88] ICO: "[A guide to data security](#)".
- [i.89] OWASP®: "[OWASP Application Security Verification Standard \(ASVS\)](#)".
- [i.90] Technical Disclosure Commons: "[Training Dataset Validation to Protect Machine Learning Models from Data Poisoning](#)".
- [i.91] NCSC: "[Device Security Guidance - Logging and Protective Monitoring](#)".
- [i.92] Towards Data Science: "[LLM Monitoring and Observability - A Summary of Techniques and Approaches for Responsible AI](#)".
- [i.93] "[API Management Overview and links](#)" from Wikipedia®.
- [i.94] The Alan Turing Institute: "[What is synthetic data and how can it advance research and development?](#)".
- [i.95] DSTL: "[Machine learning with limited data](#)".
- [i.96] NCSC: "[Secure development and deployment guidance](#)".
- [i.97] [NIST SP 800-218](#): "Secure Software Development Framework (SSDF) Version 1.1: Recommendations for Mitigating the Risk of Software Vulnerabilities".
- [i.98] NCSC: "[Vulnerability Disclosure Toolkit](#)".
- [i.99] OWASP: "[Cryptographic Storage Cheat Sheet](#)".
- [i.100] [ISO/IEC 29147:2018](#): "Information technology — Security techniques — Vulnerability disclosure".
- [i.101] CSA: "[Incident Response Checklist](#)".
- [i.102] NCSC: "[Incident management](#)".
- [i.103] CSA: "[AI Organizational Responsibilities - Governance, Risk Management, Compliance and Cultural Aspects](#)".
- [i.104] ICO: "[Security requirements](#)".
- [i.105] MITRE: "[System of Trust Framework](#)".
- [i.106] OWASP: "[Machine Learning Bill of Materials \(ML-BOM\)](#)".
- [i.107] SLSA: "[Supply-chain Levels for Software Artifacts \(SLSA\) specification](#)".
- [i.108] NCSC: "[Vulnerability Management](#)".
- [i.109] OWASP®: "[OWASP Vulnerability Management Guide](#)".
- [i.110] [ISO 2800:2022](#): "Security and resilience — Security management systems — Requirements".

- [i.111] GOV.UK: "[Guidance and tools for digital accessibility](#)".
 - [i.112] GitHub: "[AI-secure/DecodingTrust: A Comprehensive Assessment of Trustworthiness in GPT Models](#)".
 - [i.113] AI Safety Institute: "[Inspect - An open-source framework for large language model evaluations](#)".
 - [i.114] NIST: "[AI Test, Evaluation, Validation, and Verification \(TEVV\)](#)".
 - [i.115] NIST: "[Dioptra Test Platform](#)".
 - [i.116] ICO: "[Retention and destruction of information](#)".
 - [i.117] NIST: "[Cryptographic Standards and Guidelines](#)".
 - [i.118] OWASP®: "[OWASP Top 10 - 2021](#)".
 - [i.119] "[Generative AI Red Teaming Challenge: Transparency Report](#)", AI Village Defcon 2024".
 - [i.120] NCSC: "[Penetration testing](#)".
 - [i.121] Cornell University: "[Red-Teaming for Generative AI: Silver Bullet or Security Theater?](#)".
 - [i.122] NCSC: "[CHECK penetration testing](#)".
 - [i.123] [ISO/IEC/IEEE 29119-1:2022](#): "Software and systems engineering — Software testing — Part 1: General concepts".
 - [i.124] Linux® Foundation AI & Data Foundation: "[Adversarial Robustness Toolbox](#)".
- NOTE: Linux® is the registered trademark of Linus Torvalds in the U.S. and other countries.
- [i.125] Github: "[TextAttack: Generating adversarial examples for NLP models](#)".
 - [i.126] ICO: "[Generative AI second call for evidence: Purpose limitation in the generative AI lifecycle](#)".
 - [i.127] ICO: "[What do we need to know about accuracy and statistical accuracy?](#)".
 - [i.128] ICO: "[Disposal and deletion](#)".
 - [i.129] ICO: "[Guidance on AI and data protection](#)".
 - [i.130] [NIST SP 800-88 Rev. 1](#): "Guidelines for Media Sanitization".
 - [i.131] NCSC: "[Secure sanitisation of storage media](#)".
 - [i.132] ETSI TR 104 221: "Securing Artificial Intelligence (SAI); Problem Statement".
 - [i.133] ETSI TR 104 062 (V1.2.1) (2024-07); "Securing Artificial Intelligence; Automated Manipulation of Multimedia Identity Representations".
 - [i.134] ETSI TR 104 029: "Securing Artificial Intelligence (SAI); Global Ecosystem".
 - [i.135] ETSI TS 104 050: "Securing Artificial Intelligence (SAI); AI Threat Ontology and definitions".
 - [i.136] ETSI TR 104 030: "Securing Artificial Intelligence (SAI); Critical Security Controls for Effective Cyber Defence; Artificial Intelligence Sector".
 - [i.137] ETSI TS 102 165-1: "Cyber Security (CYBER); Methods and protocols; Part 1: Method and pro forma for Threat, Vulnerability, Risk Analysis (TVRA)".
 - [i.138] ETSI TS 103 485: "CYBER; Mechanisms for privacy assurance and verification".
 - [i.139] ETSI TS 104 224: "Securing Artificial Intelligence (SAI); Explicability and transparency of AI processing".
 - [i.140] ETSI TR 103 305-1: "Cyber Security (CYBER); Critical Security Controls for Effective Cyber Defence; Part 1: The Critical Security Controls".

- [i.141] ETSI TR 103 305-2: "CYBER; Critical Security Controls for Effective Cyber Defence; Part 2: Measurement and auditing".
- [i.142] ETSI TR 103 305-3: "CYBER; Critical Security Controls for Effective Cyber Defence; Part 3: Service Sector Implementations".
- [i.143] ETSI TR 103 305-4: "Cyber Security (CYBER); Critical Security Controls for Effective Cyber Defence; Part 4: Facilitation Mechanisms".
- [i.144] ETSI TR 103 305-5: "Cyber Security (CYBER); Critical Security Controls for Effective Cyber Defence; Part 5: Privacy and personal data protection enhancement".
- [i.145] ETSI TR 104 048: "Securing Artificial Intelligence (SAI); Data Supply Chain Security".
- [i.146] JTC21024: "Risk management".
- [i.147] [Directive 2006/42/EC](#) of the European Parliament and of the Council of 17 May 2006 on machinery, and amending Directive 95/16/EC.
- [i.148] [Joint programme with Digital Europe](#).
- [i.149] [Directive 2002/58/EC](#) of the European Parliament and of the Council of 12 July 2002 concerning the processing of personal data and the protection of privacy in the electronic communications sector (Directive on privacy and electronic communications).
- [i.150] [ETSI White Paper No. #34](#): "Artificial Intelligence and future directions for ETSI".
- [i.151] [ETSI White Paper No. #52](#): "ETSI Activities in the field of Artificial Intelligence Preparing the implementation of the European AI Act".

3 Definition of terms, symbols and abbreviations

3.1 Terms

For the purposes of the present document, the following terms apply:

administrator: user who has the highest-privilege level possible for a user of the device

NOTE: This can mean they are able to change any configuration related to the intended functionality.

adversarial AI: techniques and methods that exploit vulnerabilities in the way AI systems work

EXAMPLE: By introducing malicious inputs to exploit their machine learning aspect and deceive the system into producing incorrect or unintended results. These techniques are commonly used in adversarial attacks but are not a distinct type of AI system.

adversarial attack: attempt to manipulate an AI model by introducing specially crafted inputs to cause the model to produce errors or unintended outcomes

agentic systems: AI systems capable of initiating and executing actions autonomously, often interacting with other systems or environments to achieve their goals

Application Programming Interface (API): set of tools and protocols that allow different software systems to communicate and interact

Bill of Materials (BOM): comprehensive inventory of all components used in a system, such as software dependencies, configurations, and hardware

data custodian: See definition in ETSI TS 104 223 [i.1].

data poisoning: type of adversarial attack where malicious data is introduced into training datasets to compromise the AI system's performance or behaviour

Data Protection Impact Assessment (DPIA): tool used in UK GDPR to assess and mitigate privacy risks associated with processing personal data in AI systems

embeddings: vector representations of data (e.g. text, images) that capture their semantic meaning in a mathematical space, commonly used to improve the efficiency of search, clustering and similarity comparisons

evasion attack: type of adversarial attack where an adversary manipulates input data to cause the AI system to produce incorrect or unexpected outputs without altering the underlying model

excessive agency: situation where an AI system has the capability to make decisions or take actions beyond its intended scope, potentially leading to unintended consequences or misuse

explainability: ability of an AI system to provide human-understandable insights into its decision-making process

feature selection: process of selecting a subset of relevant features (variables) for use in model training to improve performance, reduce complexity and prevent overfitting

generative AI: AI models that generate new content, such as text, images or audio, based on training data

EXAMPLE: Image synthesis models and large language models like chatbots.

governance framework: policies and procedures established to oversee the ethical, secure and compliant use of AI systems

guardrails: predefined constraints or rules implemented to control and limit an AI system's outputs and behaviours, ensuring safety, reliability, and alignment with ethical or operational guidelines

hallucination (in AI): AI-generated content that appears factual but is incorrect or misleading

NOTE: This is prevalent in LLMs, which can produce plausible sounding but inaccurate responses.

harm: injury or damage to the health, or damage to property or the environment, or interference with the fundamental rights of the person

hazard: potential source of harm

hazardous situation: circumstance in which people, property or the environment is/are exposed to one or more hazards

inference attack: privacy attack where an adversary retrieves sensitive information about the training data, or users, by analysing the outputs of an AI model

Large Language Model (LLM): type of AI model trained on vast amounts of text data to understand and generate human-like language

EXAMPLE: Chatbots and content generation tools.

Machine Learning (ML): subset of AI where systems improve their performance on a task over time by learning from data rather than following explicit instructions

Machine Learning Bill of Materials (ML BOM): specialized BOM for AI systems that catalogues models, datasets, parameters and training configurations used in the development and deployment of machine learning solutions

Machine Learning Operations (ML Ops): set of practices and tools that streamline and standardize the deployment, monitoring and maintenance of machine learning models in production environments

model extraction: attack where an adversary recreates or approximates a proprietary AI model by querying it and analysing its outputs, potentially exposing trade secrets or intellectual property

model inversion: privacy attack where an adversary infers sensitive information about the training data by analysing the AI model's outputs

multimodal models: AI models that process and integrate multiple types of data (e.g. text, images, audio) to perform tasks

Natural Language Processing (NLP): type of machine learning that understands, interprets, and generates human language in a way that is meaningful and useful

predictive (or discriminative) AI: type of machine learning designed to classify inputs or make predictions based on existing data

NOTE: These models focus on identifying patterns and drawing distinctions, such as fraud detection or customer segmentation.

prompt: input provided to an AI model, often in the form of text, that directs or guides its response

NOTE: Prompts can include questions, instructions, or context for the desired output.

prompt injection: attacker exploits a vulnerability in AI models by using prompts that produce unintended or harmful outputs

Reinforcement Learning (RL): machine learning approach where an agent learns by interacting with its environment and receiving feedback in the form of rewards or penalties

Retrieval-Augmented Generation (RAG): AI approach that combines external knowledge retrieval (e.g. documents or databases) with prompts to language model generation to provide accurate and up-to-date responses

risk assessment: process of identifying, analysing and mitigating potential threats to the security or functionality of an AI system

sanitisation: process of cleaning and validating data or inputs to remove errors, inconsistencies and malicious content, ensuring data integrity and security

Software Bill of Materials (SBOM): detailed list of all software components in a system, including open-source libraries, versions and licences to ensure transparency and security

system prompt: predefined input or set of instructions provided to guide the behaviour of an AI model, often used to define its tone, rules, or operational context

threat modelling: process to identify and address potential security threats to a system during its design and development phases

training: process of teaching an AI model to recognize patterns, make decisions, or generate outputs by exposing it to labelled data and adjusting its parameters to minimize errors

web content accessibility guidelines: guidelines, as part of, internationally recognized standards for making web content more accessible to people with impairments

NOTE: They are developed and maintained by the World Wide Web Consortium (W3C) under its Web Accessibility Initiative (WAI).

3.2 Symbols

Void.

3.3 Abbreviations

For the purposes of the present document, the following abbreviations apply:

ADR	Architecture Decision Records
AI	Artificial Intelligence
API	Application Programming Interface
BOM	Bill of Materials
CI/CD	Continuous Integration/Continuous Deployment
CISA	Cyber Security and Infrastructure Agency
DPIA	Data Protection Impact Assessment
ETSI	European Telecommunications Standards Institute
GDPR	General Data Protection Regulation
GRC	Governance Risk and Compliance
GRT	Generative Read-Teaming
ICO	Information Commissioner's Office

ISO/IEC	International Organization for Standardization / International Electrotechnical Commission
LLM	Large Language Models
MFA	Multi-Factor Authentication
MITRE	MITRE Corporation
ML BOM	Machine Learning Bill of Materials
ML	Machine Learning
MLOps	Machine Learning Operations
NCSC	National Cyber Security Centre
NIST	National Institute of Standards and Technology
NLP	Natural Language Processing
OWASP	Open Web Application Security Project
RAG	Retrieval-Augmented Generation
RBAC	Role-Based Access Control
RL	Reinforcement Learning
RLFAI	Reinforcement Learning from AI
RLHF	Reinforcement Learning from Human Feedback
RMF	Risk Management Framework
RSS	Really Simple Syndication
SBOM	Software Bill of Materials
SHA-256	Secure Hash Algorithm 256-bit
T&Cs	Terms and Conditions
WCAG	Web Content Accessibility Guidelines
WORM	Write Once, Read Many

4 How to use the present document

4.1 Purpose

The present document helps implementers of ETSI TS 104 223 [i.1] by providing further information about each provision. Recommendations in ETSI TS 104 223 [i.1] are followed by AI supply chain stakeholders unless they are not appropriate. For example, because of the type of model(s) used for the AI system. It can depend on whether a stakeholder has decided to develop their own model, use, or finetune a third-party model (either directly or remotely via an API). The present document provides guidance on the scope of particular provisions.

Figure 4.1-1 highlights that the principles have been mapped to various phases of the AI lifecycle. Importantly, some of the principles and provisions are also relevant to other phases, which is clarified in the relevant clauses. An example is principle 9 (see clause 6.9), which is included under "Development", but it is also very important for the "Deployment" of an AI system. The examples provided to address the below scenarios are not exhaustive or limitative; it is possible to meet the provisions in ETSI TS 104 223 [i.1] by using other solutions, or variants of the examples provided.

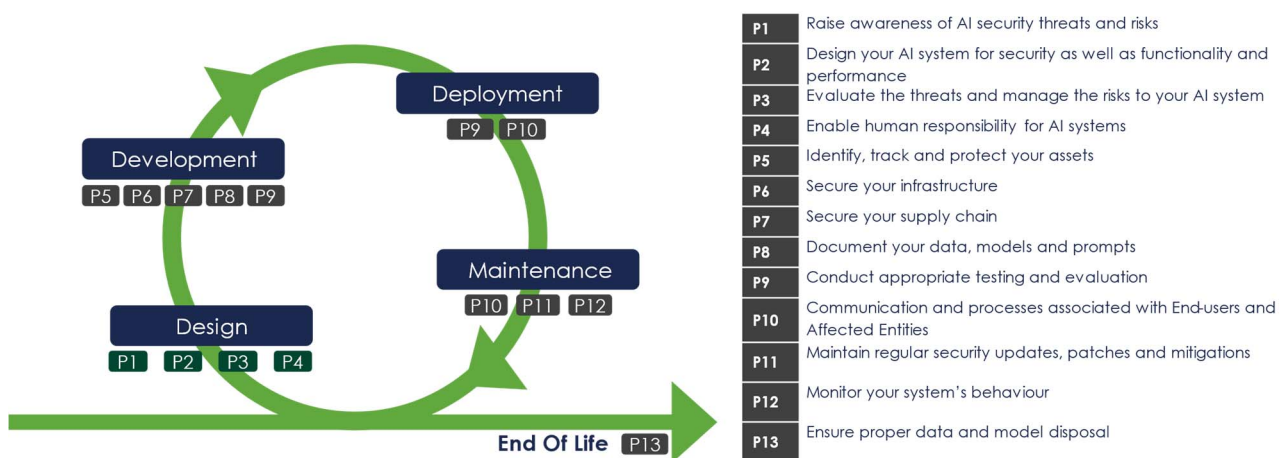


Figure 4.1-1: AI Lifecycle

4.2 Relationship to ETSI TS 104 223

ETSI TS 104 223 [i.1] sets out a detailed list of provisions based on thirteen principles. The present document can be used (when implemented) to inform the definition of test scenarios and the development of a test plan based on ETSI TS 104 223 [i.1].

5 Guidance on implementation

Clause 6 provides examples for implementing ETSI TS 104 223 [i.1] provisions based on various scenarios (outlined below) of how a stakeholder might create and use an AI system:

- **Chatbot App:** An organization using a publicly available LLM via the APIs offered by the external provider to develop a chatbot for internal and customer use. This can include:
 - 1) A large enterprise uses a publicly available LLM through an API to create chatbots for internal and customer interactions, such as answering FAQs or automating routine customer service tasks.
 - 2) A small retail business developing and using an AI-powered chatbot to handle online shopping queries, assisting customers with product recommendations and order tracking.
 - 3) A hospital developing and using a chatbot to provide general health advice and appointment scheduling, ensuring compliance with data privacy requirements.
 - 4) A local council develops and uses a chatbot to provide guidance on local planning applications and handle the applications.
- **ML Fraud Detection:** A mid-size software company selects an open-access classification model, which they train further with additional datasets to develop and host a fraud detection system. The system is designed to identify patterns of fraudulent financial transactions based solely on transactional data. It explicitly avoids linking decisions to inferred personal characteristics, behaviours, or any factors unrelated to the context of the financial transaction. The model's primary focus is on identifying fraud patterns and does not evaluate or classify individuals' social behaviour or implement social scoring. The scenario does not encompass situations that are detrimental or unfavourable treatment of certain natural persons or groups of persons that is unjustified or disproportionate to their social behaviour or its gravity, as stipulated in Article 5(1)(c) Prohibited Practices of the EU AI Act (Regulation (EU) 2024/1689) [i.3].
- **LLM Provider:** A tech company develops a new multimodal LLM capable of understanding and generating text, audio and images, providing commercial API access to developers for diverse applications, such as virtual assistants and media generation.
- **Open-Access LLM:** A small organization is developing an LLM for specific use cases. This can include:
 - 1) Developing an LLM for legal and contract negotiation use cases planning to release it as open-access and monetize via support agreements.
 - 2) A law firm using the open-access LLM to combine it with their confidential casework for legal research, enabling quick identification of relevant legal precedents and statutes.
 - 3) A rural development organization developing and using an open access LLM to offer farmers localized advice on crop management and pest control strategies.

The information in the present document will help stakeholders to protect end-users and affected entities from vulnerabilities that could result in confidentiality, integrity, or availability attacks. This includes the various threat-related examples linked to AI systems below:

- Data Poisoning, Backdoors, Model Tampering, Evasion and Supply-Chain Attacks.
- Privacy Attacks, such as Model Theft, Model Extraction, Model Inversion and Inference Attacks.
- Information Disclosure of Personal and Special Category Data, Confidential Business Information or System Configuration details.

- Prompt Injections, Excessive Agency and Training Data Extraction and Model Denial of Service.

These are the most common examples of threats, but threats will continue to evolve and new ones will emerge. For complete taxonomies, refer to OWASP AI Exchange [i.10], MITRE ATLAS [i.8] and the NIST Adversarial Attacks Taxonomy [i.9].

AI models and systems can also be misused. Although this area is out of scope of the present document, there is some crossover as AI Security underpins all aspects of AI Safety and is linked to the responsible management of AI. As a result, the present document will help reduce the risk of AI models and systems being misused. The following measures/controls set out in the present document could help to mitigate the misuse of AI systems, for example to help produce misinformation or conduct cyber attacks:

- Human Oversight Mechanisms.
- Access Control and Rate-Based Permissions.
- Threat Modelling.
- Risk Assessment.
- Documentation and Monitoring of Prohibited Cases.
- Monitoring and Logging.
- Rate Limiting.

While the primary focus of the present document is AI security, certain aspects of responsible AI management such as copyright violations, bias, unethical or harmful use and legal or reputational risks are included in specific sections of the present document. In the context of AI, these areas often stem from or are exacerbated by poor AI security practices; safeguarding them not only mitigates the misuse of AI but also strengthens trust and compliance in AI systems.

Personal data, (although a consistent theme throughout the principles), is specified only in particular circumstances. Organizations should consult official regulatory guidance for regulatory compliance where appropriate, including data protection guidance issued by the relevant data regulatory bodies. Additionally, stakeholders that adhere to ETSI TS 104 223 [i.1] will still need to ensure their compliance with other regulatory compliance requirements.

For the purposes of the present document, "regularly" denotes a frequency determined by the associated risks and operational requirements of the system. This can range from continuous, daily, or weekly actions for high-risk scenarios to quarterly or annual actions for lower risk scenarios.

Content on how stakeholders can verify conformity to each provision, is out of scope of the present document.

The present document aligns with international standardization including:

- ETSI TR 104 048 [i.145]
- ETSI TS 104 224 [i.139]
- ETSI TR 104 222 [i.2]
- ETSI TR 104 032 [i.12]
- ETSI TR 104 225 [i.13]
- ETSI TR 104 066 [i.14]
- ISO/IEC 22989:2022 [i.4]
- ISO/IEC 42001:2023 [i.15]
- ISO/IEC 25059:2023 [i.16]

6 Examples to meet AI Security Provisions

6.1 Principle 1: Raise awareness of AI security threats and risks

6.1.1 Provision 5.1.1-1

"Organizations' cyber security training programme shall include AI security content which shall be regularly reviewed and updated where necessary, such as if new substantial AI-related security threats emerge." (ETSI TS 104 223 [i.1])

Related threats/risks:

AI attack types are still being understood and evolving so staff can be unaware of unique AI vulnerabilities like data poisoning, adversarial attacks, or prompt injections, leaving the system exposed to sophisticated attacks.

Example Measures/Controls:

Establish an AI Security Awareness Training Programme that covers basic AI concepts, threats, applicable regulations, etc. Include guidance on how to monitor for threats and the escalation paths for reporting security concerns.

Chatbot App: Training on AI concepts, personal data and its regulatory implications, risks on confidential business information or system configuration, hallucinations, overreliance, and ethical use and safety; training should cover at least the national guidelines, and e.g. OWASP Top 10 for LLM applications [i.11].

ML Fraud Detection: Provide training on AI concepts and use national guidelines and e.g. the OWASP AI Exchange [i.10], to cover ML threats such as poisoning, evasion, model extraction, model inversion, inference, and supply-chain attacks.

LLM Platform: Provide training on AI concepts and threats including poisoning, prompt injections, safety and data protection and national guidelines; Cover for example OWASP AI Exchange [i.10], the OWASP Top 10 for LLM applications [i.11] and recent reports on the risks of general-purpose systems [i.18].

Open-Access LLM Model: Self-training on national guidelines and for example OWASP AI Exchange [i.10], the OWASP Top 10 for LLM applications [i.11] and recent reports on the risks of general-purpose systems [i.18].

References:

- ISO/IEC 22989: Artificial intelligence concepts and terminology [i.4]
- ICO Data Protection Audit Framework [i.19]
- ICO Generative AI- eight questions that developers and users need to ask [i.20]
- NCSC Machine Learning Principles, Part 1, 1.1 Raise awareness of ML threats and risks [i.5]
- OWASP Top 10 for LLM applications [i.11]
- OWASP AI Exchange [i.10]

6.1.2 Provision 5.1.1-1.1

"AI security training shall be tailored to the specific roles and responsibilities of staff members." (ETSI TS 104 223 [i.1])

Related threats/risks:

Without tailored AI security training, staff can lack the knowledge to address role-specific risks, leading to ineffective implementation of security measures, increased vulnerability to threats, and potential misuse or mismanagement of AI systems.

Example Measures/Controls 1:

Role-Specific AI Security Training: Provide role-specific AI security training tailored to the responsibilities of each staff category.

Chatbot App: Train engineers on secure coding and AI-specific vulnerabilities (see clause 6.1.5); for CISOs, include governance frameworks, incident response strategies, and regulatory compliance, as found in national guidance and for example the OWASP LLM Applications Cybersecurity and Governance Checklist [i.21]. For Risk Officers cover relevant frameworks such as NIST AI Risk Management Framework [i.17], AI threat modelling, and mitigation strategies; for IT Operations focus on implementing and maintaining security controls in production environments.

ML Fraud Detection: Similar approach to the Chatbot App example.

LLM Platform: Similar approach to the Chatbot App example.

Open-Access LLM Model: Similar approach to the Chatbot App example.

References:

- ICO Assessing security and data minimisation in AI [i.22]
- CISA Joint Cybersecurity Information - Deploying AI Systems Securely [i.7]
- NCSC Guidelines for Secure AI System Development [i.6]
- NIST AI Risk Management Framework [i.17]
- OWASP LLM Applications Cybersecurity and Governance Checklist [i.21]

Example Measures/Controls 2:

Incorporate Training on AI Threat Modelling and Red Teaming. Provide developers and other technical staff with training on threat modelling techniques and red teaming techniques tailored for AI.

Chatbot App: Provide training to developers and Risk owners on Threat Modelling incorporating threats and mitigations from OWASP Top 10 for LLM applications [i.11].

ML Fraud Detection: Follow the same approach as in the Chatbot App example but using MITRE ATLAS [i.8] and OWASP AI Exchange [i.10].

LLM Platform: Provide training on AI concepts and threats including poisoning, prompt injections, and data protection MITRE ATLAS [i.8] and OWASP AI Exchange [i.10]; and Generative Read-teaming (GRT) approaches including the AI Village Defcon 2024 report [i.23].

Open-Access LLM Model: Similar approach to the LLM Platform example.

References:

- OWASP Threat Modeling Cheat Sheet [i.24]
- MITRE ATLAS [i.8]
- AI Village Defcon 2024 Report: Generative AI Red Teaming Challenge [i.23]
- OWASP AI Exchange [i.10]
- OWASP Top 10 for LLM applications [i.11]

6.1.3 Provision 5.1.1-2

"As part of an Organization's wider staff training programme, they shall require all staff to maintain awareness of the latest security threats and vulnerabilities that are AI-related. Where available, this awareness shall include proposed mitigations." (ETSI TS 104 223 [i.1])

Related threats/risks:

AI systems face evolving threats, and staff who are not updated regularly on these vulnerabilities can unknowingly expose systems to risks, such as adversarial attacks or personal data leaks which will be in breach of data protection regulations.

Example Measures/Controls:

Maintain training awareness: Update training material regularly with new examples of AI threats (e.g. prompt injections, adversarial attacks) and mitigation techniques.

Chatbot App: Large organizations should conduct regular (at least annually) reviews of training material and the facility to register for changes to be updated. Smaller organizations can rely on logging new significant developments, e.g. a new version of the OWASP Top 10 for LLM Applications [i.11] and include them in knowledge sharing sessions.

ML Fraud Detection: Track and train staff on new adversarial attack patterns and data validation techniques.

LLM Platform: As in the ML Fraud Detection example. Additionally, update training with new research papers on AI vulnerabilities and share case studies on generative AI misuse and related mitigations.

Open-Access LLM Model: Update training log with risks like data memorization and unauthorized data use, using curated updates from e.g. ICO, OWASP, and others, including research, and utilizing workshop sessions.

References:

- MITRE ATLAS [i.8]
- OWASP Top 10 for LLM applications [i.11]
- OWASP AI Exchange [i.10]
- NIST Adversarial Machine Learning Taxonomy [i.9]

6.1.4 Provision 5.1.1-2.1

"These updates should be communicated through multiple channels, such as security bulletins, newsletters, or internal knowledge-sharing platforms. This will ensure broad dissemination and understanding among the staff." (ETSI TS 104 223 [i.1])

Related threats/risks:

Failure to communicate new developments through diverse channels can result in uneven dissemination of critical security information, leaving some staff unaware of vulnerabilities, mitigations, or best practices, increasing the risk of oversight and security lapses.

Example Measures/Controls:

Disseminate Regular Security Updates and Bulletins.

Chatbot App: This includes AI security bulletins, newsletters (e.g. ICO, NCSC, etc.), or messages on knowledge-sharing platforms and communication (messaging channels) to keep staff informed of the latest AI threats, vulnerabilities, and mitigations. Subscribe to e.g. the ICO, OWASP, and other curated newsletters and AI Feeds with a team member responsible in tracking changes. Everyone should contribute to team updates via team messaging channels. Attend conferences and events when possible and use free or low-cost resources such as RSS feeds or community forums to gain updates including new academic papers on emerging AI security vulnerabilities. Both the MITRE and OWASP slack channels are open to public and provide excellent information on AI Security news.

ML Fraud Detection: As before but with an automation of curated data feeds of AI Security news and content being shared and knowledge sharing sessions.

LLM Platform: As before, participation in events and conferences disseminating learnings using a knowledge sharing process and platform.

Open-Access LLM Model: As with the chatbot app with the inclusion of automated feeds with new LLM-related research.

References:

- ICO Newsletter [i.25]
- NCSC News on AI [i.26]
- OWASP Top 10 for LLM Apps Newsletter [i.11]
- MITRE Slack Channel [i.27]
- OWASP Slack Invite [i.28]

6.1.5 Provision 5.1.1-2.2

"Organizations shall provide developers with training in secure coding and system design techniques specific to AI development, with a focus on preventing and mitigating security vulnerabilities in AI algorithms, models, and associated software." (ETSI TS 104 223 [i.1])

Related threats/risks:

Without specialized training in secure coding and AI system design, developers can inadvertently introduce vulnerabilities into AI algorithms, models, or supporting software, increasing the risk of exploits, data breaches, or system failures.

Example Measures/Controls:

Provide secure coding training for engineers related to AI threats and incorporating guidelines from OWASP, NCSC, the ETSI Mitigation Strategy report [i.2].

Chatbot App: Train engineers on secure coding, including implementing input validation to mitigate prompt injections. Smaller organizations can implement this control with coding standards pointing to guidelines and using code reviews and developer mentoring as training.

ML Fraud Detection: Train developers on adversarial risks and OWASP AI Exchange [i.10].

LLM Platform: Similar to the Fraud Detection example, but with additional system design techniques to address LLM specific safety risks (e.g. model jailbreaking).

Open-Access LLM Model: Focus on secure coding techniques compliance with LLM specific threats and national guidelines for handling sensitive training data. Use the small organization approach which was described in the Chatbot App section to address resource constraints.

References:

- ETSI TR 104 222 [i.2]
- NIST Secure Software Development Framework for Generative AI and for Dual Use Foundation Models Virtual Workshop [i.29]
- OWASP Top 10 for LLM applications [i.11]
- OWASP AI Exchange [i.10]
- OWASP Secure Coding Practices-Quick Reference Guide [i.30]
- NCSC Secure Development and Deployment Guidance [i.31]
- NCSC Guidelines for Secure AI System Development [i.6]

6.2 Principle 2: Design the AI System for Security as well as Functionality and Performance

6.2.1 Provision 5.1.2-1

"As part of deciding whether to create an AI system, a System Operator and/or Developer shall conduct a thorough assessment that includes determining and documenting the business requirements and/or problem they are seeking to address, along with potential AI security risks and mitigation strategies." (ETSI TS 104 223 [i.1])

Related threats/risks:

Without assessing whether an AI system is required to meet the business requirements, systems can be unnecessary or poorly suited for their environment, leading to lack of compliance, unnecessary complexity, increased attack surface, unexpected behaviour, and security vulnerabilities.

Example Measures/Controls 1:

Conduct Business Alignment Review: Use ISO/IEC 25059 [i.16] as a quality model to review and document business requirements for the AI system to ensure that design choices align with the organization's needs and objectives.

Chatbot App: If the chatbot is for simple summarization or sentiment analysis, evaluate whether a task specific algorithm can be more appropriate than an LLM, which carries additional complexity and risk including regulatory risks.

ML Fraud Detection: When decision-making transparency is a requirement, use a simpler model (e.g. Gradient Boosting Model) with better explainability instead of a more complex black-box Deep Learning model.

LLM Platform: Review the business assessment to ensure the advanced multi-modal capabilities and complexity are required and potential risks have been considered in the business assessment.

Open-Access LLM Model: When creating a specialized LLM, evaluate complexity, data protection and security risks before deciding the route to follow. Fine-tuning a pre-built LLM is faster and simpler but can involve sharing sensitive data with the provider, raising privacy concerns and risks like unauthorized access or compliance issues (e.g. UK GDPR). Building a model from scratch offers greater control but can be complex and it requires robust internal controls, such as encrypted storage and access management, to safeguard training data. Both approaches demand careful evaluation to prevent data breaches and ensure regulatory compliance.

References:

- ICO Do we need to consult the ICO? [i.32]
- ICO Data Protection Audit Framework [i.19]
- ICO What is the impact of Article 22 of the UK GDPR on fairness? [i.34]
- OWASP Top 10 for LLM applications [i.11]

Example Measures/Controls 2:

Perform Risk Assessment: Conduct and document an AI-specific risk assessment covering data classifications, logging risks of personal data and their mitigations in DPIAs. Cover expected data volume, types of integration, the model's complexity, architecture, and number of parameters. For more information on these risk factors see the NCSC Principles for Machine Learning [i.5] and NCSC Guidelines for Secure AI system development [i.6].

Chatbot App: Focus the assessment on use of internal data and their classification, safety and abuse, reputational and legal risks if the chatbot provides misleading or inappropriate content.

ML Fraud Detection: Use a standardized assessment template, such as the ICO's AI Data Protection Risk Toolkit [i.35], to record AI risk variables and risk scores. Include in other relevant factors such as interpretability and regulatory impact, then aggregate these scores to prioritize model security. Consult finance and regulatory compliance experts to understand sector-specific compliance requirements including e.g. PCI DSS [i.36], and FCA standards [i.37] to ensure guidance. Ensure solution is compliant with the provisions of EU AI Act Explorer [i.38], 5(1)(c) in particular.

LLM Platform: In addition to using a standardized assessment template, factor in regulations, legislations and guidelines in targeted markets and cover copyright violation risks as well as emerging new risks identified in recent reports on the risks of general-purpose systems [i.18].

Open-Access LLM Model: Follow the advice in the LLM Platform example, but include the risks related to misuse of open-source and open-access components, including malicious code, data leaks, and ethical/legal liabilities of public outputs (e.g. bias or misinformation). Review responsibilities and legal liabilities related to licensing compliance, safeguarding sensitive data during model use or fine-tuning, and implementing safeguards to prevent misuse (e.g. phishing or misinformation campaigns). Leverage tools like the ICO's AI Data Protection Risk Toolkit [i.35], OWASP AI Security Center of Excellence (CoE) Guide [i.44] and MITRE ATLAS [i.8] to structure assessments.

References:

- EU AI Act Explorer [i.38]
- ICO Data Protection Audit Framework [i.19]
- ICO AI Data Protection Risk Toolkit [i.35]
- MITRE ATLAS [i.8]
- NCSC Risk management [i.33]
- International Scientific Report on the Safety of Advanced AI: Interim Report [i.18]
- NCSC Principles for Machine Learning [i.5]
- OWASP Top 10 for LLM applications [i.11]
- EU Machinery Directive [i.147]

Example Measures/Controls 3:

Integrate Risk Management and Governance Frameworks:

Embed AI system assessments within both a formal Risk Management Framework (RMF) and the AI governance structure to ensure thorough risk evaluation, consistent mitigation, and monitoring before key organizational decisions.

Chatbot App: Implement NIST AI RMF [i.17] and use it to assess the application as part of the framework's structured approach which includes defining purpose and risk objectives, risk categorization, risk assessment, control selection, implementation, monitoring, and review. Extend existing organizational governance with check lists and guidelines on how to review proposed AI solutions and criteria to escalate reviews to a full risk assessment for high-impact systems or models involving security, legal, compliance, and business units.

ML Fraud Detection: Review NIST AI RMF [i.17] to see whether it can be helpful in standardizing risk AI assessments and how to interface it with other GRC processes in the organization.

LLM Platform: See Chatbot App example.

Open-Access LLM Model: This is not applicable to this example due to the size of the team. Instead, the team documents in their Wiki page how they perform risk assessments.

References:

- NIST AI RMF [i.17]
- NIST AI RMF Playbook [i.40]
- NIST AI RMF Crosswalk Documents [i.41]
- ICO Accountability and Governance Implications for AI [i.42]
- European AI Alliance: Implementing AI Governance: from Framework to Practice [i.43]
- OWASP AI Security Center of Excellence (CoE) Guide [i.44]

6.2.2 Provision 5.1.2-1.1

"Where the Data Custodian is part of a Developers organization, they shall be included in internal discussions when determining the requirements and data needs of an AI system." (ETSI TS 104 223 [i.1])

Related threats/risks:

Failure to include the Data Custodian in discussions about AI system requirements and data needs can result in non-compliance with data governance policies, inappropriate data usage, or insufficient safeguards for sensitive data, increasing the risk of data breaches or regulatory violations.

Example Measures/Controls:

Ensure collaboration with the Data Custodian: During the design and development phases to define data requirements to identify regulatory compliance requirements. Ensure that Data Custodians are able to balance additional risks to data that come from the AI system with intended mitigations and the business need.

Chatbot App: Include data governance checklists as part of design discussions. Schedule workshops with Data Custodians, developers, and security staff to ensure ongoing alignment on data needs, compliance requirements, and data access implications.

ML Fraud Detection: Include Data Custodian reviews and feature signoffs when using personal data; provide explanations of how data will be used, risks such as memorization, extraction and inference, and options to safeguard use.

LLM Platform: Align with internal governance processes to involve Data Custodians in defining data usage and ensuring, compliance.

Open-Access LLM Model: Define who is the Data Custodian in the team and have them work with the rest of the team to perform and document DPIAs that are reviewed as part of the design process. Review the national guidelines to ensure the acting Data Custodian is performing their role in compliant manner.

References:

- ICO AI Data Protection Risk Toolkit [i.35]
- NIST AI RMF [i.17]
- NIST AI RMF Playbook [i.40]

6.2.3 Provision 5.1.2-2

"Developers and System Operators shall ensure that AI systems are designed and implemented to withstand adversarial AI attacks, unexpected inputs and AI system failure." (ETSI TS 104 223 [i.1])

Related threats/risks:

Organizations are not always successful in preventing breaches and so defence in depth requires planning for and responding to compromise. The absence of defence in depth can lead to infringement of applicable regulation.

Example Measures/Controls:

Apply Secure by Design Principles: Integrate security into the AI system's design phase by conducting threat modelling. Threat modelling covers both traditional cyber threats and AI-specific ones that might be introduced by the design choices. Incorporate standardized security controls in the system design controls to mitigate risks. Document each standardized control used in the design phase, and ensure it is integrated with specific test cases to verify its effectiveness during system testing. Ensure monitoring controls as well as incident response and recovery from failures are addressed in threat mitigation and they are documented in the system's design.

Chatbot App: Use threat modelling to identify risks specific to the chatbot app; apply controls for general application security and ones relevant to OWASP Top 10 for LLM applications [i.11], such as implementing input validation to prevent prompt injection attacks to the LLM it uses as well as preventing sensitive data exposure.

ML Fraud Detection: In addition to application security controls, incorporate controls in the design for relevant predictive adversarial AI attacks such as poisoning, evasion, model extraction, and other privacy attacks. Use OWASP AI exchange as threats and controls reference.

LLM Platform: Use both MITRE ATT&CK [i.45] and MITRE ATLAS [i.8] to perform threat modelling of both model development and operation, addressing threats from adversarial AI attacks, especially ones related to LLMs such as poisoning and safety measures to prevent jailbreaking, data extraction, and unsafe use. Review customer-facing APIs and include them in threat modelling with usage scenarios.

Open-Access LLM Model: Use a similar approach in the LLM Provider example, focusing on training but also how others might use the model and the safeguards it needs to have in place to protect them. Follow a lightweight approach to Threat Modelling as part of the design and use either MITRE ATLAS [i.8] or OWASP AI Exchange [i.10] as threats and controls library.

References:

- CISA Joint Cybersecurity Information - Deploying AI Systems Securely [i.7]
- CSA Companion Guide on Securing AI Systems [i.46], Section 2.2.1 Planning and Design
- MITRE ATLAS [i.8]
- MITRE ATT&CK [i.45]
- NCSC Secure Design Principles [i.47]
- NCSC Principles for Machine Learning [i.5]
- NCSC Guidelines for Secure AI system development [i.6]
- OWASP AI Exchange [i.10]

6.2.4 Provision 5.1.2-3

"To support the process of preparing data, security auditing and incident response for an AI system, Developers shall document and create an audit trail in relation to the AI system. This shall include the operation, and life cycle management of models, datasets and prompts incorporated into the system." (ETSI TS 104 223 [i.1])

Related threats/risks:

A lack of audit trails can lead to untraceable changes or unauthorized adjustments, complicating incident response, forensic investigations, and regulatory compliance.

Example Measures/Controls:

Automated Audit Trails for ML Operations (MLOps) and System changes: Implement automated logging for all critical operations related to model training, dataset changes, prompts and parameter adjustments. For critical systems with compliance requirements, use WORM (write-once, read-many) storage to store logs, ensuring they remain tamper-proof and accessible for audits.

Chatbot App: Use a version control system to system prompt and prompt all changes with related documentation. If the model provider's API is used for fine-tuning, ensure the training and testing datasets are logged. For RAG workflows, track the embeddings generated, log metadata about retrieved data (e.g. query terms, document IDs), and maintain versioned snapshots of smaller reference datasets used in retrieval to ensure traceability and reproducibility.

ML Fraud Detection: Use an ML Ops platform or MLOps functionality in cloud platforms to enforce versioning by tracking of all changes for models, prompts, and other experimentation. Complement MLOps with data life cycle management tools to implement similar versioning for datasets, storing details as to what data was used to train or test the system.

LLM Platform: Use the same approach described in the ML Fraud Detection example.

Open-Access LLM Model: Use open-source tools for tracking and automating workflows, such as versioning datasets, model configurations, and changes through APIs as part of regular workflows.

References:

- CSA Companion Guide on Securing AI Systems [i.46], Section 2.2.1 Planning and Design
- MLOps Principles [i.48]
- NCSC Principles for Machine Learning, Part 1, Secure Design [i.5]
- NCSC Guidelines for Secure AI System Development, Secure Design [i.6]
- OWASP AI Exchange [i.10]

6.2.5 Provision 5.1.2-4

"If a Developer or System Operator uses an external component they shall conduct an AI security risk assessment and due diligence process in line with their existing software development processes, that assesses AI specific risks." (ETSI TS 104 223 [i.1])

Related threats/risks:

Third-party components introduce risks through possible vulnerabilities in the external vendor's security practices which might not be to the same standard. This includes operating systems and libraries, container images, programming packages as well as models and datasets. Large models contain general purpose functionality, ensure that specific risks are mitigated or otherwise managed.

In the context of product safety regulation, risks to AI security include risks that can lead to hazardous situations. Hazardous situations can include risks to fundamental rights, health and safety.

For further information about risk management in the context of the EU's AI Act, see JTC21024 [i.146].

Example Measures/Controls:

Security Due-Diligence for External Components: Mandate a risk assessment process before a component (including external models) can be used, covering provenance, known risks, and when personal data is used a DPIA. Safeguard provenance by mandating in internal standards that components can only be sourced by trusted and approved sources, documenting source, version, licencing, history, and other related artifacts (e.g. Model Card for models); use checksums to verify integrity.

Chatbot App: Review documentation including known vulnerabilities and run automated vulnerability scans against application and platform packages; consult published documentation and benchmarks or conduct independent tests against the LLM model used by the chatbot. Include embedding models in the diligence if they are used to generate embeddings as part of RAG, instead of APIs.

ML Fraud Detection: As in the Chatbot app example, but with additional diligence for the third-party base model. Consult model documentation and use tools to scan for vulnerabilities such as serialization attacks. Store approved models and components in an internal repository, ensuring they are the only ones used in production. Set up alerts for changes or security notifications for the external components.

LLM Platform: Similar to Fraud Detection example but include auxiliary models that might be used. For instance, smaller models to provide embeddings API for RAG use cases and RL models for RLHF and RLFAI in the fine tuning. External datasets are of critical importance and need to be evaluated for copyright, privacy, bias, and ethical risks.

Open-Access LLM Model: Automate scans for components and models in use (foundation for fine tuning, auxiliary for testing RAG, RL for RLHF and RLFAI scenarios) and ensure they are from trusted sources. Review external datasets sourced only from reputable sources and use automated tools to detect bias, personal data, and copyright issues.

References:

- ETSI TR 104 048 [i.145]
- NIST Cybersecurity Supply Chain Risk Management [i.49]
- NCSC Supply Chain Security Guidance [i.51]
- OWASP AI Exchange [i.10]

- OWASP Top 10 for LLM applications - LLM03 Supply-Chain Vulnerabilities [i.11]

6.2.6 Provision 5.1.2-5

"Data Custodians shall ensure that the intended usage of the system is appropriate to the sensitivity of the data it was trained on as well as the controls intended to ensure the security of the data." (ETSI TS 104 223 [i.1])

Related threats/risks:

Misalignment between the intended usage of the AI system and the sensitivity of the data it was trained on can result in inappropriate data exposure, inadequate security controls, and regulatory non-compliance, leading to potential data breaches and misuse of personal data or other confidential information.

Example Measures/Controls:

Ensure Data Custodian Assurance: Require Data Custodians to review system's intended usage and the data security controls to ensure compliance and balancing these risks with business needs.

Chatbot App: Review chatbot use cases with Data Custodian, access and usage policies and ensure they are aligned with the DPIA.

ML Fraud Detection: As in the Chatbot app, but including data used for training, data memorization, inversion and inference risks, and the implications GDPR for automated fraud detection.

LLM Platform: Similar to the Fraud Detection example but integrating Data Custodian review to governance with a multi-disciplinary board including legal and data protection experts to sign off.

Open-Access LLM Model: Ensure the acting Data Custodian reviews with the rest of the team the design and plan and ensures it is compliant and aligned with the DPIA and that there are sufficient controls to mitigate personal data leakage through training data extraction and prompt injection attacks.

References:

- ICO UK GDPR Guidance and Resources [i.52]
- ICO What is the impact of Article 22 of the UK GDPR on fairness? [i.34]
- NCSC Protecting bulk personal data [i.53]
- GDPR Compliance Guidelines by EU Commission [i.54]
- ISO/IEC 27001:2022 [i.55]

6.2.7 Provision 5.1.2-5.1

"Organizations should ensure that employees are encouraged to proactively report and identify any potential security risks in AI systems and ensure appropriate safeguards are in place." (ETSI TS 104 223 [i.1])

Related threats/risks:

A lack of proactive reporting and identification of security risks in AI systems can lead to undetected vulnerabilities, increasing the likelihood of security breaches, data leaks, or misuse of AI, with potentially significant operational, financial, and reputational consequences.

Example Measures/Controls:

Support proactive reporting of security risks. Establish a clear, accessible process for employees to report potential security risks in AI systems, encourage a culture of proactive risk identification by providing training, communication channels, transparent handling, and recognition for reporting issues.

Chatbot App: Develop an incident reporting template specifically for chatbot-related risks (e.g. sensitive data leakage or inappropriate responses). Use collaborative tools (e.g. messaging channels) to establish a dedicated risk-reporting channel.

ML Fraud Detection: Train employees to identify potential risks such as biases in fraud detection or false positives/negatives. Provide an anonymous reporting mechanism for concerns and include follow-ups on how identified risks are addressed.

LLM Platform: Create a centralized risk registry for employees to log concerns about API misuse, data exposure, or unexpected system outputs. Provide regular updates on how identified risks are managed and mitigated.

Open-Access LLM Model: Provide a checklist for team members and external contributors to log risks as tickets in team's work management board. Review reported risks as part of the work and ensure resolution steps are documented and shared.

References:

- ICO Reporting Processes [i.56]
- ICO Breach identification, assessment and logging [i.57]
- NCSC Developing a positive cyber security culture [i.58]
- NCSC Responding to a cyber incident - a guide for CEOs [i.59]

6.2.8 Provision 5.1.2-6

"Where the AI system will be interacting with other systems or data sources, (be they internal or external), Developers and System Operators shall ensure that the permissions granted to the AI system on other systems are only provided as required for functionality and are risk assessed." (ETSI TS 104 223 [i.1])

Related threats/risks:

There is huge potential and interest in "agentic systems", where an AI system can decide and conduct its own actions, typically through integrations with other systems. The actions taken by an AI system are not fully predictable, and can be coerced by an attacker. This risk introduces the potential for unauthorized access, data exfiltration, and privilege escalation.

Example Measures/Controls:

Least-privilege access to data and systems accessed by AI System. Mandate a risk assessment process before a component can be used covering provenance, known risks, and evaluations. Ensure that the assessment covers all possible model states, not just the designed or expected ones.

Chatbot App: If app uses external services to enrich LLM input or drive systems (for instance a booking system), implement data minimization and granular least-privilege access policies for integration endpoints. Examine what would happen if the proposed integrations are used in the wrong order or for unintended purposes. Review against Excessive Agency as defined in OWASP Top 10 for LLM Applications and consider listing all the proposed integrations and their permissions and asking a security specialist what harm they could cause the organization if given those permissions and integrations.

ML Fraud Detection: Review and test the data inputs used (including data preprocessing and enrichment) and ensure only the required data is used. Evaluate side-effects if model outputs (predictions) are used to drive downstream services e.g. automated processes.

LLM Platform: If the system allows adding extra features, such as plugins or connections to other tools (e.g. a calendar app, an API, or an email service), test how these features work together and check for risks. For example, ensure that the system doesn't produce harmful or incorrect results when someone uses these extra features, like accessing confidential data, executing system commands, or sending misleading emails.

Open-Access LLM Model: The small organization follows a similar approach to the LLM Provider but tailored to its workflows and level of resources.

References:

- ICO Assessing security and data minimisation in AI [i.22]
- MITRE ATLAS, Privilege Escalation [i.8]

- NCSC Using a cloud platform securely - Apply access controls [i.60]
- NCSC Zero trust architecture design principles [i.61]
- NIST AI RMF [i.17]
- OWASP Top 10 for LLM Applications: Excessive Agency [i.11]

6.2.9 Provision 5.1.2-7

"If a Developer or System Operator chooses to work with an external provider, they shall undertake a due diligence assessment and should ensure that the provider is adhering to the present document." (ETSI TS 104 223 [i.1])

Related threats/risks:

Collaborating with external providers without assessing their adherence to ETSI TS 104 223 [i.1] can lead to increased vulnerabilities, lack of regulatory compliance, insecure systems, or inadequate response protocols, which could compromise the entire system's security.

Example Measures/Controls:

Security Review of External Providers: Verify the external provider's implementation of ETSI TS 104 223 [i.1] with the external provider and their overall regulatory compliance.

Chatbot App: Ask the cloud, LLM, and other providers to provide evidence of ETSI TS 104 223 [i.1] compliance.

ML Fraud Detection: Ask the cloud or other service providers to provide evidence of ETSI TS 104 223 [i.1] adherence.

LLM Platform: Follow the same approach as the ML Fraud Detection example.

Open-Access LLM Model: Ask the cloud and other service providers for evidence of ETSI TS 104 223 [i.1] adherence; In the absence of specific ETSI TS 104 223 [i.1] adherence documentation, review provider documentation to ascertain adherence and document the findings in a wiki page.

References:

- ETSI TS 104 223 [i.1]

6.3 Principle 3: Evaluate the threats and manage the risks to the AI system

6.3.1 Provision 5.1.3-1

"Developers and System Operators shall analyse threats and manage security risks to their systems. Threat modelling should include regular reviews and updates and address AI-specific attacks, such as data poisoning, model inversion, and membership inference." (ETSI TS 104 223 [i.1])

Related threats/risks:

AI systems face unique threats, such as data poisoning, model inversion, and membership inference attacks, which traditional threat models cannot account for. New threats will emerge that will need to be incorporated in threat modelling and risk management.

Example Measures/Controls:

Perform Threat Modelling including AI threats: Apply threat modelling that captures potential impacts on stakeholders including both AI and traditional cyberattacks. Document each identified threat in detail, outlining the likelihood and severity of potential impacts to the AI model and the broader system and list mitigations using standardized OWASP or MITRE controls. Both OWASP AI Exchange [i.10] or MITRE ATLAS [i.8] provide threat taxonomies and related mitigations which both types of attacks and they can be used in threat modelling. If the AI system processes or was built on personal data ICO's guidance on AI and security is useful to consult for regulatory compliance.

Chatbot App: Map threats and data flows for the app focusing on traditional and LLM threats. Include un-used functionality of the chosen models or components. For instance, if a multi-modal model is being used just for language, then model the risks if someone were to conduct attacks or abuse through giving it images.

ML Fraud Detection: Include both the inference system and the development environment to cover poisoning, model tampering, serialization attacks in addition to run-time attacks such as evasion, extraction, inversion, and inference.

LLM Platform: Follow the approach described in Fraud Detection, but with focus on generative AI risks such as overreliance, jailbreaking the LLM, and their safety, ethical, social, and regulatory consequences.

Open-Access LLM Model: The small organization performs the same type of modelling as in the previous LLM Platform example.

References:

- ICO Assessing security and data minimisation in AI [i.22]
- NIST AI RMF [i.17], AI RMF Core - Map
- NCSC Risk Management - Threat Modelling [i.39]
- OWASP Threat Modeling Process [i.62]
- MITRE ATLAS [i.8]
- OWASP Top 10 for LLM applications [i.11]
- OWASP AI Exchange [i.10]

6.3.2 Provision 5.1.3-1.1

"The threat modelling and risk management process shall be conducted to address any security risks that arise when a new setting or configuration option is implemented or updated at any stage of the AI lifecycle." (ETSI TS 104 223 [i.1])

Related threats/risks:

Failure to conduct threat modelling and risk management when implementing or updating settings or configurations during the AI lifecycle can lead to unmitigated security vulnerabilities, such as configuration errors or unanticipated attack vectors, increasing the risk of exploitation and system compromise.

Example Measures/Controls:

Conduct Threat Modelling for Configuration Changes: Perform threat modelling whenever settings or configurations are implemented or updated to identify and mitigate security risks throughout the AI lifecycle.

Chatbot App: When enabling or modifying user feedback options, assess potential risks, such as injection attacks through input fields, and implement mitigations like input validation and sanitisation.

ML Fraud Detection: When feature selection or detection thresholds change, or new geolocation risk tables are deployed, evaluate risks like adversarial evasion attacks. Develop safeguards such as monitoring for unusual patterns and conducting stress tests on configurations.

LLM Platform: When changing user authentication settings, evaluate threats for unauthorized access, and implement mitigations like rate limiting, lockouts, and enhanced logging.

Open-Access LLM Model: When modifying the model's configuration to allow community contributions or plugins, assess and mitigate risks such as the introduction of malicious code or unintended functionalities.

References:

- NIST SP 800-154 [i.63]
- NCSC Introduction to Logging for Security Purposes [i.64]
- OWASP Top 10 for APIs - API4:2019 Lack of Resources & Rate Limiting [i.65]
- OWASP Threat Modeling in Practice [i.66]

6.3.3 Provision 5.1.3-1.2

"Developers shall manage the security risks associated with AI models that provide superfluous functionalities, where increased functionality leads to increased risk. For example, where a multi-modal model is being used but only single modality is used for system function." (ETSI TS 104 223 [i.1])

Related threats/risks:

Allowing AI models to retain superfluous functionalities that are not required for the system's purpose can introduce unnecessary security risks, such as expanded attack surfaces, increased vulnerability to exploitation, and potential misuse of unused features, compromising the overall security of the system.

Example Measures/Controls 1:

Restrict Superfluous Functionalities: Limit AI model functionalities to those essential for the system's purpose to reduce the attack surface and minimize security risks associated with unused features.

Chatbot App: If the chatbot is implemented for text-based customer support, use guardrails or API blocking to disable or restrict any access to unused multimodal capabilities, such as speech-to-text or text-to-speech features to prevent unintended interactions or vulnerabilities.

ML Fraud Detection: If the system only analyses transactional data, remove or disable unnecessary model features, such as image processing or location-based predictions to minimize risks and complexity. Only enable advanced features, such as the use of RL algorithm to explore a complex space, if the security implications are fully understood.

LLM Platform: Provide different APIs for text from ones including advanced capabilities like multimodal input (e.g. image processing) for application only uses text. Provide specific voices in text to voice scenarios instead of allowing voice cloning.

Open-Access LLM Model: Since the focus is for legal advice and contract negotiation, implement safety measures to disable general purpose or other use. Provide lightweight documentation to users about the rationale and security benefits of these restrictions.

References:

- NCSC Secure Design Principles [i.47]
- OWASP AI Top 10 API Security Risks - 2023 [i.67]
- OWASP Top 10 for LLM Applications [i.11], Excessive Agency
- Building Guardrails for Large Language Models [i.68]

Example Measures/Controls 2:

Integrate Threat Modelling with AI Governance: Require completed threat models for governance approval at critical stages of the AI lifecycle, ensuring documented risk understanding and mitigation before deployment providing support and guidance and ensuring cross-discipline input (ethics, privacy, legal, etc.) to threat modelling.

Chatbot App: Establish governance policies that mandate a formal review of the threat model, including stakeholder and impact assessments, prior to each major system deployment. Provide a standardized threat modelling template using standard notation e.g. STRIDE or PASTA [i.24] with AI threats found in MITRE ATLAS [i.8] or OWASP AI Exchange [i.10].

ML Fraud Detection: Introduce the need for a threat model as part of deployment approval.

LLM Platform: As in the Chat Bot App, with the addition of a formal multi-disciplinary threat model review before approval.

Open-Access LLM Model: This is not applicable to small organizations; instead, facilitate threat model by using reusable standardized templates in free diagrammatic tools.

References:

- OWASP Threat Modeling Playbook [i.69]
- NIST AI RMF [i.17], AI RMF Core
- NCSC Risk Management - Cybersecurity Risk Management Framework [i.70]

6.3.4 Provision 5.1.3-1.3

"System Operators shall apply controls to risks identified through the analysis based on a range of considerations, including the cost of implementation in line with their corporate risk tolerance." (ETSI TS 104 223 [i.1])

Related threats/risks:

When risk tolerance is not clearly defined in the context of AI-specific risks such as data poisoning or model misuse, this could result into Inadequately prioritized controls leading to breaches, operational disruptions, or unethical decision-making.

Example Measures/Controls:

Develop a Prioritization Framework for AI Risk Controls: Use an AI-specific risk-scoring system to prioritize mitigations and controls based on the impact of threats, likelihood of occurrence, and in alignment with organizational risk tolerance. This should account for regulatory risk, including data protection, and cover AI-specific vulnerabilities, such as adversarial manipulation, model drift, and bias.

Chatbot App: Indirect Prompt Injections and Bias are rated as High when the app is used for recruitment but Low when used for summarization of internal documentation.

ML Fraud Detection: Poisoning Backdoor and Evasion attacks are prioritized as High to avoid costly fraudulent transactions.

LLM Platform: AI risks identified for general-purpose models are elevated as High including copyright violations which can result into legal liability.

Open-Access LLM Model: Follow a similar approach to the LMM Platform example.

References:

- NIST AI RMF Playbook [i.40]
- NIST AI RMF [i.17], Use Cases, Autonomous Vehicle Risk Management Profile for Traffic Sign Recognition

6.3.5 Provision 5.1.3-2

"Where AI security threats are identified that cannot be resolved by Developers, this shall be communicated to System Operators so they can threat model their systems. System Operators shall communicate this information to End-users, so they are made aware of these threats. This communication should include detailed descriptions of the risks, potential impacts, and recommended actions to address or monitor these threats." (ETSI TS 104 223 [i.1])

Related threats/risks:

Without clear communication on unresolved risks, System Operators and End-users can lack awareness, limiting their ability to apply safeguards effectively.

Example Measures/Controls:

Document and Communicate Identified Unresolved Risks: Ensure clear documentation and timely communication of any unresolved threats to all relevant stakeholders.

Chatbot App: If the app performs document summarization, the application maybe vulnerable to indirect prompt injection. Ensure system operators are aware so that they can introduce mitigations such PDF checks and reviews.

ML Fraud Detection: Notify system operators to develop real-time monitoring of inputs for suspicious patterns to cover residual evasion attack risks

LLM Platform: Document public API documentation with known risks and how to mitigate them.

Open-Access LLM Model: Inform model users about potential risks like model misuse for generating biased or harmful outputs, and document recommended safeguards such as usage guidelines or implementing content moderation mechanisms.

References:

- NIST AI RMF Playbook [i.40]
- NIST AI RMF [i.17], Use Cases, Autonomous Vehicle Risk Management Profile for Traffic Sign Recognition

6.3.6 Provision 5.1.3-3

"Where an external entity has responsibility for AI security risks identified within an organizations infrastructure, System Operators should attain assurance that these parties are able to address such risks." (ETSI TS 104 223 [i.1])

Related threats/risks:

Reliance on third parties without adequate verification could expose the AI system to unmanaged vulnerabilities.

Example Measures/Controls:

Conduct AI-Specific Security Assessments for Third Parties: Ensure third-party components and vendors undergo security assessments that specifically address AI-related risks and adherence with ETSI TS 104 223 [i.1].

Chatbot App: Request from the model provider to provide assurances addressing specific AI risks, such as model safety and secure data handling.

ML Fraud Detection: Require similar assurances similar to the Chatbot App example but from external data providers to ensure data used for training is handled responsibly and securely.

LLM Platform: See the ML Fraud Detection Example.

Open-Access LLM Model: Implement a lightweight approach by focusing on key external dependencies. For example, request a simple self-assessment checklist from any external providers (e.g. cloud hosting, pre-trained models, or APIs) to ensure they address basic AI risks such as data privacy, model integrity, and adherence to security standards. Limit reliance on complex external integrations to reduce potential vulnerabilities.

References:

- ISO 9001 What does it mean in the supply chain? [i.71]

- NCSC Supply Chain Security Guidance [i.51]
- NCSC Cloud Security Guidance [i.60] - Choosing a cloud provider
- World Economic Forum: Adopting AI Responsibly: Guidelines for Procurement of AI Solutions by the Private Sector: Insight Report [i.73]

6.3.7 Provision 5.1.3-4

"Developers and System Operators should continuously monitor and review their system infrastructure according to risk appetite. It is important to recognize that a higher level of risk will remain in AI systems despite the application of controls to mitigate against them." (ETSI TS 104 223 [i.1])

Related threats/risks:

Residual risk can be exploited by malicious actors, especially as evolving threats introduce new vulnerabilities or amplify existing ones, leading to potential breaches, disruptions, or compromised AI integrity.

Example Measures/Controls:

Establish Continuous AI Risk Monitoring Controls: Implement a regular review processes of AI developments to determine whether emerging vulnerabilities, improved mitigation techniques, or advancements in AI models necessitates updates to the risk assessment controls.

Chatbot App: Threat intelligence feeds report that prompt injection attacks, including advanced techniques like using emoticons to bypass safety measures, have suddenly gained popularity in the hacking community. As websites and applications face widespread probing for these vulnerabilities, the organization mitigates the risk by developing and deploying in-house guardrails to detect and block such attacks.

ML Fraud Detection: An updated version of the third-party model used for fraud-detection includes additional adversarial training to withstand evasions, leading to its selection, fine tuning and deploying the new version.

LLM Platform: A new government report highlights risks associated with audio generation, such as the potential for voice cloning to bypass voice authentication systems. In response, the platform develops new safeguards to prevent misuse of its audio generation capabilities.

Open-Access LLM Model: Review of a new research paper demonstrates a novel attack vector to jailbreak model, necessitating the development of additional safety features.

References:

- MITRE AI Risk Database [i.74]
- NIST SP 800-137 [i.75]
- NCSC Early Warning [i.76]
- NCSC Building a Security Operations Centre (SOC) - Threat Intelligence [i.77]

6.4 Principle 4: Enable human responsibility for AI systems

6.4.1 Provision 5.1.4-1

"When designing an AI system, Developers and/or System Operators should incorporate and maintain capabilities to enable human oversight." (ETSI TS 104 223 [i.1])

Related threats/risks:

Without built-in human oversight, AI systems will generate incorrect outputs or decisions that are difficult to interpret, verify, or override, increasing risks of data protection compliance, unintended consequences, misuse, or harmful impacts.

Example Measures/Controls 1:

Implement Mechanisms for Human Oversight: Control: Implement features that allow human operators to easily interpret, verify, and act on AI outputs, including manual release and overrides. Ensure that the design meet obligations around automated decisions and encourages meaningful human decision-making rather than passive acceptance of AI recommendations.

Chatbot App: For new features that allow the chatbot to take autonomous actions, such as scheduling appointments or order office supplies, an override control has been applied to cancel appointments or orders, because of its low to moderate impact. By contrast a feature to provide personalized packages to customers, has high reputational and legal risks, as a result manual release control has been implemented with an operator review-and-approve interface.

ML Fraud Detection: Explanation techniques such as SHapley Additive exPlanations (SHAP) or Local Interpretable Model-agnostic Explanations (LIME) are implemented, to enable operators to understand the key factors behind each decision and intervene, which grants individuals the right not to be subject to decisions based solely on automated processing that produce legal effects concerning them or that significantly affect them.

LLM Platform: Incorporate a feature that allows human moderators to review and approve AI-generated content before publication, ensuring outputs align with ethical guidelines and community standards.

Open-Access LLM Model: Provide users with tools to flag and report inappropriate or harmful AI-generated content, facilitating human oversight and continuous improvement of the model's outputs.

References:

- ICO Audit Framework toolkit on AI - Human review [i.78]
- ICO Guidance on AI and Data Protection - What is the impact of Article 22 of the UK GDPR on fairness? [i.34]
- CSA Companion Guide on Securing AI Systems [i.46], Section 2.2.4 Operations and Maintenance
- NISTIR 8312 [i.79]
- NCSC Machine Learning Principles [i.5]
- NCSC Guidelines for Secure AI system development, Secure Operation and Maintenance [i.6]

Example Measures/Controls 2:

Measure and Validate Accuracy of Human Oversight Decisions: Regularly test and measure the accuracy of human oversight decisions, validating that operators can correctly interpret and act on AI outputs and identifying areas for improvement. Assess not just individual performance but how the system supports human understanding and engagement to foster effective sociotechnical communication between the operator and AI.

Chatbot App: For a hiring version of the chatbot app, regularly review a sample of candidates who were automatically flagged as unsuitable by the AI. Assess whether human operators correctly validated or overrode these decisions. Identify patterns of misinterpretation or bias in operator actions, and use these insights to refine training, decision guidelines, or the AI's recommendation criteria.

ML Fraud Detection: Test whether human reviewers accurately validate flagged transactions, ensuring they can effectively distinguish between true fraud cases and false positives.

LLM Platform: Evaluate how well operators identify and correct biased or inappropriate content generated by the platform, using flagged examples to improve content moderation guidelines and model outputs.

Open-Access LLM Model: This is a nice to have for a small organization and in this case the company might conduct experiments to evaluate model responses and feedback from operators on how to improve as in the above.

References:

- ICO Audit Framework toolkit on AI - Human review [i.78]
- ICO Explaining Decisions made with AI [i.80]
- NISTIR 8312 [i.79]

6.4.2 Provision 5.1.4-2

"Developers should design systems to make it easy for humans to assess outputs that they are responsible for in said system (such as by ensuring that models outputs are explainable or interpretable)." (ETSI TS 104 223 [i.1])

Related threats/risks:

Without clarity and ease of use, users can not perform oversight effectively leading to failures and harm.

Example Measures/Controls:

Develop User-Friendly Human Responsibility UI: Implement UIs that display outputs, decision-making rationales, and logs clearly to make it easy for human operators to assess outputs and understand their accountability. Ensure systems are designed to encourage rigorous assessment by humans and not condition them to simply click an approve button.

Chatbot App: As the chatbot is extended to cover new cases, studies of the existing usage are analysed and consolidated in an UX library for consistent implementation of oversight and human responsibility features.

ML Fraud Detection: Include a dashboard that shows flagged transactions with explanations of the model's decision factors, providing a clear interface for review.

LLM Platform: Review explainability review and provide a web interface for API requests to include confidence scores and reasoning summaries.

Open-Access LLM Model: Offer an API with an option to receive back responses to legal queries with references referencing specific legal principles or past cases that influenced the output.

References:

- ICO Explaining Decisions made with AI [i.80]
- NISTIR 8312 [i.79]
- Multicalibration for Confidence Scoring in LLMs [i.81]
- From Understanding to Utilization: A Survey on Explainability for Large Language Models [i.82]

6.4.3 Provision 5.1.4-3

"Where human oversight is a risk control, Developers and/or System Operators shall design, develop, verify, and maintain technical measures to reduce the risk through such oversight." (ETSI TS 104 223 [i.1])

Related threats/risks:

Ineffective oversight technical measures can compromise the risk reduction effort by overburdening or failing to adequately support human reviewers.

Example Measures/Controls:

Implement Validation and Enforcement of Oversight controls: Design and implement technical measures that provide guardrails to assist human reviewers in understanding, interpreting, and acting on AI outputs.

Chatbot App: For a chatbot that can make automated triaging decisions, provide human reviewers with a summary of the chatbot's reasoning for triage decisions (e.g. key user inputs) and allow them to adjust or override decisions before escalation.

ML Fraud Detection: Provide human reviewers with flagged transactions prioritized by risk level and accompanied by explainable insights (e.g. key features influencing the fraud score), to ensure informed and efficient decision-making.

LLM Platform: Integrate API-based guardrails that automatically flag AI-generated content containing sensitive information or potential biases, providing a score as an API field for human reviewers to identify outputs requiring attention.

Open-Access LLM Model: Train the model with additional safety features to apply AI security measures such as adversarial robustness checks, and warning mechanisms for potentially misleading or harmful outputs. These measures can include applying confidence thresholds, logging flagged outputs, and ensuring model responses align with compliance policies.

References:

- ICO Audit Framework toolkit on AI - Human review [i.78]
- Building Guardrails for Large Language Models [i.68]

6.4.4 Provision 5.1.4-4

"Developers should verify that the security controls specified by the Data Custodian have been built into the system." (ETSI TS 104 223 [i.1])

Related threats/risks:

Without validation of Data Custodian controls, the system can lack necessary data protection and governance measures, potentially leading to security vulnerabilities or regulatory non-compliance.

Example Measures/Controls:

Conduct Validation of Custodian: Verify that all controls specified by the Data Custodian have been implemented correctly, with testing to validate effectiveness and alignment with data protection requirements and guidance.

Chatbot App: Ensure that data retention policies adhere to the Data Custodian's specifications by implementing automated data purging mechanisms and conducting regular audits to confirm compliance.

ML Fraud Detection: Validate that data anonymization controls specified by the Data Custodian are active and effective in protecting customer privacy.

LLM Platform: Confirm that access controls and encryption protocols for the LLM platform's API endpoints meet the Data Custodian's requirements by performing penetration testing and security assessments.

Open-Access LLM Model: Verify that the model's training data complies with the Data Custodian's guidelines on data sourcing and consent by reviewing data collection processes and conducting compliance checks.

References:

- ICO Audits [i.83]
- ICO Audits, Artificial Intelligence Audits [i.84]

6.4.5 Provision 5.1.4-5

"Developers and System Operators should make End-users aware of prohibited use cases of the AI system." (ETSI TS 104 223 [i.1])

Related threats/risks:

Without clear communication on prohibited uses, end-users can unintentionally misuse the AI system, leading to legal, ethical, or operational risks.

Example Measures/Controls 1:

Document and Train Users on Prohibited Use Cases: Clearly define and document prohibited use cases for the AI system, ensuring end-users understand limitations and restrictions. Use threat modelling to identify and inform users of all known harmful states and unmitigated risks.

Chatbot App: In document summarization use of the app, guide the user at the beginning of each conversation and online documentation not to upload classified internal documents, explaining the risks of data memorization and leaks.

ML Fraud Detection: Threat modelling has identified that AI could be abused to monitor a person's spending habits without consent. Document and communicate to operators that using the system for non-compliance-related surveillance is prohibited. Provide training on ethical boundaries and enforce compliance audits to ensure proper usage.

LLM Platform: Document prohibited use cases in use policies and T&Cs for APIs.

Open-Access LLM Model: Threat modelling highlighted that the open-access LLM could be fine-tuned or deployed to generate misinformation or manipulate public opinion. Clearly communicate in the licensing terms and documentation that using the model for misinformation campaigns or malicious automation (e.g. phishing scams) is strictly prohibited.

References:

- ICO Accountability and Governance Implications for AI [i.42]

Example Measures/Controls 2:

Monitor for Prohibited Use Cases: Implement controls to actively monitor, detect, and prevent prohibited use cases.

Chatbot App: Implement guardrails to detect and block use of personal data supported with additional automated tests and periodic audits of log.

ML Fraud Detection: Implement monitoring of patterns of unauthorized access or analysis triggering escalation alerts.

LLM Platform: Use finetuning to add safety measures blocking prohibited use cases, API guardrails to detect and prevent misuse, activity monitoring, and output watermarking to detect misuse in prohibited cases.

Open-Access LLM Model: Finetune to implement safety measures to detect and block prohibited uses, and watermarking model output to identify misuses in phishing attacks.

References:

- ETSI TR 104 032 [i.12], clause 5.3
- NIST AI RMF Playbook [i.40]
- OWASP Top 10 for LLM applications [i.11]
- OWASP AI Exchange [i.10]
- Building Guardrails for Large Language Models [i.68]
- OWASP LLM and Generative AI Security Solutions Landscape [i.85]

6.5 Principle 5: Identify, track, and protect assets

6.5.1 Provision 5.2.1-1

"Developers, Data Custodians and System Operators shall maintain a comprehensive inventory of their assets (including their interdependencies / connectivity)." (ETSI TS 104 223 [i.1])

Related threats/risks:

Without a clear understanding of AI assets and their dependencies, organizations can not be able to provide protection and risk exposing their AI systems to unauthorized access, data leakage, and vulnerability to external attacks.

Example Measures/Controls:

Establish an AI Asset Inventory: Create and maintain a centralized inventory that records all AI assets, including datasets, models, software dependencies, hardware resources, and system configurations.

Chatbot App: Document all chatbot versions customized for different use cases, training datasets, software libraries, and APIs used. Include components used for RAG and/or embeddings generation.

ML Fraud Detection: Use a detailed register of all AI models in use, listing model version, source, datasets used for training, and any updates or retraining conducted. This includes the software libraries used in the development or running of the models. As an optional safeguard use scripts and scanning tools to generate an ML BOM as machine readable inventory list

LLM Platform: Use the same approach as in the Fraud Detection example.

Open-Access LLM Model: Use the same approach as in the Fraud Detection example.

References:

- NCSC Machine Learning Principles, 2.3 Manage the full life cycle of models and datasets [i.5]
- NCSC Guidelines for Secure AI System Development, Secure Development [i.6]
- CSA Companion Guide on Securing AI Systems [i.46], Section 2.2.2 Development, Item 2.3

6.5.2 Provision 5.2.1-2

"As part of broader software security practices, Developers, Data Custodians and System Operators shall have processes and tools to track, authenticate, manage version control, and secure their assets due to the increased complexities of AI specific assets." (ETSI TS 104 223 [i.1])

Related threats/risks:

Without secure tracking, version control, and authentication, AI systems can be vulnerable to unauthorized changes, data integrity issues, and version conflicts, leading to compromised reliability and security. This is exacerbated by the rapid changes in Gen.AI tool and models.

Example Measures/Controls 1:

Implement AI Asset Tracking: Use version control and ML Ops systems to manage and track changes to AI assets including models, datasets, and related software components ensuring transparency and rollback capability.

Chatbot App: Use Git to track changes to conversation flows and training datasets, ensuring traceability when new features or intents are added. Use an internal package repository mirror is used for components.

ML Fraud Detection: Version control can track training data updates, capturing modifications to ensure traceability and data integrity whereas a model registry is used similarly for models. An internal package repository mirror is used for components.

LLM Platform: Implement the same tracking described in the Fraud Detection example, with the addition of the vendor saving model weights in a protected and isolated storage using a separate dedicated enclave as recommended by the CISA Joint Cybersecurity Information - Deploying AI Systems Securely [i.7].

Open-Access LLM Model: Implement a model registry and scripts to log versioned training datasets, model weights, and configuration files. Use hash-based verification to detect unauthorized modifications to publicly shared assets. Protect package dependency files (e.g. requirements.txt) and keep a copy of packages for each release.

References:

- See References in clause 6.5.1

Example Measures/Controls 2:

Authenticate, Authorize and Log Access to Assets: Enforce strict access controls and authentication protocols to limit asset modifications to authorized personnel only logging any access.

Chatbot App: Access to system prompt, prompt templates, datasets and embeddings used in RAG are restricted to data scientists using RBAC and MFA with updates only allowed via CI pipelines.

ML Fraud Detection: Restrict access to training datasets and models using RBAC and via APIs and CI pipelines with appropriate approvals for critical assets. All access events to sensitive datasets are logged, and alerts are triggered for access attempts from unauthorized users or unusual access patterns. Trigger alerts when models or data are updated outside the process.

LLM Platform: See ML Fraud Detection Example for this control.

Open-Access LLM Model: See ML Fraud Detection Example for this control.

6.5.3 Provision 5.2.1-3

"System Operators shall develop and tailor their disaster recovery plans to account for specific attacks aimed at AI systems." (ETSI TS 104 223 [i.1])

Related threats/risks:

Without tailored disaster recovery, AI systems can be unprepared for incidents such as data poisoning, or Large Language Model (LLM) weaponisation, leaving the organization vulnerable to prolonged disruption leading to data leakage, DoS or compromised model performance. Maintaining a reliable known good state can be challenging, particularly with continuous learning models or systems with frequent updates, increasing the risk of losing data integrity.

Example Measures/Controls:

Incorporate AI-Specific Threat Scenarios into Recovery Plans: Update disaster recovery plans to address AI-specific risks, including adversarial attacks, data poisoning, and model drift and ensure readiness for prompt recovery.

Chatbot App: Include a recovery plan to address adversarial attacks, such as prompt injections to compromise chatbot responses, by maintaining fallback mechanisms to disable automated responses temporarily while restoring integrity.

ML Fraud Detection: Develop recovery plans for data poisoning incidents, including manual backup processes to maintain operations until a tested model checkpoint with reliable detection and adversarial robustness is restored.

LLM Platform: Similar approach to Fraud Detection; consider ensuring the system remains operational during attacks by using backup infrastructure ('high availability ') or placeholders to quickly deploy tested model versions.

Open-Access LLM Model: Decide on how to roll back to a previous healthy version of the model or data if poisoning or other exploitable AI vulnerabilities are identified. Automate the process of restoring backups to speed up recovery.

References:

- CISA, JCDC, Government and Industry Partners Conduct AI Tabletop Exercise [i.86]

6.5.4 Provision 5.2.1-3.1

"System Operators should ensure that a known good state can be restored." (ETSI TS 104 223 [i.1])

Related threats/risks:

Inability to restore a known good state can lead to prolonged system downtime, data loss, or operational disruptions after failures or attacks

Example Measures/Controls 1:

Establish and Maintain a Known Good State: Regularly backup AI models, data, and system configurations, maintaining a known good state that can be restored after a disruption. Good state can contain additional internal state to help a model reach optimal performance after model warm-up. This might not be always possible, and a new model will have to be trained addressing the attack. Nevertheless, backups will help accelerate developing a mitigation.

Chatbot App: Maintain backups of system prompts, RAG data and embeddings, prompt templates, LLM API access details, and other configuration details.

ML Fraud Detection: Maintain backups of the latest clean dataset and model versions, allowing quick restoration if the current version is compromised with a backdoor attack or is discovered to have been trained on poisoned data.

LLM Platform: Retain secure versions of models, weight files, fine-tuning checkpoints, and system configurations to facilitate recovery and training datasets to quickly roll back if vulnerabilities such as poisoning attacks or unauthorized modifications are detected.

Open-Access LLM Model: Retain secure copies of model versions and training datasets to quickly roll back and release a new version if vulnerabilities such as poisoning attacks or unauthorized modifications are reported.

Example Measures/Controls 2:

Develop Recovery plans for advanced scenarios: Some incidents will include abuse and weaponisation of AI systems, especially Generative AI, to stage misinformation or other attacks abusing an organization's system. Recovering from such as attacks will require multi-disciplinary expertise and response strategies to minimize impact and harm.

Chatbot App: Design a process that addresses adversarial take-over of the chatbot to spread misinformation or exploit unused multi-modal capabilities to create deep fakes. This is used after the notification by authorities that the service has been used for the last six months for misinformation campaigns.

ML Fraud Detection: Not applicable in predictive ML.

LLM Platform: As in the Chatbot App scenario.

Open-Access LLM Model: Consider using watermark verification to help authorities and customers using the model to deal with advanced attacks.

References:

- OWASP Guide for Preparing and Responding to Deepfake Events [i.87]
- ETSI TR 104 032 [i.12], clause 5.3

6.5.5 Provision 5.2.1-4

"Developers, System Operators, Data Custodians and End-users shall protect sensitive data, such as training or test data, against unauthorized access (see principle 7 for details on securing data)." (ETSI TS 104 223 [i.1])

Related threats/risks:

Unauthorized access to sensitive data, such as training datasets, training algorithms, hyper parameters and model parameters can lead to privacy violations, data breaches, or compromised model integrity, increasing regulatory and reputational risks.

Example Measures/Controls 1:

Implement Data Encryption at Rest and in Transit: Encrypt sensitive data at rest and in transit to protect against unauthorized access, ensuring data confidentiality and security.

Chatbot App: Encrypt all user chat logs stored in cloud storage using server-side encryption with customer-managed keys. Use HTTPS with TLS 1.3 to secure communications between the chatbot and the end user, preventing data interception.

ML Fraud Detection: Encrypt transaction data stored in the database (at rest) using AES-256 encryption. Use TLS 1.3 to encrypt data transmitted between the fraud detection model and retailer APIs, ensuring compliance with PCI DSS standards. Regularly review encryption settings to ensure they meet evolving regulatory requirements.

LLM Platform: Encrypt model artifacts and training datasets stored in the development environment using industry-standard encryption algorithms. Use encrypted channels (e.g. SFTP or TLS) to transmit training data to remote servers during deployment. Periodically validate the encryption configurations compliance with the ISO 27001 standard [i.55].

Open-Access LLM Model: Store training datasets encrypted and use HTTPS or SFTP for all data transfers to prevent unauthorized access during transfers.

References:

- ICO A guide to data security [i.88]
- ISO/IEC 27001:2022 [i.55]
- NCSC Cloud Security Guidance [i.60] - Using a cloud platform securely
- OWASP Application Security Verification Standard (ASVS) [i.89]
- OWASP AI Exchange [i.10]

Example Measures/Controls 2:

Implement Strong Data Access Controls: Restrict access to sensitive data to authorized personnel only, using multi-factor authentication and role-based access controls. Programmatic updates via pipelines have strict RBAC and approvals in place to prevent abuse.

Chatbot App: Apply RBAC with strict least-privilege access to datasets used in RAG including vectors and embeddings.

ML Fraud Detection: Enforce access controls at multiple levels by restricting access to training datasets containing sensitive customer information at storage level and limiting access to the training environments. Ensure only essential team members and approved pipeline roles have permissions, with strict enforcement through RBAC and strong authentication across all layers.

LLM Platform: Same as in the ML Fraud Detection example.

Open-Access LLM Model: Same as in the ML Fraud Detection example.

References:

- See References in Example Measures/Controls 1 of this threat/risk.

Example Measures/Controls 3:

Data Access and Usage Monitoring: Implement automated monitoring tools to monitor access and usage of sensitive data including proprietary model weights, generating alerts for any unusual patterns or unauthorized attempts.

Chatbot App: Set up automated alerts for access to protected RAG data, ensuring that unauthorized access is immediately flagged.

ML Fraud Detection: Set up automated alerts for access to training /testing datasets and model weights and configuration files.

LLM Platform: Similar implementation as in the Fraud Detection.

Open-Access LLM Model: Alert for unusually high failed attempts to access and utilize cloud or platform controls to alert for high volume downloads indicating potential data exfiltration.

References:

- See References in Example Measures/Controls 1 of this threat/risk.

6.5.6 Provision 5.2.1-4.1

"Developers, Data Custodians and System Operators shall apply checks and sanitisation to data and inputs when designing the model based on their access to said data and inputs and where those data and inputs are stored. This shall be repeated when model revisions are made in response to user feedback or continuous learning. See principle 6 for relevant provisions for open source." (ETSI TS 104 223 [i.1])

Related threats/risks:

Lack of data and input sanitization can introduce biases, errors, or malicious data, leading to compromised outputs, security risks, as well as potential legal and reputational harm.

Example Measures/Controls:

Apply Data Sanitisation and Validation: Ensure that all data - both training datasets and runtime inputs - are sanitized and validated to prevent data poisoning, malicious inputs, and biases. During training, verify that datasets are clean, unbiased, and aligned with model objectives.

Chatbot App: Apply data sanitisation in incoming user prompts to reduce the risk of prompt injections.

ML Fraud Detection: Sanitize training data to remove any incorrect or malicious entries that could skew model outputs.

LLM Platform: Apply input sanitisation as part of guardrails; implement automated scripts to scan and sanitize large-scale text corpora for biases, offensive language, or duplicate entries during data preprocessing.

Open-Access LLM Model: Consider use of content filtering and other relevant data preprocessing mechanisms when training the model to check inputs and ensure data quality.

References:

- OWASP AI Exchange [i.10]
- Training Dataset Validation to Protect Machine Learning Models from Data Poisoning [i.90]

6.5.7 Provision 5.2.1-4.2

"Where training data or model weights could be confidential, Developers shall put proportionate protections in place." (ETSI TS 104 223 [i.1])

Related threats/risks:

Failure to adequately protect sensitive training data or model weights can lead to unauthorized access, intellectual property theft, or exposure of confidential information, increasing the risk of data breaches and regulatory non-compliance.

Example Measures/Controls:

See controls and examples in clause 6.5.5.

References:

- See References in clause 6.5.5.

6.6 Principle 6: Secure the infrastructure

6.6.1 Provision 5.2.2-1

"Developers and System Operators shall evaluate their organization's access control frameworks and identify appropriate measures to secure APIs, models, data, and training and processing pipelines." (ETSI TS 104 223 [i.1])

Related threats/risks:

Inadequate access control can expose sensitive data, models, and pipelines to unauthorized access, increasing the risk of data leaks, model tampering, or unauthorized modifications.

Example Measures/Controls:

Establish Role-Based Access Controls (RBAC): Implement role-based access controls to limit access to AI models, data, and pipelines based on user roles and responsibilities, enforcing the principle of least privilege. This includes research environments. Protect API endpoints and data pipelines by implementing access controls, encryption, and authentication mechanisms to prevent unauthorized access.

Chatbot App: Restrict access to system prompts, configuration data, access to the underlying LLM API and any data used for RAG. Restrict access to embeddings generation and via pipelines to relevant group of data engineers.

ML Fraud Detection: Restrict access to training data to data scientists and model tuning permissions to ML engineers, ensuring minimal access rights across roles. This should be for both interactive and API-based access.

LLM Platform: Follow the same implementation approach as in the Fraud Detection example of this control.

Open-Access LLM Model: Restrict access to training data to and model tuning permissions to developers, ensuring minimal access rights across roles.

References:

- NCSC Machine Learning Principles, Part 3: Secure deployment [i.5]
- NCSC Guidelines for Secure AI System Development, Secure your Infrastructure [i.6]
- CSA Companion Guide on Securing AI Systems [i.46], Section 2.2.3 Deployment, Item 3.1
- OWASP AI Top 10 API Security Risks - 2023 [i.67]

6.6.2 Provision 5.2.2-2

"If a Developer offers an API to external customers or collaborators, they shall apply appropriate controls that mitigate attacks on the AI system via the API. For example, placing limits on model access rate to limit an attacker's ability to reverse engineer or overwhelm defences to rapidly poison a model." (ETSI TS 104 223 [i.1])

Related threats/risks:

Externally exposed APIs increase the risk of model extraction attacks, rapid data poisoning, and abuse, potentially compromising the integrity of the AI system.

Example Measures/Controls 1:

Implement API Rate Limiting: Enforce rate limits on API requests to prevent attackers from overwhelming the system, reverse engineering the model, or rapidly injecting malicious inputs.

Chatbot App: Apply rate limit to web endpoints or own APIs to prevent using the app's API to indirectly attack the supporting model.

ML Fraud Detection: Apply rate limit to web endpoints and APIs.

LLM Platform: Apply a rate limit of X requests per minute per user to prevent bulk extraction of model responses.

Open-Access LLM Model: Not applicable.

References:

- OWASP Top 10 for APIs - API4:2019 Lack of Resources & Rate Limiting [i.65]

Example Measures/Controls 2:

Use Behavioural Analysis for API Security: Implement behavioural analysis tools to detect abnormal API usage that could indicate malicious intent, such as model extraction or poisoning attempts.

Chatbot App: Use behavioural analysis to flag repeated API calls with unusual input patterns that could indicate an attempt to exploit or overload the chatbot's system.

ML Fraud Detection: Monitor API usage for anomalies such as an unusual volume of high-risk transaction queries that can suggest adversarial testing or system probing.

LLM Platform: Use behavioural analysis with a dual LLM, moderation APIs, or anomaly detection code, to identify sudden spikes in repetitive queries that can signal reverse engineering attempts.

Open-Access LLM Model: Provide guidelines to customers on implementing behavioural monitoring of model use.

References:

- OWASP AI Top 10 API Security Risks - 2023 [i.67]

- NCSC Logging and Protective Monitoring [i.91]
- LLM Monitoring and Observability - A Summary of Techniques and Approaches for Responsible AI [i.92]

Example Measures/Controls 3:

Deploy API Gateway with Security Features: Use an API gateway with security features like throttling, dynamic rate limiting, and any other attack detection features, authentication, and logging to manage and monitor access to external-facing APIs.

Chatbot App: Use an API Gateway to manage access if the apps API is publicly accessible.

ML Fraud Detection: Use an API gateway to manage access to inference API, logging each request, applying rate limiting, and requiring OAuth for user authentication.

LLM Platform: Similar implementation to the Fraud Detection example but adding all APIs offered to customers.

Open-Access LLM Model: Not applicable.

References:

- OWASP AI Top 10 API Security Risks - 2023 [i.67]
- Wikipedia API Management Overview and links [i.93]

NOTE: Rate limiting can be relevant to other types of interfaces such as a user interface.

6.6.3 Provision 5.2.2-3

"Developers shall also create dedicated environments for development and model tuning activities. The dedicated environments shall be backed by technical controls to ensure separation and principle of least privilege. In the context of AI, this is particularly necessary because training data shall only be present in the training and development environments where this training data is not based on publicly available data." (ETSI TS 104 223 [i.1])

Related threats/risks:

Without separation and environment-specific controls, production-grade sensitive data and models can be exposed to unauthorized access via development or research workflows, leading to data leaks, tampering, unauthorized deployments, or compliance violations.

Example Measures/Controls:

Set Up Dedicated Development and Production Environments: Establish separate environments for development, testing, and production, ensuring data and models are only accessible where necessary. Restrict sensitive training data to the development environment, isolating it from production.

Chatbot App: maintain separated environments for the app with restricted access to system prompts, and API config using different LLM API endpoints for each environment. Use synthetic RAG data for development and testing.

ML Fraud Detection: Use a model registry to track and manage versions, ensuring only authorized personnel access and modify them. Leverage anonymized transaction data or synthetic datasets in development to protect sensitive customer information. Implement ML pipelines to automate the promotion of tested models to production with approvals, ensuring compliance, security, and quality standards.

LLM Platform: Similar approach as in the Fraud Detection example of this control.

Open-Access LLM Model: Use lightweight containerization to isolate development and production environments. Restrict sensitive training data to local development only and deploy approved models via automated scripts to maintain separation and minimize manual intervention risks.

References:

- Alan Turing Institute: What is synthetic data and how can it advance research and development [i.94]
- CISA Joint Cybersecurity Information - Deploying AI Systems Securely [i.7]

- DSTL Machine learning with limited data [i.95]
- NCSC Secure your development Guide [i.96]
- NIST SP 800-218 [i.97]
- Secure Software Development Framework (SSDF) [i.97]

6.6.4 Provision 5.2.2-4

"Developers and System Operators shall implement and publish a clear and accessible vulnerability disclosure policy." (ETSI TS 104 223 [i.1])

Related threats/risks:

Without a defined vulnerability disclosure policy, security vulnerabilities can go unreported, exposing the organization to delayed or missed opportunities to patch critical issues.

Example Measures/Controls:

Develop and Publish a Vulnerability Disclosure Policy: Create a vulnerability disclosure policy that details how vulnerabilities can be reported, including timelines for acknowledgment and resolution.

Chatbot App: Develop a policy using the NCSC Vulnerability Disclosure Toolkit [i.98].

ML Fraud Detection: Develop a disclosure policy to comply with the ISO/IEC 29147:2018 [i.100] Vulnerability disclosure standard.

LLM Platform: Similar approach as in the Fraud Detection example of this control.

Open-Access LLM Model: Publish a policy on the organization's website, including a contact email for reporting AI vulnerabilities and a commitment to respond within 48 hours.

References:

- NCSC Vulnerability Disclosure Toolkit [i.98]
- ISO/IEC 29147:2018 Vulnerability disclosure standard [i.100]

6.6.5 Provision 5.2.2-5

"Developers and System Operators shall create, test, and maintain an AI system incident management plan and an AI system recovery plan." (ETSI TS 104 223 [i.1])

Related threats/risks:

Without a dedicated incident and recovery plan, AI systems can be slow to recover from disruptions, resulting in prolonged downtime or compromised model performance.

Example Measures/Controls:

Develop an AI-Specific Incident Management Plan: Create an incident management plan tailored to AI-specific threats, such as discovery of data poisoning, model drift, and adversarial attacks with recovery steps to a known good state, including procedures for validating model integrity.

Chatbot App: Develop an incident plan to address adversarial inputs causing inappropriate responses and jailbreaking, including isolating malicious queries, and implementing input validation updates. If the third-party LLM API becomes unavailable or produces unreliable responses (e.g. due to downtime or compromised functionality), the recovery plan includes switching to a pre-integrated backup LLM provider or a local fallback model to maintain core functionalities. If no immediate replacement is available, the app transitions to a basic response mode using pre-scripted messages or rule-based logic.

ML Fraud Detection: For a fraud detection model, a 20 % spike in fraud signals triggers an incident investigation. The team isolated recent training datasets identified anomalies like skewed entries and rolled back to a prior model version. Containment included notifying teams and deploying enhanced validation checks.

LLM Platform: Create a comprehensive plan of handling jailbreaking, misuse, and model overload including temporarily restricting access, deploying a backup model, and conducting usage audits.

Open-Access LLM Model: Implement a plan to handle community-reported data poisoning incidents by halting any contributions, validating datasets, and rolling back to a previous model checkpoint from its model registry, validating its outputs through automated tests before redeployment.

References:

- CSA Incident Response Checklist [i.101]
- NCSC Incident Management [i.102]
- NCSC Vulnerability Disclosure Toolkit [i.98]

6.6.6 Provision 5.2.2-6

"Developers and System Operators should ensure that, where they are using cloud service operators to help to deliver the capability, their contractual agreements support compliance with the above requirements." (ETSI TS 104 223 [i.1])

Related threats/risks:

Without clear understanding of cloud service agreements, organizations might not know what security measures they can expect or demand from the cloud provider, potentially leaving AI assets exposed to gaps in data protection or compliance failures.

Example Measures/Controls:

Define and Validate Security Clauses in Cloud Contracts: Ensure cloud service agreements explicitly outline security responsibilities, compliance standards, and support provisions, including data protection, access controls, incident response, and audit capabilities. Provide detailed documentation to stakeholders to bridge knowledge gaps about the cloud provider's obligations.

Chatbot App: The developers of the chatbot app switch to a new LLM model provider hosted on the cloud, as it references encryption, data isolation and retention, limited customer data logging and a ban on model retraining with customer data in the T&Cs and cloud contract. This follows the refusal of the previous model provider to cover these requirements in the contract.

ML Fraud Detection: Review and validate cloud provider documentation.

LLM Platform: Similar approach as in the Fraud Detection example of this control.

Open-Access LLM Model: Similar approach as in the Fraud Detection example of this control.

References:

- CSA AI Organizational Responsibilities - Governance, Risk Management, Compliance and Cultural Aspects [i.103]
- ICO guidance for cloud providers obligations under NIS Regulations 2018 - Security requirements [i.104]
- NCSC Cloud Security Guidance [i.60]

6.7 Principle 7: Secure the supply chain

6.7.1 Provision 5.2.3-1

"Developers and System Operators shall follow secure software supply chain processes for their AI model and system development." (ETSI TS 104 223 [i.1])

Related threats/risks:

Without secure software supply chain processes, AI systems are vulnerable to risks like supply chain attacks, insertion of malicious components, and dependency issues, including models and datasets, which can compromise AI integrity and security. Lack of accountability across the supply-chain for jurisdiction-specific regulations can lead to data breaches and lack of regulatory compliance.

Example Measures/Controls 1:

Safeguard Provenance and Transparency: Mandate in internal standards that components can only be sourced by trusted and approved sources, documenting source, version, licencing, history, and other related artifacts (e.g. Model Card for models); use checksums to verify integrity. SBOMs can help automated creation of signed attestations of all components and their metadata. This can help safeguard transparency and provenance with non-repudiation and component tampering checks. SBOMs cover libraries and system components but not models or datasets. AI and ML BOMs are an emerging field. Monitor progress in standards such as Cyclone DX ML BOM [i.106] and introduce similar scans when tools emerge. In the meanwhile, rely on ML Ops to enforce model and dataset provenance and use file checksum to verify integrity and provenance.

Chatbot App: Publish guidelines and developer standards and use SBOMs to document the controls used to build the application applying periodic tests and audits to ensure compliance. Cover models used for embeddings in the process.

ML Fraud Detection: Same implementation to Chatbot App, but complement SBOMs with an in-house created ML-BOM documenting the model package dependencies and model cards offering a machine-readable record which is validated with a script against approved standards.

LLM Platform: Similar implementation to the Fraud Detection example but for the multimodal model, but adding datasets and auxiliary (RL, embeddings generation, and so on).

Open-Access LLM Model: Agree on the criteria and workflow to manage external components, datasets, and models (foundation if used, RL, embeddings and so on). Use scanning tools and scripts to create signed SBOMs and ML BOMs for machine readable lists.

References:

- ETSI TR 104 048 [i.145] - Data Supply Chain Security
- CISA SBOM [i.72]
- MITRE System of Trust Framework [i.105]
- NIST Cybersecurity Supply Chain Risk Management [i.49]
- NCSC Supply Chain Security Guidance [i.51]
- NCSC Machine Learning Principles, 2.1 Secure your supply chain [i.5]
- OWASP CycloneDX - ML BOM [i.106]
- OWASP Top 10 for LLM applications - LLM03 Supply-Chain Vulnerabilities [i.11]
- Supply-chain Levels for Software Artifacts (SLSA) [i.107]

Example Measures/Controls 2:

Apply Vulnerability Management: Implement a vulnerability management process that includes regular (e.g. daily) scanning of third-party components for known vulnerabilities and timely patching to mitigate risks.

Chatbot App: Use automated vulnerability scanning tools to identify vulnerabilities in third-party libraries or components. Introduce a patch management process where critical vulnerabilities are patched within 48 hours, moderate vulnerabilities within 14 days, and low-risk vulnerabilities within 30 days.

ML Fraud Detection: Use the same approach as in the Chatbot App but with additional model scanning tools and tests for model serialization attacks for base third party models.

LLM Platform: Similar approach as in the Fraud Detection example of this control.

Open-Access LLM Model: Similar implementation with the Fraud Detection example but use open-source or commercial LLM scanners to finetune the foundational model; Agree remediation timeframes for fixing defects with patching focusing on Critical and Highs and triaging Medium and Lows.

References:

- NCSC Vulnerability Management [i.108]
- OWASP Vulnerability Management Guide [i.109]
- OWASP LLM and Generative AI Security Solutions Landscape [i.85]

Example Measures/Controls 3:

Adopt Secure Supply Chain Frameworks: Follow the organization's preferred secure supply chain guidance or standard for all stages of AI development, from sourcing and integrating components to testing and deployment.

Chatbot App: Follow e.g. NCSC supply chain security guidance to assess the provider's security policies and ensure compliance with encryption, access control, and other security standards.

ML Fraud Detection: Follow the same advice as in the Chatbot App example; Additionally, ensure training datasets sourced from external vendors comply with data protection legislation, provenance checks, contractual obligations, and avoiding unlawfully collected training datasets.

LLM Platform: Complement the implementation described in the Fraud Detection example with the additional adoption of ISO 28000:2022 [i.110].

Open-Access LLM Model: This is not applicable to small organizations, who can consult the standards and NCSC guidance on supply-chain security [i.51].

References:

- NCSC Supply Chain Security Guidance [i.51]
- NIST C-SCRM [i.49]
- ISO 28000:2022 [i.110]

6.7.2 Provision 5.2.3-2

"System Operators that choose to use or adapt any models, or components, which are not well-documented or secured shall be able to justify their decision to use such models or components through documentation (for example if there was no other supplier for said component)." (ETSI TS 104 223 [i.1])

Related threats/risks:

Utilizing poorly documented or untrusted AI components because of some of their unique features introduces potential vulnerabilities, increasing risks of data leaks or unexpected behaviour.

Example Measures/Controls:

Document Justification for Untrusted Components: When using poorly documented or untrusted components, document the justification, including alternative evaluations, and the absence of better suppliers.

Chatbot App: Use a chatbot API with superior multilingual capabilities despite limited documentation, justifying the decision with performance benchmarks.

ML Fraud Detection: In the fraud detection scenario, the vendor identifies a fraud-detection model on a public model hub that outperforms all other models in false positives but lack any model cards or other documentation and the developer cannot provide more details. False positive reduction is a key market differentiator for the vendor who document the decision to migrate to the new components detailing comparative evaluations with other models.

LLM Platform: When adopting an LLM fine-tuning library with undocumented optimization techniques, document its unique benefits and conduct tests to evaluate risks against other alternatives.

Open-Access LLM Model: When using a content filtering library with limited documentation to remove harmful outputs, justify the decision with its accuracy benchmarks and document it in the project's wiki.

References:

- NIST AI RMF [i.17]
- NCSC Vulnerability Management [i.108]

6.7.3 Provision 5.2.3-2.1

"In this case, Developers and System Operators shall have mitigating controls and undertake a risk assessment linked to such models or components." (ETSI TS 104 223 [i.1])

Related threats/risks:

Utilizing poorly documented or untrusted AI components because of some of their unique features introduces potential vulnerabilities, increasing risks of data leaks or unexpected behaviour.

Example Measures/Controls:

Evaluate and Mitigate Risks for Untrusted Components: Perform a risk assessment to identify potential risks associated with unsecured AI components, specifying, and implementing mitigating controls to minimize vulnerabilities.

Chatbot App: Conduct input validation and adversarial testing on a poorly documented chatbot API to identify vulnerabilities. Mitigation includes rate limiting, anomaly detection for malicious input patterns, and monitoring for unexpected behaviour.

ML Fraud Detection: Evaluate the new model with serialization attack scans and adversarial robustness testing; this forms part of its risk assessment that also includes legal and regulatory compliance checks. Mitigation controls include real-time enhanced monitoring with anomaly detection, staggered deployment on low-risk cases in controlled environments for a period, contractual agreement with former members of the model development team to help testing.

LLM Platform: Perform a risk assessment on the fine-tuning library, testing for optimization errors, compatibility issues, and security vulnerabilities. Mitigations include using sandboxed environments for testing and implementing logging to detect anomalies during fine-tuning.

Open-Access LLM Model: Assess risks of the content filtering library, such as failure to flag harmful content or biases in filtering. Mitigation measures include integrating secondary filtering layers, testing for gaps, and monitoring flagged content for refinement.

References:

- NIST SP 800-161 Rev. 1 [i.50]

6.7.4 Provision 5.2.3-2.2

"System Operators shall share this documentation with End-users in an accessible way." (ETSI TS 104 223 [i.1])

Related threats/risks:

Utilizing poorly documented or untrusted AI components because of some of their unique features introduces potential vulnerabilities, increasing risks of data leaks or unexpected behaviour.

Example Measures/Controls:

Share Documentation with End-Users: Share documentation of non-compliant components with end-users, detailing risks, and justifications for transparency.

Chatbot App: Provide end-users with a summary of the chatbot API's undocumented features, potential risks, and implemented safeguards, along with instructions for reporting anomalies. Ensure accessibility by providing for example, a downloadable accessible PDF File.

ML Fraud Detection: Users included in the staggered deployment are informed with a WCAG 2.1 compliant document on the model choice, reasons, and mitigations and they are invited to report any suspicious system behaviour.

LLM Platform: Include documentation in the - compliant with accessibility formats - user guide detailing the fine-tuning library's benefits, risks, and mitigation strategies, and direct end-users to support channels for reporting unexpected outcomes.

Open-Access LLM Model: Document the use of the library in the readme file including the additional monitoring controls customers can use; reference the wiki pages detailing risk assessments and evaluation and how to report issues. Test and leverage repository's accessibility features or generate an accessible PDF.

References:

- See References in clause 6.7.3
- GOV.UK Guidance and tools for digital accessibility [i.111]

6.7.5 Provision 5.2.3-3

"Developers and System Operators shall re-run evaluations on released models that they intend on using." (ETSI TS 104 223 [i.1])

Related threats/risks:

Without reusable evaluations, periodic re-evaluation can be difficult, resulting into performance degradation, bias, inaccuracies, or security vulnerabilities go unnoticed in new releases.

Example Measures/Controls:

Create and Maintain Appropriate and Reusable Model Evaluation Suites: Create appropriate model evaluation suites to assess performance, accuracy, security, and potential drift when required.

Chatbot App: Consult external evaluations of LLM providers or conduct independent tests, focusing on use-cases for major model upgrades.

ML Fraud Detection: Changes in anti-money laundering legislation expand the definition of suspicious transactions to include previously normal patterns (e.g. smaller amounts). The company uses the model evaluation suite to detect any drift. Track performance metrics over time to help setting for triggering re-evaluation or re-training.

LLM Platform: Run the evaluation suite to assess a new model release and identify areas that need mitigations.

Open-Access LLM Model: Runs evaluation suites for major releases and changes to identify areas that need mitigations.

References:

- DecodingTrust: A Comprehensive Assessment of Trustworthiness in GPT Models [i.112]
- AI Safety Institute: Inspect - An open-source framework for large language model evaluations [i.113]
- NIST AI Test, Evaluation, Validation, and Verification (TEVV) [i.114]
- NIST Dioptra test platform [i.115]

6.7.6 Provision 5.2.3-4

"System Operators shall communicate the intention to update models to End-users in an accessible way prior to models being updated." (ETSI TS 104 223 [i.1])

Related threats/risks:

Lack of transparency around model updates can make it challenging for users to use a new model version, especially in evaluating data protection compliance. This creates availability and operational issues for end-users relying on certain model feature or behaviour.

Example Measures/Controls:

Provide Advance Notice of Model Updates: Notify end-users of model updates at least one month in advance, including any potential impacts on model performance or functionality and offer previews to test changes for at least a month. Ensure notices are accessible.

Chatbot App: Notify users about updates to the chatbot's third party LLM, highlighting any changes in supported languages or response behaviour, and provide access to a testing sandbox. Notices comply with WCAG 2.1 guidelines.

ML Fraud Detection: Send end-users a notice explaining upcoming changes in the algorithm and how it can impact detection and provides a test instance to allow them run evaluations before switching offering an accessible downloadable PDF.

LLM Platform: Inform end-users - using accessible notices like the ones the previous two examples- about planned finetuning updates, detailing expected improvements or deprecated features, and offer a staging environment and API version for evaluation before deployment.

Open-Access LLM Model: Publish a public notice on the repository's changelog and mailing list detailing upcoming changes, potential impacts, and instructions for testing the new pre-release version. The notice takes advantage of the repositories built-in accessibility features including WCAG 2.1 compliance.

References:

- GOV.UK Guidance and tools for digital accessibility [i.111]

6.8 Principle 8: Document Data, Models, and Prompts

6.8.1 Provision 5.2.4-1

"Developers shall document and maintain a clear audit trail of their system design and post-deployment maintenance plans. Developers should make the documentation available to the downstream System Operators and Data Custodians." (ETSI TS 104 223 [i.1])

Related threats/risks:

Without thorough documentation and audit trails, AI system operations can lack transparency, limiting traceability and increasing risks of security gaps, regulatory non-compliance, and operational inefficiencies.

Example Measures/Controls:

Develop Comprehensive System Design and Maintenance Documentation: Document the system design and post-deployment maintenance plans with audit trail of design decisions, architectural diagrams, and maintenance schedules. Ensure the documentation is accessible to downstream System Operators and Data Custodians, highlighting responsibilities, version control, and any dependencies.

Chatbot App: Document post-deployment monitoring approaches, detailing how user feedback is collected and incorporated into iterative updates. Share these protocols with Data Custodians to ensure they comply with data retention and usage policies.

ML Fraud Detection: Maintain an audit trail of ADRs (Architecture Decision Records) design choices, such as feature selection, model architecture, and testing results. Provide detailed maintenance schedules, including retraining cycles and updates to detection rules.

LLM Platform: Use the approach described in the Fraud Detection example of this control. Additionally, include an audit trail of all model fine-tuning activities and dependencies.

Open-Access LLM Model: Automate log generation from the repository of all an audit trail of all model training/fine-tuning activities and their dependencies, releases, and associated release notes.

References:

- ICO Data Protection Audit Framework Toolkits, Artificial Intelligence [i.19]
- NCSC Machine Learning Principles [i.5], clause 2.3
- NCSC Guidelines for Secure AI System Development, Document your data, models, and prompts [i.6]
- CSA Companion Guide on Securing AI Systems [i.46], Section 2.2.2 Development, Item 2.3

6.8.2 Provision 5.2.4-1.1

"Developers should ensure that the document includes security-relevant information, such as the sources of training data (including fine-tuning data and human or other operational feedback), intended scope and limitations, guardrails, retention time, suggested review frequency and potential failure modes." (ETSI TS 104 223 [i.1])

Related threats/risks:

Lack of detailed security-relevant documentation can lead to unmitigated risks from unverified data sources, misuse of the AI system, and challenges in addressing failure modes.

Example Measures/Controls:

Include Relevant Security - Information in System Documentation: Document AI systems including intended scope, limitations, known failure modes, prompts, guardrails training data sources, data retention policies, review schedules. Use **Model Cards** to capture transparent summaries of a model's capabilities, ethical considerations, performance metrics, and limitations, consider using **ML Bills of Materials (ML BOMs)** to document in machine-readable way the model and its dependencies, with component versions, and licensing.

Chatbot App: Document prompts and data sources used for RAG, intended use cases, any safeguards to prevent misuse, data retention periods, and failure scenarios.

ML Fraud Detection: Apply the implementation described in the Chatbot App to training/test data. In addition, create and maintain model cards to document the models trained.

LLM Platform: Similar implementation as in the Fraud Detection example.

Open-Access LLM Model: Creates and publish a model card. Use a standard such as OWASP CycloneDX [i.106] to generate and publish a signed ML BOM documenting model dependencies.

References:

- NCSC Machine Learning Principles [i.5], clause 2.3 Manage the full life cycle of models and datasets
- OWASP CycloneDX - ML BOM [i.106]

6.8.3 Provision 5.2.4-1.2

"Developers shall release cryptographic hashes for model components that are made available to other stakeholders to allow them to verify the authenticity of the components." (ETSI TS 104 223 [i.1])

Related threats/risks:

Absence of cryptographic hashes for model components increases the risk of tampering, unauthorized modifications, and integrity issues, compromising trust and security.

Example Measures/Controls:

Include Cryptographic Hashes for Models: Release cryptographic hashes for all models made available to stakeholders, enabling verification of component authenticity.

Chatbot App: Not applicable.

ML Fraud Detection: Provide cryptographic hashes for trained fraud detection models shared with clients, enabling them to confirm the integrity of the received models.

LLM Platform: Maintain hashes for each version internally and provide them in direct integration scenarios, e.g. a hosting by a cloud platform acting as a reseller.

Open-Access LLM Model: Generate and provide unique digital fingerprints (e.g. SHA-256 hashes) for each model version both on model hubs but also the website and Git repository, allowing developers to verify the model's integrity.

References:

- NIST Cryptographic Standards and Guidelines [i.117]
- OWASP Cryptographic Storage Cheat Sheet [i.99]

6.8.4 Provision 5.2.4-2

"Where training data has been sourced from publicly available sources, there is a risk that this data might have been poisoned. As discovery of poisoned data is likely to occur after training (if at all), Developers shall document how they obtained the public training data, where it came from and how that data is used in the model." (ETSI TS 104 223 [i.1])

Related threats/risks:

Without comprehensive documentation of training data sources and collection timestamps, organizations can be unable to determine whether their AI models have been affected by data poisoning, fraud, or insider abuse and whether they are compliant to data protection legislation. If data poisoning attacks are revealed to have occurred on public websites, producers or users of models will only know if they are affected if they have robust documentation of what data the model was trained on.

Example Measures/Controls:

Document the process of sourcing public training data: Detailing how data was collected, processed, and used in the model (e.g. pretraining, fine-tuning). Include poisoning mitigation measures such as data validation and anomaly detection. Maintain an audit trail to trace datasets if poisoning is suspected.

Chatbot App: Document any public data used in RAG scenarios.

ML Fraud Detection: Track sources of training data for fraud detection models, logging the origin and preprocessing steps to ensure traceability in the event of anomalies.

LLM Platform: Maintain an audit trail for all training and fine-tuning datasets sourced from public repositories, including documentation, evaluations, and poisoning mitigation measures like validation checks.

Open-Access LLM Model: For a community-contributed dataset, enforce logging metadata and manual filtering notes as part of the Pull Request template to enhance transparency and mitigate risks of poisoning.

References:

- ICO Data Protection Audit Framework Toolkits, Artificial Intelligence [i.19]
- ETSI TR 104 048 [i.145]
- CSA Companion Guide on Securing AI Systems [i.46], Section 2.2.2 Development, Item 2.3
- NCSC Machine Learning Principles [i.5], clause 2.3
- NCSC Guidelines for Secure AI System Development, Document your data, models, and prompts [i.6]

6.8.5 Provision 5.2.4-2.1

"The documentation of training data should include at a minimum the source of the data, such as the URL of the scraped page, and the date/time the data was obtained. This will allow Developers to identify whether a reported data poisoning attack was in their data sets." (ETSI TS 104 223 [i.1])

Related threats/risks:

Failure to record detailed metadata, such as data sources and timestamps, can hinder efforts to trace and respond to data poisoning incidents.

Example Measures/Controls:

Document metadata for all publicly sourced training data: Include the exact source (e.g. URLs) and date/time of collection. Ensure this metadata is stored in an accessible format to facilitate traceability in case of reported data poisoning.

Chatbot App: Maintain a log of URLs and timestamps for all data used in a RAG scenario.

ML Fraud Detection: Record metadata for public datasets used in feature engineering and training, ensuring URLs and collection dates are stored securely. Include supplier contact details and documentation link for commercial datasets.

LLM Platform: Capture detailed metadata for publicly sourced datasets, including data scrapers used, timestamp logs, and annotations for identified poisoning risks.

Open-Access LLM Model: Provide contributors with a template to submit dataset metadata, requiring URLs and timestamps to be included for every contribution, making required part of the Pull Request template.

References:

- See References in clause 6.8.4

6.8.6 Provision 5.2.4-3

"Developers should ensure that they have an audit log of changes to system prompts or other model configuration (including prompts) that affect the underlying working of the systems. Developers can make this available to any System Operators and End-Users that have access to the model." (ETSI TS 104 223 [i.1])

Related threats/risks:

Without an audit log for configuration changes, tracking and understanding modifications to prompts or model configurations become difficult, increasing risks of unintended system behaviour or accountability issues.

Example Measures/Controls:

Maintain an Audit Log of Prompt and Model Configuration Changes: Implement an audit log that captures all changes to system prompts and configuration settings, recording details such as change date, user ID, and a description of modifications. Provide access to System Operators and Data Custodians.

Chatbot App: Maintain logs of all prompt changes and LLM API endpoint configurations with the required metadata. Configure alerts to notify administrators when critical prompt templates implementing guardrails are modified.

ML Fraud Detection: Log changes to model parameters and weights and implement alerts when model parameters and weights are modified outside a pipeline deployment.

LLM Platform: Similar approach as in the Fraud Detection example of this control.

Open-Access LLM Model: Maintain an open audit log that records changes to prompts, model weights, and configurations, specifying contributor IDs and timestamps. Provide stakeholders with read-only access for transparency and enable alerts for critical changes.

References:

- See References in clause 6.8.4

6.9 Principle 9: Conduct appropriate testing and evaluation

6.9.1 Provision 5.2.5-1

"Developers shall ensure that all models, applications, and systems that are released have been tested as part of a security assessment process." (ETSI TS 104 223 [i.1])

Related threats/risks:

Without rigorous security assessments by qualified experts understanding AI and classic risks, AI systems can be vulnerable to exploits, unauthorized access, or data breaches, leading to potential data loss, manipulation, or misuse of the AI model.

Example Measures/Controls 1:

Implement Security Assessment Processes for All Releases: Establish a mandatory security assessment process for all models, applications, and systems prior to release, covering areas like access control, data integrity, and adversarial AI attacks. Scope testing based on the threats identified in threat models.

Chatbot App: Test against the OWASP Top 10 for LLM Apps, focusing on the indirect injection risks from PDF uploads identified in the threat model. Add tests to cover general application and platform vulnerabilities ensuring the application has good security posture.

ML Fraud Detection: Test against evasion attacks highlighted in the threat model as well as the security posture of the application, APIs, and platform using the OWASP Top 10 for Apps [i.118] and APIs [i.67].

LLM Platform: Test adversarial robustness, extraction attacks, output sensitivity, ethical violations and bias, data memorization, and jail breaking.

Open-Access LLM Model: Follow the same approach as the LLM Provider. Publish a summary of findings alongside the model to ensure transparency and improve community trust.

References:

- ETSI TR 104 066 [i.14]
- NCSC Machine Learning Principles [i.5], clause 1.4
- NCSC Guidelines for Secure AI System Development, Secure Deployment [i.6]
- CSA Companion Guide on Securing AI Systems [i.46], Section 2.2.3 Deployment, Item 3 (Release AI systems responsibly) and Annex A - Technical Testing and System Validation
- OWASP Top 10 for LLM applications [i.11]
- OWASP AI Exchange [i.10]
- OWASP Top 10 2021 [i.118]
- OWASP Top Ten for APIs [i.67]

Example Measures/Controls 2:

Use Offensive Security Assessments with Penetration Testing and Red Teaming: Use periodic penetration testing to evaluate the AI system's resilience and red teaming for models.

Chatbot App: Expand the annual penetration test to include LLM-specific threats identified in the threat model, such as indirect prompt injections through PDF uploads.

L Fraud Detection: Include evasion attacks in penetration testing, focusing on adversarial examples designed to bypass detection algorithms.

LLM Platform: Conduct red teaming exercises to evaluate the model's robustness against jailbreaking attempts, bias exploitation, and data memorization risks, leveraging multi-disciplinary teams for comprehensive assessments.

Open-Access LLM Model: For a small organization with limited resources, combine lightweight internal red teaming and invite experts from the open-source community to test for vulnerabilities, such as toxic content generation or prompt injections.

References:

- Generative AI Red Teaming Challenge Transparency Report - AI Village Defcon 2024 [i.119]
- NCSC Penetration Testing [i.120]
- Red-Teaming for Generative AI: Silver Bullet or Security Theater? [i.121]

6.9.2 Provision 5.2.5-2

"System Operators shall conduct testing prior to the system being deployed with support from Developers." (ETSI TS 104 223 [i.1])

Related threats/risks:

Without pre-deployment testing, systems can fail to meet operational and security requirements, leading to compromised performance, reduced reliability, or security vulnerabilities once deployed.

Example Measures/Controls:

Implement Comprehensive Pre-Deployment Testing: Complement development-time testing with running benchmarking emulation suites covering functional, performance, and security tests to confirm that the system meets intended requirements before deployment. These can either be in-house or independent third-party benchmarks. Integrate these tests into the development pipeline to ensure issues are identified and resolved before release.

Chatbot App: In the chatbot scenario, unit, integration, and acceptance tests are complimented with tests before deployment on performance under load, functional response accuracy, and resilience against common Top 10 for LLM attack patterns. Tests are automated at CI level.

ML Fraud Detection: In the fraud detection scenario, a similar approach is taken but with emphasis on predictive AI adversarial attacks (e.g. evasion) run against the model.

LLM Platform: For LLM releases, evaluation suits and red team exercises help ensure the model is of acceptable quality and security before deployed. Automate adversarial robustness checks and ethical evaluation frameworks in pipelines to detect biases and regressions in output trustworthiness.

Open-Access LLM Model: Apply the previous community-driven approach to early access models. Integrate tests into automated workflows.

References:

- ETSI TR 104 066 [i.14]
- NCSC Machine Learning Principles, 1.4 Analyse vulnerabilities against inherent ML threats [i.5]
- NCSC Guidelines for Secure AI System Development, Secure Deployment [i.6]
- CSA Companion Guide on Securing AI Systems [i.46], Section 2.2.3 Deployment, Item 3 (Release AI systems responsibly) and Annex A - Technical Testing and System Validation
- OWASP Top 10 for LLM applications [i.11]
- OWASP AI Exchange [i.10]

6.9.3 Provision 5.2.5-2.1

"For security testing, System Operators and Developers should use independent security testers with technical skills relevant to their AI systems." (ETSI TS 104 223 [i.1])

Related threats/risks:

Without independent security testing, systems can overlook critical vulnerabilities, especially those requiring specialized expertise, increasing the risk of undetected flaws and potential exploitation.

Example Measures/Controls:

Engage Independent Security Testers for Pre-Deployment Testing: Use independent security testers with AI expertise to validate security measures, providing an unbiased review of system security.

Chatbot App: Use external accredited Pen Testers to test the application and API including OWASP Top 10 for LLMs [i.11] in the scope.

ML Fraud Detection: In addition to the steps described for the Chatbot App, include relevant OWASP AI Exchange [i.10] items in the scope.

LLM Platform: Add for platform and API security to the steps as described for the Chatbot app and the Fraud Detection examples; employ reputable red team testers to run red team exercises against the model.

Open-Access LLM Model: Leverage the community to crowdsource red team experts for external testing; contract external experts if resources allow to complement exercises with more rigour.

References:

- NCSC CHECK Penetration Testing [i.122]
- OWASP Top 10 for LLM applications [i.11]
- OWASP AI Exchange [i.10]
- OWASP Top Ten for APIs [i.67]
- OWASP Top 10 2021 [i.118]

6.9.4 Provision 5.2.5-3

"Developers should ensure that the findings from the testing and evaluation are shared with System Operators, to inform their own testing and evaluation." (ETSI TS 104 223 [i.1])

Related threats/risks:

Without access to previous testing findings, System Operators can lack critical insights into known vulnerabilities or limitations, and mitigations that cannot be seen as adequate by operators leading to potential gaps in security coverage.

Example Measures/Controls:

Establish a Process for Sharing Testing Results: Create a standardized process for sharing all relevant testing results and evaluation findings with System Operators of testing reports and ensure timely access for all relevant stakeholders.

Chatbot App: Share: Share testing results on identified vulnerabilities with System Operators and planned mitigations.

ML Fraud Detection: Share testing findings with System Operators, highlighting a lower rate of rejecting evasion attacks when transactions involve vendors already used by the user legitimately.

LLM Platform: Provide detailed reports on testing outcomes, including risks like data memorization or output bias, and offer guidelines for secure deployment practices

Open-Access LLM Model: Publish a summary of security and robustness findings in community forums, highlighting vulnerabilities like poisoning risks and encouraging feedback for improvements. Provide more detailed reports to customers on request with detailed risks and proposed mitigations.

References:

- ISO/IEC/IEEE 29119-1: Software Testing Standards - Communication and Reporting [i.123]

6.9.5 Provision 5.2.5-4

"Developers should evaluate model outputs to ensure they do not allow System Operators or End-users to reverse engineer non-public aspects of the model or the training data." (ETSI TS 104 223 [i.1])

Related threats/risks:

Without adequate output evaluation, Operators or End-users can reverse engineer model internals or correlate outputs with sequences of predesigned inputs, enabling them to reconstruct the model or training data. This can erode commercial advantage, allow the creation of shadow models, or facilitate attacks, even affecting "open source" models if only parts, such as architecture or weights, are released.

Example Measures/Controls 1:

Conduct Security Reviews for Output Sensitivity: Engage with AI experts to evaluate model outputs and identify any information that could reveal internal model structures, ensuring sensitive aspects remain non-public.

Chatbot App: Test chatbot responses for unintended patterns that could reveal prompt templates or system configurations, adjusting output constraints to maintain operational privacy.

ML Fraud Detection: Analyse outputs for patterns revealing specific feature importance (e.g. high fraud scores tied to rare combinations), replacing explicit results with categorical risk levels to obscure decision logic.

LLM Platform: Test responses to ensure the model does not disclose memorized training data or hint at weight configurations, implementing filters to mask sensitive details.

Open-Access LLM Model: Follow the approach described in the LLM Platform example while harnessing the community and crowd sourcing.

References:

- ETSI TR 104 066 [i.14]
- ETSI TR 104 225 [i.13]
- OWASP AI Exchange [i.10]

Example Measures/Controls 2:

Integrate Adversarial Testing and Robustness Evaluation: Implement adversarial testing to assess the resilience of AI systems against attempts to reverse engineer or manipulate model outputs. This involves simulating attacks that exploit output data to uncover model weaknesses or extract sensitive information. Open-source tools like ART and TextAttack can help to develop and run these tests.

Chatbot App: Use adversarial testing to simulate prompt injection attacks aimed at eliciting unintended responses. Adjust prompt templates and input validation to reduce vulnerability.

ML Fraud Detection: Test the system against adversarial inputs designed to evade fraud detection, such as fabricated transaction patterns. Refine detection algorithms to improve resilience.

LLM Platform: Simulate jailbreaking attempts to bypass guardrails in the LLM's responses. Update output constraints and enhance filtering mechanisms to prevent such exploits.

Open-Access LLM Model: Leverage community participation to create adversarial tests for input patterns, such as queries that manipulate outputs to infer training data. Apply patches and retrain models to address these vulnerabilities.

References:

- ETSI TR 104 066 [i.14]
- CSA Companion Guide on Securing AI Systems, List of AI Testing Tools [i.46]

- Linux Foundation AI & Data Foundation: Adversarial Robustness Toolbox [i.124]
- NCSC Machine Learning Principles, 1.4 Analyse vulnerabilities against inherent ML threats [i.5]
- OWASP AI Exchange [i.10]
- TextAttack: Generating adversarial examples for NLP models [i.125]

6.9.6 Provision 5.2.5-4.1

"Additionally, Developers should evaluate model outputs to ensure they do not provide System Operators or End-users with unintended influence over the system." (ETSI TS 104 223 [i.1])

Related threats/risks:

Without evaluating for unintended influence, Operators or End-users can exploit system behaviour to align outputs with personal or malicious goals, leading to bias, misuse, or compromised trust in the system.

Example Measures/Controls:

Implement Safeguards Against Manipulation: Set controls to prevent Operators or End-users from adjusting inputs in ways that could intentionally influence the AI system's outcomes.

Chatbot App: apply prompt engineering and guardrails to constraint prevent output manipulation and abuse the app using prompt injection attacks.

ML Fraud Detection: Limit Operator access to input variables to prevent modifications that could allow compromised insiders to tamper model predictions.

LLM Platform: Implement rate limits and input validation to detect and block repeated adversarial inputs designed to manipulate model responses, ensuring consistent and unbiased outputs.

Open-Access LLM Model: Develop or encourage community source plugins to enforce input constraints, preventing users from crafting queries aimed at exploiting or biasing the model's responses.

References:

- OWASP Top 10 for LLM applications [i.11]
- OWASP AI Exchange [i.10]

6.10 Principle 10: Communication and processes associated with End-users and Affected Entities

6.10.1 Provision 5.3.1-1

"System Operators shall convey to End-users in an accessible way where and how their data will be used, accessed, and stored (for example, if it is used for model retraining, or reviewed by employees or partners). If the Developer is an external entity, they shall provide this information to System Operators." (ETSI TS 104 223 [i.1])

Related threats/risks:

Insufficient communication about data usage, access, or storage practices can lead to misunderstandings and mistrust among End-users. This lack of transparency increases the risk of misuse or unauthorized access to data, as users might not fully understand or consent to how their data is handled, potentially resulting in regulatory and reputational damage.

Example Measures/Controls:

Establish Transparent Data Usage Communication: Provide a transparent overview of data usage policies, purpose of processing, specifying whether data will be used for model retraining, third-party access, or employee review. Each processing activity needs to be logged, and a compatibility assessment needs to be made for each new purpose. Ensure end users understand all potential uses of their data and how it contributes to AI model performance or security and that documentation is in an accessible format

Chatbot App: Outline to the end user of the chatbot app how their conversations will be logged, used, and monitored. Include in a downloadable accessible PDF File.

ML Fraud Detection: Explain with clear descriptions in the end-user agreement and privacy policy how customer data will be used for online training and the protections, e.g. anonymization, in place. Document transparently regions where data is stored. Documentation is WCAG 2.1 compliant.

LLM Platform: Explain with clear descriptions in the end-user agreement and privacy policy how customer data will be used for online training and the protections, e.g. anonymization, in place. Document transparently regions where data is stored. Documentation is WCAG 2.1 compliant.

Open-Access LLM Model: Document how personal data, if any, are processed. Make this information available both as a downloadable accessible PDF file and on a dedicated, well-structured, and accessible HTML web page to ensure broad accessibility.

References:

- ICO Generative AI second call for evidence: Purpose limitation in the generative AI lifecycle [i.126]
- GOV.UK Guidance and tools for digital accessibility [i.111]

6.10.2 Provision 5.3.1-2

"System Operators shall provide End-users with accessible guidance to support their use, management, integration, and configuration of AI systems. If the Developer is an external entity, they shall provide all necessary information to help System Operators." (ETSI TS 104 223 [i.1])

Related threats/risks:

Without clear guidance, End-users might mismanage the software, leading to data leaks, vulnerabilities, or system misuse, exposing AI systems to security risks.

Example Measures/Controls:

Provide Comprehensive User Guides and Tutorials: Develop and distribute detailed user guides and tutorials that cover secure configuration, integration steps, and recommended usage practices, ensuring users understand safe operation protocols. Provide accessible notification options, such as screen reader-compatible text alerts, and customisable notifications for users with sensory impairments

Chatbot App: Provide internal system documentation app configurations including LLM API, highlighting security aspects such as secrets management, encryption, access control configuration and so on, with a downloadable accessible PDF File.

ML Fraud Detection: Internal detailed guide on deploying the model and associated services, highlighting security requirements. Documentation is WCAG 2.1 compliant.

LLM Platform: Documentation includes step-by-step instructions on integrating the model with existing systems securely, explaining the need for least-privilege access with read-only roles when invoking APIs, use of TLS, secret management for configuration credentials and so on. Documentation is WCAG 2.1 compliant.

Open-Access LLM Model: Like the LLM Platform example, but less detailed. The organization takes advantage of their Wiki's accessible features to ensure accessibility.

References:

- GOV.UK Guidance and tools for digital accessibility [i.111]

6.10.3 Provision 5.3.1-2.1

"System Operators shall include guidance on the appropriate use of the model or system, which includes highlighting limitations and potential failure modes." (ETSI TS 104 223 [i.1])

Related threats/risks:

Failing to highlight limitations and potential failure modes can lead to End-users relying on the AI system for unsupported or inappropriate tasks, increasing the risk of operational errors, data misuse, or unintended consequences.

Example Measures/Controls:

Highlight Model Limitations and Failure Modes: Clearly outline the model's appropriate uses, limitations, and potential failure modes to inform end-users of scenarios in which the model might produce inaccurate or unreliable outputs.

Chatbot App: Inform users that the chatbot can struggle with complex or ambiguous user queries and encourage fallback to human operators for critical or high-impact scenarios.

ML Fraud Detection: Inform users of scenarios (e.g. unusual market conditions) where the model can produce less accurate predictions, allowing operators to interpret results carefully.

LLM Platform: Document the inherent LLM tendency of hallucinations, providing plausible sounding but incorrect information. Highlight hallucinations in code generation and its effect on creating vulnerable code.

Open-Access LLM Model: Provide a summary of known strengths and weaknesses. Highlight the strength on legal queries and poor performance on other domains such as finance.

References:

- NIST AI RMF [i.17]

6.10.4 Provision 5.3.1-2.2

"System Operators shall proactively inform End-users of any security relevant updates and provide clear explanations in an accessible way." (ETSI TS 104 223 [i.1])

Related threats/risks:

Without proactive communication about security-relevant updates, End-users can fail to understand changes in system behaviour or associated risks. This lack of awareness can lead to improper use or failure to take necessary precautions, increasing the system's exposure to potential exploitation or security breaches.

Example Measures/Controls:

Notify Users of Security Updates: Proactively inform End-users about security update, detailing the purpose and impact of each update to promote user compliance. Ensure all users, can be informed by using accessible formats.

Chatbot App: Provide clear notifications to internal end-users when a security patch is applied to the hiring chatbot that disables Word documents and restricts to PDF files to avoid certain type of attacks. Similarly, notify internal users when e-mailing functionality is disabled following a prompt injection attack abusing the functionality. Explain how it enhances system resilience and inform them of any changes to features or functionality they need to be aware of. The notification should be available as an accessible web page or accessible downloadable PDF file.

ML Fraud Detection: Inform users of a security update applied to the fraud detection system, such as enhanced protections against adversarial evasion attacks. Explain, in a WCAG- 2.1 compliant page, how the update improves the system's ability to detect malicious patterns and reduce false negatives.

LLM Platform: Inform end-users of updates to API functionality, including improvements to safety guardrails, such as enhanced mitigation against hallucinations or bias in outputs. Provide documentation of these updates in accessible formats, such as a changelog or dedicated update page, ensuring the information is easy to understand and actionable and is WCAG 2.1-compliant.

Open-Access LLM Model: Publish detailed release notes in the repository's changelog and mailing list when security-related updates are applied to the model, such as improved robustness against poisoning attacks. Include an accessible summary page and a downloadable accessible PDF.

References:

- GOV.UK Guidance and tools for digital accessibility [i.111]

6.10.5 Provision 5.3.1-3

"Developers and System Operators should support affected End-users and Affected Entities during and following a cyber security incident to contain and mitigate the impacts of an incident. The process for undertaking this should be documented and agreed in contracts with End-users." (ETSI TS 104 223 [i.1])

Related threats/risks:

Failure to support affected End-users and Affected Entities during an incident can result in prolonged recovery times, mismanagement of containment efforts, and increased reputational damage due to inadequate communication or guidance.

Example Measures/Controls:

Establish a Documented Incident Support and Communication Process: Develop and document a support process for responding to incidents, covering steps for containment, impact assessment, and recovery, and specify the roles and responsibilities of Developers and System Operators. This should include specialist skills and AI expertise that might be require. Provide support through accessible formats, ensuring usability for all affected stakeholders.

Chatbot App: Establish a process for assisting End-users who report biased or inappropriate chatbot responses, ensuring NLP/LLM experts are available to investigate the issue. Determine whether the problem stems from model updates, adversarial inputs, or misconfigurations. Provide clear and accessible guidance to End-users on mitigating such issues and share updates as an accessible downloadable PDF file.

ML Fraud Detection: Create a support mechanism to help affected entities (e.g. financial institutions) interpret and respond to a surge in incorrect fraud predictions. Ensure that experts in interpretability techniques such as SHAP and LIME are available to explain the root cause of the issue and guide End-users on adapting thresholds or re-evaluating flagged transactions. Provide all guidance in accessible formats, such as WCAG 2.1-compliant documentation.

LLM Platform: Develop a process for handling incidents such as jailbreaking attempts or misuse of generated content (e.g. deepfakes). Engage AI ethics and adversarial attack experts to assist affected users by providing timely updates, risk mitigation advice, and recovery strategies. Share detailed incident analysis and recovery steps in accessible formats, such as a dedicated incident response webpage with downloadable PDFs.

Open-Access LLM Model: Leverage contributions from the community of researchers and developers to support affected End-users in responding to issues like data poisoning or toxic content generation. Provide detailed incident updates and remediation guidance via accessible Wiki pages, ensuring compliance with accessibility standards. Offer End-users additional support through mailing lists or forums, ensuring all communication is clear and actionable.

References:

- GOV.UK Guidance and tools for digital accessibility [i.111]
- NCSC Incident Management, Plan: Your cyber incidence response process [i.102]
- NCSC Machine Learning Principles [i.5], clause 4.3

6.11 Principle 11: Maintain Regular Security Updates, Patches, and Mitigations

6.11.1 Provision 5.4.1-1

"Developers shall provide security updates and patches, where possible, and notify System Operators of the security updates. System Operators shall deliver these updates and patches to End-users." (ETSI TS 104 223 [i.1])

Related threats/risks:

Delayed updates allow attackers to exploit vulnerabilities, increasing the risk of unauthorized access, data breaches, and compromised AI system functionality.

Example Measures/Controls 1:

Implement a Structured Patch Management Process: Establish a structured process for developing, testing, and releasing security patches for AI systems, ensuring regular updates to address known vulnerabilities with user notifications.

Chatbot App: App packages and containers are patched weekly to mitigate new library vulnerabilities being exploited to stage a breach with e-mails.

ML Fraud Detection: In addition to packages, a patching process for the model helps update the system with security-related fixes, emailing customers of forthcoming updates.

LLM Platform: Like the Fraud Detection example, with email-notifications to API customers.

Open-Access LLM Model: Like Fraud Detection example but different cadence closer to the regular monthly releases and posting updates for new updates on community forums.

References:

- CSA Companion Guide on Securing AI Systems [i.46], Section 2.2.4 Operations and Maintenance
- NCSC Incident Management [i.102], Plan: Your cyber incidence response process
- NCSC Machine Learning Principles [i.5], clause 4.3
- GOV.UK Guidance and tools for digital accessibility [i.111]

Example Measures/Controls 2:

Enable Automatic Updates Where Possible: Configure AI systems to support automatic security updates, minimizing the risk of delayed patch implementation.

Chatbot App: The organization operates automated container image patching combined with regression testing to avoid breaking changes.

ML Fraud Detection: Similar to the implementation in the Chatbot App example, as part of staged rollouts with automated regression tests and rollback mechanisms.

LLM Platform: Similar to the implementation in the Chatbot App example, as part of staged rollouts with automated regression tests and rollback mechanisms.

Open-Access LLM Model: Implement automatic updates for both open-source code and model weights. Use CI/CD tools to pull, test, and deploy new releases, ensuring updates are tested for vulnerabilities and announced on community forums to notify System Operators.

References:

- NCSC Vulnerability Management [i.108] - Put in a policy to update by default

6.11.2 Provision 5.4.1-1.1

"Developers shall have mechanisms and contingency plans to mitigate security risks, particularly in instances where updates cannot be provided for AI systems." (ETSI TS 104 223 [i.1])

Related threats/risks:

Unaddressed vulnerabilities can lead to exploitation if updates are not feasible, increasing risks of unauthorized access and system disruptions.

Example Measures/Controls:

Develop Contingency Plans for Non-Updateable Components: Establish and document contingency plans, including compensating controls, for components that cannot receive updates due to system limitations or dependencies.

Chatbot App: Deploy compensating controls such as rate limiting and enhanced logging for an outdated chatbot version integrated with a legacy CRM system that cannot be updated due to compatibility issues.

ML Fraud Detection: An earlier version of the fraud detection model for a specific market segment with some custom escalation rules is running on legacy software platform that cannot be upgraded due to incompatibilities. Instead, compensating controls such as network segmentation to isolate the system from other critical environment and use of Intrusion Detection Systems and enhanced monitoring are added to detect potential exploitations.

LLM Platform: For a legacy deployment of the platform's API gateway that cannot be updated, implement compensating controls like API request filtering, strict authentication mechanisms, and enhanced monitoring and alerting to minimize risks.

Open-Access LLM Model: Recommend compensating controls such as sandboxing to customers using an older version of the model that cannot be updated due to dependency on outdated libraries.

References:

- NCSC Vulnerability Management [i.108]

6.11.3 Provision 5.4.1-2

"Developers should treat major AI system updates as though a new version of a model has been developed and therefore undertake a new security testing and evaluation process to help protect users." (ETSI TS 104 223 [i.1])

Related threats/risks:

Inadequate testing of major updates can introduce vulnerabilities and changes in model behaviour, risking data breaches, system manipulation, or unintended behaviour.

Example Measures/Controls:

Conduct Comprehensive Security Testing for Major Updates: Perform a security assessment, including penetration testing and model red teaming, for any major AI system updates, treating each update as a new version.

Chatbot App: When switching to a new major version of the supporting LLM, use automated tests to ensure guardrails continue to be effective.

ML Fraud Detection: use adversarial robustness tests, when performing major retraining of the model with some new datasets to mitigate risks of new vulnerabilities.

LLM Platform: Use pen tests and red teaming to test a major new version with additional API endpoints.

Open-Access LLM Model: Conduct community-driven red teaming exercises for a major update, inviting experts to test for vulnerabilities like poisoning risks or model extraction attacks, and document mitigation actions.

References:

- NCSC Machine Learning Principles, 1.4 Analyse vulnerabilities against inherent ML threats [i.5]
- NCSC Guidelines for Secure AI System Development, Secure Deployment [i.6]

- CSA Companion Guide on Securing AI Systems [i.46], Section 2.2.3 Deployment, Item 3 (Release AI systems responsibly) and Annex A - Technical Testing and System Validation
- NIST AI Test, Evaluation, Validation, and Verification (TEVV) [i.114]

6.11.4 Provision 5.4.1-3

"Developers should support System Operators to evaluate and respond to model changes, (for example by providing preview access via beta-testing and versioned APIs)." (ETSI TS 104 223 [i.1])

Related threats/risks:

Without evaluation support, System Operators can overlook risks in updates, leading to potential configuration errors or unmitigated vulnerabilities.

Example Measures/Controls:

Offer Preview Access for Major Model Updates: Provide System Operators with preview access to major model updates, allowing for evaluation and adjustment to any new system behaviour.

Chatbot App: Internal System Operators should have access to test environments for new releases.

ML Fraud Detection: Similar to the Chatbot App example and for new model versions.

LLM Platform: The provider offers a quarterly beta program via API versions to help customers evaluate changes before they upgrade their applications.

Open-Access LLM Model: Early beta versions of the model are released regularly as part of a vetted beta program with confidentiality agreements and with instructions to customers' System Operators on how to host and evaluate it.

References:

- NIST AI RMF [i.17]
- OWASP AI Exchange [i.10]

6.12 Principle 12: Monitor the system's behaviour

6.12.1 Provision 5.4.2-1

"System Operators shall log system and user actions to support security compliance, incident investigations, and vulnerability remediation." (ETSI TS 104 223 [i.1])

Related threats/risks:

Insufficient logging limits incident investigation and compliance enforcement, weakening the ability to detect and respond to security incidents.

Example Measures/Controls 1:

Implement Comprehensive Logging for Security and Compliance: Establish a logging framework that captures key aspects of system behaviour, including user interactions, access events, data flows, and model outputs, ensuring logs support compliance and security standards.

Chatbot App: Log user interactions with the chatbot, including prompt metadata, such as timestamps, session Ids, and response times, to support auditability and investigations. Avoid logging full user prompts or responses unless anonymized or explicitly necessary for troubleshooting and ensure compliance with data protection regulations by implementing data minimization, retention policies, and secure storage.

ML Fraud Detection: Log user access attempts, model predictions, and flagged transactions to maintain a robust audit trail for compliance and incident analysis.

LLM Platform: Log fine-tuning activities, model deployments, and system usage metadata, ensuring sensitive operations are traceable and aligned with security policies.

Open-Access LLM Model: Maintain logs of public feedback contributions and model training iterations, capturing details such as contributor IDs, timestamps, and flagged training data to improve traceability.

References:

- CSA Companion Guide on Securing AI Systems [i.46], Section 2.2.4
- CISA Joint Cybersecurity Information - Deploying AI Systems Securely [i.7]
- NCSC guidance on Logging for Security Purposes [i.64]
- NCSC Machine Learning Principles [i.5], clause 3.2

Example Measures/Controls 2:

Ensure Secure Storage and Retention of Logs: Implement secure storage mechanisms, such as encryption (both at rest and in transit), to protect logs from unauthorized access. Define a retention policy that aligns with compliance obligations and supports security incident investigations. Avoid logging personal data unless strictly necessary; if personal data needs to be logged, ensure it is anonymized or obfuscated and managed in compliance with applicable privacy laws (e.g. UK GDPR, CCPA). Periodically review and update the retention policy to address evolving compliance and operational requirements.

Chatbot App: Store chat metadata log (in encrypted storage with access restricted to authorized personnel, retaining logs for one year to comply with operational policies.

ML Fraud Detection: Use an encrypted cloud-based storage solution for transaction logs, retaining them for X years per financial regulations, with secure deletion for older logs.

LLM Platform: Encrypt logs from training and model management systems, retaining them for a defined period based on regulatory and contractual obligations, and operational policies.

Open-Access LLM Model: Securely store logs of training contributions and flagged outputs in an encrypted repository, with RBAC granted to minimum number of people, as needed. Retain the logs for six months to enable community-driven debugging and transparency.

References:

- CISA Joint Cybersecurity Information - Deploying AI Systems Securely [i.7]
- NCSC guidance on Logging for Security Purposes [i.64]

Example Measures/Controls 3:

Establish Routine Log Analysis for Model Validation: Define a regular schedule for log analysis to assess model output consistency, detect anomalies, and verify that outputs align with desired outcomes.

Chatbot App: Regularly analyse session logs to identify trends in user escalations or repeated failed responses, which could indicate prompt misalignment.

ML Fraud Detection: Use of regular log analysis help detect poisoning patterns or drift when online learning is introduced in the fraud detection scenario for certain volatile markets where fraud patterns evolve quickly.

LLM Platform: Perform periodic log reviews to detect anomalies in API usage patterns, such as repeated requests for sensitive topics, which may signal adversarial testing or misuse.

Open-Access LLM Model: Use community feedback to support customer System Operators with log analysis to identify trends in toxic or biased outputs, correlating them with specific input patterns to address potential data poisoning.

References:

- CISA Joint Cybersecurity Information - Deploying AI Systems Securely [i.7]

6.12.2 Provision 5.4.2-2

"System Operators should analyse their logs to ensure that AI models continue to produce desired outputs and to detect anomalies, security breaches, or unexpected behaviour over time (such as due to data drift or data poisoning)." (ETSI TS 104 223 [i.1])

Related threats/risks:

Without regular log analysis, security breaches or model issues can go undetected, potentially leading to incorrect outputs or system vulnerabilities

Example Measures/Controls:

Implement Alerts for Anomalous Model Behaviour: Set up alerts to notify operators of unexpected behaviour, such as unusual outputs, abnormal input patterns, or significant deviations from historical performance.

Chatbot App: Configure alerts to flag instances where user query failures exceed a threshold, indicating potential issues with the underlying language model or prompt configurations.

ML Fraud Detection: For a fraud detection system, implement alerts to notify operators when the model starts flagging an unusually high number of transactions as fraudulent in a short period. The alert triggers a review to investigate whether the issue is due to changes in input data, a model drift event, or potential adversarial activity.

LLM Platform: Set up alerts for sudden spikes in requests targeting specific APIs, especially for sensitive topics or jailbreaking attempts, prompting a review for adversarial testing.

Open-Access LLM Model: Encourage users and System Operators send notifications via e-mail of anomalous behaviour.

References:

- CISA Joint Cybersecurity Information - Deploying AI Systems Securely [i.7]

6.12.3 Provision 5.4.2-3

"System Operators and Developers should monitor internal states of their AI systems where they feel this could better enable them to address security threats, or to enable future security analytics." (ETSI TS 104 223 [i.1])

Related threats/risks:

Lack of internal state monitoring can delay detection of security threats or model drift, increasing risks of unauthorized modifications and system failure.

Example Measures/Controls 1:

Implement Monitoring of Key Internal States: Identify and monitor critical internal states of the AI system, such as hidden layers, attention weights, or feature importance, which could provide early indicators of security threats.

Chatbot App: No action taken since model is operated by another party and existing logging is sufficient.

ML Fraud Detection: Monitor the importance of features like transaction amount, geolocation, or merchant type in the fraud detection scenario; a sudden drop in geolocation importance could indicate drift or tampering by fraudsters adapting their strategies.

LLM Platform: Critical use cases might require introspection of internal model weights, such as where patterns are identified during development and testing that indicate a failure state.

Open-Access LLM Model: The small organization needs to investigate and understand reported biases. As a result, implements this advanced type of monitoring to track shifts in attention weights during community training contributions to detect how biases are introduced.

References:

- CISA Joint Cybersecurity Information - Deploying AI Systems Securely [i.7]

Example Measures/Controls 2:

Use Secure Storage for Internal State Data: Store data from monitored internal states securely, ensuring that sensitive internal metrics are protected from unauthorized access.

Chatbot App: Encrypt and store user session data and conversation states securely in a database with restricted access.

ML Fraud Detection: The company saves the weights in a protected and isolated storage using a separate dedicated enclave as recommended by the CSI/CISA/NCSC Joint Cybersecurity Information - Deploying AI Systems Securely [i.7].

LLM Provider: Similar to the implementation in the Fraud Detection example.

Open-Access LLM Model: Use secure storage with access restrictions and encryption for any internal state information that might be generated and used during development.

References:

- CISA Joint Cybersecurity Information - Deploying AI Systems Securely [i.7]

Example Measures/Controls 3:

Track and Benchmark Model Performance Metrics: Define and monitor key performance metrics for the AI system, such as statistical accuracy, factual correctness, response time, and error rate establishing benchmarks to detect deviations.

Chatbot App: Tracking user feedback in the UI escalation rates (conversations requiring human intervention) can reveal gradual degradation of the chatbot app.

ML Fraud Detection: Monitoring prediction accuracy ensures the fraud detection model remains within the target range and detects evasion attacks; Usage spikes can indicate extraction or reconnaissance attacks.

LLM Platform: Using automated checks with semantic similarity metrics and trusted datasets, and periodic fact-checking of outputs by LLM developers can detect poisoning. Use of third-party safety APIs and custom classifiers can be used to detect biased, toxic, or inappropriate language in outputs.

Open-Access LLM Model: Monitor community-reported feedback and benchmark outputs against open datasets to detect performance issues such as reduced factual accuracy or increased bias in generated text. Use automated tools to flag significant deviations for review.

References:

- ICO What do we need to know about accuracy and statistical accuracy? [i.127]

6.12.4 Provision 5.4.2-4

"System Operators and Developers should monitor the performance of their models and system over time so that they can detect sudden or gradual changes in behaviour that could affect security." (ETSI TS 104 223 [i.1])

Related threats/risks:

Failing to monitor performance over time can hide behavioural shifts, making the system more vulnerable to degradation, attacks, and inconsistencies

Example Measures/Controls:

Implement Drift Detection to Identify Behavioural Shifts: Use drift detection tools to identify shifts in model behaviour due to changing data patterns or environmental factors, allowing proactive responses

Chatbot App: No action taken, as user prompts are not used for model training.

ML Fraud Detection: Since online training has been adopted, drift monitoring is applied, detecting shifting tactics like the use of VPNs and proxies that decrease the importance of transaction location to mask fraud. Use statistical tests to detect shifts in transaction distributions.

LLM Platform: Concept drift is actively monitored to detect and adapt to changes in language, emojis, and jargon, which could bypass content filters altering model behaviour and be exploited by attackers for jailbreak attempts. This monitoring helps identify emerging vulnerabilities and ensures regular updates to the model's safety guardrails.

Open-Access LLM Model: No action taken since no online training takes place and the model is new.

References:

- CISA Joint Cybersecurity Information - Deploying AI Systems Securely [i.7]
- NCSC Machine Learning Principles [i.5], clause 4.1

6.13 Principle 13: Ensure proper data and model disposal

6.13.1 Provision 5.5.1-1

"If a Developer or System Operator decides to transfer or share ownership of training data and/or a model to another entity they shall involve Data Custodians and securely dispose of these assets. This will protect AI against security issues that can transfer from one AI system instantiation to another." (ETSI TS 104 223 [i.1])

Related threats/risks:

Improper disposal or transfer can lead to unauthorized data recovery, risking breaches, IP loss, and non-compliance with data protection laws.

Example Measures/Controls:

Develop and Implement a Secure Transfer and Disposal Policy with Data Custodian Oversight: Establish a comprehensive policy to govern the secure transfer and disposal of training data and models. Ensure that Data Custodians oversee all actions, confirming compliance with regulatory standards, protection of intellectual property, and adherence to organizational policies. This includes compliance with GDPR when involving personal data.

Chatbot App: Define procedures for securely transferring or deleting fine-tuning datasets used in the chatbot's LLM. Data Custodians to actively approve all deletions or transfers.

ML Fraud Detection: Develop protocols for securely transferring or deleting training data and models, including sensitive customer transaction records. Require Data Custodian active approval before executing any actions.

LLM Platform: Establish a policy for securely transferring model ownership or licensing agreements, ensuring all fine-tuned and proprietary models are disposed of or transferred in compliance with contractual and regulatory obligations.

Open-Access LLM Model: For open-source contributions, include guidelines for securely deleting intermediate datasets used during training, reviewed by the nominated Data Custodian.

References:

- ICO Disposal and deletion [i.128]
- ICO Guidance on AI and data protection [i.129]
- ICO Retention and destruction of information [i.116]
- CSA Companion Guide on Securing AI Systems [i.46], Section 2.2.5 End of life
- NCSC Machine Learning Principles [i.5], Part 5: End of Life

6.13.2 Provision 5.5.1-2

"If a Developer or System Operators decides to decommission a model and/or system, they shall involve Data Custodians and securely delete applicable data and configuration details." (ETSI TS 104 223 [i.1])

Related threats/risks:

Insecure data deletion during decommissioning can lead to unauthorized access to residual data, increasing regulatory and security risks.

Example Measures/Controls:

Implement a Secure Data Deletion Policy with Data Custodian Oversight: Establish a policy for securely deleting data and models during decommissioning, specifying methods compliant with standards. Ensure Data Custodians validate all deletions to maintain regulatory compliance and traceability.

Chatbot App: Securely delete all conversation logs, prompt data, and configurations when decommissioning the chatbot system. Involve Data Custodians to verify deletion of fine-tuning data accessed via the API.

ML Fraud Detection: Enforce a policy requiring the secure deletion of all historical transaction data, model weights, and configurations during system decommissioning. Data Custodians ensure compliance with financial and privacy regulations.

LLM Platform: Notify customers before decommissioning public LLM services, ensuring all uploaded data is securely deleted after notification periods expire. Require Data Custodians to validate deletion processes.

Open-Access LLM Model: Securely delete all training and test datasets, intermediate model versions, and unreleased models when ceasing development of an open-access LLM. Require that the Data Custodians to verify the process to ensure transparency and compliance.

References:

- See References in clause 6.13.1
- NIST SP 800-88 Rev. 1 [i.130]
- NCSC Secure sanitisation of storage media [i.131]

Annex A:

Mapping from design and organization principles to SAI

CoP Principle		ETSI (TC SAI) mapping/comment
Identity	Summary text	
1	Raise awareness of AI security threats and risks.	This is not something that particularly impacts an SDO. However it is clear that a number of documents prepared by ETSI do address this in general terms, and in some domains in very specific terms. The problem statement is a key document (ETSI TR 104 221 [i.132]) as is the manipulation document (ETSI TR 104 062 [i.133]). In understanding the overall ecosystem several ETSI documents apply including ETSI TR 104 029 [i.134], the ontology in ETSI TS 104 050 [i.135], the ETSI white papers on AI [i.150] and [i.151], and ETSI TR 104 222 [i.2] all give a good grounding in understanding the problem and attack space and the viability of various mitigations. In addition ETSI is actively preparing education material on key topics and AI is planned to be covered in some detail as part of a joint programme with Digital Europe [i.148]. An initial example for AI and AI-Security was trialled with the University of Luxembourg and in 2024 and will serve as the basis of a comprehensive suite of education material to be developed in late 2025.
2	Design the AI system for security as well as functionality and performance.	The way in which ETSI (and other SDOs) develop standards assumes that developers do not favour one dimension over another. Across best practice including the practices of Secure by Default the notion of balance is well described. There might be a gap to fill across the SDOs to address this balancing act. It is noted that achieving such a balance is not specific to AI or AI-systems and is a principle that applies to all systems.
3	Evaluate the threats and manage the risks to the AI system.	The metrics for identifying and modelling threats, including the impact of AI have been addressed in the Ontology (ETSI TS 104 050 [i.135]) and in ETSI TS 102 165-1 [i.137]. The CoP recommends identifying the motivation for an attacker which has been given a great deal of qualification in ETSI TS 102 165-1 [i.137] as assessments of motivation are often inaccurate before the act. In addition the CoP recommends conducting a DPIA which is addressed in ETSI TS 103 485 [i.138] and others. In addition a number of other SDOs have addressed methods for undertaking a DPIA.
4	Enable human responsibility for AI systems.	In part this is addressed by ETSI TS 104 224 [i.139] on transparency and explicability. Whilst the principle applies to AI it can be reasonably argued that all systems with or without AI have system specific risks and that there should be full consideration of those risks. As noted above there are metrics and methods for assessing risk in ETSI TS 102 165-1 [i.137] and across many SDOs.
5	Identify, track and protect the assets.	This is addressed in the Cyber Security Controls defined in ETSI TR 103 305-1 [i.140] and updated for AI in ETSI TR 104 030 [i.136]. In particular the CSC-1 and CSC-2 apply that require organizations to catalogue hardware and software assets.
6	Secure the infrastructure.	This is key to all of ETSI's security work in a large number of its TBs. ETSI has prepared a very large number of standards that when deployed can give greater assurance of the security of the infrastructure. It can be useful for an ETSI deliverable (or interactive data driven visualization tool) to give a map to standards that apply to ensuring secure infrastructure.

CoP Principle		ETSI (TC SAI) mapping/comment
Identity	Summary text	
7	Secure the supply chain.	This is addressed specifically for the ML data supply chain in ETSI TR 104 048 [i.145] and is being more widely addressed across ETSI in various forms of standards addressing Bills of Material (so SBOMS, AIBOMs and so on). This is also being addressed in a number of other venues including FIRST with its AI Special Interest Group.
8	Document data, models and prompts.	This is the key concept of the transparency and explicability document (ETSI TS 104 224 [i.139]). It is strongly asserted there, and in the CSC documents (ETSI TR 103 305-1 [i.140] to 5 [i.140], [i.141], [i.142], [i.143] and [i.144] and ETSI TR 104 030 [i.136]), that without a clear picture of the assets in the system and how they interact that the system cannot be made secure. This is reinforced by activity in ISG ETI addressing the role of the Zero-Trust approach (captured in the ZT-Kipling method).
9	Conduct appropriate testing and evaluation.	ETSI TR 104 066 [i.14] addresses this and in addition activity in ETSI MTS-AI is considering testing and testability of AI based systems.
10	Communication and processes associated with End-users and Affected Entities.	Key to building trust in the role of AI in systems is the transparency and explicability document (ETSI TS 104 224 [i.139]) and the general rationalization of AI threats, mitigations and supply chains across all of the output of TC SAI.
11	Maintain regular security updates, patches and mitigations.	This is not specific to AI but the AI dimension of updates is included in much of the activity of TC CYBER where this dimension of system maintenance is being addressed (in particular in relation to RED, CRA and NIS2).
12	Monitor the system's behaviour.	This is quite closely tied into transparency, explicability and maintainability and addressed across the work programme of TC SAI.
13	Ensure proper data and model disposal.	The end of life of an AI system including its data and model has to ensure that the AI system cannot be reinstated in a manner consistent with requirements identified in regulation including GDPR [i.54] and the e-Privacy Directive [i.149]. If the AI system contains personal data it may be necessary to also enable partial data deletion and model update upon direction of the affected user (e.g. by enabling the right to be forgotten).

Annex B:

Bibliography

- ICO: "[Artificial Intelligence](#)".
- OWASP®: "[Machine Learning Security Top Ten](#)".

History

Document history		
V1.1.1	May 2025	Publication