

GISC 2025

Global ICT Standards Conference 2025

Enabling the physical AI era with global standards

European approach to the Physical AI and Standardization Direction

Scott CADZOW
Chair TC SAI, ETSI

ICT Standards and Intellectual Property:
AI for All

Index

- 01** The evolution and timeline of AI standards in ETSI and Europe
- 02** An overview of the approach in ETSI to standards development
- 03** A review and forecast for ETSI's standards work in AI

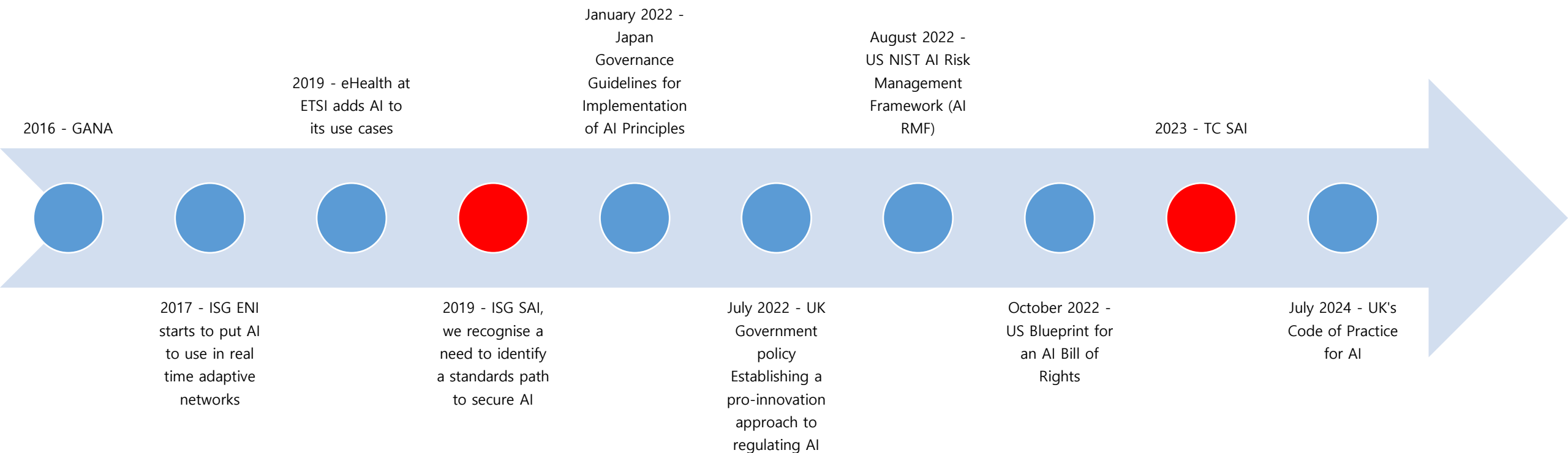
Abstract

| European standards in AI - a view from ETSI

AI Standards

An overview, primarily from the perspective of ETSI, of standards to secure and assure safe use of AI. The metric for success in standardization is proposed and outlined. The presentation addresses the ETSI approach of consensus building. It highlights the work done in the last 6 years where AI has been an active topic in ETSI and outlines the content of the standards published and the ones in development. It also identifies future areas where standards are planned and where global cooperation is essential for success

01. A timeline of AI standards and governance milestones



02. General approach to standardization in ETSI

| There are 3 distinct steps in standards development that address all 5 phases of the market

Problem identification

- What is the problem to be solved?

Solution research

- Does a solution exist? Is it viable?

Solution specifications

- The mitigations, the provisions

Different deliverable types distinguish the steps:

- Reports and Guides cover the identification and research steps
 - GRs, TRs, EGs
- Technical specifications address the solution step
 - TSs, ESs, ENs

The approval process is the only significant difference between (say) a TS, ES and EN



Design



Development



Deployment



Maintenance

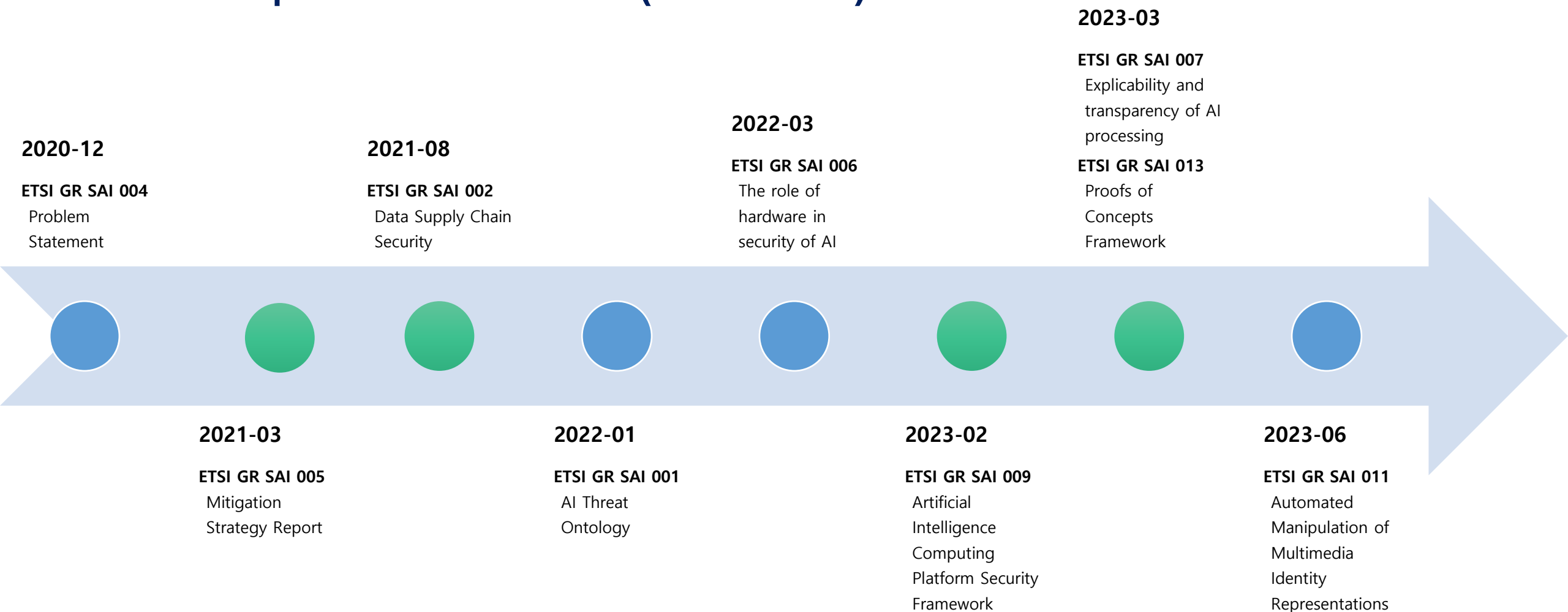


End-of-life

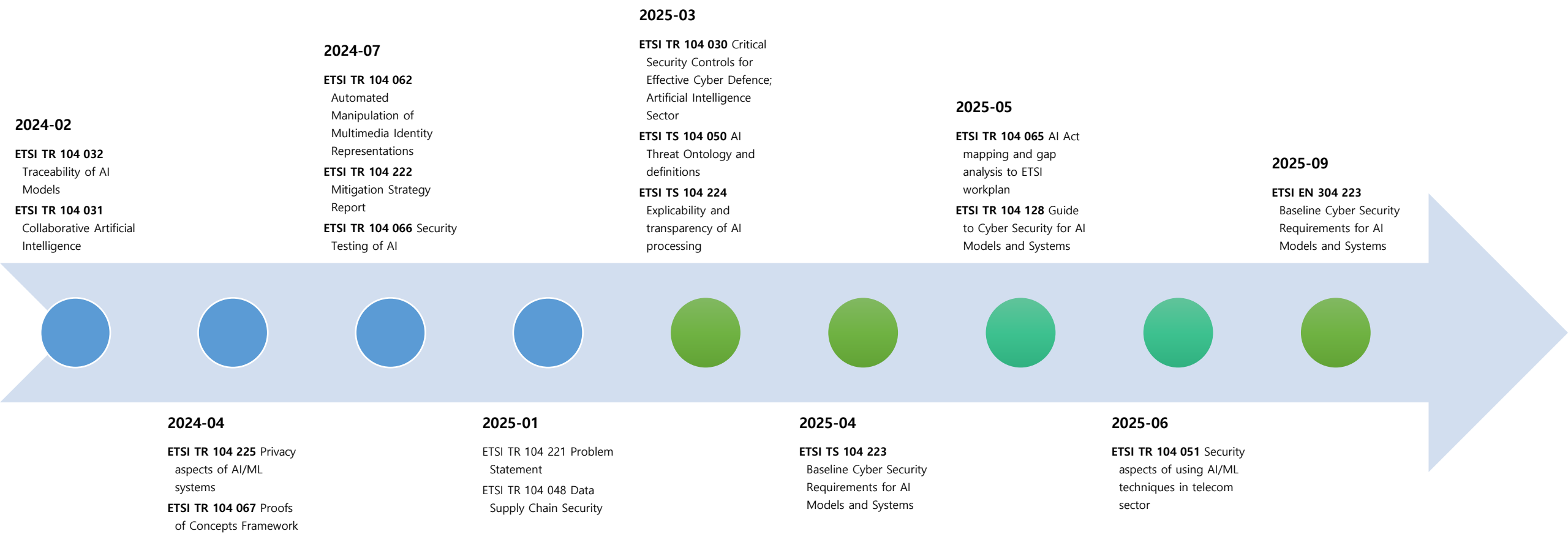
03. What does "S" in SAI mean or signify?

ETSI's Technical Committee Securing Artificial Intelligence (TC SAI) deliberately overloads the "S" in SAI. Whilst it can be first interpreted as Secure in the meaning of cyber-security and the application of the well known Confidentiality-Integrity-Availability (CIA) paradigm, it also addresses safety in the meaning of designing AI systems to be intrinsically safe (i.e. taking a similar view of AI as we do in the domain of Functional Safety in other engineering disciplines), and we also want AI to be societally relevant. This last ensures that we address how AI impacts people, societies, industries and nations. It allows us to take steps to consider the human impact of AI and how we can mitigate negative consequences of its use.

04. ETSI SAI publication timeline (ISG SAI era)



05. The wider standards picture in Europe for AI



06. The wider AI eco-system

| This is captured and updated in ETSI's TR 104 029.

Taken from Annex B: AI security relationship diagrams of TR 104 029 and expanded upon in subsequent slides



07. Collaborative standards actions in ETSI and amongst SDOs

- | Load is shared amongst SDOs and collaboration meetings, liaisons, shared memberships, work to address any potential conflict.

Standards Development Organisations

- ETSI OCG AI; TC SAI; ISG ENI; ISG ZSM
- CEN-CENELEC JTC 21
- ISO/IEC JTC 1 SC 42; IEC SEG 10
- SAC/TC 260 #AI
- ITU-T: 13 future networks; 16 multimedia; 17 Security; 20 IoT/Smart Cities; 2 operations; 15 transport

08. What is the plan from ETSI for AI standardization?

| Establish a baseline for all AI

This is captured in the baseline principles for AI security in [ETSI TS 104 223](#) (and to come EN 304 223):

13 core design principles
A total of 72 trackable provisions
5 lifecycle phases.

ETSI TS 104 223 V1.1.1 (2025-04)



**Securing Artificial Intelligence (SAI);
Baseline Cyber Security Requirements for
AI Models and Systems**



Design



Development



Deployment



Maintenance



End-of-life

09. The next steps in AI standards at ETSI

| Addressing and filling gaps -- identifying problems, developing solutions

- [TR 104 097](#) -- AI and safety -- Functional Safety in AI-dev
- [TS 104 225](#) -- Privacy -- Addressing AI privacy normatively
- [TS 104 158-1](#) (now in 4 parts) -- AI and vulnerability reporting -- AICIE
- [TR 104 159](#) -- AI and social responsibility
- [TS 104 216](#) -- Metrics and testing of AI

10. AI and safety -- Functional Safety in AI-dev

| TR 104 097 -- Guide to the secure by design considerations and assurance in AI systems development

The intent of this work item is to identify the actions of AI developers to ensure that their products and systems remain safe and secure when deployed. This work item will report on the core measures and how they work alongside the core principles of SAI. The intent is to ensure systems are both secure and safe in the deployed context. This reinforces secure by design. To also address how to give assurance of a secure and safe design.

Reinforces the design, deployment and maintenance phases



Design

Development

Deployment

Maintenance

End-of-life

11. AI and Privacy -- Addressing AI privacy normatively

| [TS 104 225](#) -- Privacy aspects of AI/ML systems

The intent of this work is to take the text of TR 104 225 and to extend it to make normative provisions addressing privacy and data protection. This will address aspects including "the right to be forgotten" and how to ensure if data is removed from an AI system that it is not recoverable.

NOTE: this will replace the TR by making normative provisions

Reinforces the design, development and end-of-life phases



Design

Development

Deployment

Maintenance

End-of-life

10. AI and vulnerability reporting -- AICIE

I [TS 104 158-1](#) -- AI Common Incident Expression (AICIE)

The intent of this work is to provide AI Common Incident Expression together with a Common Incident Scoring System that is compatible with lightweight information exchange standards.

Intended to bind AI into the same vulnerability reporting schemas as used in TC CYBER, FIRST, NIST and many others

Reinforces the deployment, maintenance and end-of-life phases



Design



Development



Deployment



Maintenance



End-of-life

10. AI and social responsibility

I [TR 104 159](#) -- Understanding and Preventing Harm from Generative AI

The intent of this work is to provide an understanding of the harm from Generative AI along with presenting the different ways to prevent that harm. This includes but is not limited to malicious code generation, deepfakes, spam messages, disinformation etc. The areas also covered are the issues of AI hallucinations, loss of confidentiality and IPR infringements. The types of methods to counter the harm from Generative AI to be considered in the report are detection, enforcement, reporting and removal.

Reinforces the design, deployment and maintenance phases



Design

Development

Deployment

Maintenance

End-of-life

10. AI and safety -- Metrics and testing of AI

I [TS 104 216](#) -- Conformance assessment for AI (EN 304 223)

In very simple terms the intent of this work is to define methods for assessing conformance to EN 304 223. In practical terms the much wider intent of this work is to develop methods that allow testing to work in order to both aid conformance and to set proper metrics.

This work is supported by many "weights and measures" experts.

Reinforces the design, development and deployment phases



Design

Development

Deployment

Maintenance

End-of-life

GLISC 2025

Global ICT Standards Conference 2025

- Thank you -

Scott CADZOW, Chair ETSI TC SAI

scott@cadzow.com

**ICT Standards and Intellectual Property:
AI for All**