



Forging Successful AI Applications
for European Economy and Society

D2.4 SUCCESS CRITERIA BY PROFESSIONAL ASSOCIATIONS IN THE EU

101177579

Project Information	
Project Number: 101177579	
Project Title: Forging Successful AI Applications for European Economy and Society - FORSEE	
Funding Scheme: HORIZON-CL2-2024-TRANSFORMATIONS-01-06	
Project Start Date: 1 February 2025	

Deliverable Information	
Dissemination Level:	PU - Public



Funded by
the European Union

This project has received funding from the EU's Horizon Europe research and innovation programme under Grant Agreement 101177579. The views and opinions expressed are however those of the author(s) only and do not necessarily reflect those of the European Union or the Agency. Neither the European Union nor the granting authority can be held responsible for them.

Work Package:	WP2 Mapping of social expectations: institutional/formal approaches to success
Document ID:	D2.4 Success Criteria by Professional Associations in the EU
Delivery Date:	31 January 2026
Type:	Report
Version:	V2
Author(s):	Linnet Taylor, Tilburg University Merve Öner Kabadayi, Tilburg University Princy Marimuthu, Tilburg University Delaram Golpayegani, Trinity College Dublin Marta Lasek-Markey, Trinity College Dublin Arjumand Younus, University College Dublin Aphra Kerr, University College Dublin Dave Lewis, Trinity College Dublin
Reviewer(s):	Marguerite Barry, Shana Cohen

History of changes	
Initial draft:	V1 - 15 December 2025
Reviewed version:	V2 – 9 January 2026
Final version:	V3 - 19 January 2026



**Funded by
the European Union**

Executive Summary

Associations of professionals play a crucial role in fostering collaboration, setting standards, and promoting ethical practices within specific industries. This study investigates the viewpoints and criteria set forth by professional bodies to gain a broad understanding of AI success from the legal and professional perspectives. The aim of the report is to discern success criteria that shape understandings of AI among international associations of professionals within the EU area such as the Council of Bars and Law Societies of Europe, European Law Institute, EurHeCa (European Health Professionals Competent Authorities), European Federation of Pharmaceutical Industries and Associations and EASE (European Association of Software Engineering). To discern the relevant baseline success criteria, the body of policies, position papers and guidelines issued by European and international legal, healthcare and engineering institutions in the window of 2018 to 2025 is reviewed. The review is primarily conducted through qualitative thematic analysis approaches. However, given the volume and diversity of the documents examined, thematic analysis is complemented with an unsupervised approach to topic modelling in order to discern new patterns in the texts associated with these international institutional stakeholders.

Across all three professional domains, a key finding is that AI success is framed as *conditional rather than absolute*. Professional bodies do not define success primarily in terms of technical performance, innovation, or efficiency alone. Instead, success is consistently linked to the preservation of human judgment, institutional accountability, and core professional values. AI is overwhelmingly positioned as a *supportive* technology that must remain subject to human oversight, ethical constraints, and established professional responsibilities.

In the legal domain, AI success is predominantly defined through legal legitimacy and rights protection. Legal professional bodies emphasise fairness, non-discrimination, transparency, explainability, procedural safeguards, and compliance with data protection law. AI systems are expected to assist rather than replace judges and lawyers, with strong insistence on human oversight and legally reasoned decision-making. Efficiency gains are acknowledged, but only insofar as they do not undermine fundamental rights, the rule of law, or trust in judicial institutions.

Healthcare professional bodies conceptualise AI success in relation to patient safety, professional autonomy, and trust. The findings show a strong emphasis on governance, risk management, data protection, bias mitigation, and AI literacy. AI is considered successful when it supports clinical decision-making, preserves the patient–doctor relationship, and aligns with medical ethics. Technological innovation is framed as secondary to safeguarding care quality, accountability, and professional responsibility in safety-critical contexts.

Engineering professional bodies exhibit more divergent expectations. Some documents, particularly those with an ethics-oriented orientation, frame AI success in terms of responsibility, harm prevention, transparency, and democratic accountability, positioning AI



Funded by
the European Union

as a technology that requires careful governance and human control. Others, especially application-focused white papers, emphasise efficiency, optimisation, accuracy, and innovation, framing success in operational and performance-based terms. This divergence reflects differences across engineering subfields, levels of public risk, and institutional roles, rather than a unified professional stance.

Interpreted through the lens of the Sociology of Expectations, the success criteria articulated by professional bodies function as performative expectations that actively shape AI governance. Concepts such as transparency, fairness, accountability, and human oversight operate as mechanisms of expectation management, helping to stabilise uncertainty about AI's future role while asserting professional authority. Across sectors, professional associations use these expectations to engage in boundary maintenance and professional mobilisation, positioning themselves as essential intermediaries between AI technologies, regulatory frameworks, and societal values.

Keywords. Professional bodies, legal bodies, engineering bodies, healthcare bodies, AI policies, trustworthy AI guidelines, thematic analysis, NVivo, topic modelling, BERTopic



Funded by
the European Union

Acronyms

AI	Artificial Intelligence
AIM-NET	Artificial Intelligence in Manufacturing NETwork
CCBE	Council of Bars and Law Societies of Europe
CEPEJ	European Commission on the Efficiency of Justice
CPME	Committee of European Doctors
EASE	European Association of Software Engineering
EFPIA	European Federation of Pharmaceutical Industries and Associations
EFCE	European Federation of Chemical Engineers
ELI	European Law Institute
EURHECA	European Health professionals Competent Authorities
EUROJUST	European Union Agency for Criminal Justice Cooperation
FBE	The European Bars Federation / Fédération des Barreaux d'Europe
GPAI	The Global Partnership on Artificial Intelligence
MHE	Mental Health Europe
UIA	Union Internationale des Avocats/International Association of Lawyers (UIA)
UNESCO	United Nations Educational, Scientific and Cultural Organization
WHO	World Health Organisation



**Funded by
the European Union**

Table of Contents

Executive Summary	3
Acronyms	5
1. Introduction	7
2. Background	9
3. Related Work	10
4. Methodology	11
4.1. General Information	11
4.2. Step 1: Manual Thematic Analysis	13
4.3. Step 2: Topic Modelling	14
4.4. Step 3: Interpreting Codes and Topics to Identify Themes	14
5. Corpus of AI Policies and Guidelines Issued by Legal and Professional Bodies	15
6. Main Themes and Findings Identified in NVivo Supported Manual Thematic Analysis	19
6.1. Analysis of the Four Main Themes in the Guidelines Published by Bodies for Legal Professionals	19
6.2. Overview of the Key Themes in the Guidelines Published by Bodies for Legal Professionals	20
6.3. Overview of Key Themes Identified by NVivo Analysis in Guidelines Published by Bodies for Healthcare Professionals	26
6.4. Overview of Key Themes Identified by NVivo Analysis in Guidelines Published by Bodies for Engineering Professionals	30
7. Themes Identified Using BERTopic	33
7.1. Overview of the key themes identified in BERTopic modeling for Legal Organizations	34
7.2. Overview of the key themes identified in BERTopic modeling for Healthcare Organizations	37
7.3. Overview of the key themes identified in BERTopic modeling for Engineering Organizations	39
8. Validation of BERTopic-based Themes through Comparison with NVivo-supported Thematic Analysis	41
8.1. Comparative Analysis for Legal Associations	41
8.2. Comparative Analysis for Healthcare Associations	42
8.3. Comparative Analysis for Engineering Associations	43
9. Key Insights based on Sociology of Expectations	43
9.1. Initial Insights on Empirical Findings	43
9.2. Applying the Theoretical Lense: SoE	44
9.3. Analysis and Discussion	45
10. Conclusion	46
Bibliography	48
Appendix	54



Funded by
the European Union

1. Introduction

This report is part of a four part study that aims to map social expectations present in institutional and formal approaches to asserting success criteria for Artificial Intelligence (AI). Specifically, this study investigates the viewpoints and criteria set forth by supranational bodies, technical organisations, academic institutions, and professional bodies in defining successful AI applications. Supranational bodies, technical organisations, academic institutions, and professional bodies collectively play a pivotal role in shaping perceptions of success in AI by establishing standards, guidelines, and ethical frameworks that influence the development, deployment, and evaluation of AI applications. These bodies serve as crucial actors therefore in shaping what comprises formal approaches to AI success.

The part of the study reported in this document is focused on the expression of AI success criteria offered in documents developed by Professional Bodies (legal, healthcare and engineering) addressing issues of AI governance. The types of Professional Bodies studied consist of parties that represent, or who are selected to be representative of, multiple different countries in developing guidelines or rules for the governance of AI.

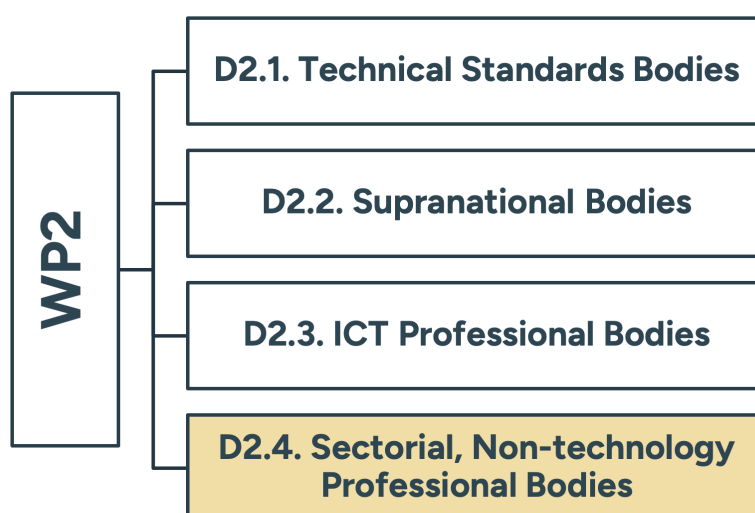


Figure 1 - Outline of the Deliverables Comprising WP2

For context, the other three parts of the study address respectively AI governance outputs from:

- **Supranational Bodies:** AI policies and guidelines from prominent international bodies, including the UN, OECD, G 20, and G 7
- **Technical Standards Bodies:** representing international and national bodies developing technical and process standards for AI.
- **ICT Professional Bodies:** While the field of ICT does not require a professional affiliation to practice, bodies representing ICT professionals internationally that have



Funded by
the European Union

engaged in guidelines or rules addressing AI governance as part of their broader remit to guide professional practice.

Overall, this study aims to discern the set of baseline criteria for AI success produced by these different classes of international institutions. This study does not aim to draw any comparison between individual countries. International perspectives are selected from bodies that can claim wide global representation as well as bodies with international representation within Europe.

The primary method of the study is thematic analysis. However, given the volume and diversity of the documents examined, this thematic analysis is complemented with an unsupervised approach to topic modelling in order to discern new patterns in the texts associated with these international institutional stakeholders. One way of addressing the research question is through the use of qualitative thematic analysis approaches (Braun & Clarke, 2021b, 2021a; Terry et al., 2017a) from social sciences. Since thematic analysis is conducted manually by domain experts, it can yield nuanced and high-quality insights, however it is inherently subjective, time-intensive, and difficult to scale. To mitigate these limitations, we adopt unsupervised topic modelling techniques from natural language processing (NLP), which enable identification of frequent topics in large text corpora, in combination with thematic analysis to interpret, validate, and provide context to the discovered topics.

The study seeks to identify and categorise recurring themes, priorities, and nuanced dimensions of success criteria outlined in the documents. It aims therefore to shed light on the evolving perspectives and priorities that shape institutional discourse on the responsible development and deployment of AI.

This study selects documents in the time window of 2018 to 2025. This represents the period to date in which international consideration of AI governance has resulted in institutional consensus, while also being of sufficient duration to discern the shifts in the criteria set forth by these influential entities.

This report specifically aims to uncover common topics and themes, both shared and divergent, within the **AI policy and guidelines issued by professional bodies**. The research question this report explores is: ***What are the prevalent themes in existing AI policies and guidelines published by associations of professionals within the EU area in 2018-2025?***

Addressing the research question, we group the guidelines and documents related to AI published by professional bodies under investigation into three sets:

- **Associations of Legal Professionals:** This category includes guidelines, white papers and positions papers issued by bar associations and associations of legal professionals such as Council of Bars and Law Societies of Europe (CCBE) and European Commission on the Efficiency of Justice (CEPEJ).
- **Associations of Healthcare Professionals:** This category includes guidelines, white papers and positions papers issued by associations of healthcare professionals such as Committee of European Doctors (CPME).



Funded by
the European Union

- **Associations of Engineers:** This category includes guidelines, white papers and positions papers issued by engineers associations in the EU.

After identifying common themes and topics from these three sets of documents, we analyse AI success factors articulated by supranational bodies through the lens of sociology of expectations.

The remainder of this report is structured as follows. In [section 2](#), we provide essential background information and key definitions. In [section 3](#), we review the relevant literature on analysis of AI usage by professional bodies, specifically legal professionals. [Section 4](#) describes our novel methodology to identify prevalent themes within AI documents through combination of topic modelling and thematic analysis approaches. [Section 5](#) presents our corpus of AI policies and guidelines issued by supranational bodies from 2018 to 2025. We present the results of our manual thematic analysis conducted on a subset of the documents in [section 6](#). [Section 7](#) details our findings from automated machine learning-based algorithms complemented with expert analysis. In [section 8](#), we validate the findings from the analysis by automated machine learning algorithms through comparison with the findings from our manual thematic analysis. In [section 9](#), we provide interpretation of the results through the sociology of expectations framework and we conclude the report in [section 10](#).

2. Background

This section establishes key definitions and provides necessary background context by clarifying the terminology used throughout this report.

Sociology of Expectations is a theoretical perspective emerging from science and technology studies in the early 2000s (Brown et al., 2003; van Lente, 2012) and informed by the Social Construction of Technology paradigm (Pinch & Bijker, 1984). This approach suggests that social expectations about technological development have the capacity to shape the direction of research, investment decisions and public discourse and can operate at micro, meso and macro levels. Social expectations can be positive or negative, they can be contested, and they can be performative (discursive) or have real impact and shape actions in the world (Kerr et al., 2020a). This perspective is discussed in more detail in section 9.

Governance as a term can include multiple forms of regulation (e.g. statutory, co-regulation and self-regulation) that operate at multiple levels, from local to regional to global. Thus governance is often taken to include vertical (state, regional, supranational) and horizontal actors (private sector). In many media and communication fields there has been a move from top down government and statutory regulation to co-regulation with a variety of non-state stakeholders over the past decades, particularly in the EU since the 2000s (Puppis, 2024). In social media and some information technology fields there has been a move from self-regulation to co-regulation. Accordingly our sample of supranational



Funded by
the European Union

documents include statutory documents, but also codes of practice, guidelines and recommendations.

Thematic Analysis is a popular method of qualitative data analysis that systematically organises datasets, and helps to identify patterns of meaning commonly referred to as themes (Squires, 2023). First described in the 1970s by Houlton – albeit the term itself had been in use even earlier – it became more prominent in the late 1990s with researchers like Boyatzis and Hayes (Terry et al., 2017b). In recent times, thematic analysis has been understood as an umbrella term for different approaches. While popular in qualitative interview data analysis, computer-assisted expert thematic analysis of legal texts, such as legislation or policy documents, appears to be less commonly employed.

Theme a broad interpretable semantic concept derived through expert interpretation, either from the topics generated by a topic modeller or directly from the text itself (**human-driven**).

Topic modelling refers to identification of prevalent topics in corpora of text using unsupervised machine learning techniques and is primarily used for uncovering topics within large sets of documents. There exist a variety of unsupervised topic models (see Churchill & Singh, 2022) for evolution of such models). Currently, BERTopic (Grootendorst, 2022) and Latent Dirichlet Allocation (LDA) (Blei et al., 2003) are the most popular topic models.

Topic a concept identified automatically by a topic modeller from text (algorithm-driven).

3. Related Work

There is a growing body of academic work examining AI governance, ethics, and policy guidelines, with most existing studies focusing on documents produced by supranational organisations, governments, and private actors. Far less attention has been paid to AI-related guidelines issued by professional and legal associations, despite their central role in shaping norms of practice, professional responsibility, and sector-specific standards. Where professional bodies are mentioned, they are often treated as peripheral stakeholders rather than independent sources of governance (see e.g. Jobin et al., 2019). This literature review therefore focuses on studies that analyse AI-related principles and guidance with relevance to professional practice, highlighting the methodological approaches used and identifying gaps that motivate the present study.

Early thematic reviews of AI ethics and governance, such as (Jobin et al., 2019) and (Hagendorff, 2020), provide foundational mappings of ethical principles including transparency, accountability, fairness, privacy, and human oversight. While these studies primarily analyse guidelines produced by governments, international organisations, and corporations, they frequently reference professional responsibility and human control as core concerns. However, professional bodies are not analysed as a distinct category, and the implications of these principles for professional practice are largely inferred rather than examined directly (see e.g. Corrêa et al., 2023)



Funded by
the European Union

More structured and systematic analyses, such as (Fjeld et al., 2020) and (Palladino, 2023), further refine these ethical principles into clusters or hierarchies that implicitly resonate with professional domains. Themes such as professional responsibility, safety, and accountability suggest relevance for professions such as law, healthcare, and engineering, yet these studies remain focused on high-level governance frameworks rather than on profession-specific interpretations. Similarly, (Birhane et al., 2024) reframe accountability as a socio-technical and institutional practice, emphasising how oversight mechanisms are embedded within power relations. Their analysis is highly relevant to professional bodies, but again does not examine professional associations as primary authors of AI guidance.

A smaller subset of studies engages more directly with sectoral or professional dimensions of AI governance. (Reuel et al., 2025), for example, shifts attention from abstract principles to concrete institutional capacities such as assessment, verification, and monitoring, implicitly highlighting the role of professional expertise in operationalising AI governance. Research on AI in law (Surden, 2018), medicine (Reddy et al., 2020), and engineering (Lu et al., 2024) similarly identifies tensions between automation and professional judgment, often emphasising human oversight, explainability, and trust as themes for AI adoption. However, these studies tend to focus on socio-technical impacts on professions rather than on the normative guidance produced by professional associations themselves.

Alongside manual thematic reviews, several studies employ automated text analysis techniques (such as topic modelling using LDA or BERT) to analyse large corpora of AI policies and ethical guidelines. Quantitative analyses, including (Roche et al., 2023), demonstrate the prominence of human-centric and rights-based language in European AI discourse, reinforcing the centrality of legal and societal considerations. While these methods are effective in identifying dominant discursive patterns, they are typically applied to broad policy corpora and rarely disaggregate professional or sector-specific documents.

Against this background, the present study addresses a clear gap in the literature by focusing explicitly on **AI guidelines published by legal, healthcare, and engineering professional associations**. Unlike supranational or governmental actors, these organisations operate at the intersection of regulation, practice, and professional identity. Their guidelines therefore offer unique insight into how AI success criteria are defined in relation to human judgment, ethical responsibility, institutional accountability, and sector-specific risk. By combining expert-led NVivo thematic analysis with BERTopic modelling, this report builds on existing methodological approaches while extending them into a comparatively underexplored domain of AI governance.

4. Methodology

4.1. General Information

The terms “**code**”, “**topic**”, “**theme**” are used extensively throughout this report. Although they may appear as synonyms, they are treated as distinct concepts and are not used interchangeably. To clarify this distinction, in the following, we provide background



Funded by
the European Union

information about how we combine **thematic analysis** and **topic modelling** in this deliverable.

In this deliverable, we consider a **document** as a formally issued textual artefact published by a legal, healthcare or engineering organisation on AI. Such documents can range from official guidelines to white papers and are collectively compiled within a corpus (see section 5).

In our methodology, we first apply **thematic analysis**, which is a popular method of qualitative data analysis that systematically organises datasets, helps to identify patterns of meaning (Squires, 2023). First described in the 1970s by Houlton – albeit the term itself had been in use even earlier – it became more prominent in the late 1990s with researchers like Boyatzis and Hayes (Braun & Clarke, 2021; Terry et al., 2017). In recent times, thematic analysis has been understood as an umbrella term for different approaches ranging from reflexive thematic analysis to codebook and coding reliability approaches. No coding reliability tests were used in this project although the combining of thematic analysis and topic modelling required negotiation between a fully qualitative and a more quantitative analysis of the data. The disciplinary backgrounds of the team and their subjectivities were crucial however to the initial manual coding, the fine tuning of the topic modelling and the final thematic analysis.

In our work, we use manual thematic analysis to inductively code a sample of our corpus of documents. The output of this process is a set of codes. Code refers to a concept that is manually identified as being both present and important within the domain. This form of thematic analysis relies on the coder's judgements and expertise. It produces contextually informed codes grounded in the documents..

However, manual thematic analysis is a time-consuming and resource-intensive process which makes it difficult to conduct thematic analysis of a large corpus of documents. Further, it does not provide insights on statistical frequency of the codes.

Addressing these specificities, we also employ **topic modelling**, which is a statistical method used to identify prevalent topics in large corpora of text. Topic modelling uses unsupervised machine learning techniques for uncovering topics within large sets of documents. There are different unsupervised topic models (see Churchill & Singh, 2022) for evolution of such models). Currently, BERTopic (Grootendorst, 2022) and Latent Dirichlet Allocation (LDA) (Blei et al., 2003) are the most popular topic models. The output of these models are clusters of keywords representing topics. **Topic** is a concept that has been identified by a computational model (e.g., LDA or BERTopic) as being present in the corpus consisting of multiple documents. In both LDA and BERTopic, a topic can be understood as a latent concept associated with a subset of documents (in our work subset of sentences within a document as explained later) and typically summarised by salient words that tend to appear when that concept is present. It should, however, be noted that



Funded by
the European Union

the underlying mathematical object defining a topic differs across LDA and BERTopic (a probability distribution over words in LDA versus a cluster of document embeddings in BERTopic).

Though topic modelling is usually applied to a set of documents, we apply it to each document separately to gain insights into each single document. Therefore, in this work topics are associated with a subset of sentences within a document.

Finally, to gain insights, we interpret the codes and topics through the lens of sociology of expectation to identify **themes**. **Theme** is a broad interpretable semantic concept derived through expert interpretation, either from the topics generated by a topic modeller or directly from the codes generated during the thematic analysis process. By illustrating the overall methodology wherein thematic analysis and topic modelling are integrated, Figure 2 shows the key concepts in the methodology and how they are related to each other.

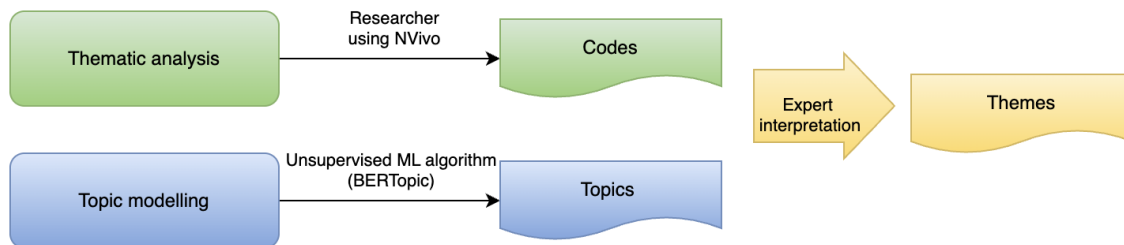


Figure 2: Overall methodology that combines thematic analysis and topic modelling to identify themes

The methodology comprises of three steps

Step 1 : Conducting thematic analysis simultaneously using Nvivo manually for the chosen documents

Step 2: Generating topics through BERTopic modeling

Step 3 : Assigning prevalent themes based on the results of the BERTopic modelling and checking for overlap between the results of Nvivo and BERTopic modelling.

4.2. Step 1: Manual Thematic Analysis

In Step 1 of the methodology, qualitative thematic analysis and quantitative topic modelling were conducted in parallel. The qualitative analysis of a subset of the corpus, including the CCBE “Considerations on the Legal Aspects of AI”, UNESCO “Draft guidelines for the use of AI in Courts and Tribunals”, CPME “Deployment of artificial intelligence in healthcare” and EFCE “White Paper on Artificial Intelligence in Chemical Engineering” were conducted with the use of NVivo software (see [section 6](#) for the findings). As explained by (Dhakal, 2022), NVivo is a CAQDAS programme that assists, rather than replaces, a human researcher. NVivo analysis is, thus, expert-led and in this project it was performed by a legal researcher.



Funded by
the European Union

Thematic analysis through NVivo is conducted by importing documents and the thematic coding is done by selecting a section of the text, assigning it to a code with a theme that is most relevant to the selected text. The assigning of the codes is done by a legal researcher. The nodes were assigned manually and were not preassigned. As a result, the analysis of the legal bodies corpus resulted in 103 codes, the engineering bodies corpus resulted in 92 codes and the healthcare bodies corpus resulted in 57 codes with several codes overlapping across all the documents. Based on these codes, 11 main themes were generated.

4.3. Step 2: Topic Modelling

Step 2 comprises three following phases:

1. **Text preprocessing:** Before application of BERTopic, we employed common NLP approaches to apply minimal preprocessing using the nltk library for:
 - Lowercasing,
 - Removing non-alphabetic tokens (e.g. punctuation marks),
 - Removing tokens that consist solely of numbers, and
 - Removing tokens shorter than two characters.

Note that we did not remove stop words or perform lemmatisation prior to embedding, given that BERTopic uses transformer-based embeddings, which require the text to be persevered for better accuracy. We removed the stop words, after generating embeddings, using the CountVectorizer when initialising BERTopic.

2. **Fine-tuning:** We finetuned BERTopic model in an iterative manner to ensure that the model identifies at least 3 topic categories for each document. The snippets used for preprocessing and topic modelling are available on GitHub under the MIT licence at https://github.com/DelaramGlp/forsee_topicmodelling.
3. **Topic interpretation and theme discovery:** This phase involved qualitative refinement of the BERTopic results. A student assistant independently analysed the generated topics and assigned themes to each topic category. Then a legal expert went through the assigned themes to determine their applicability. The proposed themes were then discussed in a joint session and consolidated into their final form.

4.4. Step 3: Interpreting Codes and Topics to Identify Themes

In Step 3, the results from the themes of Nvivo analysis and the themes generated from the topic modeller were organised, compared and discussed. The finalized themes were determined based on the comparison of both analyses.

Following inductive manual coding and automated topic analysis the research team went through a process of iterative interpretation aimed at generating, refining and naming themes. This process of meaning making involved both colour coding of topics and codes, clustering, referring back to the original texts and a structured analysis. In line with prior work on the Sociology of Expectations we also categorised the themes according to whether the themes primarily referred to actions that might be taken by actors at a micro, meso and macro levels. The multidisciplinary team then reviewed, discussed and refined the themes and clustered them into meta-themes guided by the overall research question, our

disciplinary backgrounds and our grounding in the data and the prior literature.

5. Corpus of AI Policies and Guidelines Issued by Legal and Professional Bodies

The Analysis draws upon a corpus of 22 documents collected from professional organisations in the EU that have released a guideline, position paper or framework surrounding the use of AI in their respective organisations.

The organisations chosen were of legal, healthcare and engineering fields. We collected documents based on the following criteria. The first being that AI should be mentioned in the title of the document, secondly, the guidelines/organisations should be applicable to the whole of the EU or the EU member states should be subject to those guidelines if published by an international organisation. The legal organizations were selected via a list published by the Conseil National Des Barreaux (*European Organizations / Site Name*, n.d.). This list includes European organisations, institutions, bar associations and trade unions. The engineering and healthcare organisations were chosen with the same criteria as the legal organisations, examples of these organizations include the European Chemical Engineering association and the committee of European Doctors.

Prominent healthcare and engineering organizations mentioned in the Grant Agreement such as EurHeCa and EASE were also taken into consideration for document research, that said, even though their publications were thoroughly researched, we did not come across any documents or guidelines published by these bodies with regard to AI deployment, utilization or design. Therefore, our corpus does not include any documents by these bodies. It is also important to note that the guidelines that are considered within the scope of this report are meant to be published by the associations for the “healthcare professionals” (e.g. doctors, nurses and producers), legal professionals (e.g. bar associations) and engineers. In other words, guidelines demonstrating technical standards for AI were not included in the corpus.

The second criteria was that the documents had to be published within the years of 2018 to 2025, since this is the scope of this research.

The key words used to search for the documents were the following:

Table 1 - Keywords used to search for the documents

Technical / Technological Terms	Professional Terms
1. AI / Artificial Intelligence 2. AI System 3. Generative AI/Gen AI 4. Machine Learning	1. Legal association 2. Lawyer association 3. Bar association 4. Legal professional 5. Healthcare 6. Healthcare professional 7. Healthcare association



Funded by
the European Union

	8. Doctor/Nurse 9. Health association 10. Engineer(ing) 11. Mechanical Engineer(ing) 12. Software Engineer 13. Engineer Association
Document Terms	Geographic Scope
1. Guideline 2. Paper 3. Position paper 4. White paper 5. Standard 6. Framework	1. Europe 2. European Union 3. EU

The tables below demonstrate the guidelines and documents by the legal and professional bodies included in the corpus of this report.

Table 2 - Guidelines for AI from Associations for Legal Professionals included in the corpus

ID	Document	Issuer	Type	Year
PB_L_01	CEPEJ European Ethical Charter on the use of Artificial Intelligence in judicial systems and their environment	European Commission for the Efficiency of Justice (CEPEJ)	Charter	2018
PB_L_02	CCBE Comments on the Stakeholders Consultation on Draft Artificial Intelligence Ethics Guidelines	Council of Bars and Law Societies of Europe (CCBE)	Guideline	2019
PB_L_03	CCBE Response to the consultation on the European Commission's White Paper on Artificial Intelligence	Council of Bars and Law Societies of Europe (CCBE)	White Paper	2020
PB_L_04	CCBE considerations on the Legal Aspects of AI	Council of Bars and Law Societies of Europe	Guideline	2020
PB_L_05	CEPEJ Revised roadmap for ensuring an appropriate follow-up of the CEPEJ Ethical Charter on the use of artificial intelligence in judicial systems and their environment	European Commission on the Efficiency of Justice (CEPEJ)	Guideline	2021
PB_L_06	CCBE Guide on the use of Artificial	Council of Bars and	Guideline	2022



**Funded by
the European Union**

ID	Document	Issuer	Type	Year
	Intelligence-based tools by lawyers and law firms in the EU	Law Societies of Europe		
PB_L_07	EUROJUST Artificial intelligence supporting cross-border cooperation in criminal justice	European Union Agency for Criminal Justice Cooperation (EUROJUST)	Guideline	2022
PB_L_08	CEPEJ Assessment Tool for the Operationalisation of the European Ethical Charter on the Use of Artificial Intelligence in Judicial Systems and Their Environment	European Commission on the Efficiency of Justice (CEPEJ)	Assessment tool	2023
PB_L_09	EUROJUST Generative Artificial Intelligence - The impact on intellectual property crimes	European Union Agency for Criminal Justice Cooperation (EUROJUST)	Guideline	2023
PB_L_10	ELI GPAI Principles on Co-Generated Data	European Law Institute	Guideline	2024
PB_L_11	EDPS Generative AI and the EUDPR, First EDPS Orientations for ensuring data protection compliance when using Generative AI systems	European Data Protection supervisor (EDPS)	Guideline	2024
PB_L_12	FBE (New Technologies Commission) Guidelines on How Lawyers Should Take Advantage of the Opportunities Offered by Large Language Models and Generative AI.'	The European Bars Federation / Fédération des Barreaux d'Europe	Guideline	2024
PB_L_13	UIA Guidelines on the use of artificial intelligence systems by lawyers	Union Internationale des Avocats/International Association of Lawyers (UIA)	Guideline	2024
PB_L_14	UNESCO Draft guidelines for the use of AI in Courts and Tribunals	UNESCO	Guideline	2025



Funded by
the European Union

Table 3 - Guidelines for AI by Associations of Healthcare Professionals included in the corpus

ID	Document	Issuer	Type	Year
PB_H_01	EFPIA Position on the use of artificial intelligence in the medicinal product lifecycle	European Federation of Pharmaceutical Industries and Associations	Recommendation	2024
PB_H_02	MHE Artificial Intelligence in Mental Health	Mental Health Europe (MHE)	Guiding principles	2024
PB_H_03	EMA Reflection paper on the use of Artificial Intelligence (AI) in the medicinal product lifecycle	European Medicine Agency (EMA)	Reflection paper	2024
PB_H_04	CPME Deployment of artificial intelligence in healthcare	Committee of European Doctors (CPME)	Guideline	2024
PB_H_05	CBDIO Application of AI in healthcare and its impact on the 'patient-doctor' relationship	Steering Committee for Human Rights in the fields of Biomedicine and Health (CBDIO)	Guideline	2024
PB_H_06	Health data governance in the age of artificial intelligence: policy imperatives for the WHO European Region	WHO Regional Office for Europe	Guideline	2025

Table 4 - Guidelines for AI by Engineering organisations included in the corpus

ID	Document	Issuer	Type	Year
PB_E_01	Nordic engineers' stand on Artificial Intelligence and Ethics	Association of Nordic Engineers	Guideline	2022
PB_E_02	AIM-NET White paper on Artificial Intelligence in Manufacturing	Artificial Intelligence in Manufacturing NETwork	White Paper	2023
PB_E_03	EFCE White Paper on Artificial Intelligence in Chemical Engineering	European Federation of Chemical Engineers	White Paper	2025
PB_E_04	Aerospace Formulating an Engineering Framework for Future AI Certification in Aviation	Aerospace	Guideline	2025



Funded by
the European Union

6. Main Themes and Findings Identified in NVivo Supported Manual Thematic Analysis

6.1. Analysis of the Four Main Themes in the Guidelines Published by Bodies for Legal Professionals

Based on the Nvivo supported thematic analysis, we firstly identified four general themes for AI to be considered successful by the legal associations: (i) Human centric and ethical AI, (ii) Socially functioning integration of AI, (iii) Legitimate and rights- oriented AI and (iv) Operationally efficient and responsible AI. This section will first explain these four main themes and then demonstrate the findings of the NVivo supported manual thematic analysis. In this regard, our analysis regarding these four main themes identified are as follows:

“Human-centric and ethical AI” is achieved when the AI systems are designed in a way that respects ethical considerations and takes human interests to its core. Based on the reviewed documents, we have identified a pattern that the design and implementation of AI tools should meet with 5 conditions in order to be considered ethical. These conditions are: (i) respect for fundamental rights, (ii) non-discrimination, (iii) quality and security, (iv) transparency, (which also requires explainability), (v) fairness and (vi) human control, (which also requires AI literacy).

The theme of “socially functioning integration of AI” is achieved when the utilization of AI tools by legal professionals positively affects the society as a whole. One of the indicators of societal success is the enhanced efficiency of judicial decision-making through the integration of AI technologies in the legal system. The other indicator is the ethical design of the said AI systems to ensure that the rights of the individuals are protected while increasing efficiency. This theme, although similar, differs from human-centric and ethical AI since the latter is focused more on ethics, accountability and transparency, whereas the former is concerned with public trust. This is supported by UNESCO draft guidelines *“However, the misuse of AI systems may undermine society’s trust in the judicial system.” (UNESCO, 2025)*

We have classified the theme of “legitimate and rights oriented deployment of AI” based on the discussions in the guidelines mentioning the usage of AI tools in the at different stages/parts of the judicial system including but not limited to drafting case briefs, summarizing and criminal investigations. There are certain themes mentioned under this category. The first is “supportive AI”. It is emphasized in the reviewed documents that the AI systems should not replace or undermine the judges’ decision-making power in the legal proceedings. This can be exemplified by these indications in the text *“In these cases, AI should play a supportive role only. This again underlines the need to more carefully consider the topic of the use of AI in the justice field.” (CCBE, 2019)*. “Legal reasoning” is another theme which demonstrates that any output generated by AI tools should be



Funded by
the European Union

subject to legal reasoning. Moreover, the “human oversight” theme suggests that the output generated by AI tools should be subject to critical review of lawyers. Compliance with the GDPR is another theme that the documents frequently mention which suggests that the GDPR principles should be followed during the deployment and utilization of AI tools in legal proceedings.

Based on the reviewed documents, we infer that operationally efficient and responsible adoption of AI is related to the criteria for the successful deployment of AI in the legal proceedings by humans. Our analysis based on the guidelines indicates that, for the successful implementation of AI tools in legal proceedings, AI literacy and education, establishing standards and ethical use are important. Accordingly, legal professionals are encouraged to lay down standards, follow privacy and data protection principles, establish human oversight while using AI as a supportive tool and ensure that the activities carried out by AI tools are traceable (i.e. transparency).

6.2. Overview of the Key Themes in the Guidelines Published by Bodies for Legal Professionals

The main themes identified from Nvivo manual thematic analyses for each of the legal documents are presented in the following table. This table shows the documents with the most amount of themes coded and are guidelines published from important legal organisations.

Table 5 - Overview of the key themes identified in Nvivo-supported manual thematic analysis of Associations of Legal Professionals

ID	Document	No. of themes	Key themes identified through NVivo-supported manual thematic analysis	Year
PB_L_01	CEPEJ European ethical Charter on the use of Artificial Intelligence in judicial systems and their environment	67	Legal reasoning, Machine learning, transparency, Use of AI by lawyers, Big Data, Predictive Justice, Bias and Discrimination, Human oversight	2018
PB_L_04	CCBE considerations on the Legal Aspects of AI	53	Use of AI in Courts, Operational Success, Use of AI by Lawyers, Risk mitigation, Supportive AI, human oversight, Ethical AI and Human Rights	2020
PB_L_05	CCBE Guide on the use of Artificial Intelligence-based tools by lawyers and law firms in the EU	51	Functionality of AI tools, Use of AI by lawyers, Challenges, AI literacy and education, Risk mitigation, Operational Success, Use of AI in Law firms, Obligations for lawyers and Explainability of AI systems.	2022



Funded by
the European Union

PB_L_10	UIA Guidelines on the use of AI by lawyers	21	Use of AI by lawyers, Transparency, Ethical AI, Operational success, AI literacy and Education, Usability and Confidentiality	2024
PB_L_11	UNESCO Draft guidelines for the use of AI in Courts and Tribunals	44	Use of AI in Courts, Usability, AI literacy and education, Challenges, Accountability, Fairness and Operational Success.	2025

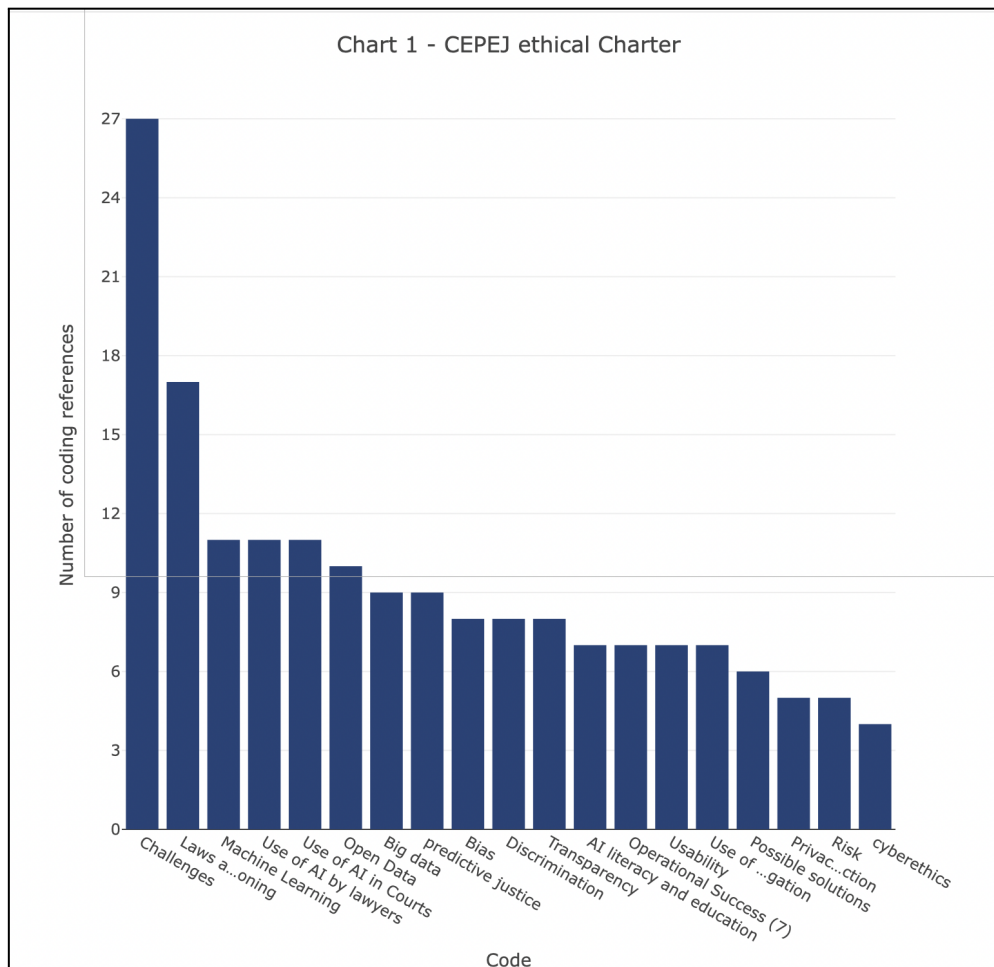


Figure 1: Nvivo Supported results for the CEPEJ ethical Charter

Figure 1 above shows the key themes in the CCBE ethical charter. The most prominent code is “challenges”, suggesting that the Charter takes a risk and challenge based approach for the deployment of AI. This is followed by “Laws and Legal reasoning” both of which indicate the contexts governing the deployment of AI. Overall, the distribution of codes suggests that the CEPEJ Ethical charter prioritises structural challenges and regulatory considerations, while also taking into account risks and ethical principles.



Funded by
the European Union

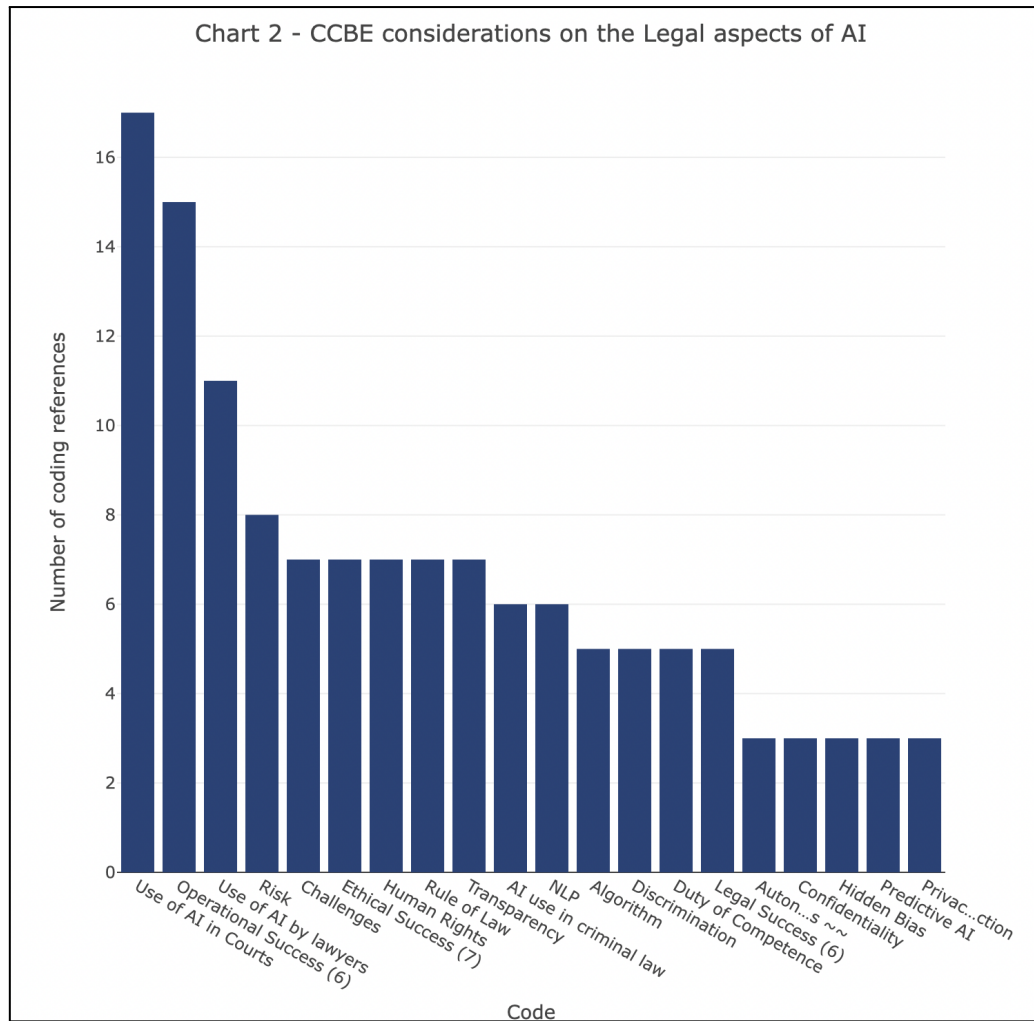


Figure 2: Nvivo Supported results for the CCBE considerations on the Legal Aspects of AI

Figure 2 above shows the themes for CCBE considerations on the legal aspects of AI. While percentage-wise the most prominent codes is “use of AI in courts”; codes such as “human rights”, “transparency”, “rule of law”, “challenges” and “risk” suggest that CCBE takes a precautionary, human-centric and risk-based approach when it comes to assessing the success of AI technologies.



**Funded by
the European Union**

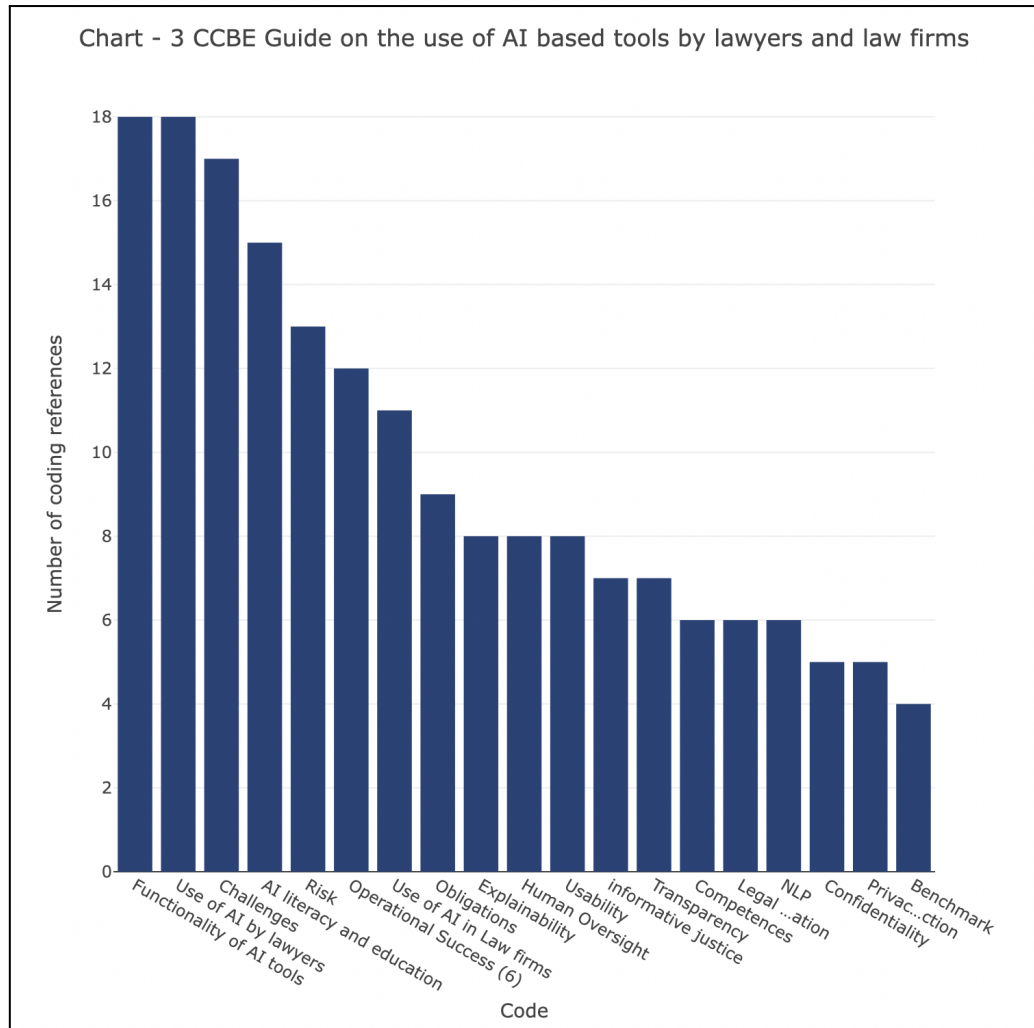


Figure 3: Nvivo Supported results for the CCBE Guide on the use of AI by lawyers and law firms

Figure 3 above presents the codes in the CCBE Guide on the Use of AI by Lawyers and Law Firms". The theme with the highest percentage coverage is "challenges". This demonstrates that the guide places considerable emphasis on identifying and managing the risks and challenges associated with the adoption of AI by lawyers and law firms. Use of AI by lawyers and functionality of AI tools highlight a strong focus on how AI technologies operate in practice. This guideline reflects a practical orientation towards assisting lawyers and law firms in navigating the responsible adoption of AI technologies.



**Funded by
the European Union**

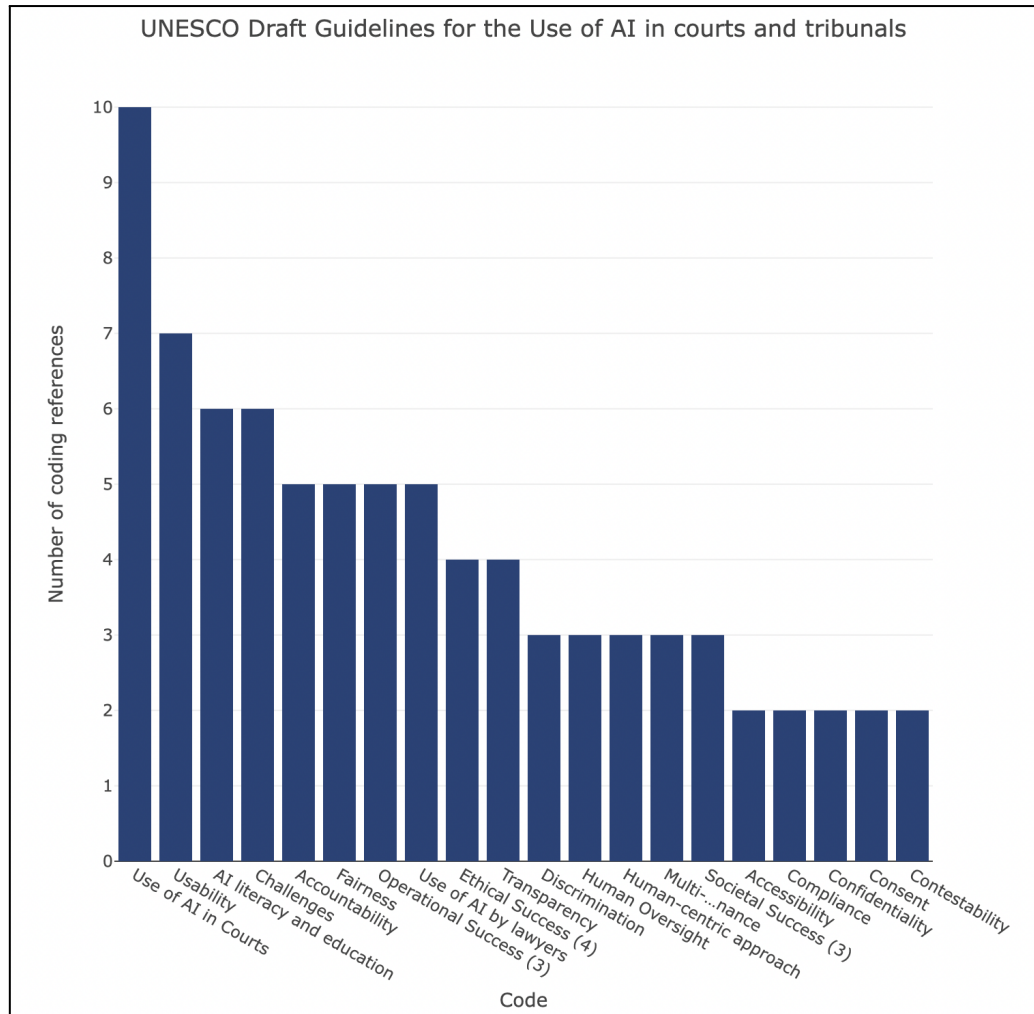


Figure 4: Nvivo Supported results for UNESCO draft guidelines for the use of AI in courts and tribunals

Figure 4 illustrates the percentage coverage of codes in the UNESCO draft guidelines for the Use of AI in Courts and Tribunals. The most prominent theme is the “use of AI in courts”. This indicates that the UNESCO draft guidelines are strongly oriented toward the institutional deployment of AI within judicial settings. The second and third most significant codes are “operational success” followed by “Usability”, which suggests a practical concern with ensuring that AI systems function effectively. The distribution of codes indicates that the UNESCO draft guidelines prioritise practical use of AI in courts.



**Funded by
the European Union**

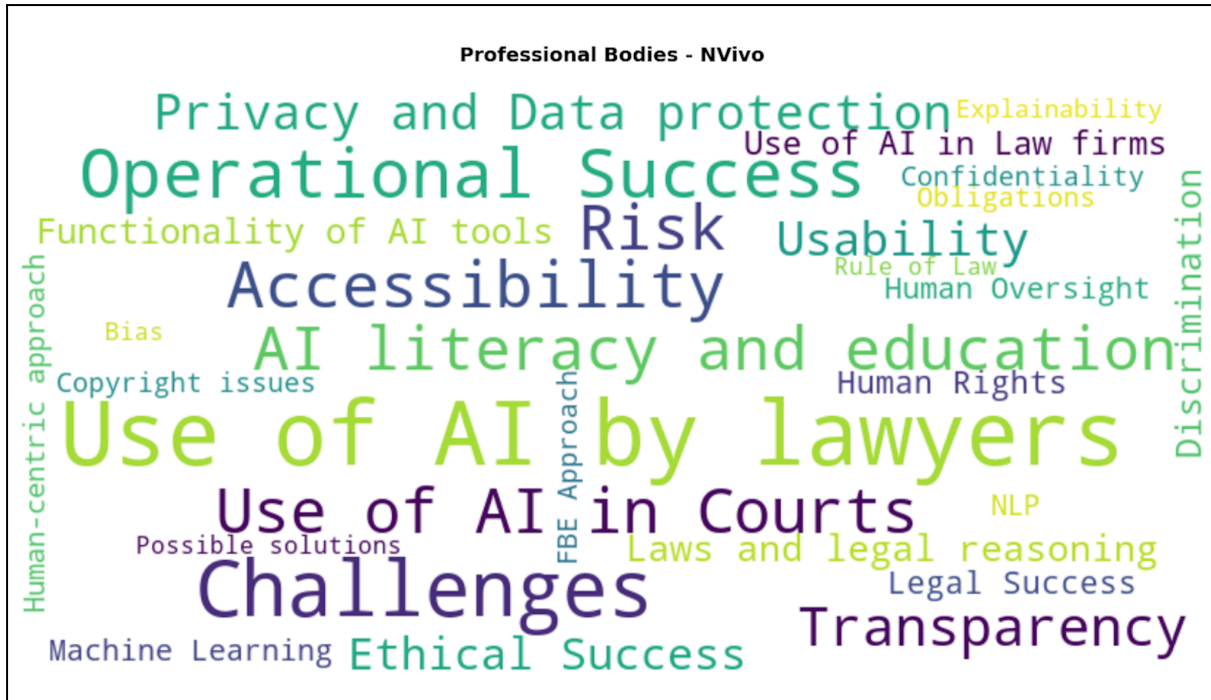


Figure 5: Word cloud of guidelines published by bodies for legal professionals

In the four NVivo charts presented above, the word cloud and the full list of themes identified in the NVivo analysis (See [Appendix](#)) were analyzed. We observed shared priorities and different points of focus across the legal professional guidelines. Across all five charts, these themes recur consistently which create a baseline for the use of AI in the legal domain. These include

- Use of AI in Courts or by Lawyers
- Operational success and Usability
- Challenges and risks
- Fairness, non-discrimination (impartiality) and transparency
- AI literacy and Education

The presence of these themes across all documents shows the shared baseline criteria for successful implementation of AI in the legal domain. They suggest that AI in legal practice must be functional, accountable and be aligned with the fundamental principles of the EU.” AI therefore is not to be treated as a technical tool but a system that requires governance and human oversight. This is supported by UNESCO as follows: *“Human critical thinking must be central in all interactions with generative AI systems”* (UNESCO, 2025). Furthermore, the documents distinguish between the use of AI in two ways, use by lawyers as individuals and use by courts as institutions. The CEPEJ and UNESCO guidelines place a significant emphasis on the institutional deployment of AI, specifically in courts and indicate that ethical principles such as fairness and impartiality (non-discrimination) should serve as a foundation for the courts to rely on (CEPEJ, 2018; UNESCO, 2025). In comparison, the guidelines by CCBE and UIA focus on lawyers as professionals and their use of AI. They discuss the responsibilities of lawyers and competences, transparency in



Funded by
the European Union

using AI tools for daily legal work. This contrast shows the difference between institutional oversight and professional responsibility for AI use. A contrast is also seen in integration of ethics and human rights at an institutional level and personal obligations for lawyers which focus more on confidentiality. Discrimination and fairness is more prominent in the guidelines for courts since they deal with the potential impact of AI on judicial outcomes. Three out of the five guidelines highlight the importance of AI literacy and education particularly in how lawyers must understand how AI tools function (CCBE *Considerations on the Legal Aspects of AI*, n.d.). This is discussed in the CCBE guide as *“For the reasons stated above, it is not only in lawyers’ best interest to understand how such tools work or in what directions they may develop in the future, but it is also important for society at large to enable as many lawyers as possible to use such tools effectively in the interest of their clients.”* Collectively all charts demonstrate that the bodies for legal professionals place emphasis on a risk-based approach, challenges related to AI, compliance to human rights, ensuring fairness, transparency and accountability. This is specified under UIA guidelines *“Lawyers should maintain transparent communication with their clients, informing them of their use of AI systems, the purposes of such use, and the precautions taken”*. Lawyers are expected to exercise competence, transparency and ethical judgement. However, to ensure ethical and legal compliance at the institutional and professional level, AI literacy and education of everyone involved in legal processes (including decision making) is identified as a significant theme.

6.3. Overview of Key Themes Identified by NVivo Analysis in Guidelines Published by Bodies for Healthcare Professionals

As also indicated in [Section 5](#), our review did not identify prominent EU-level healthcare professional associations representing healthcare professionals (such as doctors, nurses, or allied health professionals) that have published dedicated guidelines defining success criteria for the use of AI in professional practice. While AI is increasingly discussed in relation to clinical decision support, diagnostics and workflow optimisation in academic sources (See e.g. (Khalifa & Albadawy, 2024)), prominent bodies for healthcare professionals such as EurHeCa have not published any profession-specific AI guidance at this stage. That said, the findings of the thematic analysis of the existing guidelines in the healthcare sector are provided below.

Table 6 - Overview of the key themes identified in Nvivo-supported manual thematic analysis of Associations of Healthcare Professionals

ID	Document	No. of themes	Key themes identified through NVivo-supported manual thematic analysis	Year
PB_H_01	EFPIA Position on the use of artificial intelligence in the medicinal product lifecycle	37	Medicine development, risk-based approach, harmonization, innovation, regulatory oversight, patient safety, research and development, stakeholder	2024



Funded by
the European Union

ID	Document	No. of themes	Key themes identified through NVivo-supported manual thematic analysis	Year
			engagement, development optimization	
PB_H_03	CPME Deployment of artificial intelligence in healthcare	28	Efficiency, medical ethics, diagnosis and treatment, data protection, bias, AI literacy, workflow integration, doctor autonomy, patient-doctor relationship, patient interest, risk-based approach,	2024

Based on the NVivo-supported thematic analysis, four overarching themes emerge regarding how healthcare professional associations conceptualise AI success: (i) Patient-centered and safe AI, (ii) Ethical and Trustworthy medical AI, (iii) Professional Accountability and Human Oversight, and (iv) Institutional Readiness for AI.

Patient-centered and safe AI constitutes a core theme across healthcare documents. AI is considered successful when it demonstrably enhances patient outcomes without introducing unacceptable clinical risk. Codes clustered under this theme include risk management, risk-based approach, patient-doctor relationship and patient safety, reflecting the safety-critical nature of medical contexts. The references in the documents for these codes include but not limited to “...bearing in mind that a doctor must always be guided by the best interests of the patient” (CPME) and “AI governance is to have suitable and risk-based guidance for oversight that is tailored to the regulatory status and specific context of use”. Efficiency and innovation are also acknowledged with codes such as diagnosis and treatment, efficiency and development optimization (see e.g. EFPIA's statement of “Ultimately this will enable more uptake and adoption of this innovative technology...”) but these are also in accordance with patient welfare. .

A second theme concerns ethical and trustworthy medical AI, which emphasises compliance with medical ethics, non-discrimination, data protection, bias and transparency as reflected in the documents with the following references: “EFPIA acknowledges the need for appropriate transparency requirements as they foster trust and can help ensure safety of patients.”(EFPIA) and “AI systems must comply with medical ethics, data protection, privacy and security rules” (CPME). Professional associations stress that AI systems must align with established ethical principles governing medical practice and preserve patient trust. Here, fairness, explainability, and privacy operate as conditions for legitimacy rather than as performance metrics.



Professional accountability and human oversight emerges as a distinct and recurring theme. Healthcare AI is framed as a decision-support tool that must not replace clinical judgment. Codes such as regulatory oversight and doctor autonomy reflect that responsibility for diagnosis and treatment remains with healthcare professionals. This is reflected in the documents with quotations such as: *"When deploying AI, doctors should be able to question the AI regarding its decisions. Doctors should be able to interpret the output of an AI system and it should be made available, if necessary, the possibility to understand the internal functionality and the external behaviour of the AI system."* (CPME) and *"In summary, we believe that upcoming AI guidance from the EMA in conjunction with the established, well-functioning legislative and regulatory frameworks for medicines will ensure an appropriate regulatory framework for AI used in the lifecycle of medicines."* (EFPIA)

Finally, institutional readiness captures expectations relating to AI literacy, education and training and workflow integration. AI success is linked to the capacity of healthcare institutions and professionals to competently understand, supervise, and govern AI systems throughout their lifecycle. This is reflected in the documents through statements such as: *"In healthcare, this needs to be translated into competence development for healthcare professionals and systematic continuing education."*

If we look at the papers individually, the EFPIA position paper frames AI success primarily in terms of its contribution to the medicinal product lifecycle. The prominence of themes such as medicine development, research and development, development optimisation, and innovation suggests that AI is viewed as successful when it demonstrably enhances efficiency, speed, and quality. At the same time, the frequent appearance of regulatory oversight, patient safety, risk-based approaches, and harmonisation indicates that technological advancement alone is insufficient; success is also contingent on alignment with existing regulatory frameworks and on maintaining high safety standards. The inclusion of stakeholder engagement further implies that effective AI deployment is seen as a collaborative process involving regulators, industry actors, and other stakeholders, rather than a purely technical exercise.

On the other hand, the CPME document representing medical professionals, conceptualises AI success more directly in relation to clinical practice and professional values. Key themes such as medical ethics, doctor autonomy, patient–doctor relationship, and patient interest highlight that AI is considered successful only insofar as it supports the main ethical principles of medical professionalism. Practical considerations, including efficiency, workflow integration, diagnosis and treatment, and AI literacy, point to the importance of usability and competence in day-to-day clinical settings. At the same time, recurring concerns around data protection, bias, and risk-based approaches indicate an awareness of potential harms and systemic risks, reinforcing the view that AI success in healthcare is inseparable from trust, accountability, and ethical safeguards.



Funded by
the European Union

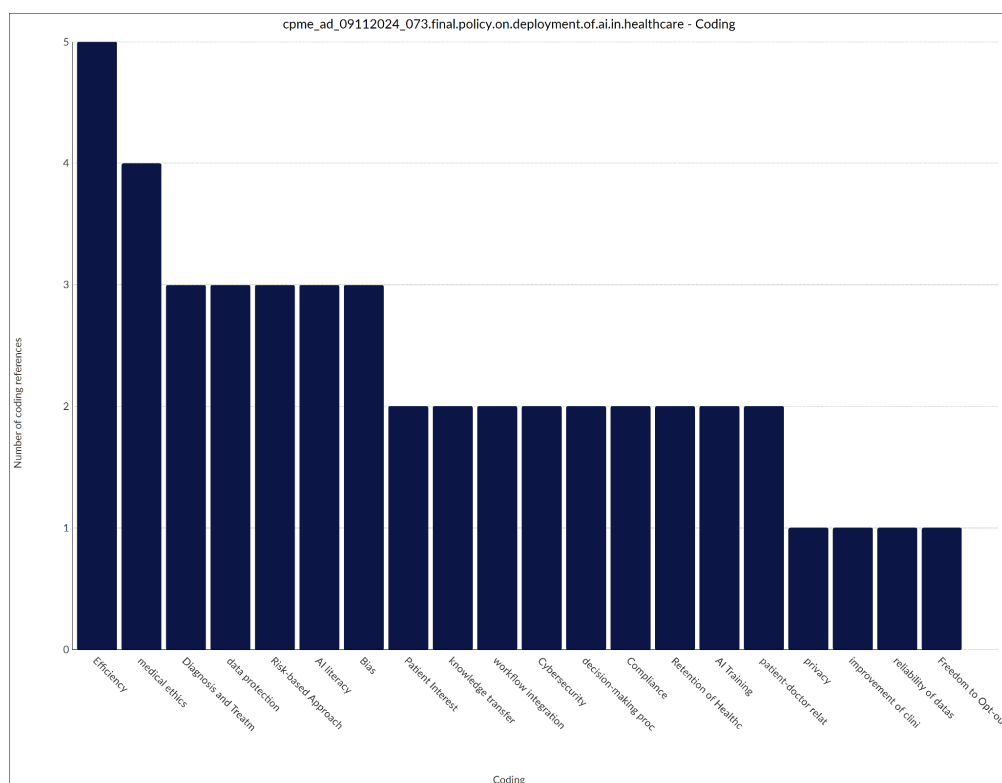


Figure 6 - NVivo Supported results for CPME Guideline on Deployment of artificial intelligence in healthcare

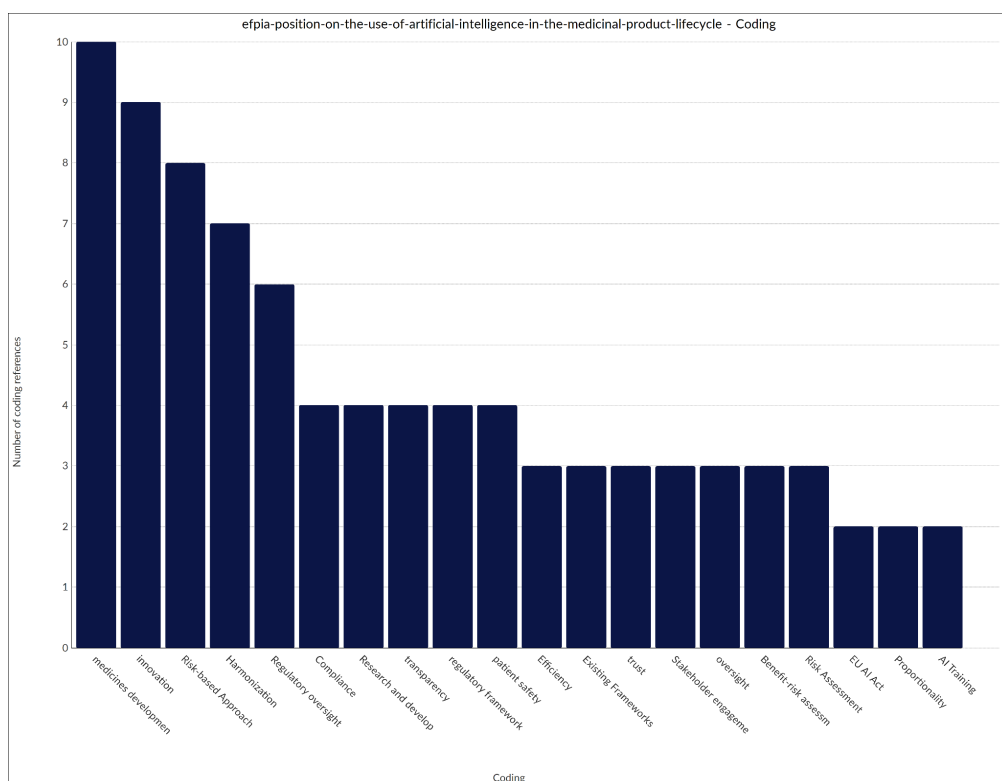


Figure 7 - NVivo Supported Results for EFPIA Position on the use of artificial intelligence in the medicinal product lifecycle



Funded by
the European Union

FORSEE has received funding from the EU's Horizon Europe research and innovation programme under Grant Agreement 101177579.

6.4. Overview of Key Themes Identified by NVivo Analysis in Guidelines Published by Bodies for Engineering Professionals

As also indicated in [Section 5](#), major EU engineering professional associations such as EASE were found to not have published AI-specific guidelines or success criteria aimed at individual engineering professionals. Although AI is increasingly used by engineers across multiple domains, only small professional associations such as the Association of Nordic Engineers and the European Federation of Chemical Engineers appear to address AI through their guiding documents or reports.

The absence of explicit guidance may reflect the diversity of engineering disciplines and professional roles, which complicates the development of common AI success criteria applicable across the profession. Another way of looking at the absence could be that engineering associations may view AI primarily as a technical tool integrated into existing design and development processes, reducing the perceived urgency for profession-level guidance.

That said, the findings of the thematic analysis of the existing guidelines are provided below.

Table 7 - Overview of the key themes identified in Nvivo-supported manual thematic analysis of Associations of Engineers

ID	Document	No. of themes	Key themes identified through NVivo-supported manual thematic analysis	Year
PB_E_01	Nordic engineers' stand on Artificial Intelligence and Ethics	62	Ethics, transparency, harm limitation, accountability, bias, democratic AI, individual responsibility, human judgement, organizational responsibility, engineer responsibility, risk mitigation	2022
PB_E_03	EFCE White Paper on Artificial Intelligence in Chemical Engineering	45	Effectiveness, accuracy, AI efficiency, innovation, AI literacy, data quality, explicability, trust, user-friendliness, transparency,	2025

The NVivo-supported thematic analysis reveals a prominent divergence in how AI success is conceptualised by the two engineering organisations. In other words, the documents demonstrate contrasting expectations about the role, risks, and desirability of AI in engineering practice. However, this heterogenous outcome can be structured around three main themes: (i) Responsible and risk-based approach, (ii) performance and functional Effectiveness and (iii) Ethical Design and Deployment.

Responsible and risk-based approach is a prominent theme in ethics-oriented and standards-focused documents. AI success is defined under codes such as harm limitation,



Funded by
the European Union

individual/organizational/engineer responsibility and risk mitigation. This is reflected in the text related to harm such as *"The notion of harm is central to the discussion of the possibilities and risks of AI from the human rights point of view. In discussions of risks and harms, the framework of human rights can often provide moral legitimacy to the expressed concerns."* (Nordic Engineers' Stand); or indications related to responsibility such as *"Nonetheless, defining clear performance criteria remains the responsibility of those developing the models."* (EFCE White Paper)

Performance and functional effectiveness appears more strongly in application-oriented and industry-facing documents such as the EFCE White Paper. Here, AI success is associated with efficiency, accuracy, user-friendliness, and innovation. This theme distinguishes engineering associations from legal and healthcare bodies, as performance metrics are more frequently treated as primary indicators of success rather than as secondary considerations. *"Effective application of AI hinges on establishing a strong business context, combining diverse data types such as time series from sensors and traceability data from production campaigns, and bridging the expertise gap between data specialists and process operators."* (EFCE White Paper)

Ethical design and deployment constitutes a bridging theme across both orientations. Engineering associations emphasise explainability, trust, transparency, ethical design and data quality as essential features of successful AI systems. These expectations locate responsibility within design processes and organisational structures, rather than solely at the point of use. Emphasis on ethical design and deployment of AI are reflected in the texts as follows: *"the Nordic countries are positioned well to be frontrunners in setting the agenda for how to address the issues of ethics in AI development and implementation."* and *"Engineers, policy makers, civil society and the general public need spaces for sustaining a living dialogue around issues of AI and ethics."* (Nordic Engineers' Stand)

Looking at the documents individually, the Nordic Engineers' stand on Artificial Intelligence and Ethics adopts a predominantly precautionary and governance-oriented framing of AI success. The dominance of themes such as ethics, responsibility (individual, organisational, and professional), accountability, harm limitation, bias, democratic AI, and risk mitigation indicates that AI is primarily approached as a source of potential societal and professional risk. Success, in this context, is not defined by performance gains or efficiency improvements, but by the extent to which AI systems are governed responsibly and kept under meaningful human control. The repeated emphasis on human judgement and engineer responsibility suggests that AI should remain subordinate to professional expertise, with engineers retaining accountability for outcomes. This document positions AI as a socio-technical challenge that requires ethical reflection, democratic oversight, and clearly assigned responsibility structures in order to prevent harm and protect public trust.

In contrast to this, the EFCE White Paper on Artificial Intelligence in Chemical Engineering presents an efficiency and productivity-based discourse surrounding AI. The prevalence of themes such as effectiveness, accuracy, efficiency, innovation, and data quality reflects a strong focus on AI's capacity to optimise processes, improve performance, and enhance



Funded by
the European Union

competitiveness within chemical engineering. In this framing, AI success is closely tied to measurable technical and operational criteria. While themes such as transparency, explicability, trust, user-friendliness, and AI literacy are present, they function primarily as enablers of fast AI adoption rather than as restrictions against its deployment. Risks and harms are less prominent and are treated as challenges to be managed in order to achieve AI's benefits, rather than as central concerns for refraining from fast adoption. .

Taken together, these documents reflect two different narratives for AI within engineering. The Nordic Engineers' document constructs AI as a potentially disruptive technology that must be carefully constrained through ethical governance and democratic accountability. In contrast, the EFCE white paper frames AI as a valuable and largely beneficial tool whose success lies in fast and effective deployment. This difference might be helpful to explain or understand why the prominent engineering professional bodies have not published guidelines demonstrating unified AI success criteria. In the existing corpus while some of the guidelines/documents emphasise caution, responsibility, and harm prevention, others prioritise efficiency, innovation, and performance optimisation. The absence of a common framework reflects underlying differences in how AI's risks and benefits are expected, valued, and governed across different areas of engineering.

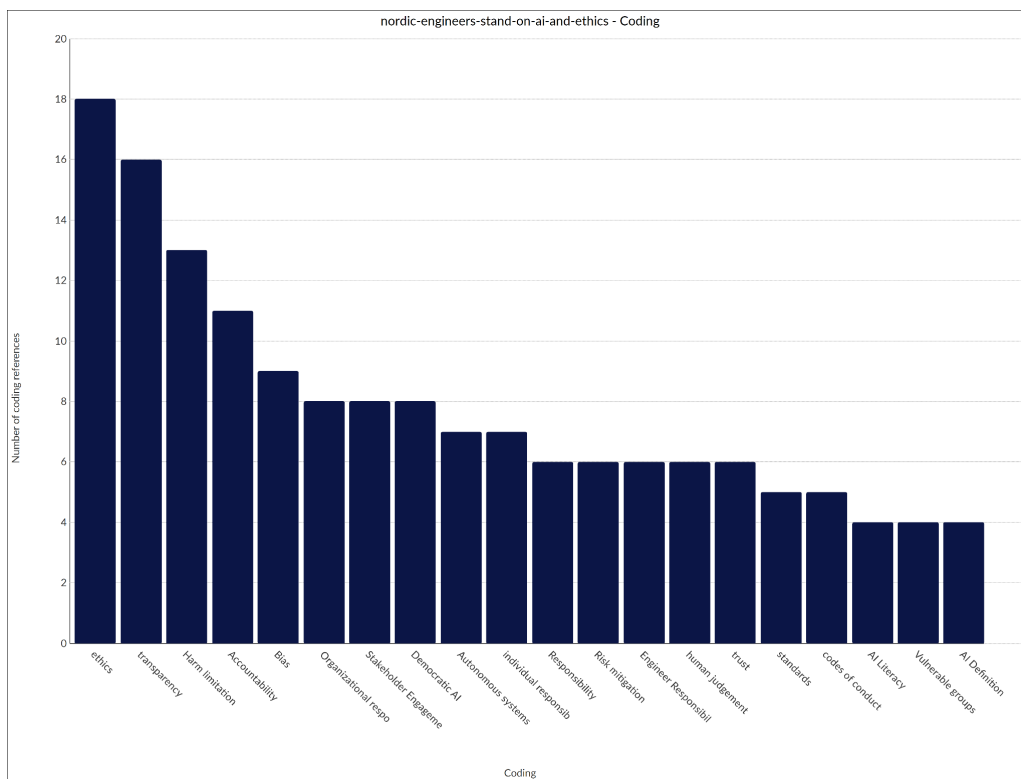


Figure 8 - NVivo Supported Results for Nordic Engineers' Stand on AI and Ethics



Funded by
the European Union

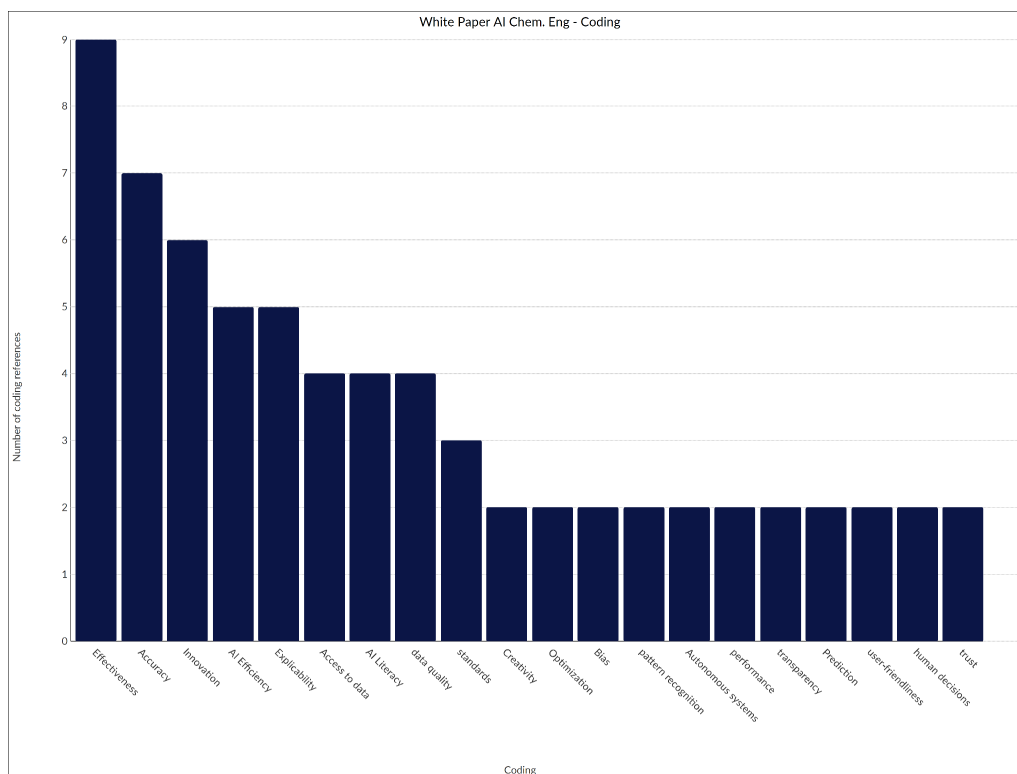


Figure 9 - NVivo Supported Results for EFCE White Paper on Artificial Intelligence in Chemical Engineering

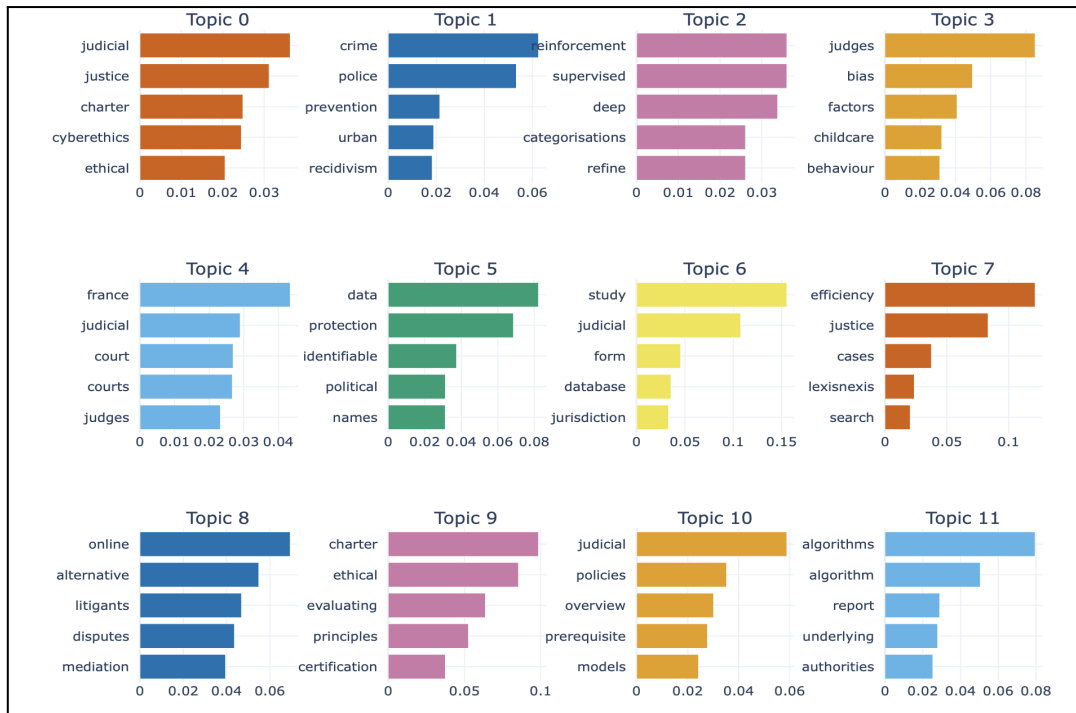
7. Themes Identified Using BERTopic

In this section we provide the themes that were identified by analysing the topics generated by BERTopic modelling. These topic clusters helped in identifying themes for the different documents by complementing the main themes identified through the manual thematic analysis conducted via NVivo by legal experts. Figure 6 shows an example of a bar chart with the BERTopic clusters for the CEPEJ ethical charter document. Table 8, 9 and 10 show the overview of the key themes identified using the BERTopic model for the legal, healthcare and engineering associations along with our analysis concerning the findings. .

Figure 10 : Example of the BERTopic clusters for the CEPEJ ethical Charter



Funded by
the European Union



7.1. Overview of the key themes identified in BERTopic modeling for Legal Organizations

The BERTopic analysis of legal organisations’ documents shows a dense and multifaceted thematic landscape, reflecting the complexity of AI’s interaction with legal systems. Across the documents, recurring themes such as ethics, rights, discrimination, bias, explainability, liability, and procedural balance indicate that AI success is primarily framed as a matter of legal legitimacy and rights protection, rather than technological advancement alone. The prominence of topics related to criminal proceedings, courts, and judicial decision-making emphasizes that legal organisations perceive AI as directly affecting core constitutional and procedural guarantees. As a result, AI is treated as a technology that must be tightly constrained by legal principles, with success measured by its ability to preserve fairness, accountability, and due process. As identified in the CCBE guidelines “... *these applications must be reconciled with the fundamental principles that govern the judicial process and guarantee a fair trial: equality of arms, impartiality, adversarial procedures, etc.*” (CCBE Considerations on the Legal Aspects of AI, n.d.)

At the same time, the presence of themes related to efficiency, legal research tools, analytics, and opportunities suggests a more pragmatic strand within the legal corpus. Here, AI success is framed in terms of professional utility, provided that safeguards such as confidentiality, security, interpretability, and competence are in place. Taken together, the BERTopic results indicate that legal organisations construct AI success through a balancing act: enabling innovation and efficiency in legal practice while preventing erosion of rights, professional responsibility, and trust in judicial institutions. That said, it can be argued based on the topic modelling that the themes related to legitimacy and rights protection prevail against the themes related to efficiency and innovation.



Funded by
the European Union

PB_L_04	CCBE considerations on the Legal Aspects of AI	29	21	Jurisdictions, Liability, Use of AI in Courts, Legal Considerations, Use of AI by lawyers, Explainability, AI in criminal proceedings and balance	2020
PB_L_05	CEPEJ Revised roadmap for ensuring an appropriate follow-up of the CEPEJ Ethical Charter on the use of artificial intelligence in judicial systems and their environment	3	1	Cyberjustice	2021
PB_L_07	EUROJUST Artificial intelligence supporting cross-border cooperation in criminal justice	2	2	Criminal Justice and Rights	2022
PB_L_06	CCBE Guide on the use of Artificial Intelligence-based tools by lawyers and law firms in the EU	39	23	Tools, Lawyers, Opportunities, Benchmark, Security, Interpretability, Competences and Analytics.	2022
PB_L_09	EUROJUST Generative Artificial Intelligence - The impact on intellectual property crimes	35	15	Malware, Criminal safeguards, Impact, Data training, Infringements, Copyright, Research mining, Policy Issues, Generative AI, Rights and Security	2023
PB_L_08	CEPEJ Assessment Tool for the Operationalisation of the European Ethical Charter on the Use of Artificial Intelligence in Judicial Systems and Their Environment	13	13	Fairness, Discrimination, Equality, Characterisation, Mitigation, Principles, Influence, Auditable, Use of AI in courts, Risks, Assessment, Bias and Competency	2023
PB_L_11	EDPS Generative AI and the EUDPR, First EDPS Orientations for ensuring data	2	4	Risks, Ethics, Regulations and Confidentiality	2024



**Funded by
the European Union**

	protection compliance when using Generative AI systems				
PB_L_12	FBE (New Technologies Commission) Guidelines on How Lawyers Should Take Advantage of the Opportunities Offered by Large Language Models and Generative AI.'	23	16	Justice, Benefits, Accuracy, Ethics, Principles, Rights, Proportionality, Standards and Privacy	2024
PB_L_13	UIA Guidelines on the use of artificial intelligence systems by lawyers	4	4	Risks, Regulations, Confidentiality and Ethics	2024
PB_L_14	UNESCO Draft guidelines for the use of AI in Courts and Tribunals	23	16	Justice, Benefits ,Accuracy, Ethics, Principles, Rights, Proportionality, Standards, Privacy, Explainability, Contestability, Governance, Transparency, Liability and Use of AI in Courts	2025

7.2. Overview of the key themes identified in BERTopic modeling for Healthcare Organizations

The BERTopic themes identified across guidelines and position papers published by healthcare organizations documents reveal a strong focus on risk mitigation, ethics, professional responsibility, and a patient-centred (or human-centered) approach. Across the documents, recurring themes include risk, governance, oversight, explainability, privacy, autonomy, and rights. This pattern suggests that AI success in healthcare is not primarily defined by performance, efficiency or innovation speed, but by the ability to integrate AI safely into existing care practices while protecting patients and supporting healthcare professionals. The repeated emphasis on the medicinal product lifecycle, regulatory compliance, and integrity further reflects healthcare's close proximity to formal regulatory regimes and safety-critical decision-making.

In addition, themes such as literacy, usability, trust, and the patient–doctor relationship indicate that AI success is closely tied to human interaction and professional practice, rather than optimization alone. Healthcare organisations appear particularly concerned with how AI reshapes human judgment in clinical processes, professional autonomy, and patient trust. The diversity of themes across documents also highlights the diversity of healthcare



Funded by
the European Union

contexts, which makes unified success criteria difficult to articulate. Overall, the BERTopic analysis suggests that healthcare organisations conceptualise AI success as a balance between innovation and ethics.

Figure 12: Word cloud of topics identified from healthcare organisations based on BERTopic results.



Table 9 - Overview of the key themes identified in BERTopic modeling for Healthcare organisations

ID	Document	No. of BERTopic clusters	No. of themes	Key themes	Year
PB_H_01	EFPIA Position on the use of artificial intelligence in the medicinal product lifecycle	8	12	Development , Risk, Regulatory framework, Oversight, Governance, Research, Collaboration and Cooperation	2024
PB_H_02	EMA Reflection paper on the use of Artificial Intelligence (AI) in the medicinal product lifecycle	14	20	Explainability, Risk, Medicinal lifecycle, Guidelines, Risk, Integrity, Compliance, Integrity, Ethical, Technical aspects, Governance and Precision	2024
PB_H_03	MHE Artificial Intelligence in Mental Health	7	13	Regulatory, Privacy, Support, Digital rights, Responsibility, Oversight, Usability, Design, and Development	2024



Funded by
the European Union

PB_H_04	CPME Deployment of artificial intelligence in healthcare	7	14	Data, Professional, Privacy, Autonomy, Literacy, Limitation, Rights, Policy, efficiency, Investment, Research, and deployment	2024
PB_H_05	CBDIO Application of AI in healthcare and its impact on the 'patient-doctor' relationship	19	22	Risk, Literacy, Application, Bias, Equity, Privacy, Trust, Quality, Standard, Guideline, Rights, Patient needs, and Autonomy	2024
PB_H_06	Health data governance in the age of artificial intelligence: policy imperatives for the WHO European Region	5	5	Training, Governance, Interoperability, Rights, and Stakeholders	2025

7.3. Overview of the key themes identified in BERTopic modeling for Engineering Organizations

The BERTopic analysis of engineering organisations sets forth a more fragmented set of expectations about AI success.

Across the documents, themes such as trustworthy AI, standards, safety, explainability, interoperability, responsibility and conformity demonstrate a concern with technical robustness, especially in areas where safety is crucial such as aviation and manufacturing. In these contexts, AI success is closely linked to certification, standardisation, and operational acceptance, reflecting engineering's traditional emphasis on verifiability, control, and compliance with technical norms.

However, the analysis of the themes also reveals a clear diversity. Some engineering documents (e.g. the Nordic Engineers' stand on AI and Ethics) emphasize on risks, responsibility, bias, and transparency, constructing AI success as dependent on ethical governance and societal accountability. Others, such as the ones focused on manufacturing or chemical engineering, emphasise availability, efficiency, innovation, and deployment readiness, framing AI success as widespread adoption, efficiency and operational benefit. This difference suggests that engineering organisations do not share a single vision of AI success; instead, expectations are shaped by sectoral application, proximity to public risk, and the institutional role of the association. As a result, AI success in engineering varies between precautionary governance and performance and efficiency-driven implementation.



Funded by
the European Union

Figure 13: Word cloud of topics identified from engineering organisations based on BERTopic results



Table 10 - Overview of the key themes identified in BERTopic modeling for Engineering organisations

ID	Document	No. of BERTopic clusters	No. of themes	Key themes	Year
PB_E_01	Nordic engineers' stand on Artificial Intelligence and Ethics	19	15	Ethics, risk, responsibility, transparency, adaptability, political support, bias, autonomous, cooperation, training, societal expectations, unpredictability	2022
PB_E_02	AIM-NET White paper on Artificial Intelligence in Manufacturing	17	19	Assessments, Monitoring, AI literacy, security, Accessibility, Cybersecurity, Sustainability, Surveillance, and interoperability,	2023
PB_E_03	EFCE White Paper on Artificial Intelligence in Chemical Engineering	19	20	Availability, confidentiality, production, challenge, collaboration, solubility, education, AI literacy, standardization, accuracy,	2025



Funded by
the European Union

				optimization, training, explainability	
PB_E_04	Aerospace Formulating an Engineering Framework for Future AI Certification in Aviation	23	26	Security, Explainability, Safety, standards, safety regulations, Acceptance, Development, Framework, Operationality	2025

8. Validation of BERTopic-based Themes through Comparison with NVivo-supported Thematic Analysis

The themes produced by Nvivo manual thematic analysis and BERTopic were compared to identify the overlaps and divergences between the two methods of producing themes. This comparison can be found in the [Appendix](#).

These tables show the documents where both Nvivo and BERTopic analyses were conducted. Considering the fact that Nvivo and BERTopic were run in parallel to each other, there were overlaps in the themes generated from each analysis.

The themes that had similar meaning but were worded differently under BERTopic and Nvivo were discussed among the team members and were later assigned to a theme that both team members agreed upon. The comparative analysis between the two shows that the themes generated by BERTopic for shorter documents were less thorough in comparison to longer documents. In addition, the topic clusters generated for all the documents by BERTopic were not the same amount as the ones generated manually by Nvivo. While the amount of themes differed in each, the key themes were similar in both analyses and the main focus of the guidelines were captured by BERTopic despite not being as thorough as Nvivo.

Below we provide the comparative analysis of NVivo and BERTopic for (i) legal, (ii) healthcare and (iii) engineering associations separately.

8.1. Comparative Analysis for Legal Associations

Across legal associations, there is significant overlap between NVivo and BERTopic around a core set of themes, including transparency, bias and discrimination, explainability, use of AI in courts, use of AI by lawyers, rights, and procedural safeguards. These recurring overlaps suggest that certain concerns are structurally central to how AI is discussed in the legal domain, regardless of analytical method. Both approaches consistently surface AI as a technology that directly affects legal reasoning, judicial processes, and fundamental rights, indicating a shared understanding that AI success in law is inseparable from legal legitimacy



Funded by
the European Union

and due process. This is supported by the UIA guidelines *“Considering that the use of AI systems can pose significant risks to human rights and may lead to unintended consequences affecting individual freedoms and protections”*.

At the same time, since the NVivo analysis is conducted manually by legal experts, it is not a surprise that it produced a much richer and more normatively articulated thematic structure. The manually conducted NVivo analysis demonstrated evaluative concepts such as (i) the main classifications based on ethical, societal, legal and operational success, (ii) human-centric approaches, (iii) human oversight, and (iv) the right to a fair trial. These themes reflect deliberate judgments and evaluations about what the written text’s actual message is and what (according to the text) should count as success, in contrast to the more quantitative findings of BERTopic. Indeed, BERTopic tends to surface more descriptive or functional clusters, such as tools, efficiency, jurisdictions, analytics, or legal research, which reflect how AI is discussed in practice but do not, on their own, articulate normative success criteria. The combined use of unsupervised topic modelling and manual qualitative analysis allows the study to move beyond just the identification of themes to examining the significance of guidelines. Both the analyses show the impact of AI on how professional bodies address responsibility and transparency. These themes help in addressing what professional bodies find imperative to forging successful AI applications within the EU and the current standards through which AI is being governed in these professions.

8.2. Comparative Analysis for Healthcare Associations

In the corpus of healthcare associations, NVivo and BERTopic again show overlaps (although not as strong as the legal associations), particularly around themes such as risk, oversight, governance, regulation, privacy, AI training, literacy, and research and development. These overlaps indicate a shared recognition that AI in healthcare is a high-risk, safety-critical domain where governance and professional competence are prerequisites for success. Both methods surface AI as something that must be carefully integrated into existing regulatory and clinical frameworks rather than treated as a standalone innovation.

However, the manual NVivo analysis reveals a substantially broader and more practice-oriented understanding of AI success than BERTopic. Manually coded themes uniquely highlight patient interest, professional autonomy, the patient–doctor relationship, clinical decision-making, ethics, trust, and workflow integration. These themes reflect healthcare professionals’ lived concerns and underscore that AI success is not only about technical performance or innovation, but about preserving care quality, professional judgment, and patient trust. BERTopic, by contrast, tends to aggregate themes at a higher level of abstraction, emphasizing on governance, policy, and deployment without consistently capturing the relational and ethical dimensions of care.

The comparison suggests that automated topic modelling captures the institutional and regulatory discourse surrounding AI in healthcare, while manual NVivo thematic analysis is



Funded by
the European Union

better suited to identifying profession-specific success criteria rooted in clinical practice. Both analyses together, the analyses reinforce the conclusion that AI success in healthcare is defined less by efficiency gains and more by safe, ethical, and trust-preserving integration into professional practice.

8.3. Comparative Analysis for Engineering Associations

For engineering organisations, the comparison between NVivo and BERTopic reveals both a higher level of divergence as well as more variation compared to the legal or healthcare domains. Even though there were many themes that overlapped including but not limited to safety, transparency, bias, responsibility, explainability, AI literacy and innovation, the amount, depth and breadth of the identified themes differ significantly.

The NVivo analysis surfaces a highly normative and responsibility-focused framing in some engineering texts, such as the highly ethics-oriented Nordic Engineers' stand. Here, the expert-led NVivo analysis uniquely captures themes such as democratic AI, harm limitation, human judgement, (classification of) institutional and professional responsibility, power dynamics, and societal expectations. These themes position AI success as contingent on ethical governance and accountability rather than deployment or performance. BERTopic, while identifying ethics and risk-related clusters, tends to represent these concerns more diffusely and places greater emphasis on adaptability, predictability, and framing.

In contrast, for application/efficiency-based documents such as the EFCE White Paper, both methods identify themes related to accuracy, data quality, optimisation, and safety, but the expert-led NVivo analysis again provides greater depth by connecting these to innovation, efficiency, responsibility, and trust. BERTopic emphasizes on availability and functionality, reflecting a discourse that takes deployment of AI at the forefront. Overall, the comparison shows that the NVivo analysis is crucial for distinguishing competing discourses of AI success within the documents produced by engineering bodies (e.g. precautionary versus efficiency-based) while BERTopic effectively highlights the technical and operational vocabulary through which these discourses are expressed.

9. Key Insights based on Sociology of Expectations

9.1. Initial Insights on Empirical Findings

Across legal, healthcare, and engineering professional bodies, AI success is predominantly framed around human oversight, institutional accountability, and the protection of core professional values, rather than stand-alone technical achievement. Legal organisations mostly emphasise on human-centric criteria such as fairness, transparency, non-discrimination, and the protection of fundamental rights, generally positioning AI as a supportive technology that should stay dependent on human oversight (see e.g. CCBE, 2020, EUROJUST, 2022). Healthcare bodies conceptualize AI success in terms of patient



Funded by
the European Union

safety, professional autonomy and trust. With an emphasis on governance, risk management, and procedural safeguards over technological performance alone. Engineering bodies exhibit diverse approaches where some of them stress ethical responsibility, harm avoidance, and democratic accountability, while others insist on AI's potential for efficiency, optimisation, and innovation.

Across these domains, the recurring emphasis on human oversight, supportive (rather than substitutive) AI, professional judgment, and institutional accountability emerges as a shared empirical pattern. Professional bodies also seem to be generally against overly simplified narratives of AI, in an attempt to refrain from eliminating human responsibility. This is similarly identified in sociological analyses of AI and future-oriented discourses on work and professional roles (Vicsek, 2020).

In legal and healthcare contexts, transparency is closely linked to individual rights, contestability, and trust, while in engineering it is more often associated with explainability, interpretability and technical reliability. Furthermore, across all sectors, AI literacy emerges as a condition for responsible AI use. Legal professionals, healthcare practitioners, and engineers are increasingly positioned as critical evaluators and decision-makers rather than passive users of AI systems.

At the same time, parts of the engineering corpus demonstrate more optimistic and efficiency-based expectations that encourage adoption and innovation. This contrast highlights internal and external variation between discourses of professional bodies, particularly when compared to the more precautionary framing observed in law and healthcare.

9.2. Applying the Theoretical Lense: SoE

Drawing on the sociology of expectations literature, these patterns can be understood as the outcome of professional bodies actively shaping anticipatory narratives about the role of AI in the future (Borup et al., 2006; Brown & Michael, 2003; van Lente, 2012b). Or in other words, the success criteria articulated by professional associations can be understood as *performative expectations* rather than as neutral or descriptive accounts of desirable AI characteristics. (Borup et al., 2006; van Lente et al., 2013) SoE scholarship highlights how expectations become particularly salient in moments when futures are contested or uncertain (Brown & Michael, 2003a).

Moreover, SoE scholarship emphasises that expectations do not merely describe futures but actively mobilise actors, coordinate activities, and manage uncertainty. As Brown & Michael, 2003 discuss, this reflects *"how the future is mobilized in real time to marshal resources, coordinate activities and manage uncertainty"*. Expectations are therefore understood as governance instruments that shape technological trajectories, institutional responses, and professional roles.



Funded by
the European Union

9.3. Analysis and Discussion

Rather than passively reflecting technological developments, legal, healthcare, and engineering organisations engage in governance by articulating desirable and undesirable futures for AI in ways that align with sector-specific risks, responsibilities, and institutional responsibilities. In domains closely tied to rights, safety, and public trust, such as law and healthcare, the expectations are observed as more precautionary (Hemphill, 2020), restricting AI through human oversight and ethical safeguards. In contrast, more efficiency-oriented expectations (e.g. engineering) can be connected to the “hype” aspect of sociology of expectations where AI’s promise of application and adaptability across industries strengthen the hypes and expectations surrounding its initial phase and increase the chance of sustaining these technologies (van Lente et al., 2013).

At the same time, these success criteria reflect a process of professional mobilisation under conditions of technological and regulatory uncertainty. Kerr et al discuss that expectations surrounding AI at the societal level require governance of the design and use of steering of AI. According to this discussion, the documents analysed here suggest that professional bodies are responding not only to uncertainty about how AI technologies will evolve, but also to uncertainty surrounding the regulatory and governance frameworks being developed by states, supranational institutions, and hybrid public–private arrangements (Kerr et al., 2020b). By articulating stable and recurrent conditions for AI success (such as rights compliance, explainability, traceability, and professional accountability) associations effectively assert that, irrespective of future legal or technical configurations, professional norms will continue to govern practice.

Across the three domains, the consistent emphasis on human oversight, supportive AI, professional judgement and institutional accountability may therefore be interpreted as a form of boundary work aimed at maintaining professional authority. By continuously positioning AI as a tool that must remain embedded within existing professional decision-making structures, professional associations articulate expectations that resist the displacement of human expertise by automated systems. These types of expectations also implicitly delimit the influence of external actors (particularly technology corporations) by asserting that legitimate AI use is conditional on compliance with industry-standard norms, standards, and ethical commitments.

Key concepts such as transparency, fairness, and accountability can function as mechanisms of expectation management shaping how AI systems are legitimised and trusted. We understand that the FAT concepts are identified by the professional associations to “manage the uncertainty” surrounding AI while also moderating the hype connected with AI-driven efficiency and innovation narratives (van Lente et al., 2013). In other words, these concepts operate as normative anchors that allow professional bodies to articulate the “promises of futures” (Brown et al., 2003; van Lente, 2012). Rather than specifying concrete technological outcomes, these concepts aim to stabilise uncertainty around the future role of AI in professional practice by rendering AI governable, contestable



Funded by
the European Union

and compliant with existing professional, ethical and legal norms even as the technology itself remains evolving.

Finally, the emphasis on AI-literacy demonstrates a signal to co-governance. From an expectations perspective, AI literacy helps ground the future of AI in professional judgment and institutional accountability, ensuring that trust in AI remains rooted in human expertise rather than technological promise. In this sense, rather than positioning themselves as passive recipients of technological innovation or regulatory mandates, professional bodies articulate expectations that foreground their role as indispensable intermediaries between AI systems, institutional systems/actors, and societal values. Through these performative expectations, professional associations seek to secure continued relevance and authority within a socio-technical field increasingly shaped by large technology firms and policy initiatives led by the European Commission.

10. Conclusion

This report set out to identify and analyse how legal, healthcare, and engineering professional bodies conceptualise baseline criteria for AI success, with particular attention to the values, conditions, and safeguards attached to the use of AI in professional practice. By combining expert-led NVivo thematic analysis with BERTopic modelling, the report aims to provide a cross-sectoral and multidisciplinary understanding of how AI success is framed by legal and professional bodies not only in terms of technological or economic criteria but as a socio-technical context

In this regard, this report sets forth that, across all three sectors, a central and consistent finding is that criteria for AI to be considered successful is conditional and not absolute.

Legal organisations frame success primarily through the protection of fundamental rights, procedural fairness, transparency, and human oversight, reflecting AI's direct impact on justice, due process, and the rule of law. Legal organisations outline success primarily through the protection of fundamental rights and human oversight, which show the impact of AI on justice and rule of law.

Healthcare bodies similarly emphasise risk-management, professional autonomy, and patient trust, highlighting the patient centered nature of clinical decision-making. In both domains, AI is predominantly positioned as a supportive tool whose legitimacy depends on its ability to preserve human judgment and institutional accountability.

Engineering organisations demonstrate a more internally diverse picture. While some bodies (particularly those with a more ethics-oriented approach) emphasize on responsibility, harm prevention, and democratic accountability, others adopt a more operational success perspective, highlighting efficiency, optimisation, and innovation. This diversity reflects sectoral variation within engineering itself, as well as differences in proximity to public risk, regulatory oversight, and societal expectations. As a result, AI success in engineering varies



**Funded by
the European Union**

between precautionary governance and performance and efficiency-based deployment, rather than a single, unified set of criteria.

That said, it should also be mentioned that further research concerning the success criteria of engineering and healthcare bodies should be conducted once there are more guidelines published by prominent bodies.

Methodologically, the combined use of NVivo and BERTopic proved its importance in capturing the complexity of the term “AI Success”. BERTopic provided a systematic overview of recurring discourses and thematic emphases, while NVivo enabled expert interpretation of normative concepts, societal values, and evaluative judgments. The overlaps between the two methods validated the key themes that have been identified, while the differences highlighted where expert insight was necessary to translate descriptive topic clusters into meaningful themes for success criteria.

Taken together, the findings show that professional bodies’ definitions of AI success operate as performative expectations in the sense articulated by the Sociology of Expectations. The emphasis on human oversight, supportive AI, and professional judgment reflects anticipatory governance that seeks to stabilise uncertainty by anchoring future AI use in existing institutional roles and responsibilities. Industry-specific framings (precautionary in law and healthcare and more efficiency-oriented in parts of engineering) demonstrate how expectations are calibrated to different risk profiles and professional areas, aligning with SoE accounts of how futures are selectively constructed and mobilised. The recurring invocation of transparency, fairness, accountability, and AI literacy further illustrates how expectations function as mechanisms of boundary work, simultaneously managing hype and asserting the position of the professional associations authority. In this way, professional associations do not merely respond to emerging AI technologies but actively shape their future governance by articulating normative conditions under which AI can be considered legitimate, acceptable, and compatible with professional practice.

Overall, the findings suggest that the absence of harmonised, cross-sectoral AI success criteria at EU level is not simply a gap in the literature, but an important demonstration of legitimate and enduring differences in professional roles, risks, and responsibilities.



Bibliography

Reviewed Documents

European Commission for the Efficiency of Justice. (2018). *European Ethical Charter on the use of artificial intelligence in judicial systems and their environment*. Council of Europe.
<https://rm.coe.int/ethical-charter-en-for-publication-4-december-2018/16808f699c>

European Commission for the Efficiency of Justice. (2021). *Revised roadmap for ensuring an appropriate follow-up of the CEPEJ Ethical Charter on the use of artificial intelligence in judicial systems and their environment*. Council of Europe.
<https://rm.coe.int/cepej-2021-16-en-revised-roadmap-follow-up-charter/1680a4cf2f>

European Commission for the Efficiency of Justice. (2023). *Assessment tool for the operationalisation of the European Ethical Charter on the use of artificial intelligence in judicial systems and their environment*. Council of Europe.
<https://rm.coe.int/cepej-2023-16final-operationalisation-ai-ethical-charter-en/1680adcc9c>

Council of Bars and Law Societies of Europe. (2020). *CCBE considerations on the legal aspects of artificial intelligence*.
https://www.ccbe.eu/fileadmin/speciality_distribution/public/documents/IT_LAW/ITL_Guides_recommendations/EN_ITL_20200220_CCBE-considerations-on-the-Legal-Aspects-of-AI.pdf

Council of Bars and Law Societies of Europe. (2022). *CCBE guide on the use of artificial intelligence-based tools by lawyers and law firms in the EU*.
https://www.ccbe.eu/fileadmin/speciality_distribution/public/documents/IT_LAW/ITL_Reports_studies/EN_ITL_20220331_Guide-AI4L.pdf

European Law Institute. (2020). *Model rules on impact assessment of algorithmic decision-making systems used by public administration*.
https://www.europeanlawinstitute.eu/fileadmin/user_upload/p_eli/Publications/ELI_Model_Rules_on_Impact_Assessment_of_ADMs_Used_by_Public_Administration.pdf

European Law Institute. (2022). *Guiding principles and model rules on digital assistants for consumer contracts*.
<https://www.europeanlawinstitute.eu/projects-instruments/instruments/eli-guiding-principles-and-model-rules-on-digital-assistants-for-consumer-contracts/>



Funded by
the European Union

European Law Institute. (2024). *ELI GPAI principles on co-generated data*.

<https://www.europeanlawinstitute.eu/projects-instruments/instruments/gpai-project-on-co-generated-data/>

European Union Agency for Criminal Justice Cooperation. (2022). *Artificial intelligence supporting cross-border cooperation in criminal justice*. Publications Office of the European Union.

<https://www.eurojust.europa.eu/publication/artificial-intelligence-supporting-cross-border-cooperation-criminal-justice>

European Union Agency for Criminal Justice Cooperation. (2023). *Generative artificial intelligence: The impact on intellectual property crimes*.

<https://www.eurojust.europa.eu/publication/generative-artificial-intelligence-impact-intellectual-property-crimes>

European Bars Federation. (2024). *Guidelines on how lawyers should take advantage of the opportunities offered by large language models and generative artificial intelligence*.

<https://www.fbe.org/wp-content/uploads/2023/06/European-lawyers-in-the-era-of-ChatGPT-FBE-Guidelines-on-how-lawyers-should-take-advantage-of-the-opportunities-offered-by-large-language-models-and-generative-ai.pdf>

Union Internationale des Avocats. (2024). *UIA guidelines on the use of artificial intelligence systems by lawyers*.

<https://rm.coe.int/uia-guidelines-use-ai-systems-by-lawyers-en-final/1680b438e1>

UNESCO. (2025). *Draft guidelines for the use of artificial intelligence in courts and tribunals*.

<https://unesdoc.unesco.org/ark:/48223/pf0000393682.locale=en>

European Federation of Pharmaceutical Industries and Associations. (2024). *EFPIA position on the use of artificial intelligence in the medicinal product lifecycle*.

<https://efpia.eu/media/tzeavw1t/efpia-position-on-the-use-of-artificial-intelligence-in-the-medicinal-product-lifecycle.pdf>

Mental Health Europe. (2024). *Artificial intelligence in mental health*.

<https://www.mentalhealtheurope.org/mental-health-europe-launches-a-study-on-artificial-intelligence-in-mental-healthcare/>

Committee of European Doctors. (2024). *Deployment of artificial intelligence in healthcare*.

https://www.cpme.eu/api/documents/adopted/2024/11/cpme_ad_09112024_073.final.policy.on.deployment.of.ai.in.healthcare.pdf

Steering Committee for Human Rights in the fields of Biomedicine and Health. (2024). *Application of artificial intelligence in healthcare and its impact on the*



**Funded by
the European Union**

patient–doctor relationship. Council of Europe.

<https://www.coe.int/en/web/human-rights-and-biomedicine/report-impact-of-ai-on-the-doctor-patient-relationship>

World Health Organization Regional Office for Europe. (2025). *Health data governance in the age of artificial intelligence: Policy imperatives for the WHO European Region*. <https://www.who.int/europe/publications/i/item/WHO-EURO-2025-11462-51234-78079>

Association of Nordic Engineers. (2022). *Nordic engineers' stand on artificial intelligence and ethics*. <https://nordicengineers.org/wp-content/uploads/2022/12/nordic-engineers-stand-on-ai-and-ethics.pdf>

Artificial Intelligence in Manufacturing NETwork. (2023). *White paper on artificial intelligence in manufacturing*. <https://aim-net.eu/wp-content/uploads/2024/01/AIM-NET-Artificial-Intelligence-in-Manufacturing-white-paper.pdf>

European Federation of Chemical Engineers. (2025). *White paper on artificial intelligence in chemical engineering*. https://efce.info/efce_media/-EGOTEC-68d1fed18a8d68470e45286e734fb51a-p-15752.pdf

Christensen, J. M., Stefani, T., Anilkumar Girija, A., Hoemann, E., Vogt, A., Werbilo, V., Durak, U., Köster, F., Krüger, T., & Hallerbach, S. (2025). Formulating an Engineering Framework for Future AI Certification in Aviation. *Aerospace*, 12(6), 482. <https://doi.org/10.3390/aerospace12060482>

Articles

Birhane, A., Steed, R., Ojewale, V., Vecchione, B., & Raji, I. D. (2024). *AI auditing: The Broken Bus on the Road to AI Accountability* (No. arXiv:2401.14462). arXiv. <https://doi.org/10.48550/arXiv.2401.14462>

Blei, D. M., Ng, A. Y., & Jordan, M. I. (2003). Latent dirichlet allocation. *J. Mach. Learn. Res.*, 3(null), 993–1022.

Braun, V., & Clarke, V. (2021a). One size fits all? What counts as quality practice in (reflexive) thematic analysis? *Qualitative Research in Psychology*, 18(3), 328–352. <https://doi.org/10.1080/14780887.2020.1769238>

Braun, V., & Clarke, V. (2021b). *Thematic Analysis: A Practical Guide*. SAGE.

Brown, N., & Michael, M. (2003). A Sociology of Expectations: *Retrospecting Prospects and Prospecting Retrospects*. *Technology Analysis & Strategic Management*, 15(1), 3–18. <https://doi.org/10.1080/0953732032000046024>



Funded by
the European Union

- Brown, N., Rip, A., & Van Lente, H. (2003). Expectations in & about science and technology. *A Background Paper for the 'Expectations' Workshop of June, 177*, 371–380. https://www.researchgate.net/profile/Harro-Van-Lente/publication/46671758_Expectations_in_About_Science_and_Technology/links/0c960527b59d75b7b6000000/Expectations-in-About-Science-and-Technology.pdf
- CCBE *Considerations on the Legal Aspects of AI*. (n.d.). Retrieved December 17, 2025, from https://www.ccbe.eu/fileadmin/speciality_distribution/public/documents/IT_LAW/ITL_Guides_recommendations/EN_ITL_20200220_CCBE-considerations-on-the-Legal-Aspects-of-AI.pdf
- Churchill, R., & Singh, L. (2022). The Evolution of Topic Modeling. *ACM Comput. Surv.*, 54(10s), 215:1-215:35. <https://doi.org/10.1145/3507900>
- Corrêa, N. K., Galvão, C., Santos, J. W., Del Pino, C., Pinto, E. P., Barbosa, C., Massmann, D., Mambrini, R., Galvão, L., Terem, E., & de Oliveira, N. (2023). Worldwide AI ethics: A review of 200 guidelines and recommendations for AI governance. *Patterns*, 4(10), 100857. <https://doi.org/10.1016/j.patter.2023.100857>
- Dhakal, K. (2022). NVivo. *Journal of the Medical Library Association: JMLA*, 110(2), 270–272. <https://doi.org/10.5195/jmla.2022.1271>
- Draft guidelines for the use of AI systems in courts and tribunals—UNESCO Digital Library*. (n.d.). Retrieved December 20, 2025, from <https://unesdoc.unesco.org/ark:/48223/pf0000393682>
- ES250132_PREMS 005419 GBR 2013 charte ethique CEPEJ WEB A5.pdf. (n.d.). Retrieved December 17, 2025, from <https://rm.coe.int/ethical-charter-en-for-publication-4-december-2018/16808f699c>
- European organizations | Site Name*. (n.d.). Retrieved December 20, 2025, from <https://cnb.avocat.fr/en/european-organizations>
- Fjeld, J., Achten, N., Hilligoss, H., Nagy, A., & Srikumar, M. (2020). *Principled Artificial Intelligence: Mapping Consensus in Ethical and Rights-Based Approaches to Principles for AI* (SSRN Scholarly Paper No. 3518482). Social Science Research Network. <https://doi.org/10.2139/ssrn.3518482>
- Grootendorst, M. (2022). *BERTopic: Neural topic modeling with a class-based TF-IDF procedure* (No. arXiv:2203.05794). arXiv. <https://doi.org/10.48550/arXiv.2203.05794>
- Hagendorff, T. (2020). The Ethics of AI Ethics: An Evaluation of Guidelines. *Minds and Machines*, 30(1), 99–120. <https://doi.org/10.1007/s11023-020-09517-8>
- Hemphill, T. A. (2020). “The innovation governance dilemma: Alternatives to the precautionary principle.” *Technology in Society*, 63, 101381. <https://doi.org/10.1016/j.techsoc.2020.101381>



- Jobin, A., Ienca, M., & Vayena, E. (2019). The global landscape of AI ethics guidelines. *Nature Machine Intelligence*, 1(9), 389–399. <https://doi.org/10.1038/s42256-019-0088-2>
- Kerr, A., Barry, M., & Kelleher, J. D. (2020a). Expectations of artificial intelligence and the performativity of ethics: Implications for communication governance. *Big Data & Society*, 7(1), 2053951720915939. <https://doi.org/10.1177/2053951720915939>
- Kerr, A., Barry, M., & Kelleher, J. D. (2020b). Expectations of artificial intelligence and the performativity of ethics: Implications for communication governance. *Big Data & Society*, 7(1), 2053951720915939. <https://doi.org/10.1177/2053951720915939>
- Khalifa, M., & Albadawy, M. (2024). AI in diagnostic imaging: Revolutionising accuracy and efficiency. *Computer Methods and Programs in Biomedicine Update*, 5, 100146. <https://doi.org/10.1016/j.cmpbup.2024.100146>
- Lu, Q., Zhu, L., Xu, X., Whittle, J., Zowghi, D., & Jacquet, A. (2024). Responsible AI Pattern Catalogue: A Collection of Best Practices for AI Governance and Engineering. *ACM Comput. Surv.*, 56(7), 173:1-173:35. <https://doi.org/10.1145/3626234>
- Palladino, N. (2023). A ‘biased’ emerging governance regime for artificial intelligence? How AI ethics get skewed moving from principles to practices. *Telecommunications Policy*, 47(5), 102479. <https://doi.org/10.1016/j.telpol.2022.102479>
- Pinch, T. J., & Bijker, W. E. (1984). The Social Construction of Facts and Artefacts: Or How the Sociology of Science and the Sociology of Technology might Benefit Each Other. *Social Studies of Science*, 14(3), 399–441. <https://doi.org/10.1177/030631284014003004>
- Puppis, M. (2024). Sociological institutionalism: Conceptualizing media governance as institution and organization. In *Handbook of Media and Communication Governance* (pp. 28–39). Edward Elgar Publishing. <https://www.elgaronline.com/edcollchap/book/9781800887206/book-part-9781800887206-10.xml>
- Reddy, S., Allan, S., Coghlan, S., & Cooper, P. (2020). A governance model for the application of AI in health care. *Journal of the American Medical Informatics Association*, 27(3), 491–497. <https://doi.org/10.1093/jamia/ocz192>
- Reuel, A., Bucknall, B., Casper, S., Fist, T., Soder, L., Aarne, O., Hammond, L., Ibrahim, L., Chan, A., Wills, P., Anderljung, M., Garfinkel, B., Heim, L., Trask, A., Mukobi, G., Schaeffer, R., Baker, M., Hooker, S., Solaiman, I., ... Trager, R. (2025). *Open Problems in Technical AI Governance* (No. arXiv:2407.14981). arXiv. <https://doi.org/10.48550/arXiv.2407.14981>
- Roche, C., Wall, P. J., & Lewis, D. (2023). Ethics and diversity in artificial intelligence policies, strategies and initiatives. *AI and Ethics*, 3(4), 1095–1115. <https://doi.org/10.1007/s43681-022-00218-9>
- Squires, V. (2023). Thematic Analysis. In J. M. Okoko, S. Tunison, & K. D. Walker (Eds.), *Varieties of Qualitative Research Methods: Selected Contextual Perspectives* (pp.



463–468). Springer International Publishing.
https://doi.org/10.1007/978-3-031-04394-9_72

Surden, H. (2018). Artificial Intelligence and Law: An Overview. *Georgia State University Law Review*, 35(4), 1305–1338.

Terry, G., Hayfield, N., Clarke, V., & Braun, V. (2017a). Thematic analysis. *The SAGE Handbook of Qualitative Research in Psychology*, 2(17–37), 25.

Terry, G., Hayfield, N., Clarke, V., & Braun, V. (2017b). Thematic Analysis. In C. Willig & W. S. Rogers, *The SAGE Handbook of Qualitative Research in Psychology* (pp. 17–36). SAGE Publications Ltd. <https://doi.org/10.4135/9781526405555.n2>

van Lente, H. (2012). Navigating foresight in a sea of expectations: Lessons from the sociology of expectations. *Technology Analysis & Strategic Management*, 24(8), 769–782. <https://doi.org/10.1080/09537325.2012.715478>

van Lente, H., Spitters, C., & Peine, A. (2013). Comparing technological hype cycles: Towards a theory. *Technological Forecasting and Social Change*, 80(8), 1615–1628. <https://doi.org/10.1016/j.techfore.2012.12.004>

Vicsek, L. (2020). Artificial intelligence and the future of work – lessons from the sociology of expectations. *International Journal of Sociology and Social Policy*, 41(7–8), 842–861. <https://doi.org/10.1108/IJSSP-05-2020-0174>



Funded by
the European Union

Appendix

1. Tables Demonstrating the Comparison of Nvivo and BERTopic for Legal Associations

Table A.1 - Comparison of Nvivo and BERTopic themes for the CEPEJ Ethical Charter

PB_L_01 CEPEJ Ethical Charter		
No. of themes	29 (BERTopic) < 67 (NVivo)	
Overlaps	7 themes: 1. Discrimination 2. Transparency 3. Predictive AI 4. Bias 5. Risk 6. Security 7. Societal Impact	
Divergences	Unique to BERTopic	Unique to NVivo
	1. social rights 2. Ethics, environment 3. Crime prevention, 4. cognitive ability 5. Legal research 6. data protection, health 7. database, efficiency 8. legal processes 9. ethics, compatibility 10. Algorithms 11. justice 12. access, openness 13. Legal reasonings 14. Digital Technologies 15. predictive software 16. Rights 17. predictive justice 18. transparency 19. recidivism, discrimination 20. professionals 21. inefficiency, policy 22. Court Decisions 23. Standards 24. Investigative tools	1. Adversarial debate 2. AI literacy and education 3. availability of data 4. Bias 5. Big data 6. Challenges 7. Civil Rights 8. Compliance 9. Confidentiality 10. cyberethics 11. Data-snooping 12. Defining Success in AI 13. CEPEJ Approach 14. Definition of AI 15. Ethical Success 16. Legal Success 17. Operational Success 18. Societal Success 19. Operational Success 20. Regulatory Success 21. Democratic Principles 22. Explainability 23. Fairness 24. Human dignity



Funded by
the European Union

	25. Legal Reasoning	25. Human friendly technology 26. Human Oversight 27. Human Review 28. Human Rights 29. Human-centric approach 30. Laws and legal reasoning 31. Legal knowledge representation 32. Machine Learning 33. Neutrality 34. NLP 35. Obligations 36. Open Data 37. Possible solutions 38. Precautionary principle 39. Predictive AI 40. Predictive criminal mapping 41. Predictive Policing 42. Privacy and Data protection 43. Proportionality 44. Reliability 45. Right to Fair Trial 46. Rule of Law 47. Strong AI 48. Supportive AI 49. Supervised learning 50. Trustworthy AI 51. Types of Justice 52. predictive justice 53. Usability 54. Use of AI by lawyers 55. Functionality of AI tools 56. Use of AI in Courts 57. Equality of Arms 58. Exploitation of conclusions 59. Principle of Impartiality 60. Use of AI in Criminal investigation 61. Use of AI in Law firms 62. Weak AI
--	---------------------	--



Table A.2 - Comparison of Nvivo and BERTopic themes for the CCBE considerations of the Legal Aspects of AI

PB_L_04 CCBE considerations of the Legal Aspects of AI		
No. of themes	21(BERTopic) < 53 (NVivo)	
Overlaps	7 themes: 1. Liability 2. Use of AI in Courts 3. Use of AI by Lawyers 4. Use of AI in Criminal law 5. Confidentiality 6. Rights 7. Explainability	
Divergences	Unique to BERTopic	Unique to NVivo
	1. Jurisdictions 2. Legal considerations 3. AI in criminal proceedings 4. Balance 5. Justice 6. Surveillance 7. Tools 8. Language Processing 9. Weak Algorithms 10. Ethics 11. Reasoning 12. Compliance 13. Legal Principles 14. Principles	1. Adversarial debate 2. Digital forensics 3. Algorithm 4. Autonomy of Devices 5. Bias 6. Blackbox Phenomenon 7. Challenges 8. Civil Rights 9. Compliance 10. Ethical Success 11. Legal Success 12. Operational Success 13. Societal Success 14. Democratic Principles 15. Discrimination 16. Duty of Competence 17. E-Discovery 18. Fairness 19. Hidden Bias 20. Human Oversight 21. Human Rights 22. Human-centric approach 23. Inherent bias 24. Neutrality 25. NLP 26. Obligations 27. Predictive AI 28. Privacy and Data protection



		29. Proportionality 30. Reliability 31. Risk 32. Rule of Law 33. Software creation 34. Strong AI 35. strong-sense interpretability 36. Supportive AI 37. Transparency 38. Trustworthy AI 39. Equality of Arms 40. Exploitation of conclusions 41. Principle of Impartiality 42. Weak AI
--	--	--

Table A.3 Comparison of Nvivo and BERTopic themes for the CCBE Guide on the use of Artificial Intelligence-based tools by lawyers and law firms in the EU

PB_L_05 CCBE Guide on the use of Artificial Intelligence-based tools by lawyers and law firms in the EU		
No. of themes	21(BERTopic) < 51 (NVivo)	
Overlaps	7 themes: 1. Bias 2. Risks 3. Obligations 4. Interpretability 5. Use of AI by lawyers 6. Transparency 7. Competences	
Divergences	Unique to BERTopic	Unique to NVivo
	1. Tools 2. Law firm 3. Lawyers 4. Opportunities 5. Benchmark 6. Security 7. Interpretability 8. Analytics 9. Argumentations 10. Capabilities	1. Accessibility 2. AI literacy and education 3. Benchmark 4. Blackbox Phenomenon 5. Challenges 6. Cloud Computing 7. Compliance 8. Confidentiality 9. Consent 10. Operational Success



	11. Measurability 12. Organisations 13. Reusability 14. Retrieval 15. Relevancy 16. Legal reasoning 17. Verification	11. Societal Success 12. Democratic Principles 13. Discrimination 14. Distributed Ledger Technology (DLT) 15. Explainability 16. Human dignity 17. Human Oversight 18. Human Rights 19. Human-centric approach 20. Legal knowledge representation 21. Machine Learning 22. NLP 23. Obligations 24. Deontology 25. Privacy and Data protection 26. Reliability 27. Rule Based Approach 28. Rule of Law 29. strong-sense interpretability 30. Supportive AI 31. Trained Model 32. Real-time Training 33. Supervised learning 34. Unsupervised Training 35. Trustworthy AI 36. Types of Justice 37. analytical justice 38. informative justice 39. Argumentation mining 40. Semantic search 41. predictive justice 42. Usability 43. Functionality of AI tools 44. Use of AI in Courts 45. Use of AI in Law firms
--	--	--

Table A.4 Comparison of Nvivo and BERTopic themes for the UIA Guidelines on the use of AI by lawyers

PB_L_12 UIA Guidelines on the use of AI by lawyers	
No. of themes	4(BERTopic) < 21(NVivo)



Funded by
the European Union

Overlaps	1 theme: Confidentiality	
Divergences	Unique to BERTopic	Unique to NVivo
	- Risks - Regulations - Ethics	- AI literacy and education - Competences - Compliance - Definition of AI - Ethical Success - Legal Success - Operational Success - Societal Success - Ethical Success - Legal Success - Operational Success - Fairness - Human Rights - Human-centric approach - Transparency - Unintended Effects - Usability - Use of AI by lawyers

Table A.5 Comparison of Nvivo and BERTopic themes for the UNESCO Draft guidelines for the use of AI in Courts and Tribunals

PB_L_13 UNESCO Draft guidelines for the use of AI in Courts and Tribunals		
No. of themes	16 (BERTopic) < 44 (NVivo)	
Overlaps	6 themes: - Transparency - Use of AI in Courts - Proportionality - Contestability - Explainability - Privacy	
Divergences	Unique to BERTopic	Unique to NVivo



	1. Justice 2. Benefits 3. Accuracy 4. Justice 5. Ethics 6. Principles 7. Rights 8. Standards 9. Governance 10. Liability	1. Accessibility 2. Accountability 3. AI literacy and education 4. Bias 5. Challenges 6. Compliance 7. Confidentiality 8. Consent 9. Ethical Success 10. Fairness 11. Human dignity 12. Human friendly technology 13. Human Oversight 14. Human Review 15. Human Rights 16. Human-centric approach 17. IP rights 18. Multi-stakeholder governance 19. NLP 20. Operational Success 21. Possible solutions 22. Precautionary principle 23. Reliability 24. Rule Based Approach 25. Rule of Law 26. Security 27. Society and Societal impact 28. Societal Success 29. Supportive AI 30. Trustworthy AI 31. Unintended Effects 32. Usability 33. Use of AI by lawyers 34. Environmental Impact 35. Equality of Arms 36. Use of AI in Law firms
--	---	---



**Funded by
the European Union**

2. Tables Demonstrating the Comparison of Nvivo and BERTopic for Healthcare Associations

Table A.6 Comparison of Nvivo and BERTopic themes for the EFPIA Position on the use of artificial intelligence in the medicinal product lifecycle

PB_H_1 EFPIA Position on the use of artificial intelligence in the medicinal product lifecycle		
No. of themes	12 (BERTopic) < 37 (NVivo)	
Overlaps	7 themes: 1. Risk 2. Oversight 3. Research & Development 4. Innovation 5. Regulatory Framework 6. AI Training 7. Transparency	
Divergences	Unique to BERTopic	Unique to NVivo
	1. Governance 2. Collaboration 3. Cooperation 4. Regulation	1. Access to data 2. AI Definition 3. Balanced approach 4. Benefit-risk assessment 5. Competitiveness 6. Compliance 7. Credibility 8. decision-making processes 9. Efficiency 10. EU AI Act 11. Existing Frameworks 12. fit-for-purpose 13. good principles 14. Harmonization 15. healthcare delivery 16. high-risk AI systems 17. integrity of data 18. Investment 19. IP Rights 20. medicines development 21. nonduplicative 22. Regulatory oversight 23. Patient Interest 24. privacy 25. Proportionality 26. regulatory sandboxes



Funded by
the European Union

		27. Safety 28. Stakeholder engagement 29. third party conformity assessment 30. Trade Secret 31. trust
--	--	--

Table A.7 Comparison of Nvivo and BERTopic themes for the CPME Deployment of artificial intelligence in healthcare

PB_H_04 CPME Deployment of artificial intelligence in healthcare		
No. of themes	15 (BERTopic) < 28 (NVivo)	
Overlaps	5 themes: 1. Research & Development 2. Privacy 3. AI Literacy 4. AI Training 5. Efficiency	
Divergences	Unique to BERTopic	Unique to NVivo
	1. Data 2. Professional 3. Limitation 4. Rights 5. Policy 6. Investment 7. Deployment 8. Healthcare	1. Access to Healthcare 2. accuracy 3. Bias 4. Compliance 5. Cybersecurity 6. data protection 7. decision-making processes 8. Diagnosis and Treatment 9. Doctor Autonomy 10. explainability 11. Freedom to Opt-out 12. improvement of clinical practice 13. interoperability 14. knowledge transfer 15. medical ethics 16. patient autonomy 17. Patient Interest 18. patient-doctor relationship 19. reliability of datasets 20. Retention of Healthcare Professionals 21. Risk-based Approach 22. trust 23. workflow integration



3. Tables Demonstrating the Comparison of Nvivo and BERTopic for Engineering Associations

Table A.8 Comparison of Nvivo and BERTopic themes for the Nordic engineers' stand on Artificial Intelligence and Ethics

PB_E_01 Nordic engineers' stand on Artificial Intelligence and Ethics		
No. of themes	15 (BERTopic) < 62 (NVivo)	
Overlaps	10 themes: 1. Accountability 2. AI Training 3. Autonomous 4. Bias 5. Cooperation 6. Ethics 7. Risk/Risk Mitigation 8. Responsibility 9. Societal Expectation 10. Transparency	
Divergences	Unique to BERTopic	Unique to NVivo
	1. Adaptability 2. Political support 3. Framing 4. Guidelines 5. Predictability	1. AI Definition 2. AI Development 3. AI Governance 4. AI Literacy 5. alignment in human values 6. Bias mitigation 7. Gender bias 8. Racial bias 9. codes of conduct 10. Complementary AI 11. cultural values 12. Democratic AI 13. Education for ethical considerations 14. Effectiveness 15. fairness 16. governmental oversight 17. Harm limitation 18. Human benefit 19. human decisions 20. human judgement 21. Human Rights 22. human-centered 23. individual responsibility



Funded by
the European Union

		24. individual rights 25. Innovation 26. integrity 27. Intelligent 28. Investment 29. large-scale data collection 30. Machine learning 31. Oversight 32. political agenda 33. Power Dynamics 34. Prediction 35. privacy 36. Regulatory framework 37. research and development 38. Engineer Responsibility 39. Institutional responsibility 40. Legal responsibility 41. Organizational responsibility 42. Safety 43. Stakeholder Engagement 44. standards 45. Sustainable AI 46. Explicability 47. Interpretability 48. Traceability 49. trust 50. utilitarian 51. Vulnerable groups 52. Workforce Diversity
--	--	--

Table A.9 Comparison of Nvivo and BERTopic themes for the EFCE White Paper on Artificial Intelligence in Chemical Engineering

PB_E_03 EFCE White Paper on Artificial Intelligence in Chemical Engineering	
No. of themes	20 (BERTopic) < 45 (NVivo)
Overlaps	9 themes: 1. Accuracy 2. AI Literacy 3. Bias 4. Collaboration 5. Data Quality 6. Explicability 7. Standard/standardization



Funded by
the European Union

	8. Optimization 9. Safety	
Divergences	Unique to BERTopic	Unique to NVivo
	1. AI Training 2. Availability 3. Challenge 4. Confidentiality 5. Education 6. Expertise 7. Functionality 8. Machine learning 9. Production 10. Process 11. Solubility	1. Access to data 2. Accountability 3. advancing industrial processes 4. AI Efficiency 5. Autonomous systems 6. cost optimization 7. Creativity 8. data integrity 9. Democratic AI 10. Effectiveness 11. error correction 12. ethics 13. Feasibility 14. human decisions 15. Innovation 16. Integration across tasks 17. Interoperability 18. Model Architecture 19. pattern recognition 20. performance 21. practical applications 22. Prediction 23. privacy 24. Process Design 25. Regulatory framework 26. Responsibility 27. Institutional responsibility 28. Organizational responsibility 29. security concerns 30. Sustainable AI 31. Technical functionality 32. transparency 33. Interpretability 34. trust 35. user-friendliness 36. Workforce Diversity

