

COVER PAGE

Title: The AI ethical network and its implementation:
embedding high-level principles into the socio-technical design of AI systems.

Author: Nicola Palladino,

Affiliation: Trinity College Dublin

Submitted to: Communication Policy and Technology Section (CPT)

EXORDO Submission ID: #2357

Introduction

“Artificial Intelligence” (AI) is a label used as shorthand for an expanding “family” of software (and hardware) systems capable of performing specific tasks by collecting, analyzing, and interpreting data, sometimes perceiving the environment in which they operate, to make decisions and take actions with a certain degree of autonomy (Russel and Norving 2003). AI is increasingly crucial in everyday life and social relations, which raises both expectations on AI’s capacity to foster human well-being as well as concerns about the risks to human autonomy and integrity (Renda 2019; Boiler 2018).

On the one hand, AI can help address the complexity of the modern world by optimizing decision-making processes and resource allocation, with relevant effects in production, transport, crisis management, environment, and health care. On the other hand, AI systems raise concerns ranging from privacy and data protection to discrimination, manipulation, misinformation, or the endangerment of democratic institutions and effects on jobs and rights in the workplace.

Stakeholders matured the awareness that the full potential of this technology is attainable only by building a trustworthy and human-centric framework, meaning that AI systems must be aligned with legitimated societal values and governed through accountable arrangements to avoid both misuse of AI applications capable of endangering people and underuse because of a lack of public acceptance (as occurred with nuclear power or GMOs)¹.

Not by chance, in the past few years, we have witnessed to a flourish of initiatives setting ethical codes and good governance principles for AI development. Attempts to map and review the AI ethics landscape have pointed out the convergence on a narrow set of principles (Jobin 2019; Floridi and Cowls 2021; Berkman Center 2020).

This convergence around an emerging AI ethical discourse appears intriguing. Discourses conceived as a “shared set of concepts, categories, ideas that provide its adherents with a framework for making sense of situations, embodying judgments and fostering capabilities” (Dryzek, 2006: 1) could play a pivotal role in norm creation and regime formation at the national and international level (Palladino 2021).

As noted discourses favor the “the convergence of multistakeholder (e.g., governments, private sector, and civil society) expectations” around a “sets of implicit or explicit principles, norms, rules and decision-making procedures” (Cogburn 2017: 27–8; Krasner 1983; Finnemore and Sikkink 1998). When discourses are produced, reproduced, and solidified in a myriad of communicative interactions and deliverables/outputs (such as statements, declarations, guidelines, and recommendations) by transnational networks of experts and practitioners, they frame an issue around salient aspects and causal linkages, enhancing or excluding possible solutions and outcomes (Palladino and Santaniello 2021). In doing so, they could lead to the development of a uniform system of norms and rules and establish a transnational mode of regulation by their incorporation into the policymaking agenda of national and international institutions (Carayannis et al. 2012; Padovani and Musiani 2010; Khagram et al. 2002).

In this view, ethical frameworks can provide a set of legitimized normative criteria from which more detailed rules and provisions can be derived, introducing a minimum of order and guidance to a messy and rapidly changing scenario. Ethical frameworks also influence the creation of coercive mechanisms of governance, informing first steps at the national and

¹ Floridi, L., Cowls, J., Beltrametti, M., et al. (2018) AI4People—An Ethical Framework for a Good AI Society: Opportunities, Risks, Principles, and Recommendations. *Minds & Machines* 28, 689–707. See also European Commission (2020) Public Consultation on the AI White Paper Final Report, <https://ec.europa.eu/digital-single-market/en/news/white-paper-artificial-intelligence-public-consultation-towards-european-approach-excellence>

regional levels when establishing mandatory requirements for AI development and deployment, as well as related punishments in the case of violations.

Nevertheless, some considerations advise caution about the possibility that the ethical turn in the AI field is effectively leading to a common, transnational regulatory framework.

Empirical investigations of AI practitioners' attitudes toward ethical approach revealed that "the sheer number of frameworks to choose from", their "abstraction" and "the lack of clarity provided by principles", which "mean different things to different people in different contexts" could cause confusion and undermine the implementation of AI principles (Morley et. Al. 2021: 6, IEEE 2021).

This study aims to shed light on these inconsistent results by introducing some innovation in the investigation of the AI ethical landscape and realizing a data-driven mapping of the frameworks proposed so far.

Unlike previous attempts, which just focused on the frequency and distribution of ethical instances arbitrarily grouped into 'silo' categories, this study pays attention to their different degrees of abstraction and interconnection.

For this purpose, content analysis has been conducted distinguishing between values (the goals to be reached in developing AI systems), principles (normative criteria guiding action), and requirements (conditions to be satisfied to fulfill principles).

Then, the result of the coding activity has been employed to perform a semantic network analysis with a twofold purpose. On the one hand, the analysis of the co-occurrences between different normative items is expected to 'validate' and get insight to refine the typology created by coding the corpus. On the other side, the analysis of the co-association between contents and documents (and their features) is supposed to shed light on the existence of a common discourse and its articulations.

Findings point out an underlying structure of normative elements in experts' and stakeholders' discourses around AI ethics giving raise to a strictly interrelated network based on a narrow set of normative items, especially principles, but still not a common normative framework. Results, indeed also suggest that principles could be articulated in more specific sub-dimensions or requirements, defining the specific meaning that they assume for different actors or contexts. Moreover, the analysis allowed to distinguish three different clusters of associated contents and documents/actors corresponding to as many patterns of norms production.

Insights from the empirical analysis will be discussed against the current tools and regulatory approaches to deal with AI ethical challenges.

2. Methodology

2.1 Data Collection

Eligible documents for this analysis have been detected according to the following selection criteria: i) documents should have an explicit prescriptive purpose, aiming at setting appropriate behavioral norms; ii) they must be drafted by a collective entity (no individual contributions allowed); iii) they must provide an overview of the ethical issues raised by AI development, deployment, and management (documents focused on just one issue, for example, privacy, have been ruled out).

The searching strategy included retrieving documents from previous similar mapping exercises (Jobin et al. 2019; Fjeld et al 2020; Shiff et. Al 2021, Greene et al. 2019) and assessing their consistency with the abovementioned criteria. Such preliminary corpus has been integrated scrutinizing the OECD, COE, and Algorithm Watch public accessible repository (respectively

<https://oecd.ai/en/dashboards;> [https://inventory.algorithmwatch.org/;](https://inventory.algorithmwatch.org/) <https://www.coe.int/en/web/artificial-intelligence/national-initiatives>) and other items mentioned in literature.

It is worth noting that many entries in the abovementioned sources have been excluded by the final corpus because of the resource was not accessible at the provided link or inconsistent with the adopted selection criteria.

2.2 Coding

The corpus obtained according to the procedure illustrated in the previous paragraph has undergone qualitative coding mixing deductive and inductive approaches (Saldaña 2013) through Nvivo software (Mayring 2019, Kaefer et al. 2015).

A manual coding approach has been preferred because the documents comprising the corpus use the same labels to indicate different concepts, or concepts with a different semantic extension; and conversely, many concepts are referred resorting to different terms. A mere computational textual analysis would not be able to deal with this polysemy issue, flawing following statistical analysis.

A coding scheme has been developed in the first instance by providing unambiguous definitions and descriptions of normative categories identified in the literature and previous mapping initiatives, in particular, the Berkman Center report, which offered a valuable systematization. In so doing the author was aimed at both providing clarity and making the most of previous classificatory initiatives, without reinventing the wheel. The initial inputs have been revisited paying attention to minimizing redundancy and maximizing hierarchical classification. Of course, the choice of which label has to be associated with which concept has been affected by an unavoidable degree of arbitrariness and inconsistency between the different sources. However, in order to minimize this inconvenience most common terminologies have been preferred when possible.

The coding scheme has been iteratively revisited during the coding operation in order to integrate new entries emerging from the analyzed texts and eventually merge or split categories.

Noteworthy, normative categories so identified have been organized hierarchically, distinguishing between values, principles, and requirements. Values indicate the general goal that actors want to fulfill in developing AI systems and applications, both in the sense of what actors want to reach *by* using AI or what must be safeguarded *while* using AI. Values are properties or states of the society potentially affected by AI systems. Instead, principles reflect normative expectations on the functioning of AI systems. They indicate the rules that must guide AI systems' actions to reach the intended values/goals, and in this sense, they are properties of AI systems themselves. Requirements could be defined as conditions to be satisfied to implement principles. In other words, they could be conceived as subdimensions of principles at a lower degree of abstraction. It is worth noting that to comply with a principle is not necessary that all potential requirements are satisfied. Requirements could refer to different paths to ensure AI system with the properties to reach the values/goals. They could be complementary (i.e inclusive design enhances bias prevention) or alternative (as in the case of data minimization and anonymization for privacy). Some of them could be more agreed than others or preferred by specific group of actors, or to be more appropriate, relevant or urgent depending by the context.

In order to reduce subjective bias, samples of coding output have been periodically reviewed by a panel of external experts who checked the conformity of coding with the definition

provided in the coding scheme, pointed out incongruences and doubtful cases and advised on the treatment of novel emerging categories.

2.3 Data Analysis

Semantic network analysis is a particular kind of social network analysis, focused on semantic elements in a text (or a corpus of texts) such as words, concepts, or themes. In this kind of network words or concepts are the actors connected to each other by “ties” their proximity or co-occurrence in the same text or portion of text. The basic assumption here is that the meaning of a text, and, in the truth, of the semantic elements themselves, emerges from the way in which semantic elements are interrelated and the structures they give rise (Segev 2021).

Several social network analysis studies attempt to identify patterns in the strength and distribution of the ties between nodes (Granovetter, 1977), or consider the position of nodes within the network, and densely connected clusters (Borgatti et al., 2013).

For this research, the coding results obtained through NVIVO has been transposed into two matrices. The first one has documents in rows and normative items in columns.

The second one was a squared matrix reporting the frequency at which two normative item co-occurred within the same paragraphs.

Both of them have been analyzed by resorting to the software Gephi, which provided basic information and metrics on the networks (degree, network diameter, graph density, etc.), visualization algorithms and clustering functions (more detail in the findings section).

3. Findings

The coding activity allowed identifying more than 80 different normative elements. After aggregating semantically overlapping items and excluding the ones with the lower frequency, this number has been reduced to 49 items, which have been employed for the semantic network analysis.

Figure 1 organizes coded normative items into an analytical schematization composed of 9 values and 6 principles articulated into 30 requirements. It is an attempt to provide an ideal-typical representation of the normative material found in the corpus. It includes only 45 items, cutting-off aggregate or less frequent normative elements. Moreover, as shown in Table 1 (where N indicates the number of documents in which the principles/requirements have been referred), it does not consider a small number of items enough relevant and frequent. However, this analytical scheme still accounts for the most part of the variance in the corpus and it represents a valid map to navigate the AI ethical landscape.

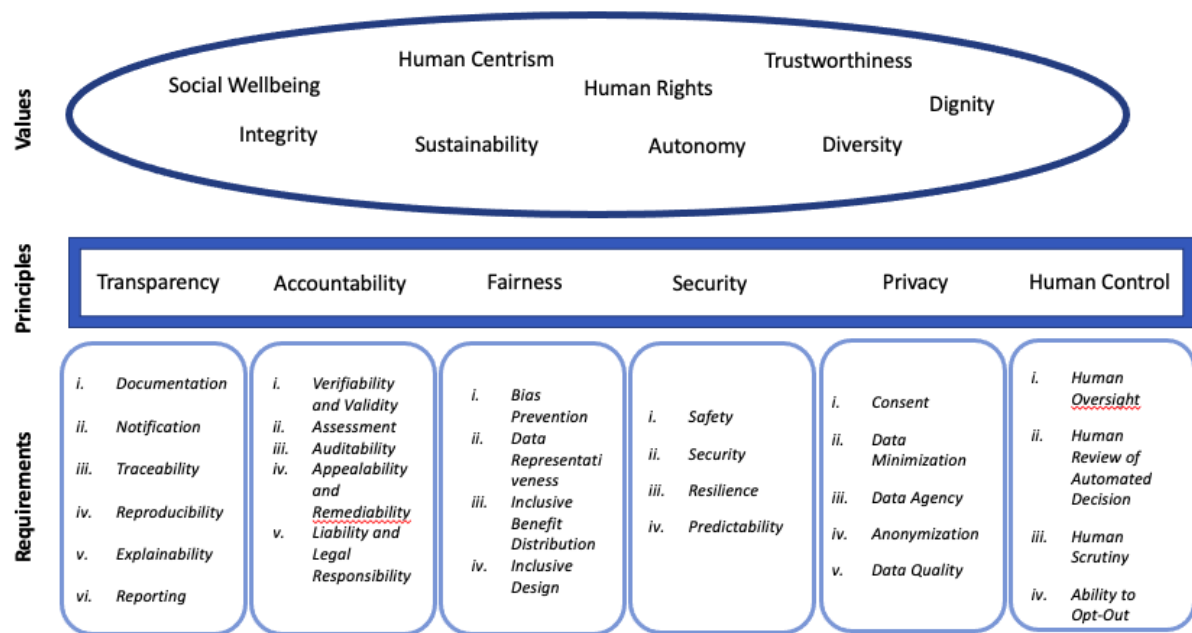


Figure 1 Normative Contents: Values, Principles, and Requirements

Table 1: Values, Principles and Requirements

		N	%		N	%		N	%
Values	Societal Wellbeing	32	52.5	Human Centricism	22	36.1	Sustainability	15	24.6
	Human Rights	31	50.8	Dignity	21	34.4	Trustworthiness	15	24.6
	Diversity	25	41.0	Autonomy	17	27.9	Integrity	12	19.7
Principle Requirement	Transparency	54	88.5	Accountability	40	65.6	Fairness	39	63.9
	Explainability	50	82.0	Liability and Legal Responsibility	31	50.8	Bias Prevention	48	78.7
	Documentation	35	57.4	Assessment	24	39.3	Inclusive Design	21	34.4
	Notification	23	37.7	Appealability and Remediability	21	34.4	Data Representativeness	14	23.0
	Traceability	20	32.8	Auditability	20	32.8	Inclusive Benefit	11	18.0
	Access to Data and Code	8	13.1	Verifiability and Validity	16	26.2			
	Reporting	7	11.5						
	Reproducibility	6	9.8						
	Privacy	46	75.4	Reliability	22	36.1	Human Control	30	49.2
	Data Agency	19	31.1	Security	37	60.7	Human Review of Automated Decision	15	24.6
	Consent	13	21.3	Safety	33	54.1	Human Scrutiny	13	21.3
	Data Quality	7	11.5	Predictability	8	13.1	Human Oversight	11	18.0
	Anonymization	5	8.2	Resilience	6	9.8	Ability to Opt-Out	7	11.5
	Data Minimization	3	4.9						
Further Normative Items									
	Multistakeholderism (principle)	19	31.1	Accuracy (requirement)	6	9.9	Intellectual Property Protection (principle)	4	6.5
	No Surveillance	4	6.5						

Table 1 also points out how discourses around AI ethics converge on very few items, namely societal wellbeing, transparency, explainability, fairness, bias prevention, accountability,

liability, privacy and security, while other items are much less referred. It is worth noting that the convergence mainly occurs around principles that are general normative criteria, whose meaning depends on the requirements in which they are articulated. In some cases (such as bias prevention for fairness, or security for reliability, explainability for transparency) the convergence occurs also, or even more, on a specific requirement. However, it should be considered that the meaning of principle relies on the combination of requirements taken into account, and this configuration may vary from one actor to another. Broadly speaking, actors share a common ethical network composed of a relatively narrow set of normative items. They also converge on some general guiding principles. But it can hardly be said they share also a common normative framework.

These findings are supported and expanded by semantic network analysis illustrated in Figure 2 and Figure 3.

To understand the graphs, it must be considered that:

- 1) the layout has been produced by the means of the 'Force Atlas' algorithm, which "simulates a physical system in order to spatialize a network. Nodes repulse each other like charged particles, while edges attract their nodes, like springs." [...] The force-directed drawing has the specificity of placing each node depending on the other nodes [...] the technique has the advantage of allowing a visual interpretation of the structure. Its very essence is to turn structural proximities into visual proximities" (Jacomy et al. 2014:2). Put otherwise, the graph provides an overview of the mutual co-associations among the nodes in the network.
- 2) the size of each node depends on the number of connections with other nodes;
- 3) the connection thickness is calculated on the number of times two items have been connected;

Figure 2 represents a bipartite network connecting documents and normative items, and it is based on a matrix reporting if a particular normative item is mentioned in a particular document (1 mentioned; 0 not mentioned). It confirms that discourses on AI ethics focus on a small number of principles and requirements (which size it's bigger inasmuch they are mentioned in a larger number of documents, and they are close to each other because mentioned in the same documents). But probably the most relevant indication we can be drawn from this graph is that all these AI ethics initiatives give rise to a unique network without any separated component. It means the discourses of the ethical initiatives analyzed tend to overlap each other to some degree. It could be said they share a common horizon, or at least a common lexicon or repertoire of normative themes and arguments, even if they could be combined into different specific configurations.

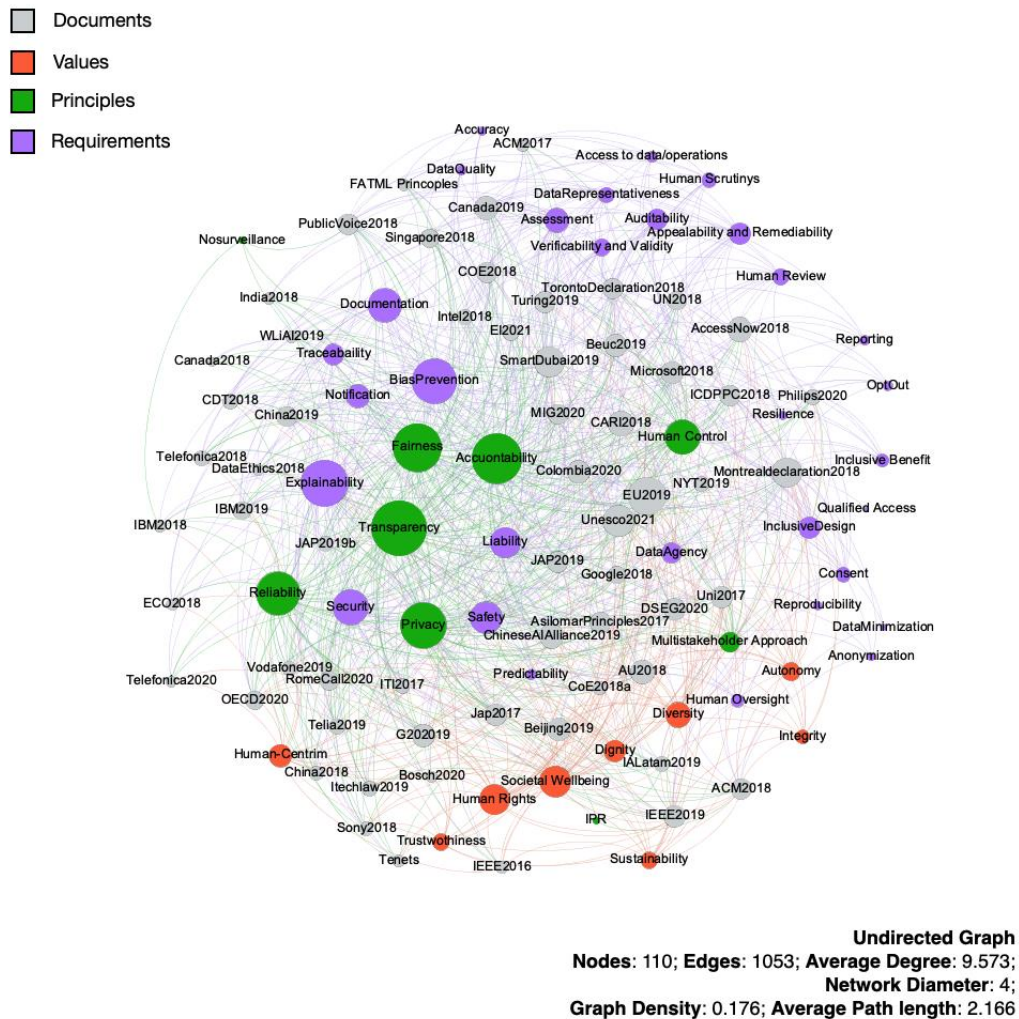


Figure 2 Documents X Normative Item Network

Figure 3 represents the graph of the normative elements coded in the corpus, as obtained from a squared co-occurrence matrix reporting how many times two normative items are mentioned in a same paragraph.

To understand the graph it should be also considered that the different colors partition the network into different communities of densely connected nodes, according to the ‘modularity’ algorithm. The algorithm identifies groups of nodes that are highly connected internally and scarcely connected to each other (Borgatti et al. 2013).

It can be said the network analysis confirms the validity of the classificatory scheme illustrated in Figure 1, both in terms of the distinctions between values, principles, and requirements and in terms of their grouping, with very few mismatches.

Modularity analysis indeed identifies seven ‘communities’. Six of them correspond almost entirely to the six sets of principles and related requirements identified while building the codebook. The remaining one gathers the normative elements previously classified as values.

The isolation of a values cluster is probably due to the fact that, as mentioned above, these general goals that should be pursued employing AI systems are usually described by resorting to other values and are generally discussed in the same preamble or other introductory paragraphs.

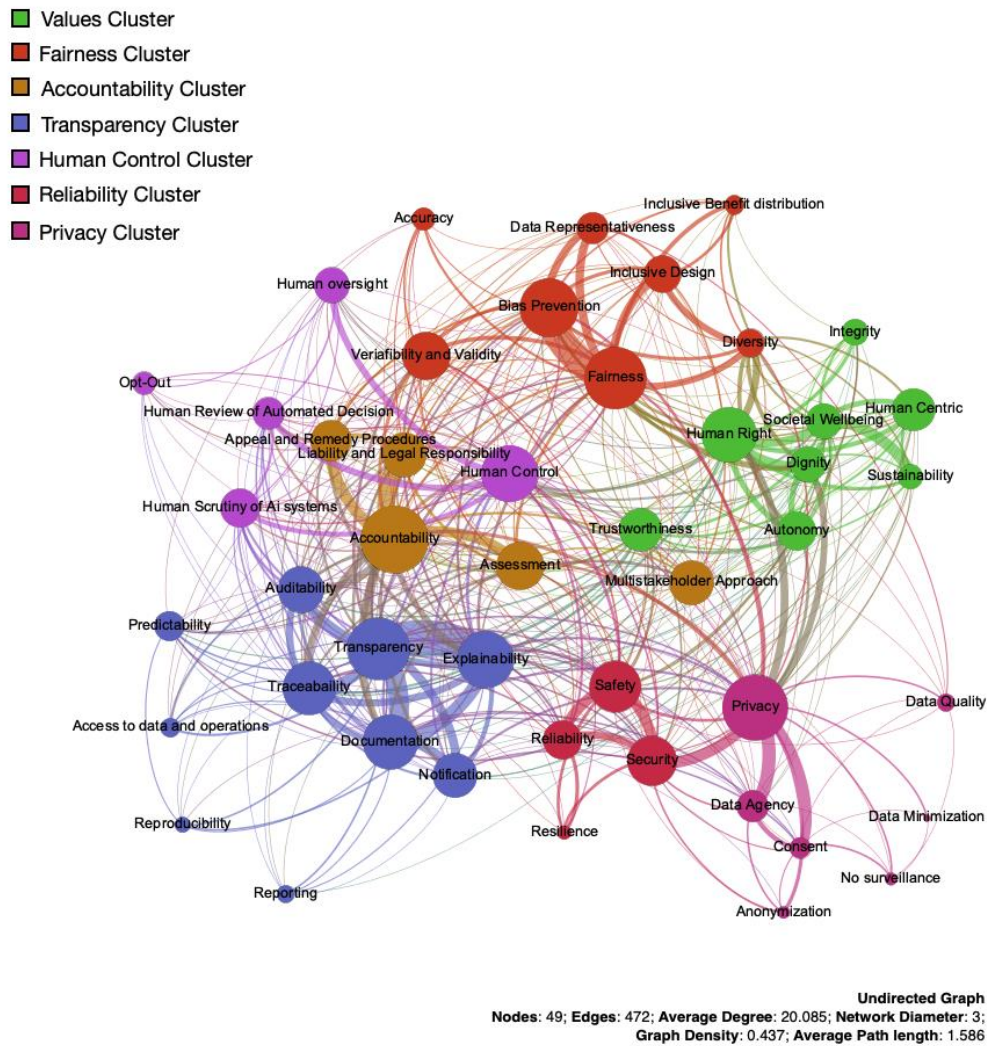


Figure 3 AI Ethical Network

The only exception in this sense is ‘diversity’ that, despite being closer to the values cluster, is part of the fairness nodes community. Furthermore looking at the thickness of concessions also human rights, autonomy and dignity seem to be particularly related to privacy.

The other mismatches also provide some useful insight. The clustering algorithm did assigned the requirements ‘Verifiability and Validity’ and ‘Auditability’ to Accountability principle as expected according to the classification adopted in this study, but rather respectively to Fairness and Transparency. In the first case because verification and validation are also called into question as processes that should be put in place to prevent biases. In the second one because traceability and documentation are often referred to as essential prerequisites for auditability. Rather than put into question the validity of the classificatory scheme adopted in this work, these mismatches seem to indicate how AI ethical principles are intertwined with each other. This latter observation is supported also by the node distributions in the layout. In particular Accountability, Transparency and Human Control look very close and their respective communities overlap in the same space. Also, Privacy and Reliability are strongly tied,

inasmuch as privacy tends to be framed according to a securitarian approach in terms of data protection in many documents.

Broadly speaking, this network points out that the principles cluster are not silos categories. On the one hand, principles are interconnected. They require the same provisions to be accomplished, or to some degree one is the pre-requisite of the other. On the other hand, actors tend to interpret principles in different ways by combining different normative materials within a common repertoire.

4. Discussion and Conclusion

The analysis conducted so far pointed out how the discourses around AI ethics have given rise to a network of recurring normative themes and arguments, according to a well-defined structure in which a set of values to be reached in developing AI system relate to six main principles which should guide the functioning of AI system, each of them articulated into a series of requirements.

Although actors converge on a relatively narrow set of normative items it can hardly be said they share also a common normative framework, leading to a convergence of expectations and the emergence of transnational norms. Rather, it could be said they share a common horizon, or at least a common lexicon or repertoire of normative themes and arguments, even if they could be combined into different specific configurations.

Values (what we want to reach with AI) are still loosely defined and badly connected with principles/requirements (how AI systems should act), causing confusion and hurdles at the implementation stage. Different values can require different principles and requirements.

Also principles could be interpreted in many ways by different actors emphasizing some dimensions over the others, or combining different elements.

Further discussion is needed. Being able to specify and hierarchize values, clarifying which normative expectations on how AI systems functioning follow, may facilitate the development of ethical approaches.

This emerging normative framework is being transposed into legislation (such as the European Artificial Intelligence Act, the American Artificial Intelligence Initiative, or the Chinese Algorithmic Recommendation), international technical standards (such as the IEEE P7000 series, ISO/IEC Joint Technical Committee SC 42 and CEN/CENELEC JTC21), and a plethora of tools supposed to deal with AI ethical challenges. While each of these further initiatives, taken individually, contributes to the concrete implementation of AI ethics, there is a risk to reproduce ambiguity and polysemy at a more operative level. Legislations, technical standards and tools referring to a common lexicon but leading to different provisions may increase the uncertainty for developer and deployer and undermine their efforts.

5. Funding And Acknowledgment

The author has received funding from the European Union's Horizon 2020 Research and Innovation Programme under the HUMAN+ COFUND Marie Skłodowska-Curie grant agreement No. 945447

6. References

- Berkman Center (2020_ Principled Artificial Intelligence: Mapping Consensus in Ethical and Rights-based Approaches to Principles for AI: <https://cyber.harvard.edu/publication/2020/principled-ai>
- Blondel, V. D., Guillaume, J. L., Lambiotte, R., & Lefebvre, E. (2008). Fast unfolding of communities in large networks. *Journal of statistical mechanics: theory and experiment*, 2008(10), P10008.
- Boiler, G., (2018) *Artificial Intelligence: The Great Disruptor*. Washington, DC: The Aspen Institute.
- Borgatti, S. P., Everett, M. G., & Johnson, J. C. (2013). *Analyzing social networks*. Sage.
- Carayannis, E. G., Pirzadeh, A., & Popescuet, D. (2012). Institutional learning and knowledge transfer across epistemic communities. New York: Springer.
- Celeste E., Palladino N., Yilma K., Redeker D. (2022) “Digital Constitutionalism and Online Content Governance”; in Celeste E., Heldt A., Iglesias-Keller C., (eds.) *Constitutionalising Social Media*, Hart Publishing, London.
- Cogburn, D. (2017). *Transnational advocacy networks in the information society*. New York: Palgrave MacMillan.
- Dryzek, J. S. (2006). *Deliberative global politics: Discourse and democracy in a divided world* (p. 45). Cambridge: Polity.
- Finnemore, M., & Sikkink, K. (1998). International norm dynamics and political change. *International organization*, 52(4), 887-917.
- Floridi, L., & Cowls, J. (2021). A unified framework of five principles for AI in society. In *Ethics, Governance, and Policies in Artificial Intelligence* (pp. 5-17). Springer, Cham.
- Floridi, L., Cowls, J., Beltrametti, M., et al. (2018) AI4People—An Ethical Framework for a Good AI Society: Opportunities, Risks, Principles, and Recommendations. *Minds & Machines* 28, 689–707
- Greene, D., Hoffmann, A.L., and Stark, L. (2019) Nicer, Clearer, Fairer: A Critical Assessment of the Movement for Ethical Artificial Intelligence and Machine Learning. In *Proc. 52nd Hawaii International Conference on System Sciences*, 2122–2131
- Haas, P. M. (2015). *Epistemic communities, constructivism, and international environmental politics*. Routledge.
- Hallensleben et al. (2020) From Principles to Practice. An Interdisciplinary Framework to Operationalise AI Ethics. Report, AI Ethics Impact Group: <https://www.ai-ethics-impact.org/resource/blob/1961130/c6db9894ee73aefa489d6249f5ee2b9f/aieig---report---download-hb-data.pdf>
- IEEE 2021, Addressing Ethical Dilemmas in AI: Listening to Engineers, <https://standards.ieee.org/initiatives/artificial-intelligence-systems/ethical-dilemmas-ai-report.html>

Jacomy M, Venturini T, Heymann S, Bastian M (2014) ForceAtlas2, a Continuous Graph Layout Algorithm for Handy Network Visualization Designed for the Gephi Software. PLoS ONE 9(6): e98679. doi:10.1371/journal.pone.0098679

Jobin, A., Ienca, M., and Vayena, E. (2019) The Global Landscape of AI Ethics Guidelines. *Nat Mach Intell*, 1, 389–399;

Kaefer, F., Roper, J., & Sinha, P. (2015). A software-assisted qualitative content analysis of news articles: example and reflections. In *Forum Qualitative Sozialforschung/Forum: Qualitative Social Research* (Vol. 16, No. 2, p. 20). DEU.

Khagram, S. (2002). Restructuring the global politics of development: the case of India's Narmada Valley Dams. *Restructuring world politics: Transnational social movements, networks, and norms*, 206-30.

Khagram, S. (2002). Restructuring the global politics of development: the case of India's Narmada Valley Dams. *Restructuring world politics: Transnational social movements, networks, and norms*, 206-30.

Krasner, S. D. (Ed.). (1983). *International regimes*. Cornell University Press.

Mayring, P.: Qualitative content analysis: demarcation, varieties, developments. *Forum Qual. Sozialforschung Forum Qual. Soc. Res.* (2019). <https://doi.org/10.17169/fqs-20.3.3343>

Morley, J., Floridi, L., Kinsey, L. et al. (2020) From What to How: An Initial Review of Publicly Available AI Ethics Tools, Methods and Research to Translate Principles into Practices. *Sci Eng Ethics* 26, 2141–2168.

Padovani, C., Musiani, F., & Pavan, E. (2010). Investigating evolving discourses on human rights in the digital age: Emerging norms and policy challenges. *International Communication Gazette*, 72(4–5), 359–378.

Palladino N. and Santaniello M. (2020) Legitimacy, Power and Inequalities in Multistakeholder Internet Governance. Palgrave MacMillan: Cham.

Palladino, N. (2021). The role of epistemic communities in the “constitutionalization” of internet governance: The example of the European Commission High-Level Expert Group on Artificial Intelligence. *Telecommunications policy*, 45(6), 102149.

Renda, A. (2019) Artificial Intelligence, Ethics Governance and Policy Challenges. Brussels: CEPS

Russell, S., and Norvig, P. (2003) Artificial Intelligence: A Modern Approach. Prentice Hall Series in Artificial Intelligence. Upper Saddle River, NJ: Prentice Hall, Pearson Education.

Saldaña, J. *The Coding Manual for Qualitative Researchers* (SAGE, 2013)

Schiff, D., Rakova, B., Ayes, A., Fanti, A., and Lennon, M. (2020) Principles to Practices for Responsible AI: Closing the Gap: <https://arxiv.org/abs/2006.04707v1>.

Santaniello, M., Palladino, N., Catone, M. C., & Diana, P. (2018). The language of digital constitutionalism and the role of national parliaments. *International Communication Gazette*, 80(4), 320-336.

Santaniello, M., De Blasio, E., Palladino, N., Selva, D., De Nictolis, E., & Perna, S. (2016). Mapping the debate on Internet Constitution in the networked public sphere. *Comunicazione politica*, 17(3), 327-354.

Santaniello M., Palladino N. (2022) "Discourse Coalitions in Internet Governance: Shaping Global Policy by Narratives and Definitions". In: Marzouki, M., Calderaro, A. (eds.) Global Internet Governance as a Diplomacy Issue, Rowman and Littlefield, Lanham, USA.