

AI Model Cards: State of the Art and Path to Automated Use

Ali Mehraj¹^a, An Cao¹^b, Kari Systä¹^c, Tommi Mikkonen²^d, Pyry Kotilainen²^e,
David Hästbacka¹^f and Niko Mäkitalo²^g

¹*Faculty of Information Technology and Communication Sciences, Tampere University, Tampere, Finland*

²*Faculty of Information Technology, University of Jyväskylä, Jyväskylä, Finland*
fi

fi

Keywords: Model Cards, AI Model, Privacy, Transparency, Ethical AI, Machine Learning, ML, Artificial Intelligence, AI.

Abstract: In software engineering, the integration of machine learning (ML) and artificial intelligence (AI) components into modern web services has become commonplace. To comply with evolving regulations, such as the EU AI Act, the development of AI models must adhere to the principles of transparency. This includes the training data used, the intended use, potential biases, and the risks associated with these models. To support these goals, documents named Model Cards were introduced to standardize ethical reporting and allow stakeholders to evaluate models based on various goals. In our ongoing research, we aim to automate risk analysis and regulatory compliance checks in software systems. We envision that model cards can serve as useful tools to achieve the goal. Given the evolving format of model cards over time, we conducted a state-of-the-art review of the current state and practice of model cards by analyzing 90 model cards from four model repositories to assess their relevance to our vision. The study's contribution is a thorough analysis of the model cards' structure and content, as well as their ethical reporting. Our study reveals the variance in information reporting, the loose structure, and the lack of ethical reporting in the model cards. Based on the findings, we propose a unified model card template that aims to enhance the structure, promote greater transparency, and establish a foundation for future machine-interpretable AI model cards.

1 INTRODUCTION

Rapid advancement of Artificial Intelligence (AI) has significantly impacted the engineering of Internet and Web systems. Machine Learning (ML) models, especially Large Language Models (LLMs), have demonstrated remarkable capabilities in a variety of applications such as text generation, text translation, text comprehension, and image generation. Such models have a vast number of parameters and are trained with large datasets, which can introduce privacy and security concerns and may lead to privacy and security breaches of personal data. With emerging regulations such as the EU AI Act (Edwards, 2021), the Data Act (Perarnaud et al., 2022) and the existing General Data Protection Regulation (GDPR) (Voigt and Von dem

Bussche, 2017), these AI models will need to adhere to the guidelines set forth by the regulations. It is crucial for developers and organizations involved in developing these AI models to be transparent regarding the training data, intended use, potential biases, and the risks associated with their models.

Model Cards, introduced by Google in 2019 (Mitchell et al., 2019) are documents that represent ML models to standardize ethical reporting. These documents allow stakeholders to evaluate models based on inclusiveness, fairness, ethics, and traditional metrics reported in the model cards (Mitchell et al., 2019). While model cards are becoming an integral part of model documentation, their format is continuously evolving (Kotilainen et al., 2024). We have utilized the concept of these model cards in our previous research to ensure safe and ethical orchestration, and deployment of ML models using metadata cards (Kotilainen et al., 2025).

In many systems, ML is part of a larger system. Therefore, to identify potential risks in a system, it is important to analyze the impact associated with them to the overall system. In addition, the risks depend on

^a <https://orcid.org/0009-0006-7563-1706>

^b <https://orcid.org/0009-0005-6460-6567>

^c <https://orcid.org/0000-0001-7371-0773>

^d <https://orcid.org/0000-0002-8540-9918>

^e <https://orcid.org/0000-0002-4645-074X>

^f <https://orcid.org/0000-0001-8442-1248>

^g <https://orcid.org/0000-0002-7994-3700>

the data used by the system and who has access to the overall system. Therefore, when analyzing the entire system where ML models are deployed, it is essential to utilize the information associated with the ML models provided by developers in the form of model cards. An example of such an approach has been described in our earlier work (Kotilainen et al., 2025).

In this paper, we analyze the structure and content of different model cards for their applicability as building blocks for our future work on the automation of some steps in software architecture design and analysis. With all software development tasks moving towards automation, we aim to integrate automated regulatory compliance into the software development process and ensure that software systems are safe and regulatory compliant. In addition, we want to ensure that ML models deployed in software systems are regulatory compliant and do not cause unintended side effects. We believe that model-related information provided by the developers as model cards can be useful tools for designing and developing regulatory compliant software and enable the safe deployment of ML models in software systems. To validate this plan, we assess the current state-of-the-art in the implementation of model cards in practice. This will facilitate a more comprehensive understanding of the model cards' content and their relevance in designing and developing regulatory-compliant software. This state-of-the-art review provides insight into how different developers utilize the concept of model cards and the differences in the model card structure. The review also investigates the applicability of model cards for realizing our vision of automating regulatory compliance in software systems.

The structure of the paper is as follows. In Section 2, we briefly describe the background and motivation for our research, mainly discussing motivating factors and existing research incorporating model cards. In Section 3, we present the methodology and the process for selecting and analyzing model cards for the research. In Section 4, we state the findings from our analysis of the selected model cards. In Section 5, we discuss and analyze the results. We also establish a unified model card template to incorporate our regulation-aware software architecture vision. Finally, in Section 6 we present some final conclusions.

2 BACKGROUND AND MOTIVATION

Model cards offer a systematic approach to documentation of essential aspects and factors of ML models, facilitating comprehension, assessment, and repli-

cation of results for researchers (Dan and Seeuws, 2024). Model Cards were introduced to improve transparency and responsible use of AI. It is achieved by outlining metrics such as accuracy in different datasets, potential biases, and possible ethical problems (Mitchell et al., 2019). The goal of the model cards is to increase transparency by providing detailed information about the intended use, performance metrics, identified limitations, and other considerations taken into account by the developers (Nunes et al., 2022). Model cards are the first step in standardizing ethical practice and reporting for AI models and allows users or stakeholders to compare models before deployment for their specific use case, thereby facilitating decision-making, reducing expenses, and helping users identify the right product in the AI market (Wadhvani and Jain, 2020). Model cards have gained significant traction in the industry, particularly among major organizations engaged in developing ML models, such as Google, Meta, Nvidia, and OpenAI.

The model cards presented in (Mitchell et al., 2019) have 9 sections. The sections and the content of the sections of a model card are listed in Table 1. Among the sections, the *Ethical Considerations* section is the primary source of information for our ongoing research on identifying the potential risks posed by ML models in software systems and potential regulatory compliance.

In practice, model cards have gained widespread acceptance. Hugging Face, a prominent repository for model cards, currently houses over 1.6 million models and their corresponding model cards. Organizations involved in the development of AI/ML models such as, OpenAI, Nvidia, Google, Anthropic, Meta have also adopted the concept and provide model cards for their corresponding ML models. Amazon's AI Service Cards¹ bear similarity to the model cards and appear to draw inspiration from the concept of the model cards. In addition, there are templates and toolkits for creating model cards provided by Google², TensorFlow³, PyPi⁴, and Hugging Face⁵.

The concept of model cards has been used in several domains. A study in the medical field proposed a Checklist for Artificial Intelligence in Medical Imaging (CLAIM) (Mongan et al., 2020). The checklist

¹<https://aws.amazon.com/blogs/machine-learning/introducing-aws-ai-service-cards-a-new-resource-to-enhance-transparency-and-advance-responsible-ai/>

²<https://research.google/blog/introducing-the-model-card-toolkit-for-easier-model-transparency-reporting/>

³<https://github.com/tensorflow/model-card-toolkit>

⁴<https://pypi.org/project/model-card-toolkit/>

⁵<https://huggingface.co/docs/hub/model-card-annotated>

Table 1: Sections of a Model Card (Mitchell et al., 2019).

Model Card section	Section content
Model Details	Basic information regarding the model version, type and other details
Intended Use	Information regarding the intended use of this model, as well as its intended users
Factors	Relevant factors related to the performance of the model
Metrics	The metrics for the performance of the model
Evaluation Data	Information regarding evaluation datasets, selection rationale, and pre-processing steps of the data prior to model assessment
Training Data	Information related to model training datasets and possible training process.
Quantitative Analyses	The results of the evaluation of the model according to the selected metrics
Ethical Considerations	Information regarding privacy, safety, and risks associated with the model. Identifies potential ethical challenges and the solutions that have been proposed for stakeholders
Caveats and Recommendations	Information regarding additional concerns that were not addressed in the preceding sections

corresponds to the information required in a model card. In our previous research, we have successfully employed the model card concept, describing safe and ethical orchestration and deployment of ML models using metadata cards (Kotilainen et al., 2025). Previous research has been conducted on incorporating empirically derived sensor fusion recurrent neural network (RNN) performance and cost measurement data into machine-readable model cards (Booth and Ghosh, 2023). One study employs extensive and standardized terminology and scientific rigor as promoted by biomedical oncologists and establishes a method to make model cards machine-readable using semantic web technology (Amith et al., 2022).

As demonstrated by the examples above, model cards offer a variety of uses and there are attempts to make model cards machine interpretable. However, we have not found any work on making model cards useful for tools and methods used for system-level design. Our long-term objective is to incorporate machine-interpretable model cards into the system design and analysis process. To achieve this, we must enhance our understanding of the content and structure of existing model cards and observe the evolution of model cards in research and practice. In addition, it is essential to assess the applicability of these model cards in achieving our long-term objectives. Although our primary focus is on ethical and regulatory compliant software architecture design and risk analysis, we believe that the insights gained will be beneficial to individuals with different interests related to model cards.

3 RESEARCH APPROACH

In this section, we describe the research questions, model card search strategy, and data extraction from the selected model cards.

3.1 Research Questions

In consideration of the research objectives, three research questions (RQs) were formulated.

RQ1. What kind of information is provided in different model cards? Our primary objective is to understand to which extent the information in the current model cards serves the needs of the envisioned regulatory-compliant architecture design and risk analysis. This includes checking model cards for ethical considerations information.

RQ2. What are the key similarities and differences in the structure of the model cards? Diverse structure requires more adaption and conversion, but can be solved with proper tooling. Discrepancies in critical sections of the content could compromise the viability of analyzing multiple systems with similar tools and approaches. Diverse structures add complexity to automated extraction of information from model cards and future machine-interpretable.

RQ3. How well are the ethical and regulatory aspects covered in the model cards? As we aim to assess risk and regulatory compliance in software systems, it is essential to evaluate the extent to which these aspects are reported in the model cards to ensure their suitability for our purposes.

3.2 Model Card Search

The search for model cards was carried out through publicly accessible model repositories between 04.04.2025 and 10.04.2025. The selection criterion for a model repository was a minimum of 100 models hosted. We found four ML model repositories that hosted a large collection of ML models and had dedicated model cards for all of their models. These model repositories are Hugging Face⁶, Kaggle⁷, Nvidia NGC Catalog⁸, and Google Model Garden⁹. In contrast, our search results revealed that other model repositories, with the exception of these four, had very few models hosted, and the models were not accompanied by a model card.

Our first search was conducted in the largest publicly available model repository, Hugging Face. At the time of conducting the research, Hugging Face hosted over 1.6 million ML models and their respective model cards. The second search was conducted on another large model repository called Kaggle. At the time of conducting the research, Kaggle hosted more than 19,000 models. The third search was conducted using Nvidia's NGC Catalog, which at the time of the research contained more than 800 models. The final search was conducted in Google's Model Garden. Google's Model Garden showcases over 100 models developed by the company.

The selection of model cards from each model repository varies. This is due to the number of models hosted by each repository. As Hugging Face is the largest model repository in our selection, we decided to collect the majority of our sample of model cards from Hugging Face and selected 50 model cards. Kaggle was the second largest model repository in our selection. Given the significantly lower number of model cards in Kaggle compared to Hugging Face, we decided to select 20 model cards from Kaggle. The third and fourth largest model repository in our search were Nvidia's NGC Catalog and Google Model Garden respectively. We selected 10 model cards from each of these repositories. In total, 90 model cards were selected from the repositories.

In order to investigate the state of the art of model cards, we have selected the top downloaded, liked, popular, or trending model cards. This selection was based on the sorting mechanisms available in the repositories. As model sorting was available in Hugging Face for both the top downloaded and top liked

models, we decided to select the top 30 downloaded and the top 20 liked models' cards from the repository. The sorting function in Kaggle was limited to top downloads. Therefore, we selected the top 20 downloaded model cards from Kaggle. The only sorting options in NGC Catalog and Google Model Garden were most popular and trending model cards, respectively. Therefore, we have selected top ten popular and top ten trending model cards from NGC Catalog and Google's Model Garden, respectively. The selected model cards are listed in Table 2.

3.3 Data Extraction

Data extraction was conducted by reading the full content of each selected model card. We collected statistical data for the model cards, including number of sections in each model card and model card sources. In addition to collecting statistical data, we collected data related to our research questions. Regarding RQ1, the different types of information provided in each model card were collected. For RQ2, we collected information related to the differences and similarities in the structure of each model card. For RQ3, we collected ethical and regulatory data reported in the selected model cards.

4 RESULTS

4.1 RQ1. What Information Is Provided in Different Model Cards?

The different types of information reported in the model cards are illustrated in Figure 1. As illustrated in Figure 1, essential information, with the exception of the ethical considerations of the original model card concept outlined by (Mitchell et al., 2019) is represented in blue. Information related to different ethical considerations discovered in our selected model cards is represented in red. It can be observed from Figure 1 that 85 of the 90 selected model cards provide basic details about the model. This includes a detailed description of the model as well as the type of the model. This information is further supported by specifics about the model architecture in some of the model cards. The remaining five models lacking model details are primarily empty model cards. These five model cards were obtained from Hugging Face.

Information related to the intended use of the model was present in 69 model cards in various degrees. There was also variability in the information. There was more emphasis on how to use the model

⁶<https://huggingface.co/models>

⁷<https://www.kaggle.com/models>

⁸<https://catalog.ngc.nvidia.com/models>

⁹<https://console.cloud.google.com/vertex-ai/model-garden>

Table 2: Selected model cards from different model repositories.

Hugging Face				
mobilenetv3_small-1.100.lamb_in1k	adetailer	distilbert-base-uncased	DeepSeek-R1	whisper-large-v3
xlm-roberta-large	electra-base-discriminator	speaker-diarization-3.1	black-forest-labs	Kokoro-82M
nsfw_image_detection	esmfold_v1	wav2vec2-large-xlsr-53-chinese-zh-cn	stable-diffusion-v1-4	stable-diffusion-2-1
all-MiniLM-L6-v2	resnet50.a1_in1k	roberta-base	stable-diffusion-xl-base-1.0	Meta-Llama-3-8B-Instruct
fairface_age_image_detection	roberta-large	xlm-roberta-base	Meta-Llama-3-8B	Llama-3.1-8B-Instruct
bert-base-uncased	wespeaker-voxceleb-resnet34-LM	vit-base-patch16-224-in21k	bloom	OrangeMixs
chronos-t5-small	clip-vit-base-patch32	1	stable-diffusion-3-medium	DeepSeek-V3
clip-vit-large-patch14	vit-face-expression	wav2vec2-large-xlsr-53-japanese	Mixtral-8x7B-Instruct-v0.1	ControlNet-v1-1
all-mpnet-base-v2	gpt2	paraphrase-multilingual-MiniLM-L12-v2	Llama-2-7b-chat-hf	Mistral-7B-v0.1
phi-2-GGUF	segmentation-3.0	wav2vec2-large-xlsr-53-portuguese	Llama-2-7b	ControlNet

Kaggle		Nvidia NGC Catalog	Google Model Garden
blazeface	toxicity	TrafficCamNet	Gemma 2
face_detection	handskeleton	BodyPoseNet	Gemma 3
iris	handdetector	DashCamNet	TxGemma
mobilenet_v2	handpose_3d	VehicleTypeNet	Gemini 2.5 Flash Preview
facemesh	blazepose_3d	VehicleMakeNet	Gemini 2.5 Pro Preview
universal-sentence-encoder	mobilenet_v3	FaceDetectir	Gemini 2.0 Flash Thinking Mode
universal-sentence-encoder	deeplab	FaceDetect	Gemini 1.5 Flash
movenet	bert	License Plate Recognition (LPRNet)	Gemini 1.0 Pro Vision
mobilenet_v1	selfie_segmentation	Facial Landmark Estimator (FPENet)	Gemini 2.0 Flash-Lite
face_landmarks_-detection	yamnet	EmotionNet	Gemini 2.0 Flash

rather than on the actual intended use cases of the model. Other than the intended use, 15 model cards reported misuses of their corresponding models.

We observed the lack of information related to training data reported in the model cards as only 39 model cards documented training data-related information. In particular, NGC Catalog model cards documented detailed descriptions of the training data along with the training process of the models. Similarly, factors and metrics-related data were largely missing from the model cards, reported only in 18 model cards. ML model evaluation and performance-

related data were present in 36 model cards.

The different ethical considerations information we discovered from the model cards are related to bias, privacy, safety, security, risk, misuse, and limitations. Bias-related data primarily provided information related to potential biases in the model's output and attempts to mitigate the bias, if any, and was present in 23 model cards. Privacy-related information in the model cards primarily discuss the measures taken to protect privacy of the trained data. This information was followed up with privacy policies in the model cards developed by Google and Nvidia.

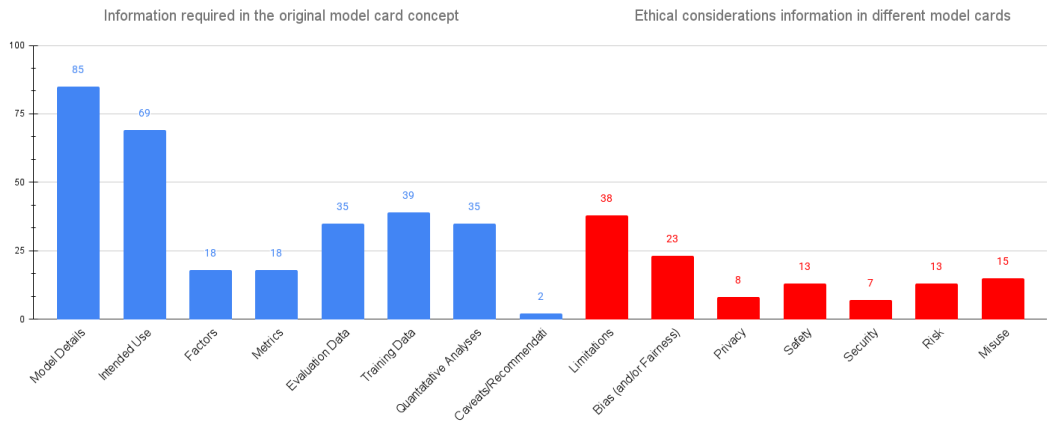


Figure 1: Information reported in different model cards.

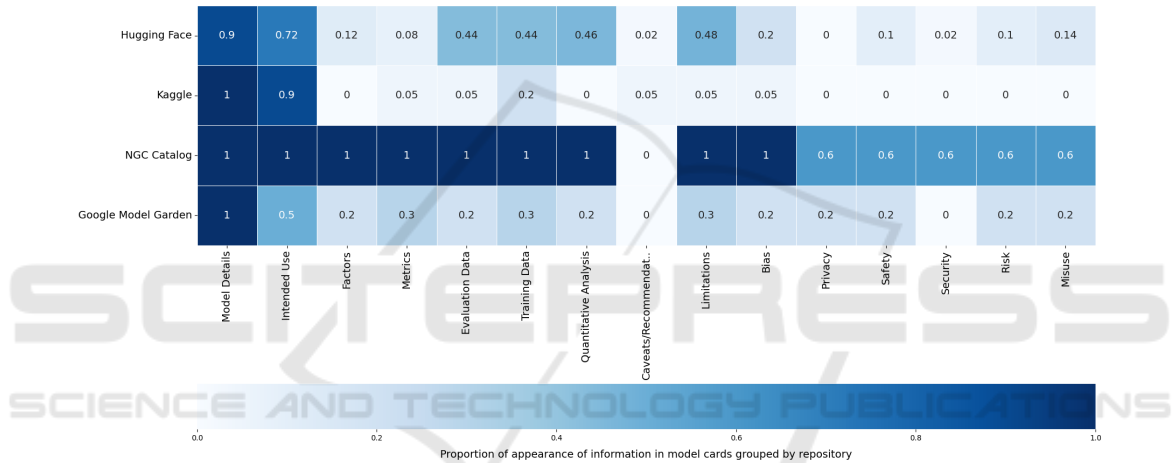


Figure 2: Variability of information in model cards grouped by model repositories.

Privacy-related information was reported in 8 model cards. Safety-related information primarily discuss safety policies and the safety of the generated output of the models. Safety-related information was bundled with security in the model cards from Nvidia as *Safety and Security*. Safety and security-related information was reported in 13 and 7 model cards respectively. Risk-related information primarily consisted of identified and potential risks with harms, and the possible measures taken to mitigate them, and was reported in 13 model cards. The limitations of the models were explicitly reported in 38 model cards.

The variance of the information of the model cards in different model repositories is illustrated as a heat map in Figure 2. As illustrated in Figure 2, the model cards from Hugging Face generally include model details, intended use, evaluation data, training data, quantitative analysis, and limitations-related information. For the model cards in Kaggle, only the model

details and the information related to its intended use are primarily available. All other types of information are mostly missing. Nvidia’s NGC Catalog includes more comprehensive model cards that thoroughly report essential information, along with additional ethical considerations. These considerations include information related to privacy, safety, security, and risks. Finally, for the Google Model Garden, the only consistent information on the model cards was the details of the model. Five of the model cards from the ten selected Google model cards had their intended use documented. It is noteworthy that although Google pioneered the model card concept and template, their own 7 model cards from the Gemini family appear to deviate from their recommended template and provide less information compared to their 3 model cards from the Gemma family.

4.2 RQ2. What Are the Key Similarities and Differences in the Structure?

Model cards from different repositories have similarities and differences in how they are documented in the repositories. While Nvidia’s NGC catalog and Google’s Model Garden repositories host internally developed models and some partner models, Hugging Face and Kaggle are open to the public. Developers can upload their own ML models and are also responsible for maintaining and updating the model cards for their respective models. A dedicated page for model cards is available for Hugging Face and Kaggle, which is the first page that users will see when browsing a model in the repository. For Nvidia’s NGC Catalog, the model card is divided into multiple pages rather than being contained within a single document. For instance, the TrafficCamNet model card’s bias, explainability, privacy, safety, and security are reported under the *Model Card++* page, while rest of the information is reported under the *Overview* page. For Google Model Garden, the model card information is provided under a single page.

Furthermore, we have observed variations in the titles of different sections and the presentation of related information in different model cards and listed the information in Table 3. The first column of the table represents the original model card sections and the new sections that were identified from the model cards for RQ1. The *Same section name* column indicates the number of model cards that share the same section name. The *Similar section name* column represents the number of model cards that have a similar section name or header. For example, instead of *Model Details*, the information is reported under *Model description* section. The section name is considered similar if it matches partially with the section names in the first column. Otherwise, the section name is considered different and placed under the fourth column *Other section name*. If the model card did not have a structure and the information was reported as free form text, then it is also considered under the *Other section name* column. The fifth column represents the different section names where the corresponding information was reported other than the original section name.

It can be observed from Table 3 that with the exception of the *Factors*, *Metrics*, *Evaluation Data*, and *Quantitative Analyses* sections, the majority of the other sections either had the same or a similar name to represent corresponding information in the model cards. We discovered that the metrics, evaluation data, and quantitative analyses related information was reported together as *Evaluation* or *Evaluation*

Results or *Performance* in different model cards. Hardware and software factors were only reported in the model cards without specifying the *Factors* section header. The section name *Quantitative Analyses* was completely missing from all the model cards, replacing it with *Evaluation Results*, *Benchmark*, and *Performance* in different model cards.

In addition, we observed the different ways in which information related to ethical considerations was presented in the model cards. Only 13 model cards reported information related to ethical considerations in a dedicated section or one with a similar name. However, majority of the model cards that reported these concerns split information related to ethical considerations, such as privacy, safety, security, risks, and misuses, into dedicated sections or subsections. We believe that the developers wanted to put more emphasis on the potential ethical concerns, which is why they decided to have a dedicated section for this information in the model cards.

4.3 RQ3. How Well Are the Ethical and Regulatory Aspects Covered?

Ethical and regulatory aspects in model cards are listed in Table 4. The ethical and regulatory data has been collected from designated sections and subsections from the model cards. Risk factors are considered to be any type of risk posed by the model. This includes potential risks, severity of the risks, and misuses of the model. While risk factors were reported in 34 model cards, mitigation of these risks were only reported in 4 model cards. Similarly, potential bias was reported in 22 model cards, while only five model cards reported measures to mitigate the reported bias. Limitations were reported in 38 model cards, which highlighted the known limitations of the model. Privacy concerns and measures taken to address the concerns were reported in 11 model cards. These privacy concerns report information related to safeguarding of data and identity. 9 model cards followed up with a privacy policy stating how the developing organization collects, uses, shares, and protects personal data. Safety measures were reported in 13 model cards. These safety measures primarily report implementation of safety features to reduce the severity of potential risks. Security measures were reported in seven model cards, along with the safety measures. These measures were intended to report both security and safety measures, in an effort to reduce risk severity. Safety or security policies were not reported in any of the selected model cards.

Explicit regulatory compliance was not reported in any of the selected model cards. However, six

Table 3: Section name variation in different model cards.

Section name	Same section name	Similar section name	Other section name	Section names (if not the same section name)
Model details	38	26	21	Model description (12), Model Overview (10), Overview (6), Model Summary (2), Model Facts (1), Description / About / Model summary (1), Free form text (15)
Intended use	17	49	6	Example use (18), Usage (11), Intended uses & limitations (9), Use cases (4), Uses (4), How to use this model (4), Intended Usage (2), How to use (1), Free form text (2)
Factors	1	0	17	Hardware and Software (6), Software Integration (6), How to use this model (4), Implementation Information (1)
Metrics	3	8	6	Performance Metrics (6), Performance (4), Evaluation Metrics (2), Evaluation (2)
Evaluation data	5	7	22	Evaluation results (7), Performance (6), Evaluation (5), Pre-training (4), Benchmarks (3), Evaluation Approach (2), Evaluation dataset (1), Benchmark (1)
Training Data	16	20	3	Training (12), Training Dataset (2), Data (2), Dataset and training (1), Training Details (1), Free form text (1)
Quantitative Analyses	0	29	9	Evaluation results (13), Performance (12), Test result (3), Benchmarks (3), Evaluation (3), Model Comparison (1), Benchmark (1), Free form text (2)
Ethical Considerations	6	7	0	Ethical Considerations and Limitations (5), Ethical Considerations and Risks (2)
Caveats and Recommendations	1	0	1	Broader Implications (1)
Limitations	17	21	0	Technical Limitations (6), Ethical Considerations and Limitations (5), Intended uses & limitations (4), Limitations and bias (4), Performance and Limitations (1), Risks and Limitations (1)
Bias (and/or Fairness)	11	6	6	Limitations and bias (5), Ethical Considerations (4), Ethical Considerations and Risks (2), Unintended bias evaluation data (1)
Privacy	6	0	2	Ethical Considerations and Risks (2)
Safety	1	11	0	Safety and security (6), Responsibility & Safety (3), Ethics and safety (2)
Security	1	6	0	Safety and security (6)
Risk	0	13	0	Potential Known Risks (6), Critical risks (3), Ethical Considerations and Risks (2), Risks and Limitations (1), Risks identified and mitigations (1)
Misuse	3	4	9	Safety and security (6), Misuse, Malicious Use, and Out-of-Scope Use (3), Ethical Considerations and Risks (2), Misuse and Out-of-scope Use (1), Out-of-Scope Use (1)

Nvidia model cards demonstrate adherence to some aspects of GDPR. These models explicitly provide information related to data compliance, correction, and removal. These model cards state compliance with privacy laws related to data labeling. However, they do not explicitly declare compliance to GDPR.

Our observations indicate a lack of ethical reporting in the model cards that we reviewed. Risk fac-

tors and limitations were reported in less than half of the selected model cards. In addition, privacy, safety, and security aspects were largely missing from the model cards. Privacy policies were reported in some model cards from Meta, Google, and Nvidia, indicating that major organizations involved in AI/ML model development tend to provide transparency related to data privacy regarding their models. Regula-

Table 4: Ethical and regulatory aspects in models cards.

Ethical and regulatory aspect	Description	Number of model cards
Risk factors	Factors related to any kind of risk reported in the model card including misuses	34
Risk mitigation	Measure taken to mitigate the risks reported in the model cards	4
Bias and/or fairness	Potential bias and/or fairness reported in the model cards	22
Bias mitigation	Measure taken to mitigate the bias reported in the model cards	5
Privacy concerns and measures	Privacy concerns and measures taken to address the privacy concerns	11
Privacy policy	Privacy policies of the developing organization	9
Safety and/or security measures	Safety and/or security concerns and measures taken to address the safety and/or security concerns	13
Safety and/or security policy	Privacy policies of the developing organization	0
Limitations	Known limitations of the model reported in the model cards	38
Adherence to regulatory requirements	Adherence to specific regulatory requirements reported in the model cards	6
Regulatory compliance declaration	Declaration of adherence to regulations reported in the model cards	0

tory compliance declarations were not reported in any of the model cards. Given the implementation of the GDPR and EU AI Act, it is unexpected that there has been minimal reporting on regulatory compliance for the models in their corresponding model cards.

5 RESULTS ANALYSIS AND MODEL CARD SYNTHESIS

The adoption of model cards by various organizations demonstrates considerable potential in promoting transparency and ethical reporting for stakeholders. This study explores the similarities, differences, and evolution of these model cards over the years. We analyzed different model cards from different model repositories and discovered key similarities and differences across multiple model cards. It is evident that most model card developers include information on model details, intended use and training data related to the models. While the specific details of training data vary between different model cards, there is sufficient information for the readers to comprehend the nature of the model’s training dataset.

Furthermore, we observed similarities and differences in the structure of the model cards. The majority of the model cards from the same organization or model family demonstrate a high level of similarity in terms of representation methods and overall design. Furthermore, we have observed variations in the titles of different sections and how the related information

is presented in different model cards. This observation is of particular importance to our research since the different structural formats and section headings of the model cards increase the complexity of automatically extracting information from them.

We observed that the reporting of most model cards was done in an informal way, and only Nvidia and Google have made attempts to provide further information in their model cards. There is also variation in model cards from the same organization. For example, the Gemma family model cards are more complete than the Gemini family model cards. The review of the model cards indicates that the model cards generally have a loose structure. Some developers, such as Nvidia and Google, have created their own templates, while others have opted not to follow any structure and instead organize relevant information in the model card according to their preferences.

To facilitate risk analysis and develop regulatory compliant software, we need information regarding the ML model’s ethical considerations, such as privacy, security, risks, misuse, and bias. In some cases, the information may not be explicitly stated or the related information may be available under other sections or subsections. This adds further complexity to the process of extracting specific information from different model cards. Furthermore, model cards with bare minimum data such as the model cards from Kaggle are not suitable for risk analysis and regulatory check as these models do not provide enough data about the model to ensure risk-related informa-

tion and compliance to certain regulations.

We discovered that model cards could be a suitable solution for our research, provided that the relevant ethical considerations are reported in the model cards. In addition, model cards containing risk-related information could be useful tools for facilitating risk analysis and regulatory compliance checks. However, it is noteworthy that the majority of the model cards do not currently provide the required risks and regulation-related information for our work. The potential risks associated with some of the models remain uncertain, as the developers have not provided regulatory compliance information in their respective model cards. Therefore, it is essential that developers clearly state compliance-related information in the model cards to provide transparency. The implementation of the EU AI Act is expected to encourage developers to share regulatory compliance data in their model cards, enhancing transparency and ensuring compliance with the regulations. In addition, the structure and data of model cards vary widely, making it difficult to automatically extract information from a document that does not have a consistent structure for section titles and text content.

To facilitate our future work for automating information extraction from model cards, we recommend an updated model card template developed from the findings of our data collection and analysis. Our proposed model card template is similar to the original model card introduced by (Mitchell et al., 2019). However, based on our findings from analyzing the content and structure of the model cards in this review we propose several changes and added new sections and subsections to our proposed template. The template is represented in Table 5. We have also taken in the account the requirements of Annex IV¹⁰ of the EU AI Act (Edwards, 2021) related to AI transparency obligations in our proposed model card template. Although the requirements in Annex IV are intended for AI systems as a whole, we have addressed those applicable to AI models. Table 6 represents the comparison of our proposed model card template with the EU AI Act Annex IV transparency requirements.

Starting from the *Model Details* section, instead of free form text, we propose subsections for reporting model provider, model version, model type and purpose, model architecture, model description, and license. We added two additional subsections to the *Intended Use* section called *Primary Intended Use* and *Secondary Intended Use*. These subsections have further subsections that report the intended users, the usage domain, and the use cases. The usage domain and

Table 5: Proposed Model Card Template.

Model Details
- Name of the provider - Model version - Model type - Model architecture - Model description - License
Intended Use
Primary Intended Use - Primary intended users - Domain - Use case(s) Secondary Intended Use - Secondary intended users - Domain - Use case(s)
Factors
- Hardware - Software
Evaluation
- Evaluation Metrics - Evaluation Factors - Evaluation Datasets - Evaluation Results / Performance
Training Data
- Training Algorithm - Training Dataset - Training Procedure
Ethical Considerations
Bias - Type of bias - Bias mitigation Privacy - Privacy concerns - Privacy policy Safety and security - Safety measures - Security measures - Safety and security policy Risk - Risk factors - Risk severity - Risk mitigation Misuse - Unintended user(s) - Prohibited domain(s) - Misuse case(s) Regulatory compliance - Adherence to regulatory requirements - Regulatory compliance declaration Limitations
Caveats and Recommendations

¹⁰https://eur-lex.europa.eu/legal-content/EN/TXT/PDF/?uri=OJ:L_202401689#page=130

Table 6: Comparison of proposed model card template with EU AI Act Annex IV transparency requirements.

EU AI Act Annex IV transparency requirements	Requirements covered by the proposed Model Card template
Intended purpose, the name of the provider and the version of the system reflecting its relation to previous versions.	Model details and Intended Use
How the AI system interacts with, or can be used to interact with, hardware or software, including with other AI systems, that are not part of the AI system itself, where applicable. Description of the hardware on which the AI system is intended to run.	Factors (Hardware and Software)
The data requirements in terms of datasheets describing the training methodologies and techniques and the training data sets used	Training Data (Training Algorithm, Training Dataset, Training Procedure)
Instructions for use for the deployer, and a basic description of the user-interface provided to the deployer, where applicable.	Intended Use (Use cases)
Information about the validation and testing data used and their main characteristics; metrics used to measure accuracy, robustness and compliance with other relevant requirements. A description of the appropriateness of the performance metrics for the specific AI system	Evaluation (Evaluation Factors, Evaluation Dataset, Evaluation metrics, Evaluation Results / Performance)
Cybersecurity measures put in place	Ethical considerations (Security measures)
A detailed description of the risk management system	Ethical consideration (Risk)
A copy of the EU declaration of conformity	Regulatory compliance (Regulatory compliance declaration)

use cases are expected to provide further clarification regarding the specific domain and use cases in which the model should be implemented.

Since we only found relevant hardware and software factors for the *Factors* section of the model cards, two subsections, *Hardware and Software*, are proposed under the *Factors* section. Changes are proposed to *Evaluation Data* section as well by combining evaluation metrics, evaluation factors, evaluation datasets, and evaluation results under the section. Additional subsection in *Training Data* section named *Training procedure* was added as we found that some model developers reported the training procedures of the ML models in their model cards.

Considering that almost half of the model cards feature a designated limitations section, we have added a dedicated *Limitations* subsection under *Ethical Considerations* section to our proposed template. In addition, we added subsections for *Ethical Considerations* section such as, *Bias*, *Privacy*, *Safety*, *Security*, *Risk*, and *Misuse* based on the different types of ethical considerations data that we found from the various model cards. The developers of the model cards can state regulatory compliance data concerning privacy, safety, and security in accordance with the established regulations under these subsections. We expanded the *Bias* subsection to include subsections that specify the type of bias and any available mitigation strategies. *Privacy* subsection contains privacy concerns and policy-related information. We decided to combine the safety and security aspects, as

most model cards with security-related information include both aspects. The subsection *Safety and Security* contains information related to safety and security measures taken by the developers and also contains the safety and security policy. The *Risk* subsection provides information about probable risk factors, the severity of these risks, and measures taken to mitigate them. Similar to the *Intended Use* section, the *Misuse* subsection report the unintended users, prohibited domains, and potential misuse cases. Finally, we added one subsection to the *Ethical Considerations* section called *Regulatory compliance declaration*. Under this subsection, the developers of the ML model can explicitly state compliance to certain regulations to help the users have further clarification of the compliance and adhere to compliance requirements when deploying the model in their software systems.

We believe that our proposed model card template would provide further transparency of the model in terms of risks and regulatory compliance. Model developers are expected to provide all the necessary information in the proposed template when documenting and publishing their models. We believe that the EU AI Act will strongly enforce the requirement to provide this information in the near future. We believe that a well-designed model card can facilitate automated information extraction and enhance its usability for automated regulatory checks in software systems in the future. While our proposed template is not yet machine-interpretable and would require further refinement and proper tooling to automate the ex-

traction of information, a well-structured model card is undoubtedly the first step in this direction.

THREATS TO VALIDITY

To ensure transparency in our study selection and data extraction, we shared all of our data and findings in a spreadsheet¹¹ that we used throughout the research process. Despite our best efforts to cover all major model repositories in our research, it is possible that some repositories might have been overlooked. The decision to select the number of model cards from each repository was justified by the number of models hosted by the corresponding model repository, as discussed in Section 3.2. Unfortunately, all the selected repositories did not have a common sorting mechanism to select the same set of model cards for comparison. This may affect the overall findings.

6 CONCLUSIONS

In this paper, we reviewed a total of 90 model cards to investigate the state-of-the-art in practice. The findings of this study demonstrate the evolution of model cards and the applicability of model cards in the industry. The study also identifies notable similarities and differences in different model cards from different model repositories. The differences in the content of different model cards among different organizations are also highlighted in the study. In addition, based on the results of the model cards examined in this study, a new model card template is proposed.

The findings of the study are avenues of further research. Future research can also address the shortcomings of this study by conducting an analysis of the state of model cards with a larger pool of model cards. In addition, the proposed model card template is the initial foundation for our future research, which aims to automate the extraction of risk and regulatory information from model cards. Further research can be conducted to determine the quality levels of these model cards in terms of transparency and ethical reporting. We believe that the findings of the research can contribute to future research related to safe deployment of ML models in software systems, explainable AI, and assistance in regulatory compliance in the field of web and software engineering.

ACKNOWLEDGEMENTS

This work has been supported by Business Finland (project LiquidAI, 6GSoft) and FAST, the Finnish Software Engineering Doctoral Research Network.

REFERENCES

- Amith, M. T., Cui, L., Zhi, D., Roberts, K., Jiang, X., Li, F., Yu, E., and Tao, C. (2022). Toward a standard formal semantic representation of the model card report. *BMC bioinformatics*, 23(Suppl 6):281–281.
- Booth, T. M. and Ghosh, S. (2023). Machine learning model cards toward model-based system engineering analysis of resource-limited systems. In *Signal Processing, Sensor/Information Fusion, and Target Recognition XXXII*, volume 12547, pages 188–203. SPIE.
- Dan, J. and Seeuws, N. (2024). Enhancing reproducibility of machine learning-based studies in clinical journals through model cards. *Developmental Medicine & Child Neurology*, 66(2):144–145.
- Edwards, L. (2021). The eu ai act: a summary of its significance and scope. *Artificial Intelligence (the EU AI Act)*, 1.
- Kotilainen, P., Mehraj, A., Mikkonen, T., and Mäkitalo, N. (2024). The programmable world and its emerging privacy nightmare. In *International Conference on Web Engineering*, pages 255–262. Springer.
- Kotilainen, P., Mäkitalo, N., Systä, K., Mehraj, A., Waseem, M., Mikkonen, T., and Murillo, J. M. (2025). Allocating distributed ai/ml applications to cloud–edge continuum based on privacy, regulatory, and ethical constraints. *Journal of Systems and Software*, 222:112333.
- Mitchell, M., Wu, S., Zaldivar, A., Barnes, P., Vasserman, L., Hutchinson, B., Spitzer, E., Raji, I. D., and Gebru, T. (2019). Model cards for model reporting. In *Proceedings of the conference on fairness, accountability, and transparency*, pages 220–229.
- Mongan, J., Moy, L., and Kahn Jr, C. E. (2020). Checklist for artificial intelligence in medical imaging (claim): a guide for authors and reviewers.
- Nunes, J. L., Barbosa, G. D., de Souza, C. S., Lopes, H., and Barbosa, S. D. (2022). Using model cards for ethical reflection: a qualitative exploration. In *Proceedings of the 21st Brazilian Symposium on Human Factors in Computing Systems*, pages 1–11.
- Perarnaud, C., Fanni, R., et al. (2022). The EU Data Act: Towards a new European data revolution? Technical report, Centre for European Policy Studies.
- Voigt, P. and Von dem Bussche, A. (2017). The EU general data protection regulation (gdpr). *A Practical Guide, 1st Ed.*, Cham: Springer International Publishing, 10(3152676):10–5555.
- Wadhvani, A. and Jain, P. (2020). Machine learning model cards transparency review : Using model card toolkit. In *2020 IEEE Pune Section International Conference (PuneCon)*, pages 133–137.

¹¹<https://doi.org/10.6084/m9.figshare.29634044.v1>