

Independent Articles

Prohibited AI Practices in Healthcare under the European Artificial Intelligence Act

Hannah van Kolfshooten 

Centre for Life Sciences Law, University of Basel, Basel, Switzerland

Abstract

The European Union's Artificial Intelligence Act introduces a novel regulatory category of "unacceptable risk," prohibiting specific AI practices that are deemed fundamentally incompatible with human rights and ethical principles. While much attention has focused on the regulation of high-risk AI systems, particularly in medical contexts, the AI Act's outright bans under Article 5 have received far less scrutiny. This paper addresses that gap by examining how these prohibitions apply to healthcare and public health, which are domains defined by rapid technological uptake, structural vulnerability, and ethically sensitive decision-making. Drawing on the European Commission's 2025 interpretative Guidelines, the paper argues that several health-related AI applications, such as emotion recognition tools, biometric categorisation systems, and technologies that influence or target vulnerable populations, may fall within the scope of the bans. It also shows that the Act's medical and safety exceptions risk weakening the vulnerability protections that the prohibitions aim to secure. By integrating legal analysis with real-world health examples, the paper offers a framework for interpreting these prohibitions and assesses how they should guide the ethical boundaries of AI in healthcare, within and beyond the European context.

Keywords: Artificial Intelligence; AI regulation and ethics; EU AI Act; health technology; digital health

1. Introduction

Imagine a hospital using an Artificial Intelligence (AI) system to analyze patients' facial expressions and decide who looks "angry" to determine how quickly they receive care. Or imagine an AI tool that infers a patient's sexual orientation from facial images to "personalize" care pathways. Under the EU Artificial Intelligence Act, both uses are categorically prohibited.

The EU Artificial Intelligence Act (AI Act), adopted in 2024, is the world's first comprehensive legal framework for artificial intelligence (AI).¹ Its impact extends beyond the EU, since all providers, importers, distributors, and deployers placing AI systems on the EU market or using them in the EU must comply with role-specific obligations.² What is often overlooked in scholarship and policy debates is that the AI Act does not only impose new requirements — it also *prohibits* certain uses of AI. Article 5 of the AI Act identifies AI practices incompatible with fundamental rights. Article 5 stands apart from the rest of the AI Act. While most provisions set conditions for "high-risk" AI, this provision classifies a small number of uses as "unacceptable" and therefore prohibited. However, several of these bans contain narrowly defined exceptions, mainly for medical or safety purposes. Understanding how these categorical prohibitions and their exceptions interact is crucial for sectors like healthcare, where both risk and necessity are high.

Email: hannah.vankolfshooten@unibas.ch

Cite this article: H. van Kolfshooten. "Prohibited AI Practices in Healthcare under the European Artificial Intelligence Act," *Journal of Law, Medicine & Ethics* (2026): 1–10. <https://doi.org/10.1017/jme.2026.10270>

In February 2025, the European Commission published guidelines on banned AI practices, clarifying the prohibitions and offering concrete examples for the healthcare sector.³ Yet the prohibitions under Article 5 remain underexplored. Studies either predate the Commission's 2025 Guidelines⁴ or treat the bans in general terms, without sector-specific analysis.⁵ In health law scholarship, the focus has been on the new rules for AI medical devices, with little attention to potential uses of AI in healthcare that the AI Act prohibits outright.⁶ Despite growing interest in the AI Act, two gaps remain. First, existing analyses treat Article 5 largely in the abstract and do not map its prohibitions onto concrete use cases in healthcare and public health. Second, the internal tension between categorical bans and the Act's narrowly defined medical and safety exceptions has not been examined from the perspective of patients and clinical practice.

This paper addresses this overlooked dimension by asking how Article 5's prohibitions, and their narrow medical and safety exceptions, should apply to healthcare and public health, and what legal and ethical implications follow for users, developers, deployers, and regulators. Drawing on the 2025 European Commission Guidelines on prohibited practices, it examines how the bans on manipulative, exploitative, or discriminatory AI practices apply in the context of health, clarifies the scope of medical exceptions, and explores the broader implications for law and ethics. Although rooted in the EU legal framework, the AI Act's concept of "unacceptable risk" speaks to a global normative debate: which uses of AI in healthcare are simply not compatible with medical ethical principles or fundamental rights? As health systems around the world move toward wider use of AI, the Act's prohibitions and underlying principles provide a useful lens for

reflecting on the boundaries of technological intervention in healthcare. Moreover, these prohibitions already influence global practice through well-established channels of regulatory diffusion, including alignment by multinational providers seeking EU market access, and public procurement requirements imposed by large European health systems.

The analysis proceeds in three parts. [Section 2](#) situates Article 5 within the broader logic of the AI Act, explaining its relationship to risk, trust, and vulnerability. [Section 3](#) applies the individual prohibitions to concrete cases in health and public health, including AI systems for emotion recognition, biometric categorization, and behavioral manipulation. [Section 4](#) concludes by reflecting on the broader implications of these prohibitions for the development and deployment of responsible AI in healthcare. The article contributes a sector-specific interpretation of Article 5 for health, clarifies the legal boundaries of the exceptions, and argues for concrete safeguards to prevent those exceptions from normalizing practices the AI Act deems unacceptable.

2. Prohibited AI in the AI Act: Trust, Risk, and Vulnerability

The “no trust, no use” narrative underpins much of the EU’s regulatory approach to AI. The core idea is that the uptake of AI depends on public trust, which in turn depends on legal guarantees that AI systems align with fundamental rights and democratic values. The AI Act, adopted in June 2024, reflects this logic by framing AI governance around the principle of “trustworthy AI.” The AI Act is often framed by EU institutions as a way to make AI systems “worthy of trust,” but in legal terms it operationalizes this narrative through risk-based obligations and fundamental rights protections.⁷

To operationalize this goal, the AI Act adopts a risk-based approach to AI regulation. This approach “should tailor the type and content of such rules to the intensity and scope of the risks that AI systems can generate.”⁸ The AI Act defines four levels of risk for AI systems: “unacceptable risk,” “high risk,” “limited risk,” and “minimal risk.” The legal requirements depend on the risk classification of the specific AI system. While AI systems in the “high risk” category have to comply with strict obligations to enter the market, the AI Act does not introduce any rules for “minimal” risk AI systems. The category of “unacceptable risk” AI systems is treated differently: it comprises a list of AI practices that are prohibited outright. According to the AI Act,⁹ they pose a threat to “important Union public interests as recognised and protected by Union law.” At the same time, the AI Act allows for narrow exemptions for these prohibited practices.

Unlike the rest of the AI Act, which will gradually apply in stages until full implementation in August 2026, the prohibitions in Article 5 entered into force already in February 2025. This reflects the severity with which the EU treats these practices: they are not merely high-risk, but they are treated as irreconcilable with the legal and ethical foundations of the EU.¹⁰ The prohibitions apply regardless of whether the AI system is purpose-built or a general-purpose AI model integrated into a prohibited use.¹¹ Developers, deployers, and users alike must ensure compliance, or risk significant penalties reaching up to €35 million or 7% of global turnover.¹²

2.1 Why Certain AI Practices Are Prohibited: Unacceptable Risk and the Role of Vulnerability

The notion of “unacceptable risk” is central to understanding why certain AI systems are prohibited under Article 5. This threshold is

not based solely on technical shortcomings or empirical performance; it is a normative judgment informed by the EU’s constitutional values, including human dignity, non-discrimination, and the protection of fundamental rights.¹³ In this sense, the prohibition of certain AI uses does not aim to improve them but rather to declare them incompatible with a rights-respecting society, regardless of technical refinement or potential benefits.

A core concern underlying several of the Article 5 prohibitions is that certain AI systems exploit, exacerbate, or institutionalize human vulnerability. While only Article 5(1)(b) refers to vulnerability explicitly (prohibiting systems that exploit a person’s age, disability, or socioeconomic status to distort behavior), the concept runs throughout the list of banned practices. It reflects a consistent concern that some individuals or groups are less able to defend themselves against the potential harms of AI.¹⁴

In particular, vulnerability sometimes acts as a legal trigger for prohibition in cases where individuals are structurally disadvantaged, dependent on public or private institutions, or placed in environments of diminished autonomy. These contexts include education, employment, law enforcement, and healthcare. In such domains, users may be subject to AI-driven decisions or interactions without meaningful alternatives or the ability to opt out, amplifying the risk of manipulation, exclusion, or harm. This concern is especially relevant in light of neuroethical and legal critiques of manipulation that bypass conscious control. As Newirth notes, AI systems that target unconscious processes threaten mental autonomy in ways existing legal frameworks struggle to address.¹⁵ Such systems may harm not only through covert influence but also by undermining individuals’ capacity for reflective decision-making, a point underscored by Galli & Novelli in calling for clearer regulatory definitions of exploitation and manipulation in contexts of structural vulnerability.¹⁶

By banning AI practices that often exploit or exacerbate these conditions, the AI Act provides a form of substantive rights protection that goes beyond procedural safeguards. Rather than merely requiring transparency or human oversight, it declares some systems entirely off-limits. For example, systems that infer emotions from biometric data are prohibited because they involve intrusive inferences into inner mental states unless they fall within the narrow medical and safety exception in Article 5(1)(f).¹⁷ Likewise, systems that score or categorize people based on sensitive characteristics, such as race, sexual orientation, or political views, are banned because they risk entrenching discrimination and treating individuals not as rights-holders, but as data subjects to be analyzed and sorted.¹⁸

This approach of banning AI practices that exploit vulnerability rather than merely regulating them represents an important normative advance. By recognizing that some technologies pose risks that cannot be adequately mitigated through transparency, oversight, or consent requirements alone, the AI Act acknowledges the limits of procedural safeguards in contexts of structural disadvantage. Patients in hospitals, elderly persons in care homes, and individuals with cognitive impairments cannot always exercise meaningful autonomy over AI systems deployed in their environments, even when formally “consented to.” Categorical prohibition in these contexts prioritizes substantive protection over regulatory flexibility. As the European Commission’s 2025 Guidelines on Prohibited AI Practices make clear, these prohibitions are intended to prevent the normalization of AI systems that undermine core principles of the EU legal order.¹⁹

However, the medical and safety exceptions contained in some Article 5 prohibitions, particularly manipulation and emotion

inference, risk undermining this protective logic. These exceptions do not apply uniformly across all prohibited practices, but where they exist, they disproportionately affect contexts of heightened vulnerability. By permitting such practices when justified by medical or safety purposes, the Act creates pathways through which vulnerability can still be exploited, often in precisely the contexts where protection is most needed. The question is not whether medical exceptions are ever justified, but whether they are defined narrowly enough to prevent the normalization of practices that the prohibitions were designed to prevent. As the following analysis demonstrates, several exceptions risk becoming loopholes that allow vulnerable populations to be subjected to the very AI practices the AI Act claims are unacceptable.

In sum, vulnerability operates both as an explicit legal criterion in Article 5(1)(b) and as a deeper normative rationale that informs the entire category of unacceptable risks. Across most of the prohibitions, the common thread is the protection of individuals whose autonomy, agency, or social position makes them disproportionately exposed to manipulation, discrimination, or intrusive inference. The next section examines how these tensions manifest when the prohibitions are applied to concrete health and public health use cases.

3. Prohibited AI Systems in the Health Sector

Healthcare is one of the leading sectors for AI adoption, with applications ranging from diagnostics and mental health support to public health surveillance. While most medical AI systems are classified as high-risk under the AI Act, some uses may fall under

the outright prohibitions of Article 5. These are less frequently examined in the health context but are especially relevant given the sensitivity of health data, structural patient vulnerability, and the power asymmetries in care settings. As shown in Table 1 below, this section explores how the prohibitions apply to concrete use cases in healthcare and public health, drawing on the European Commission's 2025 Guidelines for interpretation.

3.1. Harmful Manipulation and Deception

Article 5(1)(a) of the AI Act prohibits AI systems that use subliminal techniques or exploit unconscious processes in a manner likely to materially distort a person's behavior and cause significant harm. While this provision was initially framed in relation to consumer manipulation, such as addictive interface design or behavioral nudging in commercial contexts, it has clear implications for the healthcare sector. The potential for harm is especially acute when AI systems interact with patients who are cognitively impaired, emotionally vulnerable, or in situations of diminished autonomy: precisely the conditions under which manipulation may be most effective and least detectable.

The European Commission's 2025 Guidelines clarify that this prohibition extends to AI systems that exploit behavioral science or neuroscientific insights to influence decision-making beyond a person's conscious awareness.²⁰ In healthcare, this risk arises in applications such as mental health chatbots, AI-powered adherence tools, or systems used in behavioral therapy that employ persuasive dialogue or gamification techniques to shape patient behavior. For instance, an AI-enabled digital coach that uses reward-based

Table 1. Prohibited and Permitted AI Practices in Healthcare Under Article 5 AI Act

Article 5 AI Act	Description	Prohibited example	Permitted example
5(1)(a) – Subliminal manipulation	Use of subliminal techniques or AI exploiting unconscious processes to distort behavior and cause significant harm	A mental-health chatbot using subliminal cues (beyond user awareness) to drive medication compliance	AI that motivates behavior as part of a recognized therapeutic intervention, with transparency and safeguards (e.g., in treatment for depression)
5(1)(b) – Exploiting vulnerability	Exploiting a person's age, disability, or socio-economic situation to distort behavior and cause harm	A digital companion influencing lonely elderly patients to purchase unnecessary health services	Behavioral support tools that inform without exploiting vulnerability, and do not materially distort behavior or risk harm.
5(1)(c) – Social scoring	AI scoring individuals based on behavior or traits, used in unrelated contexts and leading to detrimental treatment	Health compliance scores used to assess creditworthiness or eligibility for public housing	Use of health risk scoring by health insurers, as long as scores are confined to the original healthcare context
5(1)(d) – Predictive criminal profiling	Predicting criminal behavior based solely on profiling or personal traits	AI in psychiatry predicting patient recidivism post-discharge based only on psychological profiling	Risk assessments based on clinical judgement and past incidents, not solely algorithmic profiling
5(1)(e) – Untargeted facial image scraping	Creating facial recognition databases via indiscriminate scraping of online images or CCTV footage	Patient identification systems in hospitals trained on datasets scraped from social media or public cameras	Systems trained on ethically sourced, consent-based image datasets for internal hospital use
5(1)(f) – Emotion inference	Inferring emotions from biometric data such as facial expressions, voice, or physiological signals	Workplace wellness AI monitoring hospital staff's stress or anxiety levels from speech patterns	CE-marked therapeutic tools emotion-related assessment in mental health or autism diagnostics
5(1)(g) – Biometric categorization	Inferring sensitive attributes (e.g., race, religion, sexual orientation) from biometric data	Diagnostic AI like DeepGestalt inferring ethnicity or sexual orientation from facial features	Classifying physical traits (e.g., skin tone for dermatology) if not used to infer sensitive attributes
5(1)(h) – Real-time remote biometric identification by law enforcement	Real-time facial recognition in public spaces by police, except under strict conditions	Law enforcement using live facial recognition in hospital waiting areas to locate forensic patients	Facial recognition for patient check-in within hospital premises; not used by or on behalf of law enforcement

nudging to steer patients toward medication compliance may cross into prohibited territory only where techniques target unconscious processing (rather than mere persuasion) and are likely to cause significant harm.

A particularly relevant and emerging area is the integration of neurotechnologies with AI, such as brain-computer interfaces (BCIs) used for rehabilitation or mental health interventions.²¹ Scholars have raised concerns that when such systems are designed to stimulate neural activity or reinforce particular behavioral patterns, without transparent user engagement, they may infringe on mental integrity or even amount to a form of technologically mediated coercion.²² In such cases, the distinction between therapeutic support and impermissible manipulation becomes both ethically and legally significant. These concerns are echoed in recent scholarship. Franklin argues that Article 5(1)(a) AIA inadequately addresses the spectrum of manipulative techniques, overlooking supraliminal strategies and harms to autonomy itself, not just to psychological health.²³

The AI Act includes a limited exception for manipulation that is “strictly necessary for medical purposes.”²⁴ However, the Commission’s Guidelines stress that this exception must be interpreted narrowly: only systems that use behavioral influence as an unavoidable component of a recognized therapeutic intervention, such as motivating a patient with severe depression to continue treatment, may qualify.²⁵ Crucially, the manipulation must be proportionate, medically justified, and accompanied by appropriate safeguards. In practice, “strictly necessary” should mean no equally effective, less intrusive alternative exists.

These limitations matter because healthcare settings are uniquely predisposed to power asymmetries and implicit trust. Patients often follow system recommendations without question, and AI that subtly shapes behavior under the guise of personalized care could undermine informed consent. As such, the prohibition of manipulative AI techniques in healthcare aims not only to protect patients from harm, but to uphold core values of dignity, autonomy, and transparency in clinical decision-making.

A common objection is that many therapeutic practices already rely on behavioral influence, motivational framing, and psychological insight. Techniques such as positive reinforcement, structured nudging, reframing, or guided suggestion are standard components of cognitive behavioral therapy and other clinical interventions. If these techniques are permissible when delivered by clinicians, it is not immediately obvious why their use should become prohibited merely because they are operationalized through AI. The answer lies in the structural differences between human-delivered and AI-mediated influence. Human therapeutic techniques operate within a relational context that is bounded, intermittent, and accountable, and in which practitioners are trained to calibrate interventions to patient agency, consent, and well-being. By contrast, AI systems can deliver influence continuously, at scale, and with forms of opacity or persistence that patients cannot detect or resist. Moreover, AI systems that target unconscious processes lack the contextual judgement, situational restraint, and ethical reflexivity that clinicians bring to therapeutic work. For these reasons, the AI Act treats AI-driven manipulation not as a digital version of standard behavioral practice, but as a qualitatively different form of influence that risks overriding autonomy in ways that conventional therapeutic methods do not. The problem is not the technique itself, but the loss of safeguards that ordinarily accompany its clinical use.

3.2. Harmful Exploitation of Vulnerabilities

Article 5(1)(b) of the AI Act prohibits AI systems that “exploit any of the vulnerabilities of a natural person due to their age, physical or mental disability, or specific social or economic situation, in order to materially distort the behaviour of that person in a manner that causes or is likely to cause that person or another person physical or psychological harm.” This prohibition recognizes that certain individuals, due to structural or situational disadvantage, are more susceptible to manipulation, particularly when AI systems are designed to influence behavior in subtle, opaque, or coercive ways. Legally, three elements must all be present: exploitation of a listed vulnerability, material distortion of behavior, and a likelihood of physical or psychological harm.

Although the provision does not single out specific sectors, its relevance to healthcare and public health is considerable. The care relationship is often characterized by cognitive asymmetry, emotional dependency, and a degree of trust that can be easily manipulated through AI-mediated interactions. In practice, this concern becomes acute in settings such as elderly care, mental health support, or digital health services targeting individuals with chronic conditions or disabilities. For instance, an AI-powered chatbot designed to provide mental health support might nudge users toward certain commercial products under the guise of therapeutic advice. Similarly, a digital companion system used in long-term care could leverage loneliness and cognitive decline to influence behavior in ways that are not transparent or clinically justified.

The Commission’s 2025 Guidelines offer several illustrative examples of prohibited conduct under Article 5(1)(b), including AI-powered toys that encourage children to engage in risky or harmful challenges and systems that offer exploitative loans to individuals in financial distress. While not health-specific, these examples highlight a broader logic that readily extends to medical contexts. An AI system that detects signs of psychological vulnerability in a patient and uses this information to promote non-evidence-based treatments or to steer users toward commercial health services would likely breach the prohibition. Similarly, targeting elderly patients with cognitive impairments through persuasive interfaces that push certain treatments or services could constitute exploitative behavior under the Act. As Galli & Novelli argue, the Act’s failure to clarify whether exploitation depends on intent or outcome creates uncertainty for health applications where patients may be systematically influenced due to their dependency or illness.²⁶

Academic literature has also drawn attention to the risks of manipulative neurotechnologies in healthcare. Studies have raised concerns about AI-driven brain-computer interfaces or neuromarketing tools that can manipulate user emotions or decisions without conscious awareness, particularly in patients with reduced cognitive autonomy.²⁷ In these contexts, the potential for material distortion of behavior, combined with the inability to resist or critically evaluate such influence, amplifies the legal and ethical concerns underpinning the Article 5(1)(b) ban.

Ultimately, the prohibition reflects a fundamental regulatory red line: where AI systems exploit reduced agency for behavioral control or influence, particularly in sensitive health contexts, they are not merely high-risk but unacceptable. This approach reinforces a commitment to safeguarding the dignity and autonomy of those who are most at risk of harm from technologically mediated manipulation. Given the structural dependency of many patients, the prohibition in Article 5(1)(b) sets a justified boundary: behavioral influence in such contexts should be allowed only when it is transparent, limited in scope, and clinically necessary.

3.3. Social Scoring in the Field of Health

Article 5(1)(c) of the AI Act prohibits AI systems that evaluate or classify the trustworthiness of natural persons over a certain period based on their social behavior, socioeconomic status, or personal characteristics, leading to detrimental or unfavorable treatment in contexts unrelated to where the data was originally generated. The European Commission Guidelines clarify that the prohibition targets AI systems creating composite profiles or scores used to determine trustworthiness across different contexts, particularly when they result in unjustified discrimination or disproportionately negative impacts.²⁸

In the health sector, AI-driven scoring systems are frequently used to assess health risks or eligibility for services.²⁹ For instance, health insurers may use AI to calculate individual risk scores based on medical records, fitness tracker data, or lifestyle factors to set insurance premiums.³⁰ According to Recital 58 and Annex III of the AI Act, such systems are considered high-risk AI systems and are generally permissible when confined to their intended health insurance context. The key reason this is allowed is that health insurers are evaluating health-related risks within the same context, namely health insurance, and are not transferring the scoring into unrelated domains. The European Commission Guidelines explicitly acknowledge that AI-based scoring systems for health and life insurance do not meet the cumulative conditions of Article 5(1)(c) when used appropriately.³¹

However, if these systems generate scores that are subsequently used outside the health insurance context, such as influencing employment screening, housing applications, or access to public benefits, they may fall under the prohibition. For example, a health behavior score that is used by banks to determine creditworthiness or by employers to evaluate a job applicant's reliability based on medical compliance or fitness data would likely be prohibited. Similarly, a public health authority's use of vaccination compliance data to influence access to unrelated social services could breach Article 5(1)(c). The recent SCHUFA case, though focused on credit scoring, highlights the broader risks of transferring profiling across domains.³² Under the AI Act, stakeholders must ensure that AI scoring systems remain strictly within their original purpose to avoid prohibited social scoring.³³ The restriction in Article 5(1)(c) is justified in health because cross-context scoring would intensify existing social and health inequalities.

3.4. Individual Criminal Offense Risk Assessment and Prediction

Article 5(1)(d) of the AI Act prohibits AI systems that are used for "making risk assessments of natural persons or predicting the occurrence or reoccurrence of an actual or potential criminal offence based solely on the profiling of a person or assessing their personality traits and characteristics." While this provision primarily targets law enforcement and security applications,³⁴ certain healthcare or health-related AI systems could inadvertently fall within its scope. According to the European Commission Guidelines, the prohibition covers any AI system that seeks to predict criminal conduct "without a direct link to a specific criminal act or suspicion and solely based on the person's profile or characteristics."³⁵

In healthcare, AI systems are increasingly used to assess patient behavior and risks. For example, AI tools in psychiatric hospitals may attempt to predict whether a patient is likely to commit criminal offenses post-release, based on psychological profiling or past behavior.³⁶ Similarly, AI systems used in addiction treatment centers could predict the likelihood of patients engaging in illegal drug-related activities based on medical and behavioral data.³⁷

Another plausible scenario is hospital-based AI systems that assess the risk of domestic violence offenders reoffending after therapeutic interventions, based largely on psychological traits and historical data.³⁸ According to Article 5(1)(d), if these systems rely solely on profiling and personality assessments, without evidence of a specific, concrete threat or event, they would be prohibited.

A further example arose during the COVID-19 pandemic, when curfews and movement restrictions were enforced across many countries. AI systems could, in theory, be used to predict which patients might violate curfew orders based on medical records, socioeconomic status, or past compliance behavior. Even though such systems might aim to support public health efforts, if the AI's purpose is to predict violations of legal measures (such as curfew-breaking) that are classified as criminal offenses and if it does so solely through profiling, it would fall under the prohibition of Article 5(1)(d). The Guidelines also clarify that administrative offenses not deemed criminal under national law would generally not be covered by this prohibition.³⁰

3.5. Untargeted Scraping of Facial Recognition

Article 5(1)(e) of the AI Act prohibits AI systems that create or expand facial recognition databases through the untargeted scraping of facial images from the internet or CCTV footage. This provision targets AI systems that process scraped images specifically for the recognition of persons. Crucially, the prohibition covers both direct usage of scraped images and the indirect use through training AI models.⁴⁰ This means that even if a healthcare provider or AI developer does not personally scrape the images but relies on datasets obtained or generated via untargeted scraping, the AI system would still fall under the prohibition.

Hospitals and other healthcare deployers have explicit legal obligations under the AI Act to verify that any AI facial recognition software or algorithms they deploy have not been trained using prohibited practices like untargeted scraping.⁴¹ Deployers must demand clear documentation from developers or providers confirming compliance. If no sufficient documentation is provided, deployers must refuse to use the system to avoid legal exposure. Non-compliance risks severe penalties, including fines up to €35 million or 7% of annual turnover.⁴²

Real-life examples where this rule is highly relevant include patient identification systems that use facial recognition to match patients with medical records during automated check-in at hospitals.⁴³ Similarly, surveillance systems in elderly care homes employ facial recognition to detect wandering or escaping patients.⁴⁴ Another emerging area is public health surveillance systems that integrate facial recognition to monitor population movement or compliance with quarantine measures during pandemics, such as those piloted in parts of Asia during the COVID-19 crisis. Even though these tools may aim to strengthen public health responses by tracking potential contacts or enforcing isolation orders, their use of facial recognition trained on datasets expanded through untargeted scraping, such as scraping CCTV footage from public areas or facial images from social media, would be prohibited under Article 5(1)(e).⁴⁵ In the context of healthcare, this prohibition seems justified in light of patients' limited practical ability to contest or meaningfully avoid systems trained on scraped facial data.

3.6. Inferring Emotions

Under Article 5(1)(f) of the AI Act, AI systems designed to identify or infer emotions from biometric data are generally prohibited. Article 3(39) defines emotion recognition specifically as processing

biometric data, such as facial expressions, vocal intonations, or physiological signals like heart rate, to directly identify or analytically infer emotions.⁴⁶ Recital 44 explicitly clarifies that this prohibition includes both identifying and inferring emotions.⁴⁷ The prohibition in Article 5(1)(f) AI Act does not refer to “emotion recognition systems,” but only to “AI systems to infer emotions of a natural person.” The AI Act stipulates a “medical exception,” by explicitly excluding “AI systems placed on the market strictly for medical or safety reasons, such as systems intended for therapeutic use.”⁴⁸ In healthcare, emotion AI could be used for detecting signs of depression in patients, assisting in suicide prevention efforts, or diagnosing autism spectrum disorders through the analysis of emotional expressions.⁴⁹

It is, however, up for discussion when an emotion AI system serves a “therapeutic purpose.”⁵⁰ The EU Commission now clarified in its Guidelines that this “medical exception” must be narrowly interpreted, necessary, proportionate, and strictly limited to therapeutic contexts. In practice, these will typically involve CE-marked medical devices explicitly intended for care.⁵¹ Examples of permissible uses include AI tools used by therapists to recognize emotional distress in patients undergoing mental health treatment, systems aiding clinicians to identify early signs of autism in children, or assistive technologies that interpret emotional cues for visually or hearing-impaired patients.⁵² Moreover, AI that detects purely physical states like pain or fatigue without inferring emotional states, such as systems used in postoperative monitoring⁵³ or to prevent driver fatigue,⁵⁴ remains explicitly outside the scope of the emotion recognition prohibition.⁵⁵ Conversely, AI systems broadly used to monitor stress levels, anxiety, anger, or general emotional well-being among hospital staff for workplace wellness initiatives remain prohibited, as they do not fulfill the strict medical exception criteria. However, this line remains blurred. As Durovic and Corno highlight, even publicly observable expressions can convey deeply private affective states. When systems infer emotions from such data without consent, they undermine the individual’s right to mental privacy and violate the dignity-based rationale behind Article 5(1)(f).⁵⁶

The prohibition has been criticized for leaving gaps: while grounded in fundamental-rights concerns, it bans emotion inference yet still allows detection of emotional expressions. Consequently, employers could use the technology to flag an employee as sounding distracted during calls, without officially claiming to know their emotional state. Moreover, since the ban is based on current technical limitations,⁵⁷ it may no longer apply once the technology improves, without tackling the deeper ethical issues of emotion tracking.⁵⁸ In the context of health, a strict reading of Article 5(1)(f) seems warranted: emotion inference should be prohibited in health because inferring sensitive attributes from biometric data is intrinsically misaligned with patient autonomy and equality, and no therapeutic aim can render that inference proportionate.

3.7. Biometric Categorization

Article 5(1)(g) prohibits AI systems that categorizes individuals based on biometric data to deduce or infer sensitive attributes, specifically race, political opinions, trade union membership, religious or philosophical beliefs, sex life, or sexual orientation. In a healthcare context, many AI systems process biometric data, such as facial images, retinal scans, or voice data, for diagnostic, therapeutic, or monitoring purposes. However, the prohibition only applies if all cumulative conditions are met: (1) The system must

be placed on the market, put into service, or used for biometric categorization; (2) It must categorize individual persons; (3) The categorization must be based on biometric data, and (4) The purpose must be to deduce or infer one of the prohibited sensitive attributes.

In contrast to Article 5(1)(f) (emotion recognition), Article 5(1)(g) provides no medical or safety exception.⁵⁹ Accordingly, the prohibition applies regardless of the sector, including healthcare, as long as all the cumulative conditions are met. The EC explains in their guidelines that, for example, classifying skin tone or eye color for medical diagnostics (like in dermatology or oncology) is permitted.⁶⁰ This is because it does not aim to deduce sensitive attributes as listed in the provision, meaning it falls outside its scope. Where a system unintentionally correlates with a sensitive attribute but is not intended to deduce it, Article 5(1)(g) may not apply.

However, (future) AI systems for health-related purposes could fall under the prohibition because they intend to deduce sensitive attributes such as race. For instance, DeepGestalt, a facial analysis AI designed to diagnose rare diseases, raised concerns due to its use of facial patterns linked to specific ethnic groups, risking racial inference.⁶¹ Similarly, AI systems such as Wang and Kosinski’s AI tool, which claimed to predict sexual orientation from facial images, would constitute a clear breach, and could, for example, be deployed in mental healthcare.⁶² Hypothetical but plausible cases include AI tools that infer religious affiliation from biometric information to customize hospital services or AI-driven risk prediction tools that use facial or voice data to deduce political views. For example, researchers have suggested the use of personal data, although not biometric data, for personalizing communication to tackle vaccine hesitancy.⁶³ Even if such systems are framed as beneficial for health outcomes, they would be prohibited under Article 5(1)(g) because they infer sensitive attributes from biometric data. The absence of a medical exemption in Article 5(1)(g) seems justified, because inferring sensitive attributes from biometric data is inherently incompatible with protecting patient autonomy and equality.

3.8 Real-Time Remote Biometric Identification

Article 5(1)(h) of the AI Act prohibits the use of AI systems for real-time remote biometric identification (RBI) of individuals in publicly accessible spaces by law enforcement authorities, except under narrowly defined and strictly necessary circumstances. These include: (i) the targeted search for victims of specific crimes (e.g., child abduction), (ii) the prevention of an imminent threat to life or physical safety (e.g., a terrorist attack), or (iii) the localization or identification of individuals suspected of serious crimes listed in Annex II of the AI Act, provided those crimes are punishable under national law by at least four years’ imprisonment.⁶⁴

In healthcare, the use of real-time RBI is generally not affected by this prohibition, as most health-related deployments take place within private or semi-public environments, such as hospitals, clinics, or secure treatment facilities. Systems used solely for clinical or operational purposes, such as patient check-in via facial recognition in a hospital lobby, or internal security monitoring in a psychiatric facility, do not fall under Article 5(1)(h), since they are not deployed by or on behalf of law enforcement and do not operate in publicly accessible spaces.

A relevant case from the COVID-19 pandemic illustrates how this boundary may blur. Vodafone, in partnership with surveillance technology firm Digital Barriers, deployed a thermal detection camera system designed to screen individuals in real time for elevated body temperatures.⁶⁵ Although primarily used in semi-

public spaces such as office buildings, the system combined thermal and HD imaging and was capable of identifying individuals and triggering alerts.⁶⁶ If such a system were to be deployed by law enforcement in publicly accessible spaces, such as airports or train stations, for the purpose of identifying potentially infected individuals or enforcing quarantine measures, it could fall within the scope of the prohibition in Article 5(1)(h).

Further grey areas arise in forensic care settings, where health-care intersects with criminal justice. If law enforcement were to use real-time RBI in a public space to locate a forensic patient subject to a court order, the use may be prohibited unless it falls under one of the explicit exceptions. For example, identification would be lawful if the individual is suspected of having committed a serious offense listed in Annex II (e.g. rape, murder, terrorism).

This provision has been the subject of extensive criticism. A central concern is that the prohibition is too narrow. First, it applies only to *real-time* identification, and thus, post hoc identification (such as reviewing footage after an event) remains permissible.⁶⁷ The European Data Protection Board and the European Data Protection Supervisor emphasized that the intrusiveness of biometric surveillance does not depend on whether the identification is real-time or retrospective, arguing that post-event processing can produce similar chilling effects on fundamental rights such as the right to freedom of assembly.⁶⁸ Civil society organizations criticize that the restriction applies solely to law enforcement authorities, thus excluding private entities that may deploy the same intrusive technologies for purposes such as private security and workplace monitoring.⁶⁹ Additionally, there are concerns that the AI Act's broad and vaguely defined exceptions for national security, military applications, and research and development, significantly weaken the effectiveness of the prohibition.⁷⁰ National security, in particular, is explicitly excluded from the scope of the regulation,⁷¹ allowing member states to deploy AI systems with biometric surveillance without being bound by the Act's safeguards.⁷² This could also have implications for AI systems for public health objectives, such as surveillance of biothreats at airports.⁷³ Member states could, in principle, frame certain biosurveillance deployments under national security exceptions, thereby sidestepping Article 5 safeguards, raising oversight concerns. Although directed at law enforcement, the logic of Article 5(1)(h) supports a strong presumption against real-time biometric identification in health settings as well.

4. Medical Exceptions and Vulnerability Protections

A consistent pattern runs through the Article 5 prohibitions: where a medical or safety exception exists, it tends to authorize uses of AI on the very populations the prohibitions are designed to protect. This means that many vulnerable patient groups, such as psychiatric patients, children with developmental disorders, and people in institutional care, may be exposed to influence, profiling, and inference techniques that the Act considers unacceptable for everyone else. The result is a structural inconsistency: the Act emphasizes vulnerability as a reason to ban certain practices, yet its exceptions apply primarily in settings where vulnerability is at its highest.

The usual justification for these exceptions is twofold. First, medical environments operate under professional standards that ostensibly mitigate risks. Second, withholding potentially useful technologies could harm patients. Neither claim addresses the central problem. Professional norms do not eliminate the power asymmetries in psychiatric wards, residential care, or acute treatment settings. Patients cannot meaningfully refuse certain interventions, and "consent" often carries little practical weight. Claims

about clinical benefit are equally indeterminate: many systems sold as "therapeutic" serve administrative or institutional interests, and Article 5 gives no criteria for separating genuine therapeutic necessity from convenience, efficiency, or commercial motives.

Some exceptions may be justified. Emotion-related tools used in autism diagnosis⁷⁴ or narrowly tailored crisis-prevention systems for severe depression⁷⁵ may provide benefits not achievable through other means. In such cases, denying access could expose patients to avoidable harm. The existence of such examples, however, does not justify the current breadth of the exceptions. An exception is only defensible when the system's influence is transparent, proportionate, and directly linked to a recognized therapeutic aim; when it avoids subliminal or covert techniques; when it does not condition care on acceptance of the AI system; and when robust evidence supports its claimed benefits. The AI Act does not require any of this. As drafted, an exception can be invoked without demonstrating necessity, without proving that less intrusive alternatives have been considered, and without any independent review of proportionality. This absence of structure is what allows the exception to drift toward institutional convenience.

The core problem is not the idea of an exception, but the lack of a principled framework to govern its use. A credible framework would require three things: demonstrable therapeutic necessity; use of the least intrusive design capable of achieving the intended aim; and clear evidence that the expected patient benefit outweighs any loss of autonomy or mental privacy. Therapeutic necessity should mean that no clinically effective, less intrusive alternative is available. Least-intrusive design should require the exclusion of subliminal techniques and limit behavioral steering to transparent, clinician-mediated interactions. The benefit-to-autonomy ratio must be shown with reference to concrete outcomes, not broad claims about innovation or institutional workload.

These requirements are not theoretical. They must be integrated into the legal architecture that already governs medical AI. Providers should be required to document therapeutic necessity and least-intrusive design in the technical file and make these claims available to notified bodies and regulators. Deployers invoking the exception should complete a fundamental rights impact assessment that directly addresses manipulation risks, patient agency, and alternatives to AI-mediated influence. Procurement contracts should include proportionality clauses, audit obligations, logs of influence techniques, and deactivation triggers. Hospital ethics committees should review not only clinical safety but the autonomy impact of the system and should be able to suspend deployment where necessary. Finally, patients should be able to decline AI-mediated interventions without penalty, and a clinically viable alternative should always be offered.

Without these conditions, the medical exceptions risk functioning as broad permission mechanisms rather than narrow derogations. They create a gap between the Act's stated aim, namely protecting individuals in situations of reduced autonomy, and its operational reality, where those same individuals may be subjected to techniques the law deems unacceptable for the general population. A prohibition regime that is premised on vulnerability must require more justification when exceptions apply in precisely those contexts.

5. Conclusions

This article has argued that the EU AI Act marks a decisive shift in the governance of emerging technologies. Beyond introducing a tiered risk framework, it draws clear normative boundaries around

certain uses of AI that are deemed fundamentally incompatible with European values. These stakes are particularly visible in healthcare, a sector characterized by structural power imbalances, heightened vulnerability, and rapid digitalization. While existing scholarship and policy debate have focused primarily on the regulation of high-risk medical AI systems, this paper has identified an equally critical yet largely overlooked regulatory dimension: the absolute prohibitions under Article 5.

By examining these prohibitions through the lens of healthcare, and drawing on the European Commission's 2025 Guidelines, this paper has shown that numerous health-related AI applications may be categorically banned. These include emotion inference tools, biometric categorization systems, manipulative chatbots, and predictive criminal profiling tools. Critically, these prohibitions do not depend on whether systems can be improved through better design, accuracy, or oversight. They reflect a principled judgment: certain technologies, when used in contexts of reduced autonomy or structural vulnerability, violate ethical and legal boundaries that cannot be justified by efficiency or innovation.

The implications for health sector stakeholders are threefold. First, compliance with the AI Act requires more than technical validation. It demands upfront normative scrutiny of whether a system's purpose is legally permissible at all. Second, actors invoking exceptions (such as for medical purposes) bear the burden of demonstrating that their systems are not only clinically justified but also narrowly tailored, proportionate, and rights-compatible. Third, the Article 5 prohibitions invite a more fundamental reflection on what kinds of technological mediation are acceptable in the care relationship, and where the line between assistance and intrusion should be drawn.

The Article 5 prohibitions represent a significant normative achievement: they recognize that certain AI practices exploit vulnerability in ways that cannot be remedied through better design or procedural oversight, and they establish substantive limits on what may be done to individuals in contexts of reduced autonomy. This is an important departure from the Act's otherwise risk-management approach and reflects a commitment to dignity-based protection that goes beyond efficiency or innovation concerns.

At the same time, the prohibitions and the exceptions disclose an ambivalence at the heart of the AI Act's approach to vulnerability. The categorical bans reflect a principled commitment to shielding individuals from forms of technological intrusion that strike at autonomy, dignity, and equality. As such, the prohibitions represent a robust model of vulnerability protection: one that does not rely on procedural safeguards or user responsibility but draws substantive boundaries around what may be done to people in conditions of dependency. By contrast, the medical and safety exceptions weaken this model by reintroducing the very logics of risk management and efficiency that the prohibitions reject. Whether they ultimately advance or undermine vulnerability protection depends on how narrowly and rigorously they are interpreted. As drafted, the Act does not ensure that exceptions will only be applied in circumstances where their use is indispensable, proportionate, and accompanied by safeguards capable of offsetting the structural disadvantages faced by patients. The current framework therefore offers strong protections in principle but leaves their realization contingent on discretionary institutional practices. If the law does not require strong proof and strict justification before allowing exceptions to the bans, it will only appear to protect vulnerable people on paper, while in reality those same people may still be subjected to harmful AI practices.

To prevent this outcome, the legal boundaries of the exceptions of Article 5 for health systems should be crystal clear and meet extra safeguards. At a minimum, deployers should be required to demonstrate therapeutic necessity, show that no less intrusive alternative exists, document proportionality and expected clinical benefit, and undergo independent review through mechanisms such as fundamental rights impact assessments and ethics oversight. These requirements should be embedded in the legal and organizational infrastructure that already governs medical AI: providers should document necessity and least-intrusive design in the technical file (and make this accessible to notified bodies and regulators), and procurement contracts should include audit obligations, logging of influence techniques, and deactivation triggers. Finally, patients must retain the ability to refuse AI-mediated interventions without losing access to care, and a clinically viable alternative should be offered.

The implications of this framework extend beyond the European Union. By designating certain AI practices as incompatible with dignity and fundamental rights, the Act sends a normative signal globally. Other jurisdictions now face a choice: align with the EU's stance by embedding similar prohibitions into their frameworks or adopt a more permissive approach where technologies banned in Europe may be deployed elsewhere. The former could establish a shared baseline of patient protection across borders; the latter risks creating regulatory havens for ethically problematic AI and deepening global health inequalities. For patients, these choices are not abstract, and access to safe, rights-compatible care may depend on geography. The effectiveness of Article 5 in healthcare will therefore depend not on the existence of its prohibitions, but on whether its exceptions are governed strictly enough to preserve the dignity-based limits those prohibitions establish.

Disclosures. The author has nothing to disclose.

Hannah van Kolschooten, PhD, LL.M., is a researcher at the Centre for Life Sciences Law, University of Basel (Switzerland).

References

1. Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence and amending Regulations (EC) No 300/2008, (EU) No 167/2013, (EU) No 168/2013, (EU) 2018/858, (EU) 2018/1139 and (EU) 2019/2144 and Directives 2014/90/EU, (EU) 2016/797 and (EU) 2020/1828 (Artificial Intelligence Act), 2024 O.J. (L 2024/1689) 1.
2. Hannah van Kolschooten and Janneke van Oirschot, "The EU Artificial Intelligence Act (2024): Implications for Healthcare," *Health Policy* 149 (2024): 105152, <https://doi.org/10.1016/j.healthpol.2024.105152>.
3. European Commission, *Commission Guidelines on Prohibited Artificial Intelligence Practices Established by Regulation (EU) 2024/1689 (AI Act) COM (2025) 5052 final* (July 29, 2025).
4. Uroš Čemalović, "Prohibited Artificial Intelligence Practices According to Article 5 of the European Union's Regulation on AI — between the 'Too Late' and the 'Not Enough,'" *International Journal of Law and Information Technology* 32 (2024): eaae023, <https://doi.org/10.1093/ijlit/eaae023>; Ros-tam J. Neuwirth, "Prohibited Artificial Intelligence Practices in the Proposed EU Artificial Intelligence Act (AIA)," *Computer Law & Security Review* 48 (2023): 105798, <https://doi.org/10.1016/j.clsr.2023.105798>; Paul Voigt and Nils Hullen, *The EU AI Act* (Springer Nature, 2024), 37–46.
5. Elif Biber, "A Close Reading of the European Commission's Guidelines on Prohibited Artificial Intelligence Practices: A Powerful Reflection of the European Approach to AI," *Journal of AI Law and Regulation* 2, no. 3 (2025): 266–73, <https://doi.org/10.21552/aire/2025/3/9>.
6. van Kolschooten and van Oirschot, "The EU Artificial Intelligence Act (2024)"; Sofia Palmieri and Tom Goffin, "A Blanket That Leaves the Feet Cold: Exploring the AI Act Safety Framework for Medical AI," *European*

- Journal of Health Law* 30, no. 4 (2023): 406–427, <https://doi.org/10.1163/15718093-bja10104>; Timo Minssen et al., “Regulatory Responses to Medical Machine Learning,” *Journal of Law and the Biosciences* 7, no. 1 (2020): lsa002, <https://doi.org/10.1093/jlb/lsa002>; Barry Solaiman et al., “Monitoring Mental Health: Legal and Ethical Considerations of Using Artificial Intelligence in Psychiatric Wards,” *American Journal of Law & Medicine* 49, nos. 2–3 (2023): 250–66, <https://doi.org/10.1017/amj.2023.30>; Hannah van Kolschooten & Astrid Pilotin, “Reinforcing Stereotypes in Health Care Through AI-Generated Images: A Call for Regulation,” *Mayo Clinic Proceedings: Digital Health* 2, no. 3 (2024): 335–341, <https://doi.org/10.1016/j.mcpdig.2024.05.004>; Sara Gerke et al., “The Need for a System View to Regulate Artificial Intelligence: Machine Learning-Based Software as Medical Device,” *npj Digital Medicine* 3 (2020): 53, <https://doi.org/10.1038/s41746-020-0262-2>; Hannah van Kolschooten, “Towards an EU Charter of Digital Patients’ Rights in the Age of Artificial Intelligence,” *Digital Society* 4 (2025): 6, <https://doi.org/10.1007/s44206-025-00159-w>; Hannah van Kolschooten, “Artificial intelligence in radiology: safeguarding patients’ rights in the digital era,” *European Radiology* 36 (2026): 5194–5196, <https://doi.org/10.1007/s00330-025-12221-9>.
7. Johann Laux et al., “Trustworthy Artificial Intelligence and the European Union AI Act: On the Conflation of Trustworthiness and Acceptability of Risk,” *Regulation & Governance* 18, no. 1 (2024): 3–32, <https://doi.org/10.1111/rego.12512>.
 8. Regulation 2024/1689, 2024 O.J. (L 2024/1689) Recital 26 (EU).
 9. Regulation 2024/1689, 2024 O.J. (L 2024/1689) Recital 46 (EU).
 10. Regulation 2024/1689, 2024 O.J. (L 2024/1689) Recital 179 (EU).
 11. Regulation 2024/1689, 2024 O.J. (L 2024/1689) Article 5(2) (EU).
 12. Regulation 2024/1689, 2024 O.J. (L 2024/1689) Article 99 (EU).
 13. Čemalović, “Prohibited Artificial Intelligence Practices According to Article 5 of the European Union’s Regulation on AI—between the ‘Too Late’ and the ‘Not Enough.’”
 14. Sue Anne Teo, “Artificial Intelligence, Human Vulnerability and Multi-Level Resilience,” *Computer Law & Security Review* 57 (2025): 106134, <https://doi.org/10.1016/j.clsr.2025.106134>.
 15. Neuwirth, “Prohibited Artificial Intelligence Practices in the Proposed EU Artificial Intelligence Act (AIA).”
 16. Federico Galli and Claudio Novelli, “The Many Meanings of Vulnerability in the AI Act and the One Missing,” *BioLaw* no. S1 (2024): 53–72, <https://teseo.unitn.it/biolaw/article/view/3302>.
 17. Regulation 2024/1689, 2024 O.J. (L 2024/1689), Article 5(1)(f) (EU).
 18. Regulation 2024/1689, 2024 O.J. (L 2024/1689), Articles 5(1)(c) and 5(1)(g) (EU).
 19. European Commission, “Guidelines on Prohibited Artificial Intelligence Practices,” § 2.
 20. European Commission, “Guidelines on Prohibited Artificial Intelligence Practices,” § 60–64.
 21. Matija Franklin et al., “Missing Mechanisms of Manipulation in the EU AI Act,” *The International FLAIRS Conference Proceedings* 35 (May 2022), <https://doi.org/10.32473/flairs.v35i.130723>.
 22. Marcello Ienca and Roberto Andorno, “Towards New Human Rights in the Age of Neuroscience and Neurotechnology,” *Life Sciences, Society and Policy* 13 (2017): 5, <https://doi.org/10.1186/s40504-017-0050-1>.
 23. Franklin et al., “Missing Mechanisms of Manipulation in the EU AI Act.”
 24. Regulation 2024/1689, 2024 O.J. (L 2024/1689), Recital 29 (EU).
 25. European Commission, “Guidelines on Prohibited Artificial Intelligence Practices,” § 65.
 26. Galli and Novelli, “The Many Meanings of Vulnerability in the AI Act and the One Missing.”
 27. Galli and Novelli, “The Many Meanings of Vulnerability in the AI Act and the One Missing”; Neuwirth, “Prohibited Artificial Intelligence Practices in the Proposed EU Artificial Intelligence Act (AIA).”
 28. European Commission, “Guidelines on Prohibited Artificial Intelligence Practices,” § 85.
 29. Azish Filabi and Sophia Duffy, “AI-Enabled Underwriting Brings New Challenges for Life Insurance: Policy and Regulatory Considerations,” *Journal of Insurance Regulation* 40, no. 8 (2021): 1–20, <https://content.naic.org/sites/default/files/JIR-ZA-40-08-EL.pdf>.
 30. Manohar Kapse et al., “Customization of Health Insurance Premiums Using Machine Learning and Explainable AI,” *International Journal of Information Management Data Insights* 5, no. 1 (2025): 100328, <https://doi.org/10.1016/j.jjimei.2025.100328>; A. Spender et al., “Wearables and the Internet of Things: Considerations for the Life and Health Insurance Industry,” *British Actuarial Journal* 24 (2019): e22, <https://doi.org/10.1017/S1357321719000072>.
 31. Regulation 2024/1689, 2024 O.J. (L 2024/1689), Recital 58 (EU); European Commission, “Guidelines on Prohibited Artificial Intelligence Practices,” § 85.
 32. Hannah B. van Kolschooten, “A Health-Conformant Reading of the GDPR’s Right Not to Be Subject to Automated Decision-Making,” *Medical Law Review* 32, no. 3 (2024): 373–91, <https://doi.org/10.1093/medlaw/fwae029>.
 33. European Commission, “Guidelines on Prohibited Artificial Intelligence Practices,” § 85.
 34. Kitti Mezei, “Artificial Intelligence in the Criminal Justice System: Exploring Risks and Opportunities Through the Lens of the European Union’s AI Act,” in *Challenges of Artificial Intelligence for Law in Europe*, ed. Marton Varju and Kitti Mezei (Springer Nature Switzerland, 2025), https://doi.org/10.1007/978-3-031-86813-9_7.
 35. European Commission, “Guidelines on Prohibited Artificial Intelligence Practices,” § 90.
 36. Leda Tortora et al., “Neuroprediction and A.I. in Forensic Psychiatry and Criminal Justice: A Neurolaw Perspective,” *Frontiers in Psychology* 11 (2020): 220, <https://doi.org/10.3389/fpsyg.2020.00220>; Gijs van Dijck, “Predicting Recidivism Risk Meets AI Act,” *European Journal on Criminal Policy and Research* 28 (2022): 407–23, <https://doi.org/10.1007/s10610-022-09516-8>.
 37. Mohit Suva and Gayatri Bhatia, “Artificial Intelligence in Addiction: Challenges and Opportunities,” *Indian Journal of Psychological Medicine* (2024): 02537176241274148, <https://doi.org/10.1177/02537176241274148>.
 38. Rongqin Yu et al., “Development and Validation of a Prediction Tool for Reoffending Risk in Domestic Violence,” *JAMA Network Open* 6, no. 7 (2023): e2325494, <https://doi.org/10.1001/jamanetworkopen.2023.25494>; Michael Mayowa Farayola et al., “Ethics and Trustworthiness of AI for Predicting the Risk of Recidivism: A Systematic Literature Review,” *Information* 14, no. 8 (2023): 426, <https://doi.org/10.3390/info14080426>.
 39. European Commission, “Guidelines on Prohibited Artificial Intelligence Practices,” § 91.
 40. Akshita Rohatgi and Tae Jung Park, “Privacy in the Public: Analysing the EU Framework to Outline Approaches for Regulating AI Personal Data Scraping,” *Computer Law & Security Review* 57 (July 2025): 106150, <https://doi.org/10.1016/j.clsr.2025.106150>.
 41. Regulation 2024/1689, 2024 O.J. (L 2024/1689), Article 29 and Recitals 78–81 (EU).
 42. Regulation 2024/1689, 2024 O.J. (L 2024/1689), Article 72 (EU).
 43. Seema Mohapatra, “Use of Facial Recognition Technology for Medical Purposes: Balancing Privacy with Innovation,” *Pepperdine Law Review* 43, no. 4 (2015): 1017–64.
 44. F. Xavier Gaya-Morey et al., “Deep Learning-Based Facial Expression Recognition for the Elderly: A Systematic Review,” *arXiv:2502.02618*, preprint, arXiv, February 4, 2025, <https://doi.org/10.48550/arXiv.2502.02618>.
 45. Vera Lúcia Raposo and Li Du, “Facial Recognition Technology: Is It Ready to Be Used in Public Health Surveillance?,” *International Data Privacy Law* 14, no. 1 (2024): 66–86, <https://doi.org/10.1093/idpl/ipad021>.
 46. European Commission, “Guidelines on Prohibited Artificial Intelligence Practices,” § 82.
 47. Riccardo Vecellio Segate, “The ‘Medical Exception’ to Emotion Detection Algorithms within the EU’s Forthcoming AI Act: Regulatory Implications for Therapeutical Smart Cobotics,” *2024 IEEE 8th Forum on Research and Technologies for Society and Industry Innovation (RTSI)* (September 2024): 408–13, <https://doi.org/10.1109/RTSI61910.2024.10761570>.
 48. Regulation 2024/1689, 2024 O.J. (L 2024/1689), Recital 44 (EU).
 49. European Commission, “Guidelines on Prohibited Artificial Intelligence Practices,” § 80.
 50. Segate, “The Medical Exception.”
 51. European Commission, “Guidelines on Prohibited Artificial Intelligence Practices,” § 256–259.

52. European Commission, "Guidelines on Prohibited Artificial Intelligence Practices," § 87; Smith K. Khare et al., "Emotion Recognition and Artificial Intelligence: A Systematic Review (2014–2023) and Research Recommendations," *Information Fusion* 102 (February 2024): 102019, <https://doi.org/10.1016/j.inffus.2023.102019>.
53. Denys Fontaine et al., "Artificial Intelligence to Evaluate Postoperative Pain Based on Facial Expression Recognition," *European Journal of Pain* 26, no. 6 (2022): 1282–91, <https://doi.org/10.1002/ejp.1948>; Gioacchino D. De Sario et al., "Using AI to Detect Pain through Facial Expressions: A Review," *Bioengineering* 10, no. 5 (2023): 548, <https://doi.org/10.3390/bioengineering10050548>.
54. Yucheng Shang et al., "Driver Emotion and Fatigue State Detection Based on Time Series Fusion," *Electronics* 12, no. 1 (2023): 1, <https://doi.org/10.3390/electronics12010026>.
55. European Commission, "Guidelines on Prohibited Artificial Intelligence Practices," p. 83.
56. Mateja Durovic and Tommaso Corno, "The Privacy of Emotions: From the GDPR to the AI Act, an Overview of Emotional AI Regulation and the Protection of Privacy and Personal Data," in *Privacy, Data Protection and Data-Driven Technologies* (Routledge, 2024).
57. Regulation 2024/1689, 2024 O.J. (L 2024/1689), Recital 44 (EU).
58. Alexandra Prigent, "Is There Not an Obvious Loophole in the AI Act's Ban on Emotion Recognition Technologies?," *AI & SOCIETY* 40 (2025): 5581–5582, <https://doi.org/10.1007/s00146-025-02289-8>.
59. Regulation 2024/1689, 2024 O.J. (L 2024/1689), Recital 44 and Articles 5(1)(f) and 5(1)(g) (EU).
60. European Commission, "Guidelines on Prohibited Artificial Intelligence Practices," p. 94.
61. Yaron Gurovich et al., "Identifying Facial Phenotypes of Genetic Disorders Using Deep Learning," *Nature Medicine* 25, no. 1 (2019): 60–64, <https://doi.org/10.1038/s41591-018-0279-0>.
62. Yilun Wang and Michal Kosinski, "Deep Neural Networks Are More Accurate than Humans at Detecting Sexual Orientation from Facial Images," *Journal of Personality and Social Psychology* (US) 114, no. 2 (2018): 246–57, <https://doi.org/10.1037/pspa0000098>; John Leuner, "A Replication Study: Machine Learning Models Are Capable of Predicting Sexual Orientation From Facial Images," arXiv:1902.10739, preprint, arXiv, February 27, 2019, <https://doi.org/10.48550/arXiv.1902.10739>.
63. Heidi J. Larson and Leesa Lin, "Generative Artificial Intelligence Can Have a Role in Combating Vaccine Hesitancy," Opinion, *BMJ* 384 (January 2024): q69, <https://doi.org/10.1136/bmj.q69>.
64. Alice Giannini and Sarah Tas, "AI Act and the Prohibition of Real-Time Biometric Identification," *Verfassungsblog*, ahead of print, *Verfassungsblog*, December 10, 2024, <https://doi.org/10.59704/a15aa6e58151c853>.
65. Jack Loughran, "Vodafone Launches Heat-Detecting Camera to Protect Offices from Covid-19," *Engineering and Technology Magazine*, May 4, 2020, <https://eandt.theiet.org/2020/05/04/vodafone-launches-heat-detecting-camera-protect-offices-covid-19>.
66. A. M. Guzman et al., "Thermal Imaging as a Biometrics Approach to Facial Signature Authentication," *IEEE Journal of Biomedical and Health Informatics* 17, no. 1 (2013): 214–22, <https://doi.org/10.1109/TITB.2012.2207729>.
67. Federica Paolucci, "Enhancing Oversight and Addressing Gaps: Assessing the Impact of the AI Act on Biometric Identification Systems," in *Next Democratic Frontiers for Facial Recognition Technology (FRT): Legal, Ethical and Democratic Implications of FRT*, ed. Natalia Menéndez González and Giuseppe Mobilio (Springer Nature Switzerland, 2025), https://doi.org/10.1007/978-3-031-89794-8_5.
68. EDPB-EDPS Joint Opinion 5/2021 on the Proposal for a Regulation of the European Parliament and of the Council Laying down Harmonised Rules on Artificial Intelligence (Artificial Intelligence Act) (June 18, 2021), https://www.edpb.europa.eu/our-work-tools/our-documents/edpb-edps-joint-opinion/edpb-edps-joint-opinion-52021-proposal_en.
69. Ella Jakubowska et al., "Retrospective Facial Recognition Tech Conceals Human Rights Abuses," *Euronews*, April 14, 2023, <https://www.euronews.com/2023/04/14/retrospective-facial-recognition-surveillance-conceals-human-rights-abuses-in-plain-sight>.
70. United Nations High Commissioner for Human Rights, Open Letter from the United Nations High Commissioner for Human Rights to European Union Institutions on the European Union Artificial Intelligence Act ('AI Act'), (Nov. 8, 2023); Rosamund Powell, *The EU AI Act: National Security Implications*, Explainer Series (Centre for Emerging Technology and Security, 2024).
71. Regulation 2024/1689, 2024 O.J. (L 2024/1689), Article 2(3) (EU).
72. Plixavra Vogiatzoglou, "The AI Act National Security Exception," *Verfassungsblog*, December 9, 2024, <https://doi.org/10.59704/292082becc7cc8e6>.
73. Tsung-Ling Lee et al., "AI and Data Surveillance: Embedding a Human Rights-Based Approach," *Journal of Law, Medicine & Ethics* 53, supp. 1 (2025): 70–74, <https://doi.org/10.1017/jme.2025.17>; Fallon J. Cochlin et al., "Unlocking Public Health Data: Navigating New Legal Guardrails and Emerging AI Challenges," *Journal of Law, Medicine & Ethics* 52, no. S1 (2024): 70–74, <https://doi.org/10.1017/jme.2024.40>.
74. Agnieszka Landowska et al., "Automatic Emotion Recognition in Children with Autism: A Systematic Literature Review," *Sensors* 22, no. 4 (2022): 1649, <https://doi.org/10.3390/s22041649>.
75. Manju Lata Joshi and Nehal Kanoongo, "Depression Detection Using Emotional Artificial Intelligence and Machine Learning: A Closer Review," *Materials Today: Proceedings*, International Conference on Artificial Intelligence & Energy Systems, vol. 58 (2022): 217–26, <https://doi.org/10.1016/j.matpr.2022.01.467>.