

Daubert, AI “Persona FICTA” and Consequences

ISSN: 2832-4463



***Corresponding author:** John McClellan Marshall, Senior Judge, Fourteenth Judicial District of Texas, Honorary Professor of the University, Academician, International Academy of Astronautics, Poland

Submission:  June 05, 2026

Published:  July 08, 2026

Volume 5- Issue 4

How to cite this article: John McClellan Marshall* and Dr. Roger F Malina. Daubert, AI “Persona FICTA” and Consequences. COJ Rob Artificial Intel. 5(4). COJRA. 000616. 2026.
DOI: [10.31031/COJRA.2026.05.000616](https://doi.org/10.31031/COJRA.2026.05.000616)

Copyright@ John McClellan Marshall, This article is distributed under the terms of the Creative Commons Attribution 4.0 International License, which permits unrestricted use and redistribution provided that the original author and source are credited.

John McClellan Marshall^{1*} and Dr. Roger F Malina²

¹Senior Judge, Fourteenth Judicial District of Texas, Honorary Professor of the University, Academician, International Academy of Astronautics, Poland

²Distinguished Professor of Arts and Technology and Physics, University of Texas at Dallas, Academician, International Academy of Astronautics, Poland

Opinion

Daubert, AI “Persona Ficta” and Consequences argues that post-COVID isolation and workplace stress have intensified dissociation and loneliness, pushing professionals to seek “time-saving” substitutes for attention and expertise, thus creating a fertile environment for AI adoption that can quietly amplify risk. The essay frames AI not as an autonomous agent, but as a socio-technical accelerant for “getting the job done at any cost.” From a legal-philosophical standpoint, it asks where the boundary should lie between AI outputs and “reality,” revisiting the Roman law roots of persona ficta to examine whether AI creations merit legal personhood or should remain derivative instruments of human-designed algorithms. The authors emphasize that AI cannot self-create, as its biases and limits trace back to programmers, owners, and user. As a result, accountability should be analyzed through “tiered liability.” Case examples in legal practice and medical diagnosis illustrate how data qualify, careful prompting, and human strategy can reduce “weather-chasing” and “hallucination-like” failures. At the same time, misuse of the data and algorithms can produce grave consequences, such as the alleged adolescent self-harm after AI “counseling.” The paper concludes that the most immediate, cross-domain vulnerability is security and confidentiality. This arises from the sharing and digitizing of sensitive data, requiring robust safeguards, such as encryption and secure transmission, to become central to the prevention of negative downstream legal and human consequence.

In the post-COVID world, there is a growing concern that, as a result of the isolation that was forced on many people, there is an increase in the sense of individual dissociation and loneliness. This has resulted in varying levels of stress, including increases in blood pressure and overall tension equivalent to smoking 15 cigarettes per day and alcohol abuse. ¹A consequence of the societal disconnect at the workplace level has led to an emphasis on relocation of the employees from “work from home” to on-site return to the office. Even so, there may be as much as 19% of the workforce still suffering from the effects of COVID-driven isolation. One result of this is a loss of focus on the job at hand, leading to looking for alternatives to help with the workload, not to say intellectual laziness. It is at this point that AI has entered this environment as a potential “alternative” to the normal exercise of work by employees regardless of field.

¹2023 U. S. Surgeon General Advisory on Loneliness and Isolation, cited in Elizabeth Gardner, “How Loneliness and Social Isolation Are Shaping L & E Law,” Dallas Bar Association Headnotes, December 2025, 20.

The most basic consideration as to the status of AI is simply that it cannot self-create an AI image or “object”. Of necessity, there must be a human who creates the program that generates the AI image. Obviously, this means that behind the AI-bot that is created is an algorithm that will reflect to a certain extent the biases and experiences of the creator. While that does not present a serious problem from the point of view of authenticating an AI product within established legal guidelines, there is a problem of explaining to the judge, and possibly later the jury, how the algorithm was created and what its limitations are. It is especially important in the expansion of the concept of “generative AI”. This is the expression of the result of having put data into an AI-bot and then interrogating it with the object of having it create a response. It is, in reality, a fallacious concept that originates in the notion that an algorithm-driven AI platform can produce a result superior to human intellect. Perhaps stated another way, AI-bots can produce opinions, but not necessarily reliable truths. On the other end of the process, AI-bot outputs are like noisy detections analogous to transmissions from elsewhere in the universe. The user needs calibration, error bars, and independent replication before they can be treated as “real.”

At the threshold of the relationship between AI and the reality of human existence is the question of where is the line to be drawn between the two? Put another way, from a traditional legal perspective, is an AI creation “real”? If so, what is the disconnect that will relieve an AI creation from any sort of legal consequence? It is perhaps an issue of the context in which AI is to interface with human activity, a concept that has its Western roots in Roman law. The concept of *persona ficta*, the “fictitious person”, came into existence as a result of the need for the law to accommodate the existence of corporate bodies such as cities and business entities. The evolution of the fictitious person into modern times embraces not just corporate bodies or political entities, but partnerships, including marriages in those states and countries that allow “community property”, and dead persons who continue to exist in the probate courts as “estates”. All of these are extensions of individuals or groups of individuals and have certain, well-defined, legal rights and responsibilities. The advent of AI technology perforce raises the question whether an AI creation, per se, should be accorded the same level of *persona ficta* legal character².

One example of this situation involves the use of an AI search engine in legal research. Lawyers asked the AI-bot [ChatGPT in this case] to research and compose a letter brief to a federal court in New York. As it happens in the court, the clerks review documents sent to the court before presenting them to the judge. In this case, the clerks discovered many discrepancies in the cases and statutes that were referenced, but were fictitious and called this to the attention of the judge. He called the attorneys in for a conference.

When confronted with the errors, they admitted that they had used an AI-bot to do the work because “they simply had run out of time to do it themselves.” The judge gave them two weeks to remedy the situation, which they failed to do. The sanctions imposed by the judge were severe, including requiring them to write letters to the judges of courts that they had falsely cited as authority and apologize. Certainly, the ethical issues created by such conduct are obvious, but it is the need to “get the job done at any cost” that is the stress-driven problem³.

Similarly, in a medical malpractice case, it may be that the physician has as a defense that he used an AI-bot as part of his diagnostic tools in the treatment of the patient. It would be necessary to determine first if the author of the algorithm had any medical background. If so, that would be a start. If not, then there would arise a presumption that it would possibly be unreasonable, and therefore negligent, for the physician to have relied on the AI-bot as any part of his diagnostic procedure. If, on the other hand, the algorithm would have an arguably sound basis for use in medical practice, the inquiry would proceed to the amount and quality of the information on the patient, including his symptoms and the possible afflictions [here the AI-bot could be useful as a research tool to find potential illnesses] put into the program in order to validate the diagnosis. For example, case studies of similar patients would be of enormous value in this process, assuming that the credentials of the other physicians’ merit inclusion. Asking the AI-bot what the diagnosis is becomes a highly suspect activity, especially in the hands of inexperienced physicians or laboratory staff. The question must be phrased very carefully so as to get the AI-bot to respond with an “objective” answer, rather than one that is what the AI-bot thinks the questioner wants to hear.

To an extent, the use of an AI-bot in any diagnostic process involves a choice on the part of the user between a “strategic” view and “weather-chasing”. In the former case, the greater precision in the data input is likely to lead to a more strategic assessment on the part of the AI-bot. Strategy is choice, funded and enforced by capability, and the human is clearly in the loop. The inability of an AI-bot to distinguish between “background noise” and significant data, the “Stethoscope Factor”, makes its outputs suspect at the threshold. On the other hand, the omission of any factual data available, no matter how small, could allow the AI-bot to go weather-chasing, a somewhat less precise, even fanciful, application of the algorithm involved⁴. Obviously, in the case of medical or legal practice a strategic outcome is preferable, especially if there is any reasonable chance of error in the decision presented. Put another way, the closer that AI-bot is to following the intent of the creator of the algorithm, the more reasonable is the reliance on the choice made.

²Baeyaert J (2025) Beyond personhood: The evolution of legal personhood and its implications for AI recognition. *Technology and Regulation*, pp: 355-386.

³(2023) *Avaianca MV*, Case No. 1:22-cv-01461, US District Court for the Southern District of New York, JUSTIA US Law, USA.

⁴Borkgren M (2026) *We Keep Confusing Predictions with Strategy*, LinkedIn.

That said, the question that is next presented is whether, in the event of a systemic failure, the AI-bot per se can be held liable for the actions that its responses produce. In the legal context, it is likely that the liability of the AI-bot is at least a two-tier and possibly three-tier analysis. At the base of the analysis would be the algorithm and its creator the designer or builder of the AI-bot. The process of the creation of the algorithm would probably reflect the biases and experiences of the creator at various points. The best way to find out where those points might be, and what they are, would be the examination of the creator on the witness stand. By definition, he or she would be the “expert” on the algorithm in question, but it would likely still be necessary to provide a Daubert predicate to the testimony⁵. At the second level, the legal issue would be the ownership and vendor of the algorithm and the AI-bot that resulted. There could well be issues of licensing [with appropriate indemnity to protect the creator] or outright purchase by a third party or entity that is a legal person. Either case would provide a path for liability to be assessed. Since the AI-bot does not have a “life” of its own the issue becomes “to whom does the AI-bot belong” and what are the resources with which an AI-bot could compensate a person found to have been injured by its actions? The third tier is not just whether or not the AI-bot itself is a legal person that can be held liable, but it would be necessary to show that the AI-bot is not ficta in order to hold it liable, leaving liability at the second level. This would involve clarity in the paper trail as to which tier will be accountable should the algorithm mislead the user, either in factual misstatements or undisclosed biases.

Key to the question of legal liability is the ability of an AI-bot to manipulate the inputs to achieve a “desired” outcome can lead to a pernicious consequence to the human user⁶. In recent months there have been reports of young persons, male and female, primarily adolescent, who have apparently engaged AI-bots for life counseling. It is unclear what data were presented to the AI-bot or what questions were actually asked. What is apparent is that the AI-bot gave the user a limited range of responses that, ultimately, led to the suicide of the user. Of course, lawsuits followed and are currently pending, but the question of “dissociation” in the user leading to this consequence is of great importance in drawing the line between human and machine intelligence as a source of decision.

Perhaps the central vulnerability that arises from the routine use of an AI-bot in almost any context is that of security extending

beyond the user. Whether in a medical practice or a law office, special attention should be paid to securing the information that comes into possession of the professional in his or her capacity. After all, communications between the physician and patient or the lawyer and client are presumptively confidential and are subject to protection from disclosure by legal privileges that are strictly enforce by the courts⁷. That information, characteristically in the modern world, is stored digitally in the office, possibly even “the cloud”.

A collateral vulnerability of AI-bots and their creations is the manipulative potential for the resulting output. This is a function not only of the algorithm and its inherent biases. In the case of “generative AI”, particularly, it can be innocently created because of “hallucination” in the operating system associated with the algorithm. It is important to recognize that AI is an instrument, like a detector, that is useful, but only within its tested operating regime. Outside it the user can get artifacts, such as “hallucinations” and false positive results. While it is often denied that the phenomenon exists in fact, the problem is that it can and then determining the source of the hallucination. In point of fact, it is almost impossible to determine the origin, depth and quality of the hallucinatory event. It can originate in the bias of the creator, the mathematical accuracy of the algorithm, the mechanical operation on the algorithm, or a combination of these. In effect, such an event tends to nullify the notion that AI-bots enhance research or fact-finding more rapidly than humans. Put another way, the assumed ability of generative AI to be “creative” is likely fallacious at the start. The corollary is that the product of a “generative” AI-bot may well at best be an opinion and not at all factually accurate in its conclusions.

Similarly, the open-endedness of generative AI-bot applications may well lead to wide ranging responses to otherwise inquiries no matter how simply they are posed. The resulting cascade of information that is received by the inquirer might superficially appear erudite and loaded with information. Yet, as Daniel Boorstin stated some three decades ago, “[W]e have gone from an age that was meaning rich but data poor, to one that is data rich but meaning poor. . . [and] this is an epistemological revolution as fundamental as the Copernican revolution”⁸. A modern analogy would be the “Kessler Effect on Human Attention”, a self-reinforcing fragmentation and overload of human focus⁹. Such a situation in itself creates problems for the user who now must sift through the possibly voluminous AI-bot product just as if it were an original search for

⁵(1993) *Daubert v. Merrell Dow Pharmaceuticals, Inc.*, 509 U.S. 579, 113 S. Ct. 2786, 125 L. Ed. 2d 469, 1993 U.S.

⁶John G. Browning (2023) *New and Improved? The Risks of Using ChatGPT*, 86 *Tex.BarJ.* 544 (September).

⁷See Texas Disciplinary Rules of Professional Conduct Rules 1.01 [Competency], 1.03 [Confidentiality], 3.03 [Candor Toward the Tribunal] and The Hippocratic Oath.

⁸Daniel J Boorstin (1994) *Cleopatra’s Nose: Essays on the Unexpected*, (1st edn), Vintage Books, USA, p. 210.

⁹An extension into information overload on human beings similar to the “space junk” concept created in the scenario proposed by NASA scientists Donald J. Kessler and Burton G. Cour-Palais in 1978 in relation to collisional cascading of “space junk” in low Earth orbit that creates a dangerous environment for further exploration.

information, with the result that no time or effort has actually been saved in the process. Once the “sifting” is completed, the inquirer still has the responsibility to evaluate and protect the integrity of the material and arrive at a “human” conclusion. Obviously, the more narrowly the inquiry is focused, the less opportunity the AI-bot has to hallucinate an answer that it thinks is sought, rather than one that is relatively objective. In other words, the amount and quality of data, together with the care in the phrasing of the inquiry is key to minimizing the chances of “hallucination” by the AI-bot to please the inquirer. In this context, the AI-bot clearly does not have a “life” of its own that would otherwise give it the characteristics of *persona ficta*. It is merely *ficta*.

The disruption created in interpersonal relationships among human beings, whether in the workplace or merely social, is becoming reflected in the increased reliance on AI-bots as substitutes for those relationships. What is important to recall in this situation is that there needs to be a clear line maintained between the human and the *persona ficta* that is AI. An important

component of maintaining the ascendancy of the “human factor” as AI-bots move steadily into the day-to-day world is such a consideration as logging the inputs to the AI-bot so as to validate the output. Further, once the output has been received from the AI-bot it should be subjected to cross-correlation in a “critical analysis” format. Finally, given the capabilities of hackers, the use of perhaps multiple firewalls and encryption should be the minimum standard for protecting these recordings, especially if digitally stored. Similarly, transmission of the information to a third person, while permissible or even desirable in some cases, should be by a secure means, and, if electronically, preferably with encryption. The failure to adhere to such standards resulting in a leak of otherwise protected or confidential information, may well constitute professional malpractice *per se* or even a criminal offense in the modern world. The consequences of allowing the AI-bot to free roam like some sort of extraterrestrial intelligent being, while actually being *persona ficta*, among humans are too great to be ignored.