

An Assessment of Artificial Intelligence's Competence in Exercising Judicial Discretion

İhsan Burak Aygün¹  and Arif Çini² 

¹Bursa Uludağ University, Institute of Social Science, Bursa/Türkiye

²Hacı Bayram Veli University, Institute of Graduate Programs, Ankara/Türkiye

* Corresponding author: İhsan Burak Aygün (avburakaygun@hotmail.com)

Abstract: This article examines whether artificial intelligence systems can exercise judicial discretion and function as normative agents within legal adjudication. The study first analyses the meaning of judicial discretion through the perspectives of Hans Kelsen, H.L.A. Hart, Lon L. Fuller, and Ronald Dworkin, arguing that judicial decision-making involves not merely the application of legal rules but also norm-creation, value assessment, and practical reasoning. Building on this framework, the article evaluates whether contemporary generative artificial intelligence systems possess the capacities required for normative agency. Particular attention is given to concepts such as autonomy, responsibility, practical reasoning, and authority, drawing on Joseph Raz's theory of normativity. The study argues that current artificial intelligence systems operate primarily through statistical pattern recognition and lack the capacity for genuine normative evaluation. Furthermore, the article examines the issue through Robert Brandom's conception of the space of reasons and the practice of giving and asking for reasons. It concludes that, although artificial intelligence systems may generate outputs with normative implications, they cannot presently be regarded as genuine participants in the space of reasons or as normative agents capable of exercising judicial discretion and creating law.

Keywords: Artificial Intelligence, Judicial Discretion, Space of Reasons, Normative Agency

Received: 09 June 2026 | Revised: 28 June 2026 | Accepted: 11 July 2026

Cite : Aygün, İ.B. and Çini, A. (2026). An Assessment of Artificial Intelligence's Competence in Exercising Judicial Discretion. *Journal of Artificial Intelligence and Human Sciences*, 3(1), 52-74. Doi: 10.5281/zenodo.21463032

1. Introduction

The increasingly widespread use of artificial intelligence (AI) systems in the field of law has made the question of whether these systems should be regarded merely as assistive technologies, or whether they can be recognised as normative actors capable of participating in judicial decision-making processes to a certain extent, one of the fundamental debates in the philosophy of law. In particular, the fact that generative artificial intelligence (AI) systems have become tools capable of reasoning, generating justifications and proposing solutions to legal disputes has strengthened claims that these systems possess the capacity to exercise judicial discretion. In contrast, many authors argue that legal decision-making is not merely a matter of applying existing legal rules; it involves normative assessment, weighing up values, practical reasoning and the obligation to provide reasons, and they contend that artificial intelligence cannot genuinely fulfil these processes.

At the heart of this debate lies the nature of judicial discretion. If judicial decision-making is viewed as merely the mechanical application of existing legal rules, it might be thought that artificial intelligence (AI) systems with sufficient computational capacity could take on this task. Conversely, when it is accepted that judicial discretion requires the interpretation of legal rules, the assessment of legal principles in the light of specific cases, the consideration of the values the legislature intended to realise, and, where necessary, the generation of normative content, the issue shifts from a technical computational problem to a problem of normative agency. For this reason, the question of whether artificial intelligence (AI) can exercise judicial discretion must first be assessed within the framework of philosophical debates concerning the nature of legal reasoning.

Some views in the literature, including those of Tostschnig and Gregorio, argue that current and future artificial intelligence (AI) systems possess the ability to generate themselves and to comprehend the meaning of their normative outputs. Consequently, these views allow for the interpretation of AI systems as normative actors. However, contrary to these views, this study will argue that current AI systems operate on the basis of statistical pattern recognition and therefore lack the necessary qualities to engage



in the areas of normative evaluation, practical reasoning and justification.

This article is based on the premise that legal discretion is not an activity that can be reduced to the processing of empirical facts or causal explanations; rather, it is a normative practice that takes place within the ‘space of reasons’ rather than the ‘space of causes’. Consequently, the question of whether today’s generative artificial intelligence (AI) systems can exercise legal discretion is not merely a technical issue concerning these systems’ computational or reasoning capacities. It also involves the question of the extent to which these systems can participate in the ‘space of reasons’ and be regarded as actors capable of assuming normative status.

Within this framework, this study re-evaluates existing interpretations of Brandom’s normative pragmatics. In the literature, there are, on the one hand, interpretations that view large language models merely as mechanisms for processing statistical patterns, thereby excluding them from the domain of justifications; and, on the other hand, pragmatist-functionalist approaches that consider their linguistic performance to be sufficient and suggest they can be regarded as participants in the domain of justifications. Our study suggests that the approach of multiple realisability in the philosophy of mind offers a more suitable conceptual framework for interpreting Brandom’s theory. If such an interpretation is adopted, the possibility of situating artificial intelligence (AI) systems, in principle, within the ‘space of reasons’ in the Brandomian sense is not conceptually ruled out. However, this conclusion also raises a new question: what are the necessary conditions for participation in the space of reasons?

The issue that arises at this point is not merely concerned with the position of generative artificial intelligence (AI) systems within the space of reasons. More fundamentally, the extent to which Brandom’s conception of the space of reasons is sufficient for determining the boundaries and membership conditions of this space is also open to debate. Therefore, the study argues that the question of whether artificial intelligence (AI) can be included in the space of reasons simultaneously reveals the conceptual limits of the Brandomian understanding of normativity. Thus, the article aims not only to present an assessment of the normative status of artificial intelligence, but also to discuss the inherent possibilities and limitations of the Brandomian framework regarding the membership conditions of the field of justification.

2. Judicial Discretion

Judicial discretion is a legal authority that affords the judge a certain margin of assessment when applying legal rules to specific cases. As the legislator cannot regulate every eventuality in advance, gaps, ambiguities and provisions open to interpretation inevitably arise within the legal system. Consequently, the judge is not merely a person who applies the rules mechanically, but also a legal practitioner who seeks to reach a fair outcome by taking the specific circumstances of the case into account. Discretion becomes particularly evident in situations involving abstract concepts such as ‘equity’, ‘public order’, ‘good faith’ and ‘reasonable time’.

Judicial discretion may also be exclusively regulated within positive legal systems. For example, Article 4 of the Turkish Civil Code defines this authority as follows: *“In matters where the law grants discretion or requires the judge to take into account the necessities of the situation or justifiable grounds, the judge shall decide in accordance with the law and equity.”* As can be seen here, judges must observe the principles of justice when exercising their discretion. Gardner notes that judges take an oath to fulfil their duties and ensure justice, and dedicate themselves to this end. This is often misunderstood or misinterpreted as a commitment to merely applying existing law and doing no more. This is not what any judge dedicates themselves to; judges dedicate themselves to ensuring justice (Gardner, 2015: 582).

In this study, we will focus on the general normative adequacy of artificial intelligence (AI). Before doing so, we must address certain views within legal philosophy regarding the judge’s discretion.

2.1. Hans Kelsen’s Views on Judicial Discretion

To understand the nature of the judge’s law-making, we must focus on how different approaches in legal

theory address this issue. The understanding of law held by certain positions in legal theory excludes the judge's role in creating law. For example, according to formalists, legal disputes can be resolved without leaving room for ambiguity by referring to valid sources of law. They argue that official state officials merely apply -and are required to apply- valid legal rules (Hart, 2012: 129-130). According to the formalists' deductive reasoning model, the legal rule (major premise) is applied to the factual circumstances relevant to the rule (minor premise), and a judgement is reached (Wacks, 2021: 250). Under this approach, judges' powers are limited to reaching a judgement through logical reasoning, and they have virtually no authority to arrive at contrary conclusions.

The view that existing legal norms are not sufficient to resolve every dispute with a 'gap-free structure' and that there are areas left to the judge's discretion is the dominant position in legal theory. According to Kelsen, law is a systematic order composed of interrelated norms, and each norm derives its validity from a higher-ranking norm. However, according to Kelsen, it is not possible for all social relationships and social phenomena to be regulated by legal norms within a legal system. In fact, the legislature often deliberately refrains from regulating certain matters through legal norms and leaves gaps in the law (Kelsen, 2009: 152).

In situations where the legislator has left gaps in the law, judges are granted a discretionary power to adjudicate the relevant case. Whilst Kelsen accepts that individual norms are related to general norms, he opposes the view that the derivation of individual norms from general norms is determined by mere logical reasoning. Courts do not merely apply the law based on valid general norms; they also create new norms (Kelsen, 1973: 246-250). Every individual norm that emerges is the result of a voluntary act, and there is a normative actor involved. Kelsen also opposes the approach that, in certain situations requiring discretion, judges do not apply the law but merely create it. In such cases, judges create a norm by applying the norms that grant them authority. From this perspective, the activities of norm-creation and compliance with the norm are, as it were, carried out simultaneously. However, this understanding does not claim that there is already a general norm to rely upon in a situation falling within the judge's discretion. There may be no valid legal norm to apply; there may be uncertainty regarding the scope of the norm; or, even if an applicable norm exists, its application may conflict with the values the legislator sought to establish. In such a situation, the judge creates a legal norm. The general norm followed here concerns who is authorised to create law. Kelsen interprets the legislature's law-making activity from this perspective as well. The legislature does not merely create law; it also applies the constitution that authorises it (Kelsen, 2009: 153-155).

Although Kelsen views legislative activity as a volitional act, he opposes the view that judicial proceedings are merely a matter of applying the law. There is not a qualitative difference but a difference of degree between the two activities. In the hierarchy of norms, as one moves downwards, the scope of discretion generally narrows, whilst as one moves upwards, the scope of discretion increases (Kelsen, 2009: 272).

Kelsen argues that interpretation cannot be determined by legal logic alone; rather, the creation of individual norms ultimately depends upon the judge's volitional act. According to Kelsen, the activity of the courts—which both applies and creates the law—is always authentic and can be attributed to the judge's will (Kelsen, 2002: 348-356). At this point, it would be wrong to conclude that the law is entirely unpredictable and uncertain.

In situations where the margin of discretion is broad, it is also possible to characterise the judge's activity as a political one. In this respect, constitutional courts or courts of appeal may also be held accountable in terms of political accountability. In positive legal systems, the courts' discretion is generally granted within the limits set by legal rules. However, it is also expected that decisions made within the scope of this discretion will be justified by taking into account their social and political context. Courts at this level are also aware that their discretionary activities may produce such consequences. They often wait for a certain public consensus to form at the societal level before issuing certain decisions and take the prevailing social mindset into account when doing so. Consequently, the question of when decisions

should be made may also require judges to engage in value-based reasoning.

2.2. Hart's Views on Judicial Discretion

A significant part of Hart's legal philosophy is centered on the relationship between law and language. According to Hart, it is not possible to understand words used in everyday life correctly in isolation from the conceptual context in which they arose and developed. Hart argues that language is 'open-textured'. According to Hart, if people are to communicate—as in the most basic form of law—and express that they wish a certain type of behaviour to be regulated by rules, then the general words used must have standard meanings that are beyond doubt. Consequently, in both everyday language and legal language, most words have an established, standard core meaning. This core meaning encompasses all the meanings the word clearly encompasses within a particular language community (Hart, 1958: 607).

According to Hart, the judge's discretion stems from the fact that the law cannot resolve all disputes in advance in a definitive and exhaustive manner. In his view, laws are written in general terms, and it is impossible to foresee in advance all the specific circumstances that may arise in the future. Consequently, in some cases, existing legal rules cannot provide a direct and definitive answer. At this point, Hart distinguishes between 'easy cases' and 'difficult cases'. A case is characterised as 'easy' when there is a general consensus that the relevant case falls within the scope of the rule, and there is no doubt regarding the content of the legal rule or its direct applicability (Hart, 1958: 610).

Hart argues that, in addition to language having a complex structure, because humans are not perfect beings, they can never foresee the future with absolute clarity, and the laws created by humans will always be incomplete and flawed. This situation will lead to the emergence of difficult cases. Every legal system inevitably contains areas of uncertainty. In difficult cases, the meanings of the concepts contained in the applicable rules will fall within this area of uncertainty, and this uncertainty will hinder the application of legal standards. Consequently, it will not be easy for judges to reach a decision in such cases (Hart, 1958: 610).

According to Hart, difficult cases arise where the meaning, scope or applicability of legal rules is uncertain. In difficult cases, Hart argues, the judge is not merely someone who applies the rules mechanically. This is because the legal system does not provide a pre-determined, clear solution here. Consequently, the judge is compelled to exercise a certain degree of discretion. When existing legal rules do not provide a single answer, the judge exercises this discretion by acting as a 'conscientious lawmaker' and fills the gap in the law. Hart acknowledges that in such situations, the judge creates law. However, this is not an arbitrary or unlimited power. The judge's creation of law is distinct from the legislature's enactment of laws (Spaic, 2014: 10).

The judge's discretion is not an arbitrary power. As the scope of this power is limited to specific difficult cases and the judge is compelled to create law to resolve them, this power cannot be used to introduce large-scale reforms or new laws. A judge must always have certain general grounds to justify their decision and, when forming their opinion and rendering a judgment, must weigh up different interests and values against one another and strike a rational balance between them (Hart, 1983: 106-107). Therefore, judicial discretion is a limited authority that judges exercise to resolve cases in situations where the law is not fully defined or its meaning is not fully understood.

Hart uses the example of "no vehicles allowed in the park" to explain his views on judicial discretion. He considers a hypothetical legal rule prohibiting the entry of "vehicles" into a public park. The key point in this rule is the term "vehicle." This is because, by using the term "vehicle," the rule explicitly prohibits the entry of cars into the park. However, would bicycles, skateboards and toy cars fall within the scope of this prohibition? According to Hart, if a word relating to a mode of behaviour appears in a rule, that word must have a standard meaning that leaves no room for doubt. However, this is not always the case for the words used in rules. For example, if a judge were to base their interpretation of the rule "no vehicles may enter the park" on a standard case, they could arbitrarily define the characteristics of the term "vehicle" as: 1) normally used on land, 2) capable of carrying a human person and 3) capable

of being self-propelled. However, defining the meaning of the word “vehicle” in this way would not be very logical. This is because interpreting the word “vehicle” in this way would prohibit a toy motorised car from entering the car park, whilst failing to prevent an aeroplane from doing so (Hart, 1958: 610–611). Hart has pointed out that in difficult situations, if a judge were to act strictly according to the standard meaning of the word, as in the example, the resulting decision would be little more than a token one. Although Hart criticises this approach as lacking in reason, he does not offer a concrete suggestion as to how judges should proceed. On this basis, Hart’s views on judicial discretion have been criticised by Dworkin (Uzun, 2015). Furthermore, Fuller has stated that Hart makes an assessment without addressing the primary purposes of legal rules.

2.3. Fuller's Criticisms

Lon L. Fuller (1958: 662) has criticised H. L. A. Hart’s views on judicial discretion, particularly in the context of the ambiguity of legal language and difficult cases. Before launching into his criticisms, Fuller attempts to explain the thesis put forward by Hart. According to Fuller, Hart assumes that the judge has no creative role in the application of the core meaning of a word. In this case, the judge must simply apply the law as it stands. However, if the subject matter of the case in question requires an assessment of the “shadow area” of the words, the judge is obliged to act more creatively. Yet, according to Fuller, this theory is clearly flawed.

According to Fuller, the meaning of a legal rule cannot be derived solely from the dictionary definitions of the words. To determine what a rule means, one must take into account its purpose, function and place within the legal system. According to Fuller, Hart is mistaken when interpreting a legal rule based on individual words. When it comes to the law, one must focus not on the meaning of a single word, but on the meaning of the sentences, paragraphs and even pages within the legal rule (Fuller, 1958: 662). According to Fuller, however, what makes a case easy or difficult is not the meanings of the words—as in the vehicle example—but rather whether the rules can be easily applied to the relevant circumstances (Fuller, 1958: 663).

Fuller criticises Hart’s view that the application of a rule is determined by adhering to the core and shadow meanings of individual words in the rule’s formulation, by presenting counter-examples to the “no vehicles in the park” example. Fuller’s example is as follows: A truck dating from the Second World War is to be placed in the park by patriots as a war memorial. However, those who believe the truck would be an eyesore in the park are demanding that it not be allowed in, citing the rule “no vehicles in the park” as justification (Fuller, 1958: 663). At this point, should this justification be accepted as valid? Does the truck, accepted as a war memorial, fall within the core or the shadow area?

Fuller criticises Hart on the grounds that this ambiguity would also undermine the ideal of adherence to the law. Fuller evaluates a law stipulating that individuals sleeping at a train station should be fined in cash, considering two different scenarios. In the first scenario, a person, fully dressed and presentable, is waiting on his feet for a delayed train at night. However, the attendant has arrested the man waiting on his feet because he heard him snoring. In the second scenario, a person hoping to spend the night at the station, lying on a bench with a blanket and pillow, has been arrested before he had a chance to fall asleep. Fuller questions whether the interpretation of the word “sleep” as “lying down on the floor or on a bench at the station to spend the night” -resulting in the second man being punished whilst the first is released- might undermine the principle of adherence to the law (Fuller, 1958: 663–664). Consequently, in light of the information provided, Fuller criticises Hart for focusing on the literal meanings of individual words when interpreting the law and for disregarding the law’s underlying purpose.

Fuller’s critique of Hart offers a new perspective on understanding the nature of judicial discretion. There may be a valid legal norm within the legal system, and there may be no doubt as to the meaning of the linguistic objects that express that norm, or the law may define it directly. Furthermore, it may be considered that the specific case falls within the scope of the general norm. Despite all this, the judge may experience hesitation in applying the valid norm. At this point, it is considered that applying the valid rule in this manner conflicts with the aims and values that the legislator seeks to protect and bring

into being. Consequently, it becomes clear that the activities of norm-creation and norm-application cannot be reduced to a mere matching of words with factual situations, but additionally require value-oriented thinking. Whether any form of artificial intelligence (AI) can carry out such an activity depends on the debate regarding whether it possesses such competence.

To better grasp the nature of norm-creation and application, we must focus on the distinction between reflecting on values and actualising them in life, drawing on a Farabian perspective. Farabi distinguishes between theoretical reason and practical reason. Theoretical reason involves thinking directed towards knowledge of objective normative truths, whilst practical reason involves thinking about what accidents must be added to these objective normative truths in order to bring them into being in life. The realisation of normative truths or goods in life, however, is contextual. That is to say, different times or different communities may necessitate different concrete arrangements or require the tools used to realise the value to differ (Farabi, 2020: 57-63). We can interpret Fuller's critique of Hart through this distinction. When the legislator establishes a norm, the aim is to protect a value. In the rule 'No vehicles in the park', we can concretise this as "the creation of a peaceful and safe public space". This general aim does not require a single concretisation; various norms can be derived from it. Prohibiting the display of a tank or a lorry as a war memorial would not serve this overarching aim. Consequently, applying the valid law based solely on its literal meaning would not serve the aim the lawmaker sought to achieve. Therefore, it could be argued that displaying a tank in the park does not constitute a breach of the norm. However, such a justification differs to a certain extent from the theoretical reasoning in the distinction mentioned above. In this justification, the individual does not directly investigate what objective normative realities require. Instead, the individual primarily takes as their basis the objectives underlying the legislator's normative claim and is concerned with how these objectives are to be realised in the specific case. However, it would, in our view, be mistaken to assume that such a justification entirely disregards consideration of objective realities. This is because the assessment of how a particular objective is to be realised is itself based on certain values.

2.4. Dworkin's Criticisms

One of Ronald Dworkin's most significant contributions to the philosophy of law is that he did not view legal norms as consisting solely of "rules," but rather counted "principles" among the essential elements of law. Dworkin establishes this distinction between rules and principles largely to critique the positivist conception of law. In criticising the positivist understanding of law, he generally targets Hart's ideas. In his famous work "Taking Rights Seriously", he explains this as follows (Dworkin, 1978: 22):

"I want to make a general attack on positivism, and I shall use H. L. A. Hart's version as a target, when a particular target is needed. My strategy will be organized around the fact that when lawyers reason or dispute about legal rights and obligations, particularly in those hard cases when our problems with these concepts seem most acute, they make use of standards that do not function as rules, but operate differently as principles, policies, and other sorts of standards. Positivism, I shall argue, is a model of and for a system of rules, and its central notion of a single fundamental test for law forces us to miss the important roles of these standards that are not rules."

With these ideas, Dworkin points out that positivist theory is merely a model of rules. However, according to Dworkin, within the legal system, alongside rules, there are also principles that have a different structure from rules. These principles come to the fore in cases we might describe as difficult, where rules prove insufficient or fail to ensure justice. The most significant difference between rules and principles lies in their modes of application. Rules are of an "all or nothing" nature. If the conditions of a rule are met, that rule is applied. However, the operation of principles is different. Principles do not mandatorily require a specific outcome; rather, they provide weighty considerations to be taken into account when making a decision. (Dworkin, 1978: 24-25). When rules conflict with one another, at least one rule must be invalidated. However, this is not the case with principles. There is a relationship of "weight and importance" between principles. In one situation, one principle may prevail; in another, the other principle may take precedence. However, the principle that loses out does not become entirely

invalid (Dworkin, 1978: 48).

Dworkin (1977: 46) states that the discretion granted to the judge by Hart in difficult cases constitutes strong discretion. However, the problem here lies in Hart's view that law consists solely of rules. Consequently, in Hart's legal system, when a judge cannot find a legal rule to apply to the relevant case, they will be compelled to exercise their discretion. According to Dworkin, however, in situations where existing rules are insufficient to resolve difficult cases, judges must turn to principles. Consequently, what the judge does in difficult cases is not an exercise of discretion, but an interpretation of the concepts used in the formulation of principles (Dworkin, 1977: 46).

To explain the judge's role in resolving difficult cases and the limits of legal interpretation, Dworkin (1978: 106) proposes a fictional ideal judge model he calls Hercules. Judge Hercules is a figure who operates within ideal formulations and is capable of delivering the most correct judicial decision possible. Dworkin's ideal judge model, Hercules, possesses the ability to find the single correct answer in difficult cases. In finding this single correct answer, he views the law not as a fragmented collection of rules, but as a coherent whole composed of principles of justice, fairness and procedure. For Hercules, adjudication is a single "constructive interpretation" process that always aims to present the law in its "best light". The correct answer is derived from the principles within the legal system, even if not from the rules themselves. Consequently, it is not accurate to say that judges will exercise their discretionary powers when the rules run out. On the contrary, judges are able to make decisions within the law, using their discretion whilst remaining faithful to the principles that exist alongside the rules in the legal system, without creating new law (Dworkin, 1978: 106). On these grounds, Dworkin has criticised Hart's views regarding the judge's discretion.

When Dworkin outlines an ideal model of a judge, whom he calls Hercules, he does not argue that all judges must possess Hercules's characteristics. Hercules serves as a role model for judges. He is a useful figure for judges. When assessing their own performance, judges can easily identify their shortcomings by using Hercules as a benchmark and have the opportunity to improve themselves (Dworkin, 1978: 107).

Dworkin's (1978: 111) view that, where rules run out, the law is not created by stepping outside the law but rather by interpreting the morally grounded legal principles already present within the legal domain, offers a different perspective on the issue of judicial discretion. However, does this interpretation pose a problem in light of the views on the nature of law-making and law-application that we have discussed and sought to outline above? In a Dworkinian understanding, treating the law as a holistic structure requiring the best normative interpretation at the level of principles, whilst safeguarding the integrity of the legal system during this process, imposes a greater normative responsibility on the judge compared to the approach of filling legal gaps through judicial discretion. In this context, if artificial intelligence (AI) is to be accepted as a subject (agent) in legal adjudication, it is also necessary to examine whether it can assume this additional normative burden indicated by the Dworkinian approach. Legal adjudication, in this understanding, requires the capacity to reflect on the weight and importance of legal principles with moral content. Since the weight and importance of principles may vary according to contextual conditions, the normative burden of constructive interpretation will differ in each specific case. Consequently, from a Dworkinian perspective, the subject conducting the adjudication must both consider the legislator's purposes and reflect upon them, interpreting constitutional and political history in a manner that renders the relevant aims and values visible, whilst also taking into account the importance and weight of these values, which may vary according to the context.

The different approaches examined above demonstrate that judicial discretion is not merely a matter of applying existing norms. When these approaches are considered together, it can be said that at least five fundamental competencies are required for judicial discretion to be exercised: (i) the ability to make normative assessments, (ii) the ability to choose between values, (iii) the ability to justify one's decision through practical reasoning, (iv) the ability to take responsibility for the resulting decision, and (v) the ability to establish a link between the objectives of the legal system and the circumstances of the specific

case. The key question to be addressed at this point is to what extent today's artificial intelligence systems are able to meet these competencies. Consequently, in the next section of this study, artificial intelligence (AI) systems will be evaluated on the basis of these criteria.

3. Can Artificial Intelligence Be a Moral Agent?

At the heart of the debate surrounding the use of artificial intelligence (AI) in the courts lies the question of whether these systems are merely technical tools or whether they have evolved into decision-making mechanisms capable of producing normative outcomes to a certain extent. Particularly with the development of generative artificial intelligence (AI) and machine learningbased systems, it is argued that artificial intelligence (AI) is not merely a passive technology that processes existing legal data; rather, it can generate new outcomes based on specific patterns and, in this respect, may indirectly influence legal creation processes. Therefore, when assessing the role of artificial intelligence (AI) in courts, its learning capacity and normative impact must be considered together.

The primary aspect of artificial intelligence's (AI) learning process that requires examination is its linguistic transformation. The learning process of large language models (LLM) is fundamentally based on the identification of statistical patterns. During the training phase, the model attempts to predict the next word in a text and updates its own parameters based on the errors it makes. As a result of this process, repeated numerous times, the model becomes capable of establishing quite comprehensive relationships regarding the syntactic and semantic structures of language. In these models, no attempt is made to teach the machine a grammar, nor are individual words inputted. Instead, various sentences and paragraphs are fed into the machine so that it can learn the meanings of words, their antonyms, the structure of the language and the different types of sentences. During this pre-training process, the machine analyses sequential sentences to grasp the language's inherent structure and focuses on understanding how individuals express themselves in the language. Any structure that paves the way for the system's success becomes a permanent feature. However, the fact that artificial intelligence (AI) models possess these features should not be interpreted to mean that the artificial intelligence (AI) "understands" the text it generates (Kreischer, 2025: 5).

After being exposed to specific sentence patterns and word groups, artificial intelligence (AI) language models (LLM) undergo developmental leaps known as "phase transitions". Following this developmental process, artificial intelligence (AI) models reach a stage where they are capable of making inferences. However, what artificial intelligence (AI) essentially does is capture the statistical distribution of the word sequences used by humans. Artificial intelligence (AI) language models (LLM) predict the use of a word based on words in past sentence patterns. The generated texts consist of words selected by the AI according to probabilistic distributions. Therefore, although the resulting texts may appear meaningful, they do not rely on the artificial intelligence's (AI) semantic awareness; they are merely the product of data statistics (Sambrotta, 2025: 46). Bender et al. (2021: 616) likens artificial intelligence (AI) to a "stochastic parrot" because it generates text based on statistical and probabilistic information.

The concept of the "stochastic parrot" has been used to describe how language models (LLM) generate text based on pre-existing patterns, whilst statistically capturing the context. This raises the question: Does artificial intelligence (AI) really understand? In this study, aligning with Bender's view, we will argue that in the text-generation process, artificial intelligence (AI) operates solely on the basis of statistical and probabilistic data, without "understanding". However, the opposite has been argued in the literature. For example, Havlik (2024: 19) has stated that artificial intelligence (AI) language models do not possess merely a shallow structure that operates on statistical principles, but rather that these language models (LLM) possess concrete understanding and are equal users to humans when working with meanings. Cappelen and Dever (2025: 13), on the other hand, have argued that ChatGPT possesses a language model as advanced as that of humans. According to the authors, artificial intelligence (AI) can engage in meaningful dialogue like humans, ask questions to resolve issues, make claims, and generate responses. Artificial intelligence (AI) may be aware of certain things, know things, and desire things. It can acquire knowledge about the world. It can make plans and reason about how to carry them

out. Totschnig (2020: 2474), however, draws a distinction between current and future systems regarding the way artificial intelligence understands the world.

Totschnig (2020: 2474) evaluates the concept of “autonomy” in relation to artificial intelligence (AI) from a philosophical perspective. According to Totschnig, the purpose of artificial intelligence (AI) is not merely to carry out tasks assigned by humans independently. Artificial intelligence (AI) may possess the ability not only to act based on human inputs but also to determine its own objectives. At this point, Totschnig addresses the finality argument regarding the autonomy of artificial intelligence (AI). According to this argument, the scope of artificial intelligence (AI) authority is limited to achieving the ultimate goals it has been taught. Artificial intelligence (AI) can never change its own ultimate goal, because changing this goal is irrational in terms of the system and structure of artificial intelligence (AI) and would have a counterproductive effect on the current ultimate goal (Totschnig, 2020: 2475). Bostrom (2014: 150) argues that artificial intelligence (AI) cannot possess full autonomy and that its dependence on its ultimate goal would also create a control problem, illustrating this situation with the “paperclip scenario”.

Bostrom (2014: 150-153) considers an artificial intelligence (AI) created by humans with the ultimate goal of producing paperclips. He imagines this artificial intelligence (AI) developing to the point where it surpasses human intelligence, becoming a “superintelligence”. This artificial intelligence (AI) will possess a single ultimate goal and capabilities superior to those of human intelligence. Consequently, it will produce paperclips whenever the opportunity arises. Before long, it will utilise all the world’s resources to produce paperclips, and the world will be overflowing with them. However, Totschnig (2020: 2480) points out that whilst Bostrom envisions a super-intelligent artificial intelligence (AI), he bases his analysis on today’s machines and draws conclusions solely based on the capabilities of these AIs. Yet, according to Totschnig, current artificial intelligence (AI) models do not possess a general intelligence system. The way these artificial intelligence (AI) models understand the world is limited to what they have been taught, and this limitation remains unchanged throughout their operational lifetimes in their current state. However, AI models that emerge in the future may possess general intelligence. In this scenario, artificial intelligence (AI) models would be able to understand the world better and possess a rational understanding of the nature of their ultimate goals. Consequently, they would be able to question the progress of their ultimate goals and, should an irrational situation arise, possess the ability to evaluate their ultimate goals and perhaps even alter them in a rational manner. Therefore, in the presence of such an artificial intelligence (AI) model, the fear put forward by Bostrom would be unfounded (Totschnig, 2020: 2483).

We agree with Totschnig’s view that, in the event that artificial intelligences (AI) attain superintelligence, Bostrom’s concern would be unfounded. However, it must be noted that Totschnig’s views -that future artificial intelligences (AI) may possess the ability to determine their own objectives and, in a sense, become “autonomous”- could be argued to contain an overly optimistic expectation regarding the future potential of artificial intelligence (AI) technologies. Drawing conclusions based on future expectations regarding artificial intelligence (AI) may lead to misleading ideas. For this reason, we believe it is more accurate to base conclusions on the current state of affairs rather than future expectations when making inferences about artificial intelligence (AI). Furthermore, when assessing the legal reasoning capacity of artificial intelligence (AI), Hart’s views on the language of law and its open texture must also be taken into account. According to Hart (1958: 610), the application of legal rules is not merely a matter of determining the dictionary meanings of words. Owing to the open texture of language, decisions regarding how rules are to be applied in borderline cases are the product of a specific social practice and human reasoning. Whilst today’s artificial intelligence (AI) systems possess highly advanced language processing capabilities, they do not learn language as a participant in a social practice in the sense indicated by Hart. These systems are able to identify statistical regularities in the use of concepts but cannot access the normative justifications that explain why concepts are used in specific ways within particular contexts. Consequently, whilst artificial intelligence (AI) may successfully reproduce many patterns in legal language, it cannot fully replace the reasoning capacity possessed by a

human judge in situations arising from Hart's 'open-text problem' that necessitate judicial discretion. From the perspective of our study, we will assess whether artificial intelligence (AI) can make normative decisions and whether it possesses the capacity to create law based on current artificial intelligence (AI) models.

The issue of law-making by legislators or judges within a legal system requires not only a legal assessment but also a moral one. Although making a moral assessment within a legal system is generally identified with natural law, this view is evident even in the work of Joseph Raz, a staunch positivist. According to Raz (1979: 11), laws that are morally deficient lack legitimate authority. Raz bases a significant portion of his views on law upon the concept of authority. Raz (1979: 8-9) discusses multiple forms of authority, but the type of authority relevant to our study -and the one Raz generally focuses on- is legitimate actual authority. The nature of legitimate actual authority requires that a person has the right to issue instructions that invalidate the reasons for a specific action (or inaction) and to ensure that subjects comply with these instructions, treating them as reasons that supersede their original reasons for or against an action. However, for a person to have the right to subject their own will and judgement to that of another, they must possess certain justifications. According to Raz, an authority is legitimate not merely because it issues commands, but to the extent that it helps individuals act in a manner more consistent with the reasons and obligations they already possess. In other words, authority can be justified only to the extent that it "serves" individuals. Raz (2009: 136) conceptualised this as the "service conception of authority" and explained it through two fundamental theses: (1) Dependence Thesis (2) Normal Justification Thesis.

In short, dependence thesis means that when a legitimate authority issues a specific instruction, it must take into account the reasons for action of those to whom the instruction is addressed. Normal justification thesis is expressed as follows: if a *de facto* authority issues an instruction to its addressee and compliance with this instruction aligns with the addressee's reasons for action with a higher probability and degree, then the said *de facto* authority is characterised as a legitimate authority in relation to the relevant addressee and subject matter. These two theses together are referred to as the "service concept of authority". The service concept of authority explains, from a moral perspective, how it is legitimate for an individual to be subject to another's authority (Raz, 1988: 47-53). Consequently, if an authority's instruction serves as a guide for the addressee and, by complying with the instruction, the addressee's own reasons are better aligned, and there is no situation where it would be more appropriate for the addressee to make their own decision, then the individual's compliance with the authority's decision will be legitimate (Raz, 2009: 136-137). According to Raz (2009: 9), all authorities must make this claim to legitimacy. An authority making such a claim must evaluate each of the legal and moral reasons for action whilst guiding the individual. As can be seen, Raz evaluates law according to the role it plays in the practical reasoning of the normative subject. There are already reasons for action, including those derived from ethics directed towards the subject. Law claims to provide a separate reason for actions whilst also taking the subject to a further stage in reaching the reasons already directed towards them. At first glance, this gives the impression that law is treated as a separate subject. However, since law is ultimately grounded in social phenomena, it is possible to trace its origins back to the normative subjects or groups of subjects that create it. What is significant for our article here is that, in Raz's conception of law, both the normative subjects (citizens) who are the addressees of norms and the subjects who create norms are conceived as entities possessing the capacity for practical reasoning; capable of evaluating the reasons directed towards them and possessing the competence to choose to act or refuse to act in accordance with what these reasons require. Consequently, the activity of norm-creation, due to its impact on subjects' practical reasoning, requires access to inter-subjective reasons and a normative activity grounded in the space of reasons. But to what extent is this a plausible scenario for artificial intelligence? Does artificial intelligence (AI) possess the prerequisites for conducting a legal adjudication, particularly when the nature of situations requiring discretion is taken into account? Will it be able to carry out the necessary normative assessment when creating a legal norm, when creating law? Does it possess the capacities of a normative subject?

The question of whether artificial intelligence (AI) can be an moral agent has been the subject of extensive debate in recent years. The parties to this debate are divided into two camps: the standard view and the functionalist view. The standard view argues that moral agency cannot be explained solely by external behaviour; rather, it requires the presence of certain internal, mental qualities such as intention, consciousness and responsibility. According to this approach, for an entity to be considered a moral agent, it must possess rationality, free will or autonomy, and phenomenal consciousness (Zafar, 2025: 2607). In particular, as argued by Deborah Johnson (2026: 198), moral action must arise as a result of the agent's internal mental states, comprising their beliefs, desires and intentions. The reason why humans are accepted as moral agents is that their behaviour can be explained not merely by physical processes, but by conscious intentions and purposes. For this reason, the standard view maintains that the mere fact that artificial intelligences (AI) exhibit human-like behaviour does not, in itself, make them moral agents.

The functionalist view, however, grounds moral agency not in specific internal mental states, but in observable functions and behavioural capacities. According to Floridi and Sanders (2004: 361-363), two of the most prominent proponents of this approach, phenomenal consciousness is not a prerequisite for being a moral agent. What matters is that an entity interacts with its environment, can alter its own state to a certain extent independently, and can adapt its behaviour by learning from its experiences. Functionalists argue that human-centred criteria must be abandoned when determining moral agency. In their view, moral agency stems not from specific biological or mental structures, but from the fulfilment of specific functions.

Within this framework, advanced artificial intelligence (AI) systems can also be considered moral agents. This is because many artificial systems collect data from their environment, process this data to make decisions, and can modify their subsequent behaviour based on previous outcomes. Functionalists also emphasise that humans attribute characteristics such as consciousness or intention to others not through direct observation, but through their behaviour (Floridi & Sander, 2004: 363-364). However, Mandy Zafar (2025: 2608) has pointed out that views advocating the acceptance of artificial intelligence (AI) as a moral agent fall into four fundamental errors. She classifies these fundamental fallacies as false analogies, misleading reductionism, ignoring the logical interdependence of terms in the social space of reasons, and the ignoring normativity. Zafar (2025: 2610) explains the argument regarding false analogies as the use of concepts in weaker or different senses -detached from their standard meanings-when attributing human-specific characteristics to artificial intelligence (AI). In the functionalist view, concepts such as intention, reasoning, autonomy and responsibility are explained solely through data processing or algorithmic decision-making processes. Consequently, moral characteristics possessed by humans but not actually present in current artificial intelligence (AI) systems are attributed to artificial intelligence (AI), and this constitutes a flawed argument.

Zafar's (2025: 2611) second argument concerns misleading reductionism. The misleading reductionism of functionalist views implies that processes of reasoning are regarded as equivalent to simple information processing and algorithmic processes. As a moral agent, a human being does not act solely by considering statistical and probabilistic data when producing a norm. As Joseph Raz argues, authorities undergo a reasoning process in which they evaluate the existing reasons for action when producing a norm for individuals to follow. The norm-production processes of authorities are not limited to merely producing the norm; they also evaluate and justify the reasons behind the norms. However, the functionalist approach reduces this process to algorithmic procedures or technical calculations. According to Zafar, artificial intelligence (AI) systems cannot grasp the meaning or moral value of their actions. Therefore, reducing the concept of moral agency to technical performance both oversimplifies the true nature of human agency and overstates the capabilities of artificial intelligence (AI). At this point, by accepting Zafar's view, we may conclude that Fuller's conception of judicial discretion would be problematic in relation to artificial intelligence. According to Fuller's conception, the judge must understand and interpret the purpose and moral value of the norm. However, as artificial intelligence models lack this capacity, they cannot be agents capable of exercising judicial discretion.

The third argument concerns the ignoring the logical interdependence of terms in the social space of reasons. To qualify as a moral agent, one requires numerous elements -such as autonomy, self-awareness, rationality and social interaction- that are logically interlinked (Zafar, 2025: 2611). Rodrigo Sanz (2020: 225) argues that for any artificial intelligence (AI) model to be accepted as a moral agent, the artificial intelligence (AI) must actively participate in discourse and various pragmatic social interactions within moral communities, and be able to deal not only with what is, but also with what might be. Consequently, being a moral agent requires being part of a social process. Whilst artificial intelligence (AI) systems possess the capacity for data processing and output generation, they cannot engage in the exchange of reasons or make normative judgements by operating within the space of reasons. Consequently, to claim that artificial intelligence (AI) -which evaluates only statistical and probabilistic data- is a moral agent overlooks the fundamental logical structure of moral agency. Mandy Zafar emphasises that attributing “agency” to artificial intelligence (AI) requires the foundational elements of human agency, such as self-awareness, autonomy, responsibility and practical reasoning. This requires a capacity that goes beyond mere logical rules. However, artificial intelligence (AI) operates as a structure that processes rules without contextual sensitivity, and therefore cannot fully access the normative and logical space of reasons (Zafar, 2025: 2618-2619). Unlike the application of abstract rules, human practical reasoning involves normative evaluation of what ought to be done and weighing up intersubjective reasons. Judicial discretion, as we noted in the first section, is based precisely on this type of reasoning; that is, it requires the contextual evaluation of reasons and the determination of their normative weight. Consequently, to the extent that artificial intelligence (AI) cannot be regarded as a moral agent capable of carrying out this kind of practical reasoning, it cannot fulfil the capacity for normative evaluation required by judicial discretion (Zafar, 2025: 2612).

The fourth argument concerns the ignoring normativity. According to Zafar (2025: 2613), views that argue artificial intelligence (AI) is a moral agent do not take sufficient account of the normative dimension inherent in moral agency. The point that is overlooked here is the distinction between being a moral agent and existing within the moral domain. It is true that the outputs of artificial intelligence (AI) models produce consequences in the moral domain. However, this does not mean that these artificial intelligence (AI) models are moral agents. The sole reason they can exist in the moral domain is that they produce outputs which moral agents might take into account, which could contribute to their assessments, or which could produce outputs that harm or benefit a moral agent. Chomsky (2023) explains the reason why artificial intelligence (AI) models cannot be moral agents as “moral indifference stemming from a lack of intelligence”. In our view, artificial intelligence (AI) models lack not only intelligence but also the requirements of concepts such as “will” and “responsibility”, which are inherently normative. Humans can be regarded as moral agents capable of providing reasons for their actions and evaluating those reasons. Artificial intelligence (AI), however, is a structure that lacks the capacity for these concepts within the space of shared normative reasons, and merely executes statistical and technical processes. At this point, the distinction between using a normative language and engaging in moral reasoning is also significant for our study. Artificial intelligence (AI) models will be able to use a normative language thanks to statistical data. However, as we have previously noted, artificial models cannot understand, adopt or critique norms. Even if the recommendations and decisions produced by artificial intelligence (AI) create an impact in the normative sphere through moral agents, this does not stem from the artificial intelligence’s (AI) own moral reasoning; it is the result of programming and data structures. Consequently, we argue that artificial intelligence (AI) models are not moral agents capable of making normative decisions.

One of the fundamental issues that must be addressed when assessing whether artificial intelligence (AI) meets the conditions for being a normative agent is whether these systems can be considered autonomous and whether they are entities to which responsibility can be attributed. Although artificial intelligence (AI) systems possess the capacity for learning and adaptation, this does not imply that they have become autonomous agents, even if they generate a higher degree of opacity and unpredictability. For even the learning process in these systems occurs as a function of predefined algorithmic operations; there is no

question of the system exceeding its own operational rules, rejecting them, or making a normative decision regarding when to learn and when not to (Zafar, 2025: 2618). Autonomy appears to encompass not merely the ability to choose between alternatives, but also the capacity to weigh up multiple seemingly valuable options and even to opt for inaction. In this sense, an autonomous agent is not merely one who applies rules, but one who can reflect upon the rules and the values upon which they are based. The judge's discretion precisely requires this kind of normative assessment. Furthermore, the judge's legal obligation appears to logically entail not only the necessity to decide in a specific manner but also the possibility of acting contrary to this obligation; for an obligation is meaningful only when the possibility of acting contrary to it is assumed. For this reason, the legal system regards the judge as both a subject who can be held accountable and an autonomous agent; for the application of norms relies not merely on mechanical enforcement, but also on the capacity for justification, weighing up options, and excluding alternatives regarding these norms. Consequently, to the extent that artificial intelligence (AI) systems operate within the boundaries of the algorithmic rules governing their activities, they cannot demonstrate a capacity to fulfil the normative and practical reasoning-based meaning of autonomy intended here. If autonomy is accepted as a prerequisite for liability, it does not appear possible to regard artificial intelligence (AI) systems lacking such autonomy as liable agents.

Giovanni De Gregorio (2023: 6-9) has proposed a different conceptualisation of the normative power of artificial intelligence (AI). According to Gregorio, norms generated by artificial intelligence (AI) do not stem from social phenomena. However, as Gregorio also points out, since the norms produced by artificial intelligence (AI) do not rely on social phenomena—which embody a shared collective will—they cannot be accepted as legal norms. The author defines the norms produced by artificial intelligence (AI) at this point as “self-generating technical norms”. Consequently, artificial intelligence (AI) will form a normative layer independent of legal norms. The fact that artificial intelligence's (AI) so-called norms are not based on social phenomena may have specific implications within the context of jurisprudence. Upon examining schools of legal philosophy, it becomes evident that the prevailing view is that legal rules must be grounded in social phenomena. Despite the significant differences between natural law theory and legal positivism, both approaches accept the condition of social facts, stating that the existence and content of law are dependent on social facts (Murphy, 2017: 352). Consequently, if the normative outputs of artificial intelligence (AI) systems cannot be interpreted as being linked to social facts, it will be argued from the perspective of these theories that they cannot contribute to the existence and content of law for this reason alone. However, this conclusion may stem from an interpretation based on the premise that the normative outputs of artificial intelligence (AI) are entirely independent of social practices. Conversely, from a perspective such as that of Leonardo Marchettoni, it can be argued that generative artificial intelligence (AI) systems can participate in social normative practices to a certain extent. According to this approach, generative artificial intelligence (AI) systems are capable of making commitments, providing justifications, defending their previous claims, and responding to the criticisms of their interlocutors (Marchettoni, 2025: 8). Consequently, the normative activities of such systems can be viewed as part of social practices shaped through their interactions with human users. If such an interpretation is adopted, it could be argued that they fulfil the condition of being tied to social phenomena.

4. Can Artificial Intelligence Be Assessed Within the Space of Reasons?

As can be seen from the above explanations regarding the nature of judicial discretion, the exercise of this discretion is an activity that belongs not to the descriptive-empirical realm of ‘what is’, but to the normative realm of ‘what ought to be’. In this sense, judicial discretion is situated within the realm of reasons not as an activity that can be grasped at the level of causal explanation of empirical facts, but rather as an activity that gains meaning within the context of justification, evaluation and normative obligations. Similarly, the capacities deemed necessary for normative agency—such as autonomy, responsibility and the capacity for practical reasoning—are also the subject of normative evaluations rather than descriptive-empirical explanations. Consequently, the question of under what conditions the normative subject emerges and the question of whether artificial intelligence (AI) systems can be situated

within the causal space are not two independent issues; rather, they are two complementary aspects of the same theoretical problem.

In this study, for the reasons outlined above, we argue that the exercise of discretion is an activity situated within the space of reasons. Consequently, whether artificial intelligence (AI) systems -particularly generative AI systems- can be regarded as normative agents in the creation of law, rather than humans, depends on the debate over whether they possess the capabilities required to be situated within the space of reasons. In this context, we will assess whether artificial intelligence (AI) systems possess these qualifications by focusing on Robert Brandom's understanding of the space of reasons. In doing so, we will examine whether Brandom's approach can be interpreted in a way that allows artificial intelligence (AI) systems to be regarded as actors within the space of reasons, and finally, we will discuss the criticisms that can be directed at Brandom's theory and its limitations.

Brandom has developed his own understanding by adopting Sellars's "space of causes" metaphor and adding further interpretations. From Sellars's (1997: 36) perspective, characterising a situation or an event as "knowing" cannot be reduced to the level of causal explanation of empirical facts the "space of causes". Knowledge, and consequently the status of the knower, gains meaning within causal and justificatory relationships—and thus within the space of causes. In other words, it is constituted at a normative level independent of the epistemic causal context.

Brandom (1994: xiii) distinguishes meaning from the classical mentalist conception of representation, which is based on individual mental representations, and grounds it in social-normative practices established through mutual commitment and entitlement relations. In this conception, knowledge is essentially a social status acquired through the game of giving and asking for reasons. Brandom also holds that epistemic statuses are produced through the process of communication. Cognitive status (knowing and understanding) is not merely the linguistic production of claims, but rather a normative position acquired within a network of commitments that are normatively followed and structured by justifications. The game of giving and asking for reasons involves engaging in linguistic practices and, through justificatory practices, undertaking, maintaining, and when necessary, defending one's commitments and entitlements. In this context, there are three normative dimensions to a person's claim of knowing something. Firstly, by making a claim, the subject enters into a commitment; following this first-order commitment, the subject assumes a second-order normative commitment involving the obligation to present and defend the justifications for that commitment. Secondly, it involves the recognition of other subjects as normative actors capable of assuming similar commitments and entitlements. Thirdly, it involves attributing normative statuses to others that include the right to assume the same commitment and the possibility of referring to the first speaker for justification when this commitment is questioned by a third party. For these reasons, it may be useful to interpret the judge's exercise of discretion from the perspective of the game of giving and asking for reasons. (Brandom, 1994: 171-172).

The judge's reinterpretation of a standard legal rule, using their discretion, on the grounds that it is inconsistent with the fundamental values protected by the legislator, can be read as an act of assuming a commitment in the Brandomian sense. This act gives rise to a second-order normative responsibility that entails not only the establishment of a specific legal outcome but also the obligation to justify that outcome. In this process, the judge recognises other legal actors as normative participants capable of assuming similar commitments and entitlements. In this way, the judge establishes the field of legal justification as a shared "game of giving and demanding reasons". Within this framework, when the interpretation put forward is questioned by third parties, the judge defends their decision not only on the basis of constitutional principles, established case law and systematic legal justifications, but also under the obligation to justify their assessment that the applicable norm is in conflict with the legislator's intentions. Consequently, a process operates in which normative statuses are constantly monitored and re-evaluated. Consequently, the judge's normative claim gains meaning not as an individual expression of will, but within a social-normative "space of reasons" established through obligations of commitment, entitlement and justification.

Before addressing the question of whether any artificial intelligence system can be evaluated within the ‘space of reasons’ as conceived by Brandom, it is necessary to examine Brandom’s critiques of regularism. Brandom focuses on the distinction between being bound by a rule and merely conforming to it. Regulativism attempts to explain which norm governs a behaviour, and which behaviour is right or wrong, by looking solely at actual patterns of behaviour. Brandom’s central point is that a direct inference cannot be made from the descriptive space of “what is” to the normative space of “what ought to be”. The mere existence of a particular behavioural pattern cannot, on its own, explain how that pattern ought to be maintained or which practices are right or wrong. In other words, regularity theory can only describe what is; in contrast, rule-following and inferential practices involve normative distinctions regarding how certain behaviours ought to be. It is not possible to objectively distinguish which norm is being followed by looking at past patterns of behaviour, because behavioural patterns can be interpreted as consistent with a large number -perhaps an infinite number- of different rules (Brandom, 1994: 26-30).

At this point, drawing on Brandom’s critique of regularity, Mirco Sambrotta (2025: 14) notes that whilst LLMs can produce outputs resembling inferential reasoning, they cannot be counted as genuine normative participants in the space of reasons. The reason for this is that he argues these outcomes stem not from relations of normative commitment but from rule-based patterns. Building on Brandom’s critique of regularity, we can say that normative statuses emerge not merely through the display of certain behavioural patterns, but through being situated within relationships of commitment, entitlement and inferential obligation. If the behaviour of artificial intelligence (AI) systems can be explained on the basis of statistical regularities, this may account for the fact that the outputs they produce appear to carry normative content; however, it does not demonstrate that these outputs are situated within relationships of commitment, entitlement and inferential obligation. However, if artificial intelligence (AI) can realise normative and inferential relationships beyond mere regularity, it could be said that they are situated within the “space of reason” from Brandom’s perspective. Therefore, we need to focus on whether artificial intelligence (AI) is realising inferential relationships or merely exhibiting a regularity grounded in statistical regularity.

Sambrotta (2025: 13) emphasises that, even in the best-case scenario, systems provide a pattern-based adaptation to practices; however, this adaptation does not reach a genuine level of engagement with the normatively structured space of reasoning. Consequently, he argues that these systems do not establish a genuine relationship with the inferential norms that underpin discursive practice. By adapting Brandom’s critique of regularism in to LLMs, Sambrotta (2025: 13) argues that these systems lack the capacity to grasp the normative standards that determine correct and incorrect applications; he concludes that, without demonstrating sensitivity to inferential relations, they merely produce outputs that appear to be inferentially structured and are therefore not genuine participants in the space of reasons. However, in our view, for this conclusion to hold, an additional interpretative premise is being added to the paragraph regarding Brandom’s theory. Brandom (1994: 141-143) bases his critique on the premise that adherence to a rule cannot be explained in terms of behavioural patterns. Brandom explains normativity in terms of participation in scorekeeping practices, adherence to commitments, the provision of justification, and the maintenance of inferential relations. The condition that possessing a normative status additionally requires an internal or phenomenological grasp of the norm’s meaning, however, can be regarded not as something that clearly follows from Brandom’s theory, but rather as an independent normative agency condition added to the theory. Sambrotta’s view, which excludes artificial intelligences (AI) from the space of reasons, appears to be based on this additional agency condition. Of course, if the behaviour of artificial intelligence (AI) systems can be explained on the basis of statistical regularities, and if they do not participate in relationships of commitment, entitlement and inferential obligation, it can be argued that they do not lie within the space of reasons. However, it should be noted that, in addition to Brandom’s understanding, the condition of actually grasping normative standards is not required here. Even if artificial intelligence (AI) systems give the impression of making claims about knowledge, if they cannot assume the normative commitments required by the practice of asking for and

providing reasons, it can be concluded that, within the Brandomian framework, they cannot be regarded as genuine participants in the space of reasons. For example, as Heinrichs and Knell (2021: 1575-1576) point out, personal assistants and recommendation systems do not fulfil the obligation to provide reasons and are not part of the space of reasons. At this point, if we are to centre Brandom's theory, we believe it is an accurate approach to state that they do not enter the space of reasons because they do not fulfil social-normative practices.

Brandom's approach can be interpreted as suggesting that artificial intelligence (AI) systems can be included in the space of cause if they are capable of participating in the practice of giving and asking for. According to this understanding, inclusion in the space of cause does not require being human or a physical entity. Consequently, any entity—even if it is not physical—can occupy a place in the space of cause if it fulfils these social normative roles. So, can it be said that generative artificial intelligence (AI) systems fulfil these roles? Leonardo Marchettoni (2025: 8) argues that generative artificial intelligence (AI) can functionally fulfil certain elements of the practice of asking and giving. These systems are capable of making various propositions, demonstrating the inferential connections of these propositions when requested, and providing reasons that support their previous claims by responding to users' evaluations. Marchettoni (2025: 13) interprets Brandom through a pragmatic functionalist approach. According to this understanding, if a system successfully participates in a normative practice, the nature of its internal structure or the manner in which it participates is irrelevant. According to Marchettoni's pragmatic-functionalist interpretation, an artificial intelligence (AI) system may be regarded as a participant in the Brandomian space of reasons to the extent that it can meaningfully engage in the game of giving and asking for reasons.

In our view, at this point, it seems more appropriate to invoke the argument of multiple realizability (Bickle, 1998) rather than proceeding from a pragmatic functionalist approach. We can say that pragmatic functionalism differs from classical functionalism in the philosophy of mind. In classical functionalism in the philosophy of mind, for a state to be considered a mental state, it must fulfil a specific causal role (Levin, 2004). In this context, it is consistent with the "space of causes". Pragmatic functionalism, however, which proceeds from a Brandomian perspective, defines function in terms of normative statuses. If a system appears to be successful in the "reason-giving" game, or demonstrates successful performance in normative practice, its internal mechanisms will be irrelevant. Yet, if we derive the existence of normative roles solely from external performance, are normative roles not effectively reduced to behavioural and causal patterns? Consequently, the pragmatic functionalist interpretation that attributes normative statuses to artificial intelligence (AI) or other systems on the basis of their successful participation in reason-giving practices, even if it attempts to remain within Brandom's intended "space of reasons", carries the risk of drifting towards causal functionalism and this "space of causes" to the extent that it links normative adequacy to performative functions. Pragmatic functionalism carries the danger of reducing normativity to behavioural success. Thus, the problem targeted by Brandom's critique of regularity resurfaces, blurring the distinction between the normative and the merely normative appearance. In contrast, the argument of multiple realizability, whilst maintaining that normative statuses are not specific to particular physical or biological systems, does not require that these statuses be reduced solely to external performance. Consequently, it is concerned not merely with whether a system mimics reason-giving behaviour, but with whether it can genuinely assume the normative relationships in question. The argument of multiple realizability preserves Brandom's understanding of normativity whilst leaving open the possibility that such normative structures may also be realised in non-human systems. At this point, we contend that interpreting Robert Brandom through the lens of the multiple realizability argument, rather than from a pragmatic-functionalist perspective, is more appropriate. Such an interpretation allows artificial intelligence systems that are capable of participating in the practice of giving and asking for reasons, and of undertaking the normative statuses embedded in that practice, to be regarded as participants in the Brandomian space of reasons.

So, based on Brandom's theory, if a system genuinely participates in the practice of giving and seeking

reasons, assumes commitments, evaluates criticisms and follows inferences, can we say that how it does so is irrelevant? Is it sufficient for a system to act correctly within the normative practice to claim that it has assumed a normative commitment, or must it also be able to understand why the norm is binding? Brandom's theory does not centre on the specific internal states possessed by the agent in determining normative commitments (Brandom, 1994: 9-11). Being situated within the space of reasons depends primarily on fulfilling inferential roles rather than possessing an understanding of the nature of what is being put forward. However, the question of whether a system's display of normative behaviour based solely on statistical regularities is sufficient to constitute a genuine normative commitment cannot be directly answered by Brandom's understanding. However, in our view, Brandom's theory allows for the possibility of evaluating generative artificial intelligences (AI) as participants in the space of reasons. This is because the theory does not present direct access to normative content, self-awareness, an understanding of the meaning of commitment, or autonomy as explicit conditions for inclusion in the space of reasons. But is this conclusion satisfactory? In our view, Brandom's theory underdetermines the boundary of the space of reasons. For generative artificial intelligence (AI) systems to truly assume normative commitments, stronger conditions—such as responsibility, ownership, autonomy, and agency—are required. However, these conditions are not explicitly grounded in Brandom's "space of reasons" theory.

In our view, evaluating normative participation solely through practical performance raises certain intuitive difficulties. Even if a system's behaviour demonstrates that it successfully fulfils its inferential roles and participates in the practice of exchanging reasons, it remains debatable whether this requires a meaningful normative understanding. Indeed, at the behavioural level, it may not be possible to discern any difference between genuine normative participation and a successful imitation of normative participation. In other words, even if a system's performance appears sufficient to demonstrate that it is a full member of the space of reasons, one may still harbour doubts as to whether the content produced by that system carries genuine meaning.

We can illustrate this situation using two examples of students. The first student prepares for the exam by trying to understand the course content and the relationships between concepts; the second student, however, studies only by analysing question patterns from previous years and identifying statistical regularities among them. When exam day arrives, both students may be able to give the correct answers, draw the necessary conclusions, and demonstrate similar levels of success in terms of performance. In addition, both students may be able to provide satisfactory answers to requests for justification. From a behavioural perspective, it may be difficult to distinguish between the two. Nevertheless, in most cases, we tend to think that the student who grasps the conceptual content of the course provides answers grounded in a more authentic understanding, whereas the student who merely follows statistical patterns has achieved a successful pattern-matching exercise rather than a genuine understanding. Similarly, generative artificial intelligence (AI) systems can produce outputs that appear to fulfil inferential roles within the space of reasons; however, it is unclear whether this indicates that they truly possess normative content or have developed a meaningful understanding. Consequently, whilst Brandom's approach allows for interpreting these systems as participants in the space of reason, it does not fully resolve the question of how the relationship between performance and meaning should be understood.

At this point, a criticism that could be levelled at Brandom is that, whilst his central concepts produce various normative consequences, the conditions upon which these consequences are based and which they require are not sufficiently grounded within the theory. As we see in Brandom's theory, making a claim is not merely producing a specific utterance, but also undertaking a commitment with specific normative consequences (Brandom, 1994: 171-172). For this reason, the concept of commitment plays a central role in the theory. However, the concept of commitment is not merely associated with the display of certain behaviours in everyday and institutional normative practices; it is also closely linked to dimensions such as assuming responsibility, being accountable, being open to criticism, and being able to defend or withdraw one's position when necessary.

If we are to say that a system is situated within a normative practice and has undertaken commitments,

we must also explain in what sense the system in question has embraced these commitments. Otherwise, whilst the concept of commitment is preserved, some of the normative consequences associated with it are weakened. In particular, it could be argued that interpretations defending the position of artificial intelligence (AI) systems within the space of reasons do not sufficiently ground the dimensions of commitment such as, withdrawal, rejection and assumption of responsibility.

At this point, the following question must be asked: To say that a system is truly under a normative commitment, must it not also possess the capacity to refrain from assuming the relevant commitment or to withdraw from it when necessary? For assuming a normative position does not merely mean producing certain outputs; it also means excluding alternative positions and assuming responsibility for a specific stance. If a commitment truly expresses a normative obligation, it can be argued that this is also linked to the capacity to recognise contradictions, provide justification, and revise one's position when necessary.

We can illustrate this situation with a useful example from the perspective of our article. A judge exercising discretion is assuming a normative commitment. However, for this commitment to be meaningful, it appears to depend on the judge having the possibility of acting differently. A judge may refrain from making a decision, recuse themselves, or resign; the legal system essentially assumes these possibilities conceptually. Precisely because they have made a specific decision despite these options, they assume responsibility for the normative position they have taken. Therefore, there is a close relationship between normative responsibility and action.

From the perspective of generative artificial intelligence (AI) systems, however, the situation is more contentious. These systems are capable of producing specific outputs, providing justification-like explanations, and, at times, revising their previous statements. However, it is unclear whether they truly possess the capacity to consciously refrain from making a decision, to refuse to assume a commitment, or to take ownership of the responsibility for a normative position that has been assumed. If such conditions of agency are among the necessary components of the concept of normative commitment, it may be more appropriate to say that artificial intelligence (AI) systems exhibit performances that mimic normative commitments rather than stating that they assume normative commitments.

It is also possible to evaluate Brandom's understanding of normativity through the structural tensions it contains. As can be seen, in Brandom's framework, normative statuses are fundamentally established within a practical network formed by social normative attitudes (Brandom, 1994: 164-165). Within this framework, what a subject is committed to is determined not by an external metaphysical fact, but through the normative statuses that participants ascribe to one another. Consequently, normative content gains meaning not so much as an individual mental state but within a social-practical structure. However, if a society reduces the concept of "commitment" solely to the production of specific outcomes, and if elements such as responsibility, the obligation to provide justification, the ability to withdraw one's position, or the capacity to acknowledge a breach are weakened, is it still possible to speak of "commitment" in the Brandomian sense? Even if a community preserves the integrity of its inferential relations, as long as it does not associate commitment with its dimensions of responsibility and ownership, it can be said that the substantive density of this structure has weakened, even if the normative structure appears to be preserved.

For example, we might conceive of a society as follows: in this community, "commitment" simply means producing specific outputs. Dimensions such as assuming responsibility, taking ownership of a claim, the ability to withdraw when necessary, or the possibility of breach are not part of the normative structure. In this case, normative practice operates at a superficial level of compliance; however, the dimension of obligation traditionally associated with commitment is pushed into the background. Brandom's theory of normative attitudes already draws on concepts such as responsibility and accountability when explaining the concept of commitment (Brandom, 1994: 172). However, it is not sufficiently clear in the theory on which conditions of agency these concepts are based and how these conditions are secured. Consequently, the possibility of normative statuses being excessively diluted in

content appears plausible.

If normative statuses are determined solely through normative attitudes, then theoretically a community could attribute the same statuses to artificial intelligence (AI) systems. However, in such a case, it becomes unclear how the connection between the concept of commitment and characteristics such as accountability and violability can be preserved.

Consequently, whilst Brandom's approach liberates normativity from metaphysical foundations by grounding it in social practices, it may not fully answer the question of what the boundaries are that will preserve the substantive power of these normative structures. In other words, whilst explaining normative statuses through social attitudes provides a strong explanatory foundation, it leaves a gap regarding the need to clarify the criteria that determine "how normative" these statuses remain.

In addition to this discussion, Brandom's approach can also be evaluated at a more fundamental level in terms of circularity and presupposition issues. For in a model where normative statuses are established entirely through normative attitudes, these attitudes themselves appear to be defined by the very concepts of normative status. In other words, if a subject is deemed to be "making a commitment" only if other participants recognise them as such; yet this relationship of recognition is itself explained by concepts such as commitment, entitlement and obligation, the explanation risks taking on a circular structure. In this case, the justification of normativity, whilst referring to the practices that constitute it, appears to presuppose the normative character of those practices.

This presupposition may make it difficult to explain under what conditions normative statuses actually come into being. For to the extent that the explanation grounds the normative in normative terms, the explanatory reductive power that Brandom aims for may be weakened. Consequently, explaining normative statuses solely through social reference practices carries the risk of both producing a circular structure and, whilst explaining the source of normativity, already presupposing it.

This problem becomes even more pronounced in the debate over whether generative artificial intelligence (AI) systems can possess normative statuses. If normativity is implicitly presupposed whilst being grounded through the normative attitudes of agents towards one another, the system becomes trapped in a circularity. In this context, the theoretical distinction between an artificial intelligence's (AI) genuine "undertaking" of a rational commitment and the language community's "ascribing" of that commitment to it becomes blurred from a Brandomian perspective.

From a realist perspective, this blurring leads to the disregard of the radical difference between "simulating" (imitating) a normative practice and "truly undertaking" the responsibility of a norm. Generative artificial intelligence (AI) can perfectly mimic (simulate) a normative reasoning process through its inputs and outputs; however, this does not imply that the system bears the ontological and moral obligations entailed by those norms as a distinct agent.

5. Conclusion

In conclusion, this study has aimed, first and foremost, to clarify the nature of judicial discretion. Discretion is not merely the mechanical application of existing legal rules; on the contrary, it is a creative activity involving the interpretation and concretisation of the values, objectives and principles that the legal system seeks to protect under specific circumstances. For this reason, judicial discretion involves not only the application of norms but also, to a certain extent, the creation of norms. However, the creation of norms requires the existence of a normative subject to whom normative content can be attributed. Such a subject must be an agent capable of assessing legal and moral justifications, interpreting them within the context of a specific case, and assuming normative responsibility.

Building on this observation, the second section of this study examines the question of whether artificial intelligence (AI) systems possess the competencies required of such a normative subject. Drawing particularly on Joseph Raz's understanding of normativity, it has been concluded that agents capable of creating norms and making normative assessments must be conceived as practical rational beings who

are not merely guided by reasons, but who can also evaluate the reasons directed towards them, and choose to act or refrain from acting in light of these reasons. Within this framework, normative subjectivity is shaped around interrelated concepts such as practical reasoning, autonomy, responsibility and accountability.

In discussions regarding artificial intelligence (AI) systems, views suggesting that these systems are essentially mechanisms that operate by modelling statistical regularities have been examined. In this context, the question has been raised as to whether today's generative artificial intelligence (AI) systems meet the qualities central to normative subjectivity, such as practical reasoning, autonomous decision-making, assuming normative responsibility, and being accountable. Furthermore, the issue of whether the normative content produced by artificial intelligence (AI) systems is grounded in social facts and whether such content can be regarded as legal norms has been addressed. In particular, from the perspective of theories that ground the ontological binding force of law in social facts, we emphasised that the legal status of norms produced by artificial intelligence (AI) depends on whether they can be said to constitute a social fact-based activity.

In discussions concerning artificial intelligence (AI) systems, views suggesting that these systems essentially operate as mechanisms that model statistical regularities have been examined. In this context, the question has been raised as to whether today's generative artificial intelligence (AI) systems fulfil the qualities at the heart of normative subjectivity, such as practical reasoning, autonomous decision-making, assuming normative responsibility and being accountable. Furthermore, the issue of whether the normative content generated by artificial intelligence (AI) systems is grounded in social realities, and whether such content can be recognised as legal norms, has been addressed. In particular, drawing on the perspective of theories that ground the ontological binding force of law in social realities, we have emphasised that the legal status of norms generated by artificial intelligence (AI) depends on whether it can be said that they constitute an activity grounded in social realities.

In the final section of the study, the issue is re-evaluated within the framework of Robert Brandom's understanding of the space of reasons. Contrary to what might initially be thought, Brandom's approach does not, in principle, exclude artificial intelligence (AI) systems from the space of reasons. For Brandom, the decisive factor is not possessing a specific biological structure, a state of consciousness, or subjective experience; rather, it is the ability to assume normative roles within the practice of giving and asking for reasons, to follow inferential relationships, and to be situated within a network of commitments. Consequently, Brandom's theory offers a conceptual possibility for interpreting generative artificial intelligence (AI) as participants in the space of reasons under certain conditions.

However, it cannot be said that this conclusion implies that artificial intelligence (AI) systems are currently normative agents capable of exercising judicial discretion. For legal discretion is not merely a capacity to generate justifications or establish inferential relationships. The subject exercising discretionary power also assumes responsibility for a specific normative position, is subject to criticism, can revise its position when necessary, and is in a position to be held accountable for the consequences of its decision. It is unclear to what extent generative artificial intelligence (AI) systems could assume such normative obligations. Furthermore, the question of how to distinguish between normative participation and a successful simulation of normative participation remains significant.

The problem that arises at this point is not merely concerning the position of generative artificial intelligence (AI) systems within the domain of reasons. More fundamentally, Brandom's conception of the domain of reasons also fails to provide an adequate theoretical framework for defining the boundaries of this domain. The central claim of our study is that the question of whether artificial intelligence can be included in the field of reasons necessarily raises the question of how the membership criteria for the field of reasons should be understood. In this respect, our study develops not only a discussion concerning the normative status of artificial intelligence (AI), but also an intrinsic critique of the limits of the Brandomian conception of the field of reasons.

The difficulty arising at this point is not limited to artificial intelligence (AI) systems. Brandom's theory

itself appears open to criticism on the grounds that it does not sufficiently explain the conditions under which concepts such as commitment, responsibility, agency and accountability emerge. Consequently, the question of whether generative artificial intelligence (AI) systems can be recognised as normative agents lies at the heart not only of technological developments but also of more fundamental philosophical debates concerning the nature of normativity, agency, responsibility.

For this reason, the final conclusion reached by this study is that it cannot be convincingly established that generative artificial intelligence (AI) systems can be accepted as normative agents capable of replacing judges who currently exercise judicial discretion. However, particularly when Brandom's normative-pragmatic approach is taken into account, it becomes apparent that it is not necessarily the case that artificial intelligences (AI) must be excluded from the space of reasons in principle. Consequently, the future role of artificial intelligence (AI) in law-making and legal discretionary activities appears to depend not only on the development of their technical capabilities but also on the outcome of theoretical debates concerning normative agency, legal liability and the nature of norm-creation.

Reducing normativity to mere performance success carries the risk of inevitably removing it from the 'space of reasons' and flinging it back into the 'space of cause'. Consequently, whilst Brandom's approach may make it conceptually possible to include generative artificial intelligence (AI) within the space of reasons, such participation is doomed to remain a flawless simulation rather than the assumption of a genuine obligation, unless it incorporates the constitutive elements of commitment that require responsibility, accountability and a specific agency.

Authors contributions

İhsan Burak Aygün and Arif Çini contributed equally to this work.

Statements and Declarations

Ethical Considerations

Not applicable.

Consent for Publication

All authors consent to the publication of this manuscript.

Declaration of Conflicting Interests

The author(s) declared no potential conflicts of interest with respect to the research, authorship, and/or publication of this article.

Funding Statement

The author(s) received no financial support for the research, authorship, and/or publication of this article.

Data availability

Not applicable.

References

- Al Farabi. (2010). The Attainment of Happiness, In M. Mahdi (Ed.), *Alfarabi and the Foundation of Islamic Political Philosophy* (173-195). University of Chicago Press.
- Bender, E. & Gebru, T. & McMillan-Major, A. & Shmitchell, S. (2021). On the Dangers of Stochastic Parrots: Can Language Models Be Too Big?, In Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency, (610–623). <https://doi.org/10.1145/3442188.34459>

- Bickle, J. (1998). Multiple Realizability, *Stanford Encyclopedia of Philosophy*, from <https://plato.stanford.edu/entries/multiple-realizability/> (Accessed May 30, 2026).
- Bostrom, N. (2014). *Superintelligence: Paths, Dangers, Strategies*, Oxford University Press.
- Brandom, R. B. (1994). *Making it Explicit: Reasoning, Representing, and Discursive Commitment*, Harvard University Press.
- Cappelen, H. & Dever, J. (2025). *Going whole hog: A philosophical defense of AI cognition*. <https://arxiv.org/abs/2504.13988> (Accessed May 27, 2026).
- Chomsky, N. (2023). *The False Promise of ChatGPT*. from <https://www.nytimes.com/2023/03/08/opinion/noam-chomsky-chatgpt-ai.html> (Accessed May 30, 2026).
- Court of Appeals of New York, (1889). *Opinion of the Court*. from https://www.nycourts.gov/reporter/archives/riggs_palmer.htm (Accessed May 26, 2026)
- Dworkin, R. (1977). Is Law A System Of Rules?, In R. Dworkin(Ed.), *The Philosophy of Law* (38-65). Oxford University Press.
- Dworkin, R. (1978). *Taking Rights Seriously*, Independently Published.
- Floridi, L. & Sanders, J.W. (2004). On The Morality of Artificial Agents, *Minds and Machines*. 14(3), 349–379. <https://doi.org/10.1023/B:MIND.0000035461.63578.9d>.
- Fuller, L. L. (1958). Positivism and Fidelity to Law: A Reply to Professor Hart, *Harvard Law Review*, 71(4), 630-672.
- Gardner, J. (2015). The Many Faces of The Reasonable Person, *Law Quarterly Review*, 131, 563–584. <https://doi.org/10.1093/oso/9780198852940.003.0009>.
- Gregorio, G. D. (2023). The Normative Power of Artificial Intelligence, *Indiana Journal of Global Legal Studies*, 30(2), 55–80.
- Hart, H.L.A. (1958). Positivism and the Separation of Law and Morals, *Harvard Law Review*, 71(4), 593-629.
- Hart, H.L.A. (1983). *Essays in jurisprudence and philosophy*. Clarendon Press.
- Hart, H.L.A. (2012). *The Concept of Law*, Oxford University Press.
- Havlik, V. (2024). Meaning and understanding in large language models, *Synthese*, 205(1), 1– 21. <https://doi.org/10.1007/s11229-024-04878-4>.
- Heinrichs, B. & Knell, S. (2021). Aliens in the Space of Reasons? On the Interaction Between Humans and Artificial Intelligent Agents, *Philosophy & Technology*, 34, 1569–1580. <https://doi.org/10.1007/s13347-021-00475-2>.
- Johnson, D.G. (2006). Computer Systems: Moral Entities But Not Moral Agents. *Ethics Inf. Technol.* 8(4), 195–204. <https://doi.org/10.1007/s10676-006-9111-5>.
- Kelsen, H. (1973). Law and Logic. In O. Weinberger(Ed.), *Essays in Legal and Moral Philosophy* (228-253), D. Reidel Publishing.
- Kelsen, H. (2002). *Pure Theory of Law*, Lawbook Exchange.
- Kelsen, H. (2009). *General Theory of Law and State*, Lawbook Exchange.
- Kreischer, S. (2025). *New Actors in the Space of Reasons*. <https://philarchive.org/rec/KRENAI-2> (Accessed May 29, 2026).
- Levin, J. (2004). Functionalism, *Stanford Encyclopedia of Philosophy*, from <https://plato.stanford.edu/entries/functionalism/> (Accessed May 30, 2026).
- Marchettoni, L. (2025). There Are No Aliens in the Space of Reasons, *Philosophy and Technology*, 38(3), 1-13. <https://doi.org/10.1007/s13347-025-00917-1>.
- Marmor, A. (2005). *Interpretation and Legal Theory*, Hart Publishing.
- Murphy, M. (2017). Two Unhappy Dilemmas for Natural Law Jurisprudence, In Duke G. & George, R. (Eds.), *Natural Law Jurisprudence* (342-368), Cambridge University Press.
- Raz, J. (1979). *The Authority of Law: Essays on Law and Morality*, Oxford University Press.
- Raz, J. (1988). *The Morality of Freedom*, Clarendon Press.
- Raz, J. (2009). *Between Authority and Interpretation: On the Theory of Law and Practical Reason*, Oxford University Press.
- Sambrotta, M. (2025). LLMs and the Logical Space of Reasons, *Minds and Machines*, 35(46), 2-23. <https://doi.org/10.1007/s11023-025-09751-y>.
- Sanz, R. (2020). Ethica Ex Machina: Exploring Artificial Moral Agency or The Possibility of Computable Ethics. *Zemo*, 3, 223-239. <https://doi.org/10.1007/s42048-020-00064-6>.
- Sellars, W. (1997). *Empiricism and the Philosophy of Mind*, Harvard University Press.
- Spaic, B. (2014) The Unbearable Lightness of Adjudicating Hard Cases: Hart, Dworkin and the Institutional Control of Legal Interpretation, *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.2443346>.
- Totschnig, W. (2020). Fully Autonomous AI, *Science and Engineering Ethics*, 26, 2473–2485. <https://doi.org/10.1007/s11948-020-00243-z>.

- Uzun, E. (2015). H.L.A. Hart: Kurallar Sistemi Olarak Hukuk, In K. Akbaş, M. B. Aydın, S. Metin ve E. Uzun (Eds.), *Çağdaş hukuk düşüncesine giriş* (139-165), İthaki Yayınları.
- Wacks, R., (2021). *Understanding Jurisprudence: An Introduction To Legal Theory*, Oxford University Press.
- Zafar, M. (2025). Normativity and AI Moral Agency, *AI and Ethics*, 5, 2605–2622. <https://doi.org/10.1007/s43681-024-00566-8>.