

PART II: HOW IHE MMEGHARI ANYA (THIRD-TIER PERCEPTION OR ILLUSION) IMPACTS REASON AND CONSCIOUSNESS: IMPLICATIONS FOR HUMAN-AI INTERACTION

<https://dx.doi.org/10.4314/ezumezu.v3i1.2>

Submission: April 10, 2026

Accepted: June 30, 2026

Innocent I. ASOUZU

University of Calabar, Calabar, Nigeria

Email: frasouzu@unical.edu.ng

ORCID NO.: <https://orcid.org/0009-0001-4214-9263>

Abstract

This essay explores the human dimensions of AI safety, moving beyond purely technical requirements. It examines the influential role of ihe mmeghari anya (third-tier perception or illusion) in human-AI interaction. By applying some basic presuppositions of a realist ibuanyidanda approach to analyse the internal workings of human consciousness, I show how the interaction between reason and consciousness is mirrored in the interplay between positive and negative missing links. The investigation further addresses the ontological status of AI and its implications for post-humanist discourse, scientific methodology, and the broader concept of truth. I argue that AI's ontological nature is best captured in its "Swift but Slow-motion Progression." Because our cognitive frameworks largely dictate our relationship with reality, I examine the philosophical preconditions for AI safety by linking these frameworks to codes, algorithms, and machine learning. Ultimately, I recommend integrating complementary logic into human-AI interactions rather than relying solely on deterministic binary calculus. By understanding AI's unique progression, I show how this type of progression can be fruitfully harnessed in the process of "reverse-engineering," an AI-inspired type of noetic propaedeutic; a measure needed to mitigate the challenges posed by ihe mmeghari anya.

Keywords: Ontological boomerang effect, complementary logical reasoning, AI safety, Ezumezu Logic, positive missing link, negative missing link, AI-induced noetic propaedeutic.

The Method and Principles of Ibuanyidanda: Providing the Groundwork for the Harmony of Reason

Approaching reality with a divisive exclusivist type of mind-set turns out to be one of the greatest challenges we have to deal with in human-AI interaction, and, by extension, issues dealing with AI safety. Such issues arise the moment human reason is not grasped as a unified faculty; and where technical, practical, and theoretical functions of reason are treated as disjointed independent faculties that can be in conflict with each other. One of the major reasons the human subject approaches the world and reality generally with a divisive exclusivist type of mind-set is the failure of the human subject to manage all existential situations that are tension-laden, ambivalent, and beclouded by the phenomenon of concealment (ihe mkpuchi any). On account of this difficulty, the human subject always seeks to address issues in a unilateral exclusivist mode, believing that this is the most strategic course of action. Following this route is always destined for failure. In addition to the challenges our tension-laden ambivalent situations pose, there is still a more radical challenge deriving from what I refer to as ihe mmeghari anya (third-tier perception or illusion). Ihe mmeghari anya (third-tier perception or illusion) ensues from the omnipotent aura that surrounds the AI agent. In this form, AI has the capacity to simulate a third-tier category between the real and ideal divide; and one that has all it takes to complicate further human experience of the world.

Confronted with these constraining existential conditions, human experience of reality can be manipulated and distorted, to the point that it can impact, shrink, and even blur the boundaries of truth between the real and ideal divide. In this way, the *mmeghari anya* (third-tier perception or illusion) includes all forms of manipulation of real-life situations deriving from the AI agent; so much so that in this way new identities are capable of being generated, and even fused beyond recognition. Here, one only needs to think of what we call **Deepfake** as a typical depiction of such radical deceptions and distortions of reality. Since the *mmeghari anya* (third-tier perception or illusion) can challenge the ontological boundaries of truth between the real and the ideal divide in many matters of digital derivation, it poses some of the greatest challenges to the application of reason ever in the AI age. This has to do with its capacity to influence radically the way we perceive reality as presented through AI and allied technologies. Besides, the *mmeghari anya* (third-tier perception or illusion) has all it takes to amplify the ambivalent character of all existential realities that are beclouded by the phenomenon of concealment (the *mkpuchi anyi*), including that of AI itself. Impacting our perception of the world in this way, what is real and ideal, what is ideal and fake, can very easily get fused into a web of indecipherable mix-up. Where the *mmeghari anya* (third-tier perception or illusion) is active, the tendency is to mismanage reason such that technical, theoretical, and practical reason are perceived as disharmonised faculties. In all matters concerning AI, the task is to uphold the harmony between reason and consciousness, so that technical reason is not given undue preference over theoretical and practical reason.

Addressing such matters adequately falls within the realm of the method and principles of *Ibuanyidanda* philosophy, as these strive to provide the means needed to guide the human subject successfully through the constraints that all existential situations pose. The method and principles of *Ibuanyidanda* philosophy strive to accomplish this by showing how harmony can be achieved, first within the subject itself, and then between the subject and the totality of reality. The aim is to make it clear that a harmonised relationship with reality can only be achieved when the faculty of reason remains unified in its operations, reflecting a mind-set that is not exclusivist, disharmonious, or polarising. Basically, it is an attempt to show how the human subject can address the world credibly with a complementary harmonised type of mind-set. It entails working out the theoretical preconditions for a radical shift in the type of mind-set with which we relate to the world and reality generally. It is a shift from solving problems with an exclusivist, polarising type of mind in favour of a mind-set of mutually complementary harmony. We are dealing here with a typical case of the application of reason in a more comprehensive way devoid of polarisation and exclusiveness. Human reason is not a fragmented faculty that operates in the mode of discrete quantities, to pursue set goals in an exclusivist, disjointed manner. Its true harmonised character is fully revealed when it operates in relationship with the totality of consciousness with which it builds real and total harmony. Understanding this type of harmony entails a full analysis of the internal workings of human consciousness itself to determine how it imposes its categories on reason and the totality of human actions, dictating their modes of operation.

Expanded Inquiry into the Nature of Reason in its Relationship to Consciousness

I seek to accomplish this type of inquiry by giving a detailed account of the way what I call *akara obi* or *akara mmụọ* (forms of the mind) inform consciousness as it performs its harmonising function. The expressions *akara obi* or *akara mmụọ* are derived from the Igbo language of Nigeria, whose nearest English language equivalent translates to **the transcendent categories of unity of consciousness**. They are transcendent categories and not transcendental categories because they aid the mind or consciousness to transcend or stay above the impositions of the senses in performing the complementary act of *ima-onwe-onye* (self-consciousness). Thus, some known *akara obi* or *akara mmụọ* or transcendent categories of unity of consciousness include: (Bolded in English language):

1. In the Igbo language, we encounter these transcendent categories as **absoluteness**, which is accessible in the experience of *arūsi/alūsi*, i.e., the inviolable absolute entity, Chukwu, i.e., the eternal being, Chineke, i.e., the omnipotent creator.
2. Reality in its **insufficiency, relativity, fragmentation, historicity and world immanent predetermination** is a lived experience that is evident in such expressions as: *uwa ezuo*ke, i.e., the world is incomplete or insufficient; *onye ka ozuru?* i.e., who is perfect?
3. **Totality as dimension of universal comprehensiveness** is expressed as: *ozurumba*, *ikwu na ibe*, *nta na imo*, *mkpochakōta*, *obiōha*.
4. In the Igbo language, we encounter the category of **unity** in the experience of *ibuanyidanda*, i.e., no task is insurmountable in mutual dependence; *njikō ka*, i.e., togetherness is the greatest virtue; *Igwe bu ike*, i.e., togetherness is strength.
5. In the case of **future reference**, we aver in the Igbo language: *ihe ukwu ikpe azu*, i.e., the future harbours far more than we can imagine; *nke iru ka*, i.e., greatest events are in the future; *eche di ime*, i.e., tomorrow is pregnant with surprise; *onye ma echi?* i.e., who knows what tomorrow will bring? (cf. ASOUZU 2007a, 235ff; 2004, 291ff; 2005, 298ff; 2007b, 323ff; 2013, 58ff).

These transcendent categories virtually show the modes of expression of reason in self-conscious individuals. Informed by these transcendent categories, as its mover, the faculties of reason operate as a unified faculty. In this mode, these transcendent categories drive our actions and interactions, determine the way we see, judge, and define the world, the goals we set and pursue, the values we assign to things, our choices and aversions and all those things that characterise us as human. For this reason, our success or failure in any encounter with the world, and reality generally, depends on how we react to the demands of these transcendent categories at any given moment. In other words, in what direction they steer our actions, determine our judgement and understanding, constitute our minds, and steer our reason that is the direction these activities are bound to follow. In this way, they constitute the latent content of human consciousness as the forms or sources from which it derives its action. This is why they are called *akara obi* or *akara mmụọ* (forms of the mind). They are activated in tune with the demands of any given situation. By relying on them as the latent content of a unified consciousness, actors will seek to affirm those experiences they perceive as aggregable and reject those they consider objectionable.

In this interplay of affirmation and negation, acceptance and rejection, each category upholds its constitutive character in relation to the others: This shows that they constitute a unified whole experience at any given moment, since they are the transcendent categories of the same unified human consciousness. These transcendent categories are not the voice of conscience. Rather, in the transcendent categories of unity of consciousness, all human powers find harmony in their dynamism: This is how the mind, the will, the soul, understanding, the heart, thinking, judgement, reasoning, and all thinkable ways these powers are denoted, find expression as situations demand. This is why we can convey the same meaning by recourse to any of these faculties to express similar and identical experience within given contexts. This is how we can make recourse to the mind, the soul, reason, understanding, heart, etc to denote or express the same or similar meaning within given contexts. In all and at all moments, it is the totality of the human subject that performs such a unified operation as expression of harmonised consciousness; without division and non-exclusivist. This is why, in Igbo language, we say *obim gwaram* (my heart or mind tells me), *ọ ihe m na eche* (that is what I am thinking), etc., as interchangeable ways of expressing the same meaning even if the heart or mind is taken as the thinking faculty in the sentence. In this way, these transcendent categories serve a unifying, regulative and constitutive function as they direct and dictate which way the human subject expresses these faculties at any given moment in view of conforming to the demands of any given situation. It is in this unifying and harmonising form that they can steer

the human subject to transcend all the challenges arising from our tension-laden, ambivalent existential situations that are beclouded by the phenomenon of concealment (ihe mkpuchi anya); just as they compel the ego to be above the impositions of ihe mmeghari anya (third-tier perception or illusion). In playing their constituting roles, they compel the ego to perceive reality as a consistent act of critical self-examination driven by the need to affirm all existent realities as missing links. *This is why I insist, for example, that **Ibuanyidanda philosophy is a critique of ibu anyi dada. In other words, it is a consistent self-criticism deriving from the challenges the uncritical usage of such descriptive statements as ibu anyi dada pose to cognition.** This is important if we remember how imprecise descriptive statements can be, and how they can impact language and thinking. In the critique of the descriptive statement ibu anyi dada, I seek to show how synthetic analytic concepts can be derived from pure synthetic ones that are experience bound. In this case, how we can derive the synthetic analytic concept Ibuanyidanda from such descriptive statements as ibu anyi dada (no task is insurmountable for dada the ant), igwebuike from igwe bu ike (there is strength in huge numbers), umunnabuike from umunna bu ike, i.e., the kindred is strength, njikoka from njiko ka (unity is the highest virtue etc. (ASOUZU 2013, 10-29, 69-72). This difficulty is especially evident when a language form depends heavily on description to express reality. We encounter this in many African languages that depend heavily on idioms and wise sayings to express reality. It is by recourse to the transcendent categories that human consciousness is able to perform such synthetic-analytic tasks as it rids the mind of confusion and constraints. That is to say, in the activity of the transcendent categories we observe the attempt of human consciousness striving to rid the mind of ambivalences associated with concepts and ideas derived from the harmful contamination of sense experience. Hence, not even the idea of ibuanyidanda is above such scrutiny. On the contrary, such are the types of ideas that deserve consistent scrutiny in their application.*

AI within the Context of the Transcendent Categories and Affirming Reality as Missing Links

As this matter concerns AI directly, one can say that the problem of human relationship with AI arises the moment the constituting and regulative functions of these transcendent categories are suppressed or even undermined. This can easily happen when the human subject, in an encounter with AI, mismanages its ambivalent characteristics, such that AI is perceived as a surrogate deity and the perfect tool for solving all human needs in a unilateral technological mode. It is on account of this unilateral perception of AI that stakeholders are bound to fail: Its ambivalent character having been concealed from them, they fail to realise the full implications of what it is like to relate to a technology which, by its nature, is multifunctional. Thus, when focusing only on those things that gratify the ego and the senses in relation to AI, commensurate efforts are not invested in understanding some of its most severe harmful consequences. As a result of this negligence, there is often a tendency to misuse AI, even against the interests of its human creators. This is when actors seeking to resolve challenges in the most selfish ways possible undertake measures that can even work against their own interests, imagining that they are wise and smart. Such approaches backfire because acts of consistent self-interest, within any complementary framework, equate to acts of anti-self-interest. Hence, acting in the most unilateral selfish ways is beclouded by error. This error of judgement can be corrected the moment stakeholders seek the harmony of reason and consciousness by affirming all existent realities as missing links in tune with the promptings of the transcendent categories. In affirming existent realities as missing links that are in mutual complementary relationship, the human subject is doing nothing other than choosing those things it considers aggregable and rejecting those it considers objectionable. Thus, demonstrating that its actions comply fully with the demand of a harmonised reason. In this interplay of affirmation and negation,

acceptance and rejection, actors come to choose those things they think are in tune with their overall well-being and reject what they perceive as objectionable; often with respect to the same object at any given time and place.

One thing stands out, though: The capacity to choose, in the face of the ambivalence to which reason and consciousness are exposed, reveals that realities, as missing links, present themselves to the mind in both positive and negative modes at once. Hence, when one affirms that anything that exists serves as a missing link of reality, it is very important always to remember this ambivalent mode of presentation of data to the mind. This is a very critical fact as it relates to the way we interact with AI in the digital age. We have to bear it in mind always when we read the metaphysical principle of Ibunya philosophy. What this means is that **the metaphysical principle of Ibunya philosophy seeks to grasp all existential realities in the ambivalence of their constitution as positive and negative missing links for the joy of being.** When, therefore, I posit that anything that exists serves as a missing link of reality, I seek to grasp this hybrid, paradoxical and ambivalent nature of all existential situations. It is in this mode that all existent realities, including AI, are presented to human consciousness; something that can be accounted for only in a positive and a negative mode.

Mirroring the Mode of Interaction of Reason with the Transcendent Categories through the Interplay between Positive and Negative Missing Links

In all given situations, human reason always strives to grasp reality in the most natural, varied, and complementary way. It can do this only if it conforms to the way the transcendent categories process the data presented from sense experience. Such data is presented, not in the most straightforward mode, but in an ambivalent, varied and hybrid mode. This mixed mode of presenting data takes most active form in the modes of positive and negative missing links. This is why, when I affirm that anything that exists serves a missing link, I have in mind this varied mode of both positive and negative missing links as they stimulate the mind in its operation. What this means is that there are positive missing links, just as there are negative missing links; as these together constitute reality in a mode of consistent interplay. One can therefore say that missing links encapsulate the varied mode of self-revelation of existent realities in basic positive and negative forms. It is by relating to reality in the mode of positive and negative missing links that the transcendent categories of unity of consciousness recognise them either as agreeable or disagreeable, objectionable or acceptable. Those that are agreeable are immediately categorised as positive missing links by the mind. Those perceived as objectionable are treated as negative missing links and categorised accordingly. This is the way human consciousness relates to individual things, to collectives, institutions, systems, judgements, thoughts, categories, values, norms and all modes of determination. This is why in a given system, the diverse parts play diverse functions in tune with their modes of determination. Even when they fail in their operation, they still serve as a missing link in their negativity. Since all missing links are geared to the joy of being, the task in any given situation is to try to understand in what new way something that has failed within a given system can still be useful in view of grasping the joy of being, its ultimate goal. In other words, since all missing links are directed towards the joy of being, a negative missing link can still be useful as a means of grasping the joy of being. What this shows is that in its negativity it still retains a transforming and revitalising character. This is why:

[A] system would perform its function well if all its parts can be revitalised as to perform their functions well. Here even those units that negatively work against the life of the system have indirect positive functions to play since they help us understand what is lacking within a given situation. We say that such units stay in a negative complementary relationship to each other and serve

negative missing links. This is a useful negativity. The reason is simple: without the awareness of such negative facts, we would forever remain ignorant of why a system does not attain its desired objective. That is to say, through knowing what hinders its progress, we increase our knowledge even more positively about the whole system. In this way, the negative coordinates of a system enable the mind to grasp what is expected of certain vital aspects of the system. In all cases therefore, both negative and positive missing links of reality are directed towards one ultimate authentic source of joy in a mutual complementary manner. (ASOUZU 2004, 280)

When, therefore, I affirm that anything that exists serves a missing link of reality for the joy of being, I hold that a unit acquires meaning only within a comprehensive complementary web of relations that constitute being. It is the most natural way reason, as technical, practical, and theoretical faculties, reveals itself as harmonised and differentiated; as it is constituted by the transcendent categories. What becomes clear in this mode of expression of reason is that consciousness perceives reality in the most varied and natural mode. In the process, it assigns meaning as situations demand and in tune with the informing categories (akara obi or akara mmụọ) that determine its action. This interplay of positive and negative missing links in given situations does not mean that all units are equally valuable, nor that all contribute harmoniously. Rather, it means that no unit is self-sufficient; each plays a role, theoretically, practically, technically, directly or indirectly, in revealing the structure of reality and the demands human reason places on them towards complete revelation of meaning. In other words: Every unit of reality, whether contributing constructively or destructively, helps to disclose what is needed for the flourishing of the whole. This is the ontological basis that sustains reality as it is processed by the transcendent categories of unity of consciousness (akara obi or akara mmụọ). In this way, positive missing links are perceived as those units that fundamentally and directly enhance the life, coherence, and proper functioning of a system. They are those units that promote authentic harmony, facilitate the system's flourishing, and reveal how complementarity ought to operate. We say that they are desirable as positive missing links. Conversely, negative missing links seek to negate the way the system should function properly and fundamentally, but retain their revelatory characteristics.

In other words, in the interplay of positive and negative missing links, we realise that negative units have an indirect positive epistemic function; just as they reveal, for example, the necessity to explore the truth content of what is presented to consciousness, or to seek the values that ought to be pursued or the best means to accomplish specific tasks. One can even say that negative missing links play diagnostic technical functions as they stimulate technical reason to address matters in the most appropriate way. For this reason, even if a harmful element disrupts the functioning of the system as a negative missing link, it nevertheless exposes hidden weaknesses as it serves an epistemic function and demands correction in a normative direction. At such moments, we see the harmony of the technical, practical, and theoretical functions of reason at play as they are unified by the transcendent categories as situations demand. It is through this process that our idea of the world is broadened technically, theoretically, and practically. In given situations, therefore, we are thinking of useful negativity that may be epistemically useful but morally objectionable. Here, the transcendent categories of unity of consciousness are actively directing human reason in the interplay between positive and negative missing links in view of guiding the subject successfully through the constraints imposed on it by all existential challenges. The end of this interplay remains the realisation of the joy of being. By affirming all existent realities as serving a missing link, I am not saying, for example, that evil is good. It means the recognition of evil, dysfunction, or conflict potentials as these become moments of revelation in the process of being always seeking higher

forms of self-actualisation and legitimisation. What this means, in clear terms, is that relevance does not imply approval and ontological inclusion does not imply moral equality. Fundamentally, therefore, complementarity does not collapse into a simple form of harmony where everything is forced to be good. On the contrary, it is in complementarity that the dynamism driving the diverse functions of reason is experienced in human consciousness as a unified whole. What *ibuanyidanda* rejects is absolutising missing links and treating negative realities as isolated or self-sufficient modes, ignoring their role in revealing systemic weaknesses. This is important if we remember the truth and authenticity criterion of *ibuanyidanda*, which states: “never elevate a world-immanent missing link to an absolute instance” (ASOUZU 2007a, 197).

Therefore, understanding the dynamics of missing links of reality entails understanding the epistemic role of negativity, the relational and teleological structure of complementarity; the varied nature of reason and the harmonising character of the transcendent categories of unity of consciousness. In this way, *ibuanyidanda* is a framework for understanding the revelatory role missing links play and one that does not allow them to vanish into meaninglessness. It means the capacity to understand the difference between ontological inclusion, i.e., that everything has a relational role, and normative valuation, i.e., that not everything is good or is to be preserved. This is the noetic complementary challenge those forms of human reason face, which focus only on pursuing reason in a unidirectional mode; as in those cases that pursue the primacy either of technical, theoretical or practical reason alone, forgetful of the comprehensive character of reason in its operation. It is for this reason that *Ibuanyidanda* ontology is teleologically constituted to denote the constitutive character of the transcendent categories of unity of consciousness in a future-related mode. This is why our capacity to perceive the world as missing links has to be reflected in the type of mind-set with which we engage reality, and, by extension, in the logic with which we steer AI at all times and places. Hence, perceiving the world as missing links has great implications always for the way we understand and relate to the notion of truth. This is very important today, where our notion of truth is highly challenged when many have come to view AI as the future. That AI is not the future can be illustrated better by trying to understand closely its ontological status as a world-immanent missing link.

Upholding the Ontological Status of AI in its Ambivalence: Problematic Idea of AI as the Colossus

The many advantages and benefits of human-AI interaction notwithstanding, most public conversations about it seem to be loudest when they concern its most severe negative implications. Some of these challenges become more apparent the more we navigate the boundaries of Artificial Narrow Intelligence (ANI) or Weak AI towards Artificial General Intelligence (AGI) and Artificial Super Intelligence (ASI). Of all the challenges AI is projected to present, one of the most dreaded is that of the possibility of human agents losing control over AI operations, and such that AI, operating beyond the demands of human natural intuition, takes full control. This is the type of situation Dario Amodei, Anthropic CEO, was referring to when he describes the different steps that can lead to such a dreadful scenario: “[t]he technology will first present bias and misinformation, as it does now. Next, it will generate harmful information using enhanced knowledge of science and engineering, before finally presenting an existential threat by removing human agency, potentially becoming too autonomous and locking humans out of systems” (ROGELBERG 2026). In other words, the fear is that of AI becoming too autonomous; and to the extent of “locking humans out of systems”. The much cited Paperclip Maximiser thought experiment that was proposed by Swedish philosopher Nick Bostrom in his 2003 paper is often taken as a typical example to illustrate this extreme case of AI running riot (BOSTROM, 2020). In this thought experiment, I see AI being presented as a sort of **Colossus** standing over the world and humanity, capable

and ready to unleash mischief at will. These are the type of projections that reinforce the impression of an AI that has absolute power, and one that can very easily mislead into assuming that AI is a surrogate deity. If this actually is what AI is all about, then there are bound to be reasons to worry about AI getting out of control. However, AI, as a human creation, remains what it is: A relative world-immanent tool without absolute characteristics. Hence, hypothetical models of the kind Bostrom presents serve a very important purpose: They serve as a missing link of reality. In this form, they serve a revelatory function as to how not to relate to AI in its world-immanence. In other words, models of this kind are basically saying: The moment we absolutise world-immanent relative existents, we are bound to contend with intractable difficulties in our relationship to the object in question. This is at least what we learn from the truth and authenticity criterion of ibuanyidanda philosophy that is enshrined in the command: “never elevate a world immanent missing link to an absolute instance” (ASOUZU 2007a, 197).

Hence, models of this kind, as apocalyptic as they sound and as much as they have the capacity to inject fear into us, serve a missing link of reality. The question, however, is whether the characteristics being imputed into AI such that it appears absolute are ontologically tenable? One can say that such projections in the negativity of what they connote uphold their value and should be read as cautionary notes if the benefits deriving from progress in the deployment of AI should not be counterproductive. One thing is certain though: the case of AI running loose and turning against its creators is thinkable. However, the case of AI gaining absolute control as to be self-constituting and an all-legitimising existent is not ontologically tenable. By its nature, AI is a world-immanent relative existent without absolute characteristics. That is its ontological status. This observation is important in view of addressing such extreme projections of AI running out of control and taking absolute control, which can even lead to the ultimate disruption of order, as the “Paperclip Maximiser” seems to project. AI’s destructive potentials, as a relatively world-immanent existent, are in the measure severe as they are absolutised. Hence, for example, the more comprehensive and universal its outreach, the more an AI model seeks self-definition, the more its world-immanent destructive potentials are enhanced. The more fragmented an AI model is perceived and is constituted, the less harm it is capable of causing. AI is human creation; and the moment it is absolutised as to negate its relative world-immanent constitution, that is the moment its lethality is enhanced and the more its inherent character as a tool for positive transformation is diminished. AI remains a very powerful tool of transformation with near inexhaustible potential that can be deployed either as a positive or a negative missing link of reality. However, the moment we absolutise its character and see in it something that is eternally self-constituting, that is the moment the *mmeghari anya* (third-tier perception or illusion) tightens its grip to give us all sorts of ideas concerning AI. Unmasking the world-immanent character of AI is bound to be a consistent challenge.

We can then understand why it is questionable whether the best strategy in AI development and deployment is one that pushes for limitless optimisation of AI capabilities in view of maximising profit, and one that does not care much for the full philosophical implications of such actions. Whenever the motives driving the development and deployment of AI are to optimise it beyond all thinkable limits, AI is bound to transit from a mere tool towards nebulous objectives. In this form, it can even be elevated to an object of adoration in order to exploit some of its potentials. This has always happened throughout human history, where the human subject reacts to novel and fascinating experiences with a near religious sense of awe. Ours cannot be different. Even if perceptions about the character of AI are still divided in some quarters, there are justifiable indications that AI has created an aura of the sacred in many quarters also and one that beckons adoration; something that can always be exploited. For many, it is something shrouded in mysteries beyond comprehension. For this reason, those who are directly responsible for the development, training, and programming of AI seem to be enjoying the status of seers or diviners mediating between this near divine object (AI) and an

outside that is stuck in wonder. With this, an incredibly high margin of knowledge asymmetry between these insiders and the general public can very easily be created, and one that can very easily also be manipulated and exploited. Where adequate measures are not taken to address this situation urgently, the few insiders may be tempted to further mystify AI's capabilities, widening this asymmetry all the more. By hyping the mysteries of AI, the insiders create mass hysteria that bewilders the outsiders, which bestows selfish advantages on insiders for nefarious objectives. This has all it takes to exploit and marginalise those who remain in the dark. This is why the mystification of AI as a prohibitively expensive venture is bound to have a very negative impact on its development, especially when the impression is created that it is an undertaking that requires tons of money only the very affluent can afford. Such an impression is bound to stifle creativity and discourage talented people from joining AI research and development.

However, with the entry of DeepSeek into the AI industry, the cost of building AI models comparable to the more established ones has been demystified. It was built relatively cheaper than comparable Western AI models (THIAGARAJAN & THIND 2025)). While AI development still requires significant resources, this shift proves that the large language models (LLMs) Club is no longer restricted to the ultra-wealthy. This is why all cases that are inspired by the near inexhaustible potentials of AI will remain reasons to worry. Such can always be reasons to misuse AI for the wrong purposes. This is why I insist that the **therapeutic dimension of Ibanyidanda philosophy** is of utmost importance in matters that tend to bestow advantage in a way that human agents can pursue their interests in a one-dimensional absolute mode. The goal is to shed light on the dangers of such unilateral actions that can trigger the ontological boomerang effect of Ibanyidanda philosophy. For this reason, in dealing with AI acquiring an ibanyidanda type of logical reasoning is indispensable as a veritable tool towards superseding the impositions arising from the concealment all ambivalent situations are capable of presenting, and most especially as this concerns a technology that has the potential to always conceal its true nature and characteristics. In exposing and overcoming such concealment, stakeholders are better in a position to see the world in a clearer and more nuanced way. By so doing, a favourable condition for a balanced relationship with AI can be created. Such a balanced and enlightened condition is very high in demand today in a world where AI has the capacity to twist human consciousness in measures never seen before. By acquiring a complementary type of mind-set, and learning to defuse the tension that the ambivalent characteristics of AI can create, stakeholders will be more in a position to reap the benefits of AI fully (ASOUZU 2013, 74ff).

Ontological Status of AI and the Future-Related Character of Complementary Truth (Ihe ukwu kpe azu)

At no time has the notion of truth come under severe scrutiny as today, where many see in AI the future. This is all the more complicated by the impositions of the mmeghari anya (third tier perception or illusion), made possible by AI and allied technologies on the way we perceive reality. When people say AI is the future, there is every indication that they mean it. What becomes evident by such assertions is the unshaken confidence many have come to repose in AI, and AI-driven technologies. By affirming AI as the future, they seem to imply also that AI is the truth. This attitude immediately raises some doubt about the way we conceive truth; that is, if we remember that the transcendent category of future-relatedness proclaims truth as something belonging to the future when it states *nke iru ka* (the greatest event or occurrence is in the future). In that case, that which authenticates ultimately, in the sense of ultimate truth, is future-related. Going by the dictates of this transcendent category, AI, as something world-immanent, is definitely not the ultimate authenticating future. When the transcendent category of future-relatedness states *nke iru ka*, i.e., the greatest events are in the future, it is another way of saying that what is true is future-related and beyond world-immanence. In other words,

when many proclaim AI as the future, they thereby put the truth claim of this transcendent category to severe test. We are dealing here with a very important character of truth by reason of which it has a future-related dimension that legitimises and authenticates all. It is by reason of this future-relatedness that truth drives all human activities and inspires research and innovation. Remove this dimension of truth and world-immanence becomes self-constituting. It is this notion of truth that is future-related, which has to contend with the omnipotent character of AI by reason of which AI always seeks to impose itself as the driver of progress, the ultimate legitimiser of all things that seek to be true and authentic. In AI, many tend to see that which determines the future; so much so that it seeks to enthrone itself as the future towards which everything is directed. The moment this happens, the very thing that gives truth its authentic character is bound to be blurred and suppressed; and so much so that the legitimising function of truth is inadvertently transferred wholly and entirely to AI (a world-immanent existent). Even if we tend to suppress and forget its legitimising force, truth, as something future-related, resists suppression and has an inherent character of persistence, by reason of which it ultimately legitimises and constitutes. In this connection, Bertrand Russell cautions in the ninth of his ten commandments: “Be scrupulously truthful, even if the truth is inconvenient, for it is more inconvenient when you try to conceal it” (RUSSELL'S SOCIETY NEWS, 1987). In other words, truth cannot be suppressed easily. It always resurfaces no matter how long it takes; as the case of Herostratus stands to testify. In spite of the efforts made to suppress this infamous act, the truth about it remains indestructible. That is the character of truth which cannot be suppressed. Similarly, the triumph of the truth that Nikola Tesla's patents were prior and Guglielmo Marconi used those ideas to build the first commercially successful radio system remains relevant and highly instructive in today's artificial intelligence era. As against all forms of fanaticism, I would even say, do not fight fanatically for truth because truth has a way of asserting itself: Not even the unrestricted urge to enthrone AI as the future will withstand the untruth enshrined in such desires; hence, *nke iru ka*, i.e., greatest events are in the future; “*ihe ukwu kpe azu*, i.e., the future has the capacity to reveal the greatest of events; and by extension the full character of truth. This future-related character of truth is deeply ingrained in human consciousness and cannot be suppressed, as it continually reminds us that behind AI is the human person who is responsible and accountable. In all, and more importantly still, there is a need to be driven by truth itself if we must reap fully the benefits of AI, which is not the future. This attitude to truth as the ultimate mover and legitimiser of all aspirations is very important in the way we relate to the world and most especially in the way scientific investigations are conducted. It is this notion of truth that drives scientists towards new insights, and it is such that steers the mind to strive beyond all attainable boundaries as to break new ground. This notion of truth is negated the moment we assume that scientific propositions are, in themselves, self-authenticating, and that AI is self-legitimising and self-constituting. Attempts at negating this fact are bound to result in stifling human creativity and to follow blindly human craving to satisfy dreadful curiosities in the name of doing science or in the name of seeking optimisation of human potentialities.

Even then, *ihe mmeghari anya* (third-tier perception or illusion) has a way of suggesting that the contrary is true, and as such it fires our passion to steer AI usage to extreme limits. This is why whatever is deployed to suppress the self-evident character of the future-related nature of truth is commensurate to what is needed to evoke the ontological boomerang effect always. In its future-related mode, truth belongs to the intersubjective memory of humanity and drives reality in a way that should make practitioners of science and technology aware that they are mere agents, not ends in themselves. The same is valid to the way AI is conceived. It is not an end in itself but something that draws its legitimisation from the future. Here, Heraclitus reminds us that *logos* asserts itself despite concealment. The same idea is later voiced by St. Augustine that truth participates in eternity and therefore resists annihilation:

“[w]e have, as I recollect, concluded that Truth cannot perish, for the reason that should not only the whole world pass away but even Truth itself, it would still be true that the world and Truth had perished. But nothing can be true without Truth. In no sense, then, can Truth perish” (1910, 63). In other words, scientists are not merely drivers of truth in the sense of truth seekers; they are all the more driven by truth itself that is future-related. Failure to observe this fact easily leads to the point where science forgets this future-reference and seeks truth in world-immanent missing links. **The moment this happens, the danger is always to go against the truth and authenticity criterion of ibuanyidanda that is enshrined in the command: “never elevate a world-immanent missing link to an absolute instance” (ASOUZU 2007a, 197).** Where we insist on elevating world-immanent missing links to an absolute instance, it is bound to have very catastrophic consequences in the way we use and deploy AI.

This is why the harmony between truth and science, and between science and AI, must always be upheld. Upholding this harmony is a necessary condition for intelligent use of AI. Where we are forgetful of this fact, we would be constrained to misuse AI for purposes contrary to its intent and contrary to our own interests too, believing that we are smart and wise. These are the types of missteps that can automatically trigger the ontological boomerang effect. Hence, there is always the need to be committed to the harmony between truth and science and between AI and science by recourse to a unifying type of logic of complementarity. It belongs to the nature of this logic to make insightful the type of inherent necessary relationship between AI, its human creator and the way these relate to the totality of reality.

Some Challenges of Ihe mmeghari anya (third tier perception or illusion) on the Way we do Science

The higher the stakes are, in the same measure does ihe mmeghari anya (third tier perception or illusion) tighten its grip on our relationship with the world and reality generally. This is why some of the ideas being pursued by variants of post-modernist movements are one of the places where this phenomenon of ihe mmeghari anya (third-tier perception or illusion) is bound to exert some of its most radical challenges. AI has a way of evoking a strong, reassuring feeling that humanity is at last in possession of a magical tool that can be easily deployed to find lasting answers to most age-old questions that have hitherto defied solution. Typical examples that go to reinforce this feeling abound; the widely reported case of AI revolutionising archaeology and history by decoding previously unreadable ancient texts comes readily to mind. We have this illustrated by the case of the charred Herculaneum scrolls, where researchers using AI and X-ray technology have successfully read text inside the charred, unopened scrolls. These are feats that were once considered almost impossible. Emboldened by AI capabilities, we see renewed efforts being invested in projects whose goals appear very exciting and almost like things beyond the ordinary. I have in mind here some of those projects that aim at transforming and even manipulating human physical and cognitive capabilities beyond our imagination. Thus, enormous resources are being invested in achieving higher levels of physical, cognitive, and emotional functioning in human beings, with the aim of turning the human individual into a radically new creature. Such are the ideas driving variants of post-humanist and transhumanist yearnings that exploit the idea of machines being smart to merge biological and digital intelligence. We see this given in the ideas behind Cyborg, then, Brain-Computer Interface research and technology that seeks to use technology like Neuralink to augment human brains directly with ASI level processing. Some of the major motives behind these ideas is the deep-seated, often tech-driven desires to fundamentally alter and even eliminate some biological limitations of the human condition. Other variants concentrate on attaining digital immortality, for example, in the form of theoretically “uploading” human consciousness into a digital medium, with the hope of escaping biological decay. Behind most of these ventures is the assumption that with AI and its promises, the time has at last come to have some of these dreams fully realised.

Thinking in this manner is often deeply rooted in the impatience with which some approach human frailty; something some consider unnatural. The question will remain whether human frailty is indeed an unbearable burden that must be eliminated at all costs. This must not be the case if we remember that insufficiency, relativity, fragmentation, historicity, and world-immanent predetermination, as transcendent categories, are inherent moments of authentic existence. This transcendent category of human insufficiency captures the fundamentally fragile condition of human nature by proclaiming: *uwa ezuoke*, i.e., the world is incomplete. Hence, in tune with this transcendent category, insufficiency is an inherent moment of human relative existence that must not be wished away, but which we must contend with reasonably. What is particularly interesting with some of these transhumanist ideas is the audacity and confidence with which some pursue their goals beyond all set boundaries of scientific investigation. In other words, stakeholders hardly care if they are venturing into areas outside their subject matter. One thing is certain: What has hitherto sustained scientific investigation and given it its special character is the humility with which scientists go about their work. This humility is deeply embedded in the character of science itself, which operates in the form that each discipline has specific goals as its subject matter, and something that determines its method. With the advent of AI, there seems to be the urge to venture into areas outside the subject matter of individual sciences. It is a situation in which many give the impression that there are no longer any rules guiding research. It is a situation where, by recourse to technological methods, attempts are made to venture into areas reserved to other sciences, in the course of which perplexing conclusions can be drawn. The challenge of doing science in a way that unites and does not divide remains one of the greatest challenges in the AI era. Those scientific methods are more likely to give credible answers that speak clearly on the nature of their subject matter and not those that seek to blur the content of their subject matter altogether. There is a need to give a clearer view concerning what is being investigated and the method used to investigate it in order to avoid confusion.

When scientists seek to give purely empirical answers to metaphysical questions and are thereby emboldened by AI promises, this is bound to blur the type of answers they seek. A typical example is the desire to theoretically “upload” human consciousness into a digital medium, with the hope of escaping biological decay. Even if such cravings are now being pursued under transhumanism, digital immortality and futurism, they share a lot in common with archaeology or, better still, “neo-archaeology” and with those sciences that delve into ultimate realities. Ancient Egyptians were once obsessed with the same idea of achieving immortality through mummification. In matters of this kind, it is always important for scientists to state more clearly what their subject matters are, and with which methods they are pursuing their goals. What types of questions are they asking, and what types of answers are they expecting? Is not giving the impression that answers derived from such transhumanist yearnings are the ideal answers, and that every possible answer is of that kind, an attempt to misuse the promises of technology to sow seeds of confusion and raise false hopes? This is still worse when we remember that we are dealing here directly with one of the most fundamental human concerns as this relates to issues of immortality. With AI and its near inexhaustible potentials, the tendency is to follow its tempo in addressing matters, and such that gives the impression that all pressing and complex matters can be resolved with the finality of technology. This way of approaching research and technology is bound to cast deep shadows over scientific methods that do not often focus on their subject matter but instead encroach on other areas. The question will always remain: How viable will giving a technical answer to questions of ultimate realities be, even by recourse to the best available technologies? Not even by recourse to AI in a world-immanent manner?

The need to approach matters with an open and complementary type of mind-set while doing science in the age of AI is highly recommended. For this reason, one of the best approaches in matters of this kind should be: *Each on his lane but all in mutual complementary relationship*. In this way, researchers come to appreciate the value of pursuing their subject matters with a collaborative complementary type of mind-set. This is the approach that values openness, tolerance, and, above all, the type of humility that does not seek to polarise or divide unnecessarily. AI has all it takes to widen our knowledge and broaden our perspective; however, where what gives it its dynamism is compromised, it is bound to shrink our knowledge of the world all the more. Here, the inherent moment of concealment arising from the activities of the *mmeghari anya* (third-tier perception or illusion) can be so severe that researchers easily derail by asking the wrong questions. Where this happens, they are bound to receive the wrong answers which they celebrate as knowledge. The *mmeghari anya* (third-tier perception or illusion) has the capacity to make us stubborn in pursuing those goals we think serve our interests most, but which in actual fact have the capacity to negate such interests. It is due to the imposition arising from the *mmeghari anya* (third-tier perception or illusion) that the human subject can persist in its errors, believing that it is wise and smart. The question will always remain whether the goals driving some of these projects derive from reason. How realistic, reasonable and responsible are they? Just as AI has come to be viewed as a typical instance of the summit of Aristotle's rational man, one of the greatest problems we have to contend with in an AI era is that of striving to encase this ideal type of reason in a perfect human body. When this happens, the saying *mens sana in corpore sano* (a healthy mind in a healthy body) is bound to equate to *mens perfecta in corpore perfecto* (a perfect mind in a perfect body).

Some transhumanist extremist ideas evoke the impression that human insufficiency is a burden or an irritant that must be eliminated by all means. Attitudes of this kind can easily lead to the assumption that being and its attributes are not complementary but irreconcilable opposites. In other words, essence and accidents exist in diverse regions of being and are as such not complementary. These are the types of ideas that enhance the tendency to treat accidental insufficient qualities with disdain, and that can easily motivate eliminating them without much care. We see this concretely given when human subjects desire to eliminate those peoples, ideas, and institutions they view as irritants, as dispensable, disposable and inconsequential; most especially in asymmetrical situations of power imbalance. Those goals, actions, and desires remain realistic, reasonable and responsible, always striving to uphold the harmony between being and its attributes; and thereby enhance a mutually complementary relationship between all existent realities. Remove this decisive dimension of our relationship with the world, and reality generally, we become those monsters we all dread. Where the complementary character of reality is obscured, its ambivalent nature is enhanced, and this can blindfold actors into taking very extreme and questionable measures in seeking solutions. Being awake to the beclouding, ambivalent character of the world, and especially of AI and related technologies, is very important in view of credibly answering the most difficult life challenges. As straightforward as issues of human insufficiency may appear, confronting their intricate nature remains a great challenge in the age of AI: the irony is that, notwithstanding our thirst for perfection, we wear our insufficiency like an old cloak that refuses to go away. This is why any sincere attempt to fully understand the depth of human insufficiency can be inherently destabilising and is even bound to lead to a state of intellectual paralysis or cognitive impasse. More often than not, it is this incapacity to come to terms with the depth of human finitude that has misled many into trying to carve out a niche for themselves with a view to generating an ideal human creation of their own making and phantasies.

Addressing Uwa Ezuoke (human insufficiency), Post-humanist Ideas and Some Implications of the Ontological Status of AI

Granted, some of these transhumanist-inspired scientific research methods have often been deployed successfully to alleviate human suffering, as we can see clearly in cases of people with artificial heart valves, cochlear implants, advanced robotic prosthetics, enhancing or restoring functions like sight, movement, etc.; those projects stand to sustain their benefits that are not just undertaken as mere proof of concepts. The same can be said with regard to projects motivated by pure egoistic concerns, and such that seek only to absolutise world-immanent existents. What it takes to absolutise human relative existence as to make meaning a world-immanent commodity is what it takes to obscure meaning indefinitely. To make humans better is fundamentally a good idea, provided we are not tempted by the high hopes raised by AI to focus on aims that degrade what it means to be human. If AI has to contend with the difficulty of remaining a relatively world-immanent tool, one of the greatest challenges it poses is raising the false hope that, with it, we can solve most life challenges with the ease of computation, including the question of human physical immortality. This question remains very relevant if we consider the enormous efforts and resources invested to push research to the edges in view of finding ways to overcome biological limitations like ageing and in view of prolonging life indefinitely. Since we are dealing here with fundamental human problems, pursuing such projects has to contend with equally radical challenges. One of these concerns the rapid changes imposed on human consciousness and perception of reality by AI itself. They are such rapid changes that can impair our perception of the world completely, and in such a way that can make us susceptible to the *mmeghari anya* (third-tier perception or illusion) very easily: By recourse to AI and allied technologies, any attainable results can soon be overtaken and rendered obsolete, to such an extent that the human subject has difficulty contending with what it means to be real. This is why, within such a charged context, human consciousness can become so overwhelmed that meaninglessness becomes a habit. Operating under such tense conditions, and with the attendant distorted mind-set it has the capacity to impose, can be very challenging and disorienting in view of tackling fundamental problems of life and existence. Consequently, one of the central challenges most post-humanist ideas, which are obsessed with perfectibility of human conditions, have to contend with is that of navigating the tension between two extremes: the desire for AI to provide a definitive, absolute structure to our existence, in which case AI becomes a technological surrogate for the absolute, and the crushing sense of emptiness or hopelessness that arises when we realise that such a machine-led world lacks inherent human value. Concretely, it is still the challenge to successfully navigate the tension between what it means to be happy and the excesses of world-immanent impositions. A world of human relative happiness is desirable but not one bestowed by world-immanent form of absolute perfection or perfectibility.

Hence, even if AI has the capacity to transform our relative existential conditions immensely, what is needed is the capacity to come to terms always with the fact that human insufficiency is an inherent moment of authentic existence. Internalising this lesson is possible only with a mind-set that is complementary and one that has the capacity to understand existent realities as missing links that are in mutual complementary relationship for the joy of being. In the face of most of these world-immanent absolutistic concerns about the perfectibility of human potentialities, it may be beneficial not to forget, all too easily, some of the major fallouts of Eugenics and its catastrophic consequences during the Second World War. The *mmeghari anya* (third-tier perception or illusion) has a way of awakening the feeling of absolute certainty that such cannot happen in our highly enlightened and technologically advanced age. How wrong we can be if we remember how issues of differences are being vigorously mismanaged daily; where people whose lifestyles are different are being persecuted with all subtle means thinkable. That human beings forget all too easily in most important matters can be attributed

to the constraining impact of our tension-laden existential situation, which is now complicated all the more by *ihemmeghari anyi* (third-tier perception or illusion). This is all the worse when our most cherished interests are at stake; things we seek to pursue blindly, often not minding some of the most severe consequences. This is why some of these challenges fall under what I call the “triadic forgetfulness” (ASOUZU 2007b, 55) with which the human subject has to contend, especially in the way we do science. To these belong a) the Forgetfulness of the tension-laden human ambivalent situations, b) Forgetfulness of the subject matter of philosophy, and c) Forgetfulness of ethnocentric intrusions (ASOUZU 2007b, 55).

Some of these post-humanist ventures, as glorified as they may sound, often address particular needs. This is the case with such concerns that strive towards human immortality. Definitely, such special concerns serve a missing link of reality. As such, they should not be discouraged. However, would it not be more beneficial to invest more time and energy in matters that have long-term benefits for humanity and such that are geared towards the overall improvement of mutual coexistence between diverse peoples of the world? The desire to theoretically upload human consciousness into a digital medium, for example, as much as it addresses high-tech archaeological needs, as we had this given by ancient Egyptians who were obsessed with mummification in view of immortality, may not be a priority in a world where we have more pressing needs. Is it not an irony that in a world blessed with enormous resources to take care of the needs of her citizens, a sizable number of the world population still languish in abject poverty and still live under catastrophic inhuman conditions? Based on the 2025 State of Food Security and Nutrition in the World (SOFI) report, an estimated 8.2 per cent of the global population, or approximately 673 million people, experienced hunger in 2024. This means roughly 1 in 12 people worldwide faced hunger in 2024 (FAO, IFAD, UNICEF, WFP, WHO, 2025). That is a staggering number. Besides the area of nutrition, many in the world today do not have access to basic necessities of life in terms of health care, social and economic empowerment, functional social institutions, etc. This is still the case in some very affluent and developed societies and economies where the divide between the rich and poor continues to widen daily. Besides our inability to address some of the most basic human needs like feeding and housing in a world replete with the most advanced ways of addressing issues of this kind, many in diverse regions of the world are experiencing untold hardship stemming from human failures and negligence. The same can be said with regard to diverse forms of extremist tendencies, wars, religious fanaticism, intolerance, and the resurgence of dangerous forms of nationalist tendencies; the same is true with regard to recourse to economic, political, and religious policies geared towards marginalising nations and communities. We have the same given in the unending disregard of humans for their environment and the attendant dangers this poses to human continued existence. That is to say, it is an irony that at no time is the world almost unable, or should we say unwilling, to address its most basic needs adequately as of today; our high level of technological sophistication notwithstanding. This is still the case even though the technical know-how and the means for producing food and other necessities of life have improved significantly. In the face of all these fundamental challenges, one can say that there is basically nothing wrong with investing in how to conserve our mortal remains eternally after death as some transhumanists strive to accomplish: After all, some inquiries that may appear secondary today have turned out, in some cases, to yield surprising unintended future benefits. In most of these cases, we are dealing with issues that demand choosing well. Languishing in despair and succumbing to our finitude is not a choice; doing something to confront it with courage is a choice that has to be handled imaginatively. This is especially the case when we are dealing with niche goals.

In all, it is a matter of getting our priorities right when it matters most. This is why Igbo say: *onye unọ ya na agba oku anaghi achụ oke*, i.e., A man whose house is on fire does not chase rats. Rather than prioritising transhumanist research that pursues extreme goals like

seeking human immortality, or those that serve mere niche human needs, would it not be more beneficial to channel more resources towards enhancing the living conditions of humans, animals, and the ecosystem? Furthermore, is it not better to prioritise the psychological and ethical enhancement of our attitude towards better engagement with our current realities in the world? What this means is that there is a need for a thorough overhaul of the type of mind-set with which we address reality. The aim is to make it possible for the human subject to understand better why human beings are unable to do those things they identify as good and praiseworthy, but insist on doing those things they condemn and despise as evil. Where we invest heavily in addressing technological matters without trying also to understand the reasons, we fail to render the assistance we know is needed. We can hardly do much to improve on our impoverished human conditions: Talk less of rendering the help we know is needed. This is why, to address human needs adequately, there is need to address the constraint of our tension-laden, ambivalent existential situations that are beclouded by the phenomenon of concealment (ihe mkpuchi anya) and complicated by ihe mmeghari anya (third-tier perception or illusion). Here AI has a lot of important roles to play in matters that have to deal with cognitive enhancement towards improving human intelligence and empathy to solve global conflicts. This is why the therapeutic dimension of Ibuanyidanda philosophy or noetic propaedeutic can serve as a foundational tool needed to address most complex post-humanist cravings. This type of noetic propaedeutic aims at imparting the type of transformative complementary type of mind-set basic to most human activities. The tension between world-immanence and absoluteness is an enduring one, and one that will always present itself in a world challenged by our tension-laden, ambivalent existential situations that are beclouded by ihe mkpuchi any (phenomenon of concealment) and now complicated by ihe mmeghari anya (third-tier perception or illusion). How we handle this existential challenge will always go a long way in determining the type of world we seek to build and live in. This tension can only be navigated successfully with a mind-set that is mutually complemented and one that can be acquired in the process of noetic propaedeutic.

Capturing and Harnessing the Penetrative Transformative Character of AI in its Progression

One thing is certain, as a tool for radical transformation, the powers of AI can always be harnessed in addressing most intricate human problems, including those that have to deal with a distorted type of mind-set as caused by ihe mmeghari anya (third-tier perception or illusion). The task is to study the AI's mode of operation and understand the nature of its powers in order to use it when necessary. When carefully observed, AI exhibits a peculiar characteristic that is hardly perceptible even to a careful observer: This is what I refer to as its inherent swift-but-slow-motion, penetrative, transformative mode of progression. This type of progression can be captured by the dictum *festina lente* (make haste slowly). Hence, as ***the powers of AI increase, in the same measure does it accelerate its fast but slow-motion type of progression; and in the same measure does this type of penetrative, transformative progression impact reality.*** This is why even under ideal circumstances, this type of swift-but-slow-motion type of progression is not always fully perceptible even to a careful observer. This mode of progression can be useful but also misused, since it is paradoxical. These are the types of conditions under which some of those apocalyptic projections and scenes associated with AI are thinkable when human encounters with AI remain paradoxical in most critical matters. One can say that it is due to this, its mode of progression, that many do not fully realise that some of those projected apocalyptic scenarios associated with AI's impact are already taking place somehow. This is why it may be too late should we continue waiting for the worst-case scenarios, where AI, hanging over the world like a Colossus, locks the human agent out completely and takes over control. That some of the worst forms of the projections associated with AI are already taking place is easy to expound: To start with, AI impacts our perception of reality swiftly, and yet

slowly and steadily, and in a very penetrating, transformative mode. This is the mode in which AI impacts both positively and negatively. Even if we are not fully aware of this form of progression, cases abound to demonstrate that AI is impacting our world in ways that have almost an apocalyptic outlook: We see this given when AI has been used in some of the worst forms of manipulation and blackmail that have jeopardised, and even ruined, lives and livelihoods. Even in its present form as ANI, it has been deployed to manipulate identities, to invade privacy, misused to deceive and even drive people to extreme forms of mental illness, and in some cases, even to untimely death and suicide. We are often witnesses to situations where, through the use of AI and automation, relationships have been ruined, life savings wiped out, and people rendered hopeless due to the deployment of AI in ways that are nearly incomprehensible. Cyber-attacks and cyberbullying are repeated occurrences possible only by recourse to AI and allied technologies. Hence the question of ASI lording it over its human creators or AI locking humans out of the system is not something futuristic; it is already happening in many cases where AI, even as ANI, is already being used as a Deepfake to manipulate and deceive in highly sophisticated ways.

What this indicates is that AI is bound to pose problems, but not of the type that would have a universal, comprehensive, annihilating impact. As a relative human creation, deeply embedded in fragmentation, and deriving its identity therefrom, it can impact relatively also; and in the worst-case scenario, as a fragmented system collapses at various times and places; not such that it can constitute a universal comprehensive threat to the world. One can confidently say that the dangers AI poses are not necessarily those deriving from machines acting autonomously and quite independent of human control, as many think. It is a danger where humans become careless and overconfident as to mismanage and misuse the positive transformative characteristics of AI for the wrong purposes. They are typical existential conditions that can be occasioned by the *mmeghari anya* (third-tier perception or illusion). In such cases, humans are deceived into misusing and misapplying AI's swift and yet slow and steady penetrating transformative potentialities for inordinate ambitions driven by pursuit of profit and other vain cravings. Worst still is the threat posed by those who can manipulate such potentialities for sadistic, nihilistic purposes. In matters of this kind, AI does not need to act as an independent agent without the instrumentality of human subjects for such extreme scenarios to be possible. This is why we have to think of the possible, but credible, massive deployment of AI for economic, political, military and other forms of sabotage. By any stretch of the imagination, deploying AI in these ways can cause incalculable harm, anxiety, social unrest, and upheaval, with outcomes that are not always easy to predict. Would we not be constrained to think, somehow, of activities of Artificial Super Intelligence running riot in such cases? This is why such projected threats of AI getting out of control are human problems that require a human answer. They are not purely machine-induced problems. It is a problem that has much to do with the mind-set with which we approach AI, the mind-set with which we use AI, and over and above all, the mind-set deployed in machine language, in machine learning and training. Such extreme catastrophic situations are conceivable the moment the inherent complementary relationship among AI, science, and human consciousness is put into jeopardy, suppressed, or even completely abnegated. What this means is that much effort is needed to harness AI's potentialities in view of transforming the human subject positively, such that it learns to accomplish tasks in AI's swift, and yet slow and steady, transformative penetrating mode. This can be accomplished if there is a way to reverse-engineer this positive character of AI and have it reflected in the type of mind-set with which humans and AI interact.

This is important because, in the final effect, AI is a human creation that is designed for human overall flourishing. Where both operate as mutually complementing units, they are bound to benefit each other tremendously from their peculiarities. Ultimately, the goal of establishing AI is good. It is geared towards the advancement of human overall well-being, not

its negation. Therefore, in final effect, what drives AI is something positive, in most cases, and not its own negation. Hence, in all things and at all moments the primary motive driving AI is not to negate, but to uphold and reinforce its usefulness in sustaining human flourishing. In other words, for AI to achieve its benefits for humanity, its human creator has to think of what it takes to enhance some of the positive characteristics of AI. This can happen only where the mind-set needed to embrace AI is harmonising and complementary. Where such a mind-set is bifurcating and exclusivist, this will automatically reflect in the operations of AI. For this reason, recourse to a complementary type of logical reasoning is inevitable in all matters dealing with AI integration, and most especially as this relates to machine languages of all types, to machine learning and to training. It is by recourse to such a complementary logical reasoning that we can guarantee that the harmony between humans and machine is upheld always.

Commoditisation of AI, Harmony of Differences and Acquisition of A complementary Type of Mind-set

We learn from a recent report that: “Oracle cofounder and CTO Larry Ellison has put his finger on what he sees as the fatal flaw in today's AI race: every major model, from ChatGPT to Gemini to Meta's Llama, is trained on essentially the same publicly available internet data. During Oracle's fiscal Q2 2026 earnings call in December, Ellison argued this shared foundation is rapidly turning cutting-edge AI into a commodity product with razor-thin differentiation” (THE TIMES OF INDIA 2026). What is striking about this report is the tendency to misuse AI as a mere commodity and the resultant lack of differentiation such a situation can cause. This is not surprising where the goal driving AI is utility-based and profit-driven. When the goal is utility-based, the danger is often given to misuse the advantages AI offers to maximise profit and even as a veritable means for edging out our competitors. Ellison's apprehension corroborates this fact we see in the AI industry these days, where commoditisation has led to a tense situation where AI optimisation is pursued almost as an end in itself, such that it is geared towards edging out those that are perceived as threatening the interests of other stakeholders. In such a situation, the tendency is to ignore caution, to the point that hardly anyone would bother if such optimisation obsession leads to AI acting in ways that are against its human creators. Insofar as AI is driven by utility and profit, the danger of commoditisation and absolutisation tendencies will continue to be a threat that can easily trigger the ontological boomerang effect of Ibunya philosophy, whose outcome is easy to predict. Where the goal driving AI optimisation is utility and profit-based, stakeholders would not mind deploying the most sophisticated strategies they can afford to put their interests first at the cost of those of other competitors they identify as threatening their interests. They will still do so even if this leads to negating the very high ideal that all cherish in the advancement of AI technology. This is one of the reasons the often talked about use of AI to promote overall human well-being can very easily be forgotten, and it is bound to diminish in proportion to the extent that stakeholders are blindfolded in the pursuance of their personal interests. It is not completely their fault: We are dealing here directly with a typical case of why and how the human subject can mismanage the challenges our tension-laden, ambivalent existential situations present and how this can impact decision-making processes. In most existential situations of life, stakeholders tend to put their interests first against an outsider that is perceived as threatening this interest: Because human existential situations are ambivalent, tension-laden and are beclouded by *ihē mkpuchi anya* (phenomenon of concealment), and complicated by *ihē mmeghari anya* (illusion), actors tend to act in the most unilateral exclusivist ways possible in view of securing their interest first. In doing so, they imagine that they are acting to the best of their interest, whereas this is not always true. For this reason, they would not mind abnegating those high ideals they cherish most in an attempt to pursue and uphold their interests first against an outsider that they imagine threatens such interests. In the

course of doing so, they run into a collision course with others who equally have their interests to protect. Where all stakeholders behave in this exclusivist mode, they are bound to trigger the ontological boomerang effect of Ibuanidanda philosophy, that is the proverbial ill wind that blows nobody any good. The reason is that within the framework of mutually complementary units, what any of the units constituting the whole undertakes to undermine the interest of any of the other units is bound to boomerang, and so much so that all can easily become losers. The situation Larry Ellison is describing points exactly to such situations where stakeholders approach reality with a divisive exclusivist type of mind-set that is guided by a pre-deterministic type of logical reasoning. In such situations, where stakeholders are blindfolded by the unbridled pursuit of their interests, they are often most willing to negate the high ideals they cherish and praise to high heavens, yet keep doing the very things they criticise and condemn. Addressing this matter adequately entails recognising the central role that mechanisms that regulate differences play in all contentious situations of life. This is where recourse to complementary logical reasoning can be very useful in resolving conflicts of interest. It is the type of logical reasoning that makes us appreciate the invaluable role that differences play in human interpersonal relationships: Differences can be assets rather than a liability when we know how to manage and harness them. They must not be a cause of disharmony and conflict. Thus, complementary logical reasoning is most adapted towards addressing the tension that differences can create and towards harnessing the benefits differences can bestow. By resorting to complementary logical reasoning, stakeholders learn to appreciate the need to affirm all existent realities as missing links in a mutually complementary relationship. They are thereby enriched, knowing that differences are opportunities rather than disadvantages.

It is exactly this invaluable insight that is negated by unilateral actions, which are made attractive by the concealment deriving from approaching the world with a divisive exclusivist type of mind-set. Where this type of situation persists, it indicates that stakeholders are not fully aware of the constraints and impositions that can arise from the tension-laden ambivalent character of all existential situations. Where such concealment dominates, the tendency is to pursue a type of absolutisation of individual yearnings and such that can yield “razor-thin differentiation,” as Larry Ellison rightly observes. That is to say, the tendency is to negate differences when our most cherished interests are at stake, and to opt for the maximisation of individual egoistic plans and strategies. However, as a world-immanent missing link, the moment differentiation is removed from matters dealing with AI development, that is the moment it loses its dynamism. As a world-immanent relative missing link, AI’s life is its ability to amplify differentiation as an aspect of authentic relative existence. It is only in this way that it can resist absolutisation and homogenisation. This is why it is very doubtful if any attempts at homogenising AI in all its operations are in the best interest of the industry and its usage generally.

To start with, a world of a universal network of computers working in unison is bound to assume a totalising character, and one that easily evokes the image of an absolute Colossus that has the power to annihilate. How such a world is possible in the face of differences that characterise the use of data, as Larry Ellison pointed out, is difficult to conjecture. Aspiring to such a high standard of unified machines is not even desirable since it does not in any way reflect the character of the human subject that, in the singularity of its constitution, seeks its autonomy. A situation, for example, where all military systems of the world are standardised and connected to the same universal network of computers, though thinkable, is not realistic, since such a scenario would negate the fundamental human instinct to uphold autonomy in differentiation. Until such a level of standardisation and uniform control is achievable, the fact of AI taking over control and lording it over its human creator is bound to remain a mere hypothetical meta-theory of cautionary importance. Besides, such high-level standardisation is

not opportune. What this demonstrates is that the dimension of differentiation, as an inherent aspect of relative world-immanent predetermination, belongs to the character of AI as a tool at the service of humanity. This its character is an inherent dimension of the transcendent categories of unity of consciousness that reflects its human origin. Just as the internal workings of human consciousness, at all moments, strive to reconcile the tension between world-immanent predetermination and the tendency towards absolute determination of all relative missing links, in the same way should this be reflected in the mode of operation of very sophisticated means necessary to address human needs. Hence, it belongs to the character of AI to explore differences better and how these can be integrated within any framework of action and interaction.

Where ChatGPT, Gemini and Meta's Llama, among others, give the same answers to the same prompts, without differentiation and in a pre-deterministic mode, they would hardly evoke those surprises and semblance of creativity that fascinate us. Basically, therefore, complementarity does not admit of negation of differences. It sees differences as an inherent dynamic moment of authentic existence. Where differences are negated, there is always the tendency towards the type of standardisation that ultimately leads to absolutisation. Such is thinkable were all AI systems to be collapsed into one undifferentiated, all-controlling system, and all-legitimising Super system. Naturally, such a scenario would have severe consequences on the way we understand human freedom and autonomy. As relative subjects that always strive to uphold some sort of differences to make life meaningful, it is doubtful whether human fundamental instinct of self-preservation would allow such unification and harmonisation of all modes of AI into an all-controlling and legislating Super AI. One thing should be noted, though: it is the fact that difference is an inherent moment of authentic existence, as the transcendent categories have shown. Therefore, in the face of differences, what is needed is a formidable control mechanism that has the capacity to address or regulate the tension such existential conditions can create. It is such a control mechanism that can help address such differences that appear irreconcilable. This is important because, where tension persists and each party pursues its interests first at the cost of everyone else's, all are bound to lose their interests due to the inherent moment of complementary necessity that unites all actors within any given framework of action and interaction. We can defuse such tension by recourse to the insight provided by the method and principles of mutual complementation: This is what makes it clear to stakeholders that they can realise their interests optimally only if they recognise that others too have interests to protect. Recourse to such a control mechanism becomes effective only when stakeholders have actually learnt to internalise the value of complementary logical reasoning. It is only by appropriating such logical reasoning that actors can be compelled to understand the value of affirming all existent realities as missing links that are in a mutually complementary relationship. By so doing, they break the circle of mutual negation. Where they fail to be committed to such a regulative mechanism, they risk working against their own interests. It is this type of complementary logical reasoning that is a necessary condition for addressing issues arising in human-AI interaction.

The principles of Ibuanyidanda Philosophy: Machine Language, Machine Learning, and Machine Training

If acquiring a complementary type of mind-set is a necessary condition for relating creditably with AI and allied technologies, doing so is not an end in itself. Its benefits can be amplified if such a mode of complementary logical reasoning can be transferred to machine operations, so as to reflect this way of relating to the world in the way AI operates. This can be achieved if there is a way to ensure that this complementary reasoning is reflected in how codes and programs that operate machines are written and structured. One of the major aims is to make AI learn how to affirm missing links as mutually complementary. It thereby also learns that consistent self-interest is tantamount to anti-self-interest and must be avoided at all times. That

is to say, AI learns that consistent acts of self-interest are bound to trigger the ontological boomerang that brings no one any good. Hence, implementing the idea of the ontological boomerang effect and its implications in the way programs and codes are written is a very useful component of the search for multi-layered solutions for AI safety. This strategy is important insofar as it aims at upholding the integrity and safety of AI as a condition to perform its actions in a way that AI does not harm itself nor overreach as to turn itself against its human creators. Where this can be achieved, as in the case of the human subjects who understand the implications of the ontological boomerang effect, there is the likelihood that efforts being invested in making AI more adaptable to human needs, without being a liability, are bound to be enriched all the more. Even if AI is not expected to operate at the same level as self-conscious human persons are capable of due to its mechanistic characteristics, following this route will at least enable it to mirror or approximate the way the human subject processes data. The aim is to mirror the complementary logical reasoning evident in the way human subjects address the world and to give this in the way machines operate. Naturally, this does not imply that, by doing so, we make AI human. On the contrary, the goal is to humanise its operations so as to align them with the goals set by its human trainers, in ways that make those goals most adaptable to human needs such that they are technologically useful, practically meaningful, and theoretically realistic.

In the application of complementary logical reasoning, I seek to strike a balance between innovation and ethical application. This is one of the surest routes needed so that, in the operation of AI, opposites and diversities are understood as missing links and not as exclusivist categories at war with each other.

Issues Arising from Transferring Complementary Logical Reasoning to Machines and Ezumezu Logic

One of the greatest difficulties of having the ideals of a complementary way of reasoning reflected in the way machines operate is that relating to the complex character of human behaviour and the mechanistic nature of AI operation. Bridging this gap can prove difficult but not insurmountable if we remember that this is what humans have been doing all along in human-AI interaction. One of the major differences is that the processing power and speed with which AI accomplishes tasks has increased and continues to increase tremendously. In the case of ANI, we encounter machines grappling with complex mathematical calculations that are based on binary statistical probabilistic logic that human subjects input. Even at this level, the type of nuances obtainable in human reasoning is reflected in the way AI operates. Hence, adopting a complementary reasoning is nothing other than attempting to improve on this type of formal pre-deterministic logical calculus already in vogue. What this means is that AI can hardly attain the type of self-consciousness reserved only for humans just because AI has learnt to see all existent realities as missing links. There is a big difference between *seeing all existent realities as missing links* and *affirming all existent realities insightfully and consistently as missing links of reality* that are in a mutually complementary relationship for the joy of being. The capacity to affirm all existent realities as missing links in a mutually complementary relationship for the joy of being is a typical human act that can hardly be imitated by AI. One can say that attaining the point where AI learns to avoid the ontological boomerang effect of Ibunanyidanda philosophy is very crucial in machine learning. Even if implementing such complementary logical reasoning in the way AI operates may appear difficult, there is every indication that pursuing this goal is realistic if we take the case of ubiquitous Generative AI tools that rely largely on binary probabilistic statistical logic for their training as a point of reference. One of the things that is fascinating about this technology is its ability to show reservations about certain matters in line with its training. These generative AI tools do not always follow all prompts; on the contrary, they show some reservations in matters they do not consider agreeable. This mode of differentiated logical operation needs to be improved upon

by exploring the options offered by a complementary logical reason, most especially as it opens up the possibility of AI avoiding the ontological boomerang effect of Ibuanyidanda philosophy. In all matters, AI must be able to confront reality in its pre-determination and still be in a position to address the relative absolute tension that all existential situations are exposed to. That is to say, adopting such complementary logical reasoning in machine language aims to reflect the way the mind reacts to the interplay between the relative and the absolute divide between the pre-deterministic and probabilistic statistical. In other words, in dealing with AI, the ultimate aim is to approximate the character of the mind that is not only absolutising but relativising, all at the same time, and in tune with the dictates of the transcendent categories of unity of consciousness. Even then, there is no guarantee that, by adopting this route, we will ever be in a position to eliminate all the challenges and risks that AI poses. On the contrary, just like in all human affairs, there is no guarantee ever that the human subject can be the master of his existential situations completely: What more, machines that are mechanistic in their character. This is all the more the case if we remember the future-related character of reality, which is filled with surprises. What is being aimed at is the capacity to manage the challenges our tension-laden existential situations pose, which will make a relatively harmonious coexistence possible. Hence, in dealing with AI, as in all human affairs, the aim is to reduce risk to the barest minimum, and not to eliminate risks completely, bearing in mind the existential constraints to which the human subject is exposed at any given moment.

The major task, then, is to actually translate such complementary concepts into workable machine language and code capable of controlling the way AI functions. Where this is successfully implemented, we can then guarantee that AI fully complies with the dictates of the complementary logic that derives from the way its programmers reason and write code. When AI is trained in this way, its points of view are likely to broaden, so that, in its operation, it can assess differences more accurately. By so doing, it would be in a position to see the connection between its action and the need to avoid triggering the ontological boomerang effect in a way that impinges even on its own operations. Here, I highly commend the efforts of the Nigerian Igbo logician, Prof. Jonathan O. Chimakonam, who, in what he calls Ezumezu Logic (CHIMAKONAM 2019), has endeavoured to formalise a variant of complementary logic in tune with some of the basic presuppositions of Ibuanyidanda philosophy. His is a major milestone that serves as a powerful stepping stone for future development in this area. However, addressing issues of this kind remains a collaborative complementary task where philosophers, logicians, mathematicians, software and hardware engineers, and all stakeholders mobilise their resources to seek better solutions in the way we relate to AI. Where this complementary effort can be guaranteed, we can then confidently say that realising some of the goals of having a complementary logic reflected in machine language, machine learning, and machine training is something worth pursuing as a teamwork of determined theoreticians and engineers. What this shows is that pursuing this matter is not a purely technical issue; it is something that has a deeply philosophical dimension that has to be borne in mind always. This is important due to the special character of Ibuanyidanda logic that has an inherent future-related dimension.

Special Character of Ibuanyidanda Logic in its Future-relatedness (Nke Iru Ka)

One thing about complementary reasoning grounded in complementary logic, as pursued in Ibuanyidanda philosophy, is its ontological grounding that is future-related. This future-relatedness is exactly the character of reality we refer to as **nke iru ka** (the greatest events are in the future, or the future determines ultimately) in the Igbo language. This naturally has some implication for machine language, machine learning and training in the application of logic. By reason of its future-relatedness, the goal pursued expressly in ontology by far exceeds what is capable of being apprehended fully and represented, at any given moment, in codes, symbols, and mathematical formulas. At the conflux of logic and ontology is the very character of reality

that is incomplete and imperfect; yet striving to completion in a future-related mode. This is why, even if the answers we give at any given moment appear sufficient or even insufficient, we can still feel deep down in us the urge to strive even higher in view of giving the best touches to the questions we pose and the answers we seek. It is this urge that lies at the root of innovation and attempts to break new ground, as we strive beyond boundaries at any given time and seek new answers through the application of reason in scientific research and technological innovation. The situation in our relationship with AI, and most especially in view of replicating complementary logical reasoning in machine language, machine learning and training, cannot be different. Hence, attempts to translate and denote a transcendent, complementary, comprehensive, future-related idea into codes, mathematical, and symbolic logical forms are bound to encounter some challenges. This has much to do with the character of reality itself, which, in its consistent unfolding, can stretch the application of language to its limits. Even if what is intended cannot be fully expressed in language form at any given moment completely, striving towards the future in view of exploring the best possible opportunities remains something desirable. We know this due to the instigations of the transcendent category of future-relatedness, deeply embedded in human consciousness, as this drives all language forms and proclaims: *ihe ukwu ikpe azu*, i.e., the future harbours far more than we can imagine; or again *nke iru ka*, i.e., greatest events are in the future. Hence, by following this future-relatedness, it is immediately revealed to the inquirer that any stage achievable at any given moment, no matter how insignificant, how novel and revolutionary, can still be exceeded and improved upon always. Where we fail to recognise this fact, we would be constrained to celebrate world-immanent achievements as the summit of human creativity. This is one of the most radical challenges we have to face in the use of AI, where we seek immediate, complete, and quick answers, such that any results that do not reflect the way AI addresses issues are bound to be suspect. Worst still is when we are deceived into assuming that AI is the future itself. Where we suffer such illusions and cannot contend with the limited character of symbols, algorithms, graphs, codes, and mathematical formulas based on which machine language, machine learning and training rely to express meaning fully at any given time, we are bound to forget the ambivalent challenges of all existential realities, including AI and related technologies. It is one of the cardinal duties of *ibuanyidanda* logic to expose this concealment and point at the dangers it presents. In its future-relatedness, logic based on ontology can only represent reality partially and in an incomplete form; so also the symbols, algorithms, graphs, codes and mathematical formulas deployed to program and deploy AI.

In other words, recourse to complementary logic to address issues will not guarantee an error-free use of AI. It can only improve on the binary statistical and probabilistic logical method in vogue. However, it belongs to the character of this *ibuanyidanda* logical reasoning that is future-related to demonstrate that matters of machine language and AI are not purely issues of computation and symbolic logical representations. They are issues that have a very fundamental philosophical dimension that must always be borne in mind. One can say that in its future-related mode, complementary *ibuanyidanda* logic propels all modes of relations in a universal, complementary, and comprehensive way beyond all thinkable forms of valuation. This is the *ibuanyidanda* logic that derives its dynamism from the principles of *ibuanyidanda* philosophy, which legitimise both towards the inside and to the outside as they point to the limited nature of all modes of Logic. This is how these principles seek to sustain all forms of complementary logic; be these two valued, three valued, four valued, etc. In other words, the principles of *ibuanyidanda* show that all forms of logic are world-immanently structured; as such, they need to be superseded and legitimised always. This is why they are grounded in a New Complementary Ontology, one that is aimed at ushering in a New Complementary Turn altogether.

Confronting the Uneven Reciprocal Influence between AI and its Human Creators

The speed at which AI processes data with near accuracy is so amazing that, by every indication, AI now has the capacity to outpace humans in most computational tasks involving the deployment of syllogistic calculus to attain optimal results. This is at least what becomes obvious from the way ubiquitous generative AI tools like ChatGPT, Gemini, Grok, Copilot, Claude, etc., generate human-like text based on their training data. What comes after these ANI based tools is anyone's guess. Hence, in human-AI interaction, it is now almost a race of who overtakes whom, and such that impacts heavily on the way humans perceive and relate to the world generally. First and foremost, it has to be remarked that there is no indication whatsoever that AI can actually perform tasks like humans in the strict sense of the word, since it takes more than processing data with speed and accuracy to become human. Yet, the mutual interaction between AI and its human creators hardly proceeds without some unintended consequences that ought to be remembered always. This subsists in the undeniable mutual reciprocal influences between machines and their human creators: Just as humans strive to optimise machines to reap the benefits they have to offer, machines in turn impact human behaviour in quite radical, unexpected ways. In this mode of reciprocity, there is an asymmetry in the reciprocal relationship between AI and its human creators where AI is gradually transforming humans into machine-like agents. This is, at least, what becomes obvious, going by human projections about what the future portends for AGI and ASI. What this means concretely is that in the process of transferring their capabilities to machines, humans inadvertently inherit some of the characteristics peculiar to machines. The danger is that when humans are so accustomed to AI, they may be constrained to attribute capabilities to AI which it does not possess, believing that this is normal. On account of this type of uneven reciprocal relationship, and the astonishing speed with which AI is crossing human paths, the projections about how long it will take for AI to surpass human intelligence continue to narrow daily. These are the types of transformations that inspire certain post-humanist ambitions that seek to capitalise on these advancements to enhance human capabilities in extraordinary ways with the aid of AI. It remains to be seen how best such post-humanist ideas can better serve human interests. Therefore, addressing the uneven reciprocal relationship between humans and AI, one that increasingly gives AI an advantage over its creators and sometimes leads humans to behave like machines, remains a major challenge in our era. Equally pressing is the question of how prepared we are to resist viewing AI as something inherently human simply because it performs tasks that resemble human activity. More broadly still, the difficulties arising from this asymmetrical reciprocity between humans and AI will continue to shape our concerns in the age of intelligent machines: As we become fascinated by AI's operations, we risk gradually imitating and internalising, often unconsciously, its rigid, data-driven, and hyper-efficient tendencies. Once conditioned in this way, humans may become so immersed in the world of AI that they lose sight of the dangers it poses. Worse still, this immersion can erode or obscure some of the very qualities that define us as human beings. For this reason, our interactions with AI must be balanced in ways that safeguard what is essential to human identity. The deliberate cultivation of qualities such as compassion, empathy, wisdom, patience, tolerance, moral discernment, and other capacities needed to navigate the complexities of life becomes crucial. These human attributes are decisive in managing the uneven reciprocity between humans and AI and in ensuring that technological advancement does not encroach upon the foundations of our humanity.

Hence, the question will always remain how we can adequately address the tension in the internal workings of human consciousness that are responsible for narrowing the gap between what it takes to be human and what AI denotes as an agent. Where we are able to address such tensions credibly, we would all the more be in a position to process data we feed into machines in a way that does not polarise our relationship with the world and reality

generally. If, on the other hand, we are absorbed into the life of AI without due regard to the exclusivist type of logical reasoning that drives it, we would be constrained to view the world in a sterile, polarising, and exclusivist mode at all times. A situation of this kind makes it imperative to always review the type of mind-set deployed in dealing with matters relating to AI so as to ensure that such a mind-set has all it takes to understand the complementary character of reality and one that understands the special character of all existent realities as missing links that are in mutual complementary relationship to each other.

At times, we fail to fully realise the implications of using an exclusivist, polarising mind-set when addressing practical issues in life. This attitude to life can be very harmful in the way we view reality generally and most especially in the way we relate to other human individuals concretely. Yet, we are not always to blame when we fail in this respect, because attitudes of this kind can at times be attributed to long years of indoctrination, education, and socialisation, in which stakeholders are accustomed to defining existence mostly in terms of competition. Worst still is when such a spirit of competition devolves into using technology to seek solutions in the most exclusive ways. When this happens, the human agent always seeks to address the world with a divisive type of mind-set in a way that suggests that existence is nothing other than competition, where stakeholders have to devise strategies to win at all costs; even at the cost of abnegating some of those values they hold in very high esteem. We see this type of situation in our daily struggle, most especially when resources are scarce: This is when human fundamental instinct of self-preservation can fire our passion to always get what we want at the cost of those we identify as outsiders, as weak, and dispensable; those who threaten our interests. Not even what should be a sophisticated AI environment is spared this type of egoistic challenge. In situations of this kind, stakeholders pursue their most cherished interest blindly, even at the risk of abnegating some of those ideals they cherish and praise to high heaven. This is why an internal control mechanism is a nonnegotiable precondition for AI use. It is an internal control mechanism that can ensure AI-human interaction remains profitable for all stakeholders. Such an internal mechanism of control subsists in acquiring a complementary type of mind-set. However, acquiring such a mind-set does not come by chance; it is something that has to be internalised in the rigorous process of noetic propaedeutic.

Invaluable Nature of Complementary Logical Reasoning and AI Safety

One of the fundamental mistakes humans make in human-AI interaction is the tendency to forget that AI is an integral part of, or an extension of, human conscious acts. Without humans, there is no AI. Hence, to act as if AI and humans are discrete quantities without any form of relationship whatsoever will continue to be an error of judgement that can becloud clear reasoning always. This danger arises when we tend to see AI as a completely discrete, autonomous agent, forgetting the complementary relationship between AI and its human creators. The *mmeghari anya* (third-tier perception or illusion) can deepen this tendency to isolate AI's existence from our lives, so much that we imagine AI as a sort of humanoid existing completely outside of the region of human conscious act. When we approach the world in this way, inculcating the necessity to affirm all existent realities as missing links can become a difficult task. The major task is how to inculcate the intrinsic complementary relationship between AI and its human creators such that this mode of relationship can be profitable across the board.

As humans, we know how difficult it can be to internalise the need to affirm all existent realities as missing links that are in a relationship of mutual complementary service for the joy of being; what more, having this reflected in the ways codes, graphs, logarithms, mathematical formulas, logical symbols, hardware, and all requisite software are structured in view of guiding the activities of machines. It then means that more effort is needed to uphold complementary logical reasoning across the board and as a viable means of cementing the relationship between humans and machines. Here, a lot of patience, perseverance, and

optimism are demanded because any gains recorded in this direction, no matter how minimal, are bound to resonate tremendously in the way we deploy AI for the good of humanity. It is in this way that we can demonstrate how invaluable complementary reasoning and acquiring a complementary type of mind-set can be in confronting life's challenges generally, including the way we relate to machines. In matters of this kind, we are dealing directly with providing credible guardrails for our own conducts as humans and how those regulative principles we cherish can be extended to tools we rely on to make life more meaningful. AI does not bestow a complementary type of mind-set, it is humans who have to work for and towards this. Because we are accustomed to seeing AI as a discrete entity besides its human creator, we often imagine that the transformations resulting from AI usage are something deriving from the outside. This is the type of deception ihe mmeghari anya (third-tier perception or illusion) is capable of creating. It is a type of attitude that needs to be corrected in human-AI interaction. When we say that AI is conditioning or transforming human behaviour, we have to bear in mind always that there is no AI out there performing these tasks as such. We are dealing directly with human engineers and companies defining set objectives they feed into machines to calculate and execute. From here derive those powerful impacts that condition our behaviour. In the final analysis, all forms of AI transformation stand to mirror the intentions of its human agent. What this demonstrates is that all talk about AI safety has an internal human dimension that has to be borne in mind always. We are dealing here with the type of mind-set or logical reasoning with which machines are programmed. Where the mind-set is divisive and the logic binary, we are bound to contend with an AI that has the capacity to present some of those threats we all dread. If, on the other hand, the mind-set with which AI is programmed and the codes we feed into machines are complementary, then AI is bound to be far safer in its deployment. Therefore, all matters of AI safety are not machine-induced; they also have a noetic dimension that must be borne in mind always. This is why noetic propaedeutic, as the therapeutic dimension of Ibuanidanda philosophy, is an integral part of all efforts to enhance AI safety.

“Reverse-Engineering” AI Progression in the Process of an AI-Induced Type of Penetrative Noetic Propaedeutic

Noetic propaedeutic, as the name indicates, is the act of pre-training the mind in complementary reasoning. It can be acquired both passively and actively. It can be acquired actively in the rigorous process of self-imposed training of the mind in the act of reasoning in a complementary way. It can be imparted passively in the process of educating others in the act of complementary reasoning. This happens mostly during early childhood education, where the child is trained to affirm missing links as mutually complementary. By appropriating some of the positive characteristics of AI, we can enhance and make the act of noetic propaedeutic more efficient. This is where the inherent swift-but-slow-motion penetrative, transformative mode of progression that defines AI success in its ontological constitution comes into play. This positive characteristic of AI needs to be appropriated in the process of noetic propaedeutic of all forms, and for the enhancement and improvement of overall human wellbeing, most especially human cognitive and emotional needs. One can then say that the apex of noetic propaedeutic, whether as a self-imposed act or as an act of teaching others, in the AI era lies in humans seeking to replicate, through the practice of complementarity, the type of penetrating progression that marks AI success. By imitating AI in the form of reverse engineering, humans can enhance their communicative skills and conditions, improve their emotional and cognitive potentialities as they seek to communicate in a complementary way. It is an attempt to reset and enhance human cognitive and emotional intuitions by understanding the secrets behind AI success and having this reflected in the way humans communicate generally. Again, it is an attempt at absorbing consciously, but in a positive way, AI's penetrative, swift-but-slow-motion type of progression in the ways humans think and execute tasks. In such situations, the

human subject strives to make complementary thinking and complementary living second nature in an impactful, penetrating and sustained way. In other words, it is an attempt to imitate the penetrating thoroughness with which AI accomplishes tasks, but in a human way that is insightful, consistent, and compassionate. It is in this process that the human subject acquires a complementary comprehensive type of mind-set that, in its thoroughness, is penetrating, swift, yet slow and steady. This is why I designate such a condition as a type of AI-induced penetrative, transformative type of noetic propaedeutic. When this is applied in concrete situations of human interaction, it will translate to a state of quick, thoroughgoing, penetrative, and an all-embracing type of complementary intuition or in-feeling capability; something the human subject carries into the world in encounter with all existent realities; and most especially in the practice of virtues. If this can be successfully implemented, then the human subject is bound to think always in a complementary mode that is in tune with the demands of our ever-changing existential conditions, which are determined by the tempo of AI development.

In all, Ibunyaandanda philosophy seeks to raise awareness of the need for this type of transformation in our digital age powered by AI. It is an awareness towards ushering in a radical change of orientation, far-reaching change of mind-set and lifestyle that internalises the lesson firmly that every mode of existence, human or nonhuman, plays a vital role in the totality of existence. Where this type of transformation is achievable, we can start thinking of a New Complementary Turn where all modes of relations are perceived in a complementary mode. That is to say, all modes of determination are perceived as forming a complementary whole, i.e., every component, ideas and ideas of ideas, things and things of things, units and units of units, categories and categories of categories, entities and entities of entities; spiritual, material, visible and invisible, AI, complementary logic, binary logic, etc. More concretely still: Every component, from humans (individuals and societies) to the fauna (animals) and flora (plants), as well as fungi, microbes, the soil, water, and the sunlight and geothermal energy that fuel the intricate ecosystems they form, etc., plays a vital role in the totality of existence. In other words, our world is a single, magnificent web: Every creature and element matters; they are missing links in mutual complementary service for the joy of being. Where this type of understanding is reflected in the way we approach reality generally, there is bound to be a paradigm shift from competition to cooperation, from profit thinking to human flourishing, from “we against them mentality” to a mentality of mutual support and cooperation. Having this type of awareness internalised and reflected in human-AI interaction, and most especially in the way machine language, machine learning, and training are structured, is bound to help in improving greatly the way humans conduct themselves in our digital technological age sustained by Artificial Intelligence. This lesson is at no time more relevant and urgent than in the digital age where AI plays a central role in almost all areas of our lives. It is still more urgent when we think of the frightening scenario hinted at by Amodei, Anthropic CEO of AI, “becoming too autonomous and locking humans out of systems” (ROGELBERG 2026). Basically, AI is a technology whose goal is good as it is geared towards enhancing human flourishing. However, this goal can be attained only if the mind-set underlying human-AI development and use is also adequate enough to address human needs and not work against them. What we often forget is that there is no AI without a human agent. Behind every AI is the human agent who dictates the code and algorithms that keep it running. Where machines are trained with an exclusivist type of logic deriving from an exclusivist and polarising type of mind-set, they are bound to be polarising and exclusivist in their operations, so much so that they can, at our unguarded moments, even become too autonomous as to even lock humans out of the system.

Conclusion

Grounded in the basic presuppositions of Ibunyaandanda philosophy regarding reason and consciousness, viable solutions to the question of AI safety revolve around the mindset with which we address reality generally. Once the important philosophical dimension of AI safety

is ignored, AI safety is reduced to a mere technical alignment problem, which mistakenly frames containment as the only effective safety strategy. Such reductionism easily fuels the illusion of AI as an overwhelming Colossus. This manifests, for example, in nightmarish scenarios like Nick Bostrom's "Paperclip Maximiser". Anxieties that project AI as an absolute, uncontrollable power. By its very nature, however, AI remains a world-immanent, relative tool that can never ontologically achieve absolute autonomy or function as a surrogate deity. Consequently, the fear of AI 'locking humans out' distracts from the genuine danger: human mismanagement. Current AI development suffers from an overreliance on pre-deterministic, binary Aristotelian logic, rendering the technology relatively divisive and exclusivist: In this way, very vulnerable to biases. To steer AI away from these adversarial, exclusivist characteristics, there is a need for developers to integrate a complementary logical framework into machine learning, programming languages, and algorithmic design. This integration would allow AI to mirror the unified, harmonised character of human reason and consciousness and safely navigate the harmony of differences. The specific mindset embedded within codes and algorithms inevitably dictates how AI operates. The most radical shift occurs when AI is programmed through a complementary logic to recognise all entities as mutually dependent missing links within a unified whole. Such a framework inherently guards against the ontological boomerang effect of *ibuanyidanda* philosophy wherein any entity acting out of pure, exclusive self-interest triggers a systemic rebound that ultimately negates its own goals. Ultimately, addressing this critical noetic dimension is one of the foundational challenges of the machine age, as it constitutes the centre of gravity of the concern of *ibuanyidanda* (complementary reflection) philosophy in AI safety conversations. Because of the intrinsic, complementary relationship between humans and their technological creations, humanity must reverse-engineer the ontological character of AI, embracing its relativity to maximise its benefits. By taking the bull by the horns, human creators can mirror AI's structural capacities to enhance their own cognitive, emotional, and virtuous skills. True safety is achieved when humans internalise the deep, systemic, and progressive precision characteristic of AI's ontological status, transforming a technical threat into a catalyst for human flourishing.

Declarations

*The author declares no conflict of interest or ethical issues for this work.

Relevant literature

1. ASOUZU, Innocent. [Gedanken über die religiöse Problematik der Gegenwart im Licht der Theologie der Religionen: Eine Reflexion im Anschluss an Wolfhart Pannenberg], 1986. Peter Lang Verlag: Frankfurt am Main.
2. _____. [The Method and Principles of Complementary Reflections in and Beyond African Philosophy], 2004. University of Calabar Press: Calabar.
3. _____. [The Method and Principles of Complementary Reflection In and Beyond African Philosophy], 2005. LIT Verlag: Münster.
4. _____. [Ibuanyidanda: New Complementary Ontology. Beyond World-Immanentism, Ethnocentric Reduction and Impositions], 2007b. LIT Verlag: Münster
5. _____. [Ibuaru: The Heavy Burden of Philosophy: Beyond African Philosophy], 2007a. Lit Verlag: Münster.
6. _____. "Ibuanyidanda and the Philosophy of Essence". [50th Inaugural Lecture of University of Calabar], 2011. University of Calabar Press: Calabar.

7. _____. [Ibuanyidanda (Complementary Reflection) and Some Basic Philosophical Problems in Africa Today: Sense Experience,” Ihe Mkpuchi Anya” and the Super-maxim], 2013. Lit Verlag: Münster.
8. CHIMAKONAM, Jonathan O. [Ezumezu: A System of Logic for African Philosophy and Studies], 2019. Springer Publishers: Cham.
9. FAO, IFAD, UNICEF, WFP and WHO. [The State of Food Security and Nutrition in the World 2025 – Addressing High Food Price Inflation for Food Security and Nutrition], 2025. FAO, IFAD, UNICEF, WFP, WHO: Rome.
10. NICK, Bostrom “The Ethical Issues in Advanced Artificial Intelligence”. [Machine Ethics and Robot Ethics, Wendell WALLACH & Peter ASARO Ed.], pp69-75, 2020. Routledge: Abingdon.
11. ROGELBERG, Sasha. “‘I’m Deeply Uncomfortable’: Anthropic CEO Warns that a Cadre of AI Leaders, including Himself, should not be in Charge of the Technology’s Future.” [[Fortune](#)], February 19, 2026.
12. RUSSELL SOCIETY NEWS. “Bertrand Russell's Ten Commandments.” [Bertrand Russell Society, Inc.], May 1987. No54.
13. ST. AUGUSTINE. [The Soliloquies of St. Augustine, Rose Elizabeth CLEVELAND, trans.], 1910. Little, Brown, and Company: Boston.
14. THE TIMES NEW OF INDIA. “Oracle Cofounder Larry Ellison on the Biggest Problem that All AI Models, including ChatGPT, Gemini, Grok, Llama Have.” [[The Times New of India](#)], January 27, 2026.
15. THIAGARAJAN, Ramu & THIND, Aman. “Navigating DeepSeek’s disruption: Opportunities and challenges in AI advancement.” [[State Street](#)], March 2025.